

An Analysis of Non-Iranian Persian Language Learners' Active Vocabulary in Writing: Lexical Diversity Based on Guiraud's Index

Vol. 17, No. 2, Tome 92
pp. 33-75
Summer 2026

Seyed Akbar Jalili¹, Rezamorad Sahraei^{2*}  & Amir Zand-Moghadam³ 

Abstract

The purpose of this research was to investigate the active vocabulary in written texts of non-Iranian Persian Language learners based on Lexical Diversity (LD). Using one of the formulaic measures of LD called Guiraud's Index (GI), LD of the written texts of the subjects of the present study was calculated and its relationship with variables such as nationality, gender, age, first language and university education was determined. First, according to the principles and rules of corpus transcribing, and using the LancsBox software, Types and Tokens of the texts of 251 learners from four nationalities, who participated in the final exams of the Persian Language Center of IKIU, were extracted and counted. After that, using GI, LD of each subject's writing was calculated and the research hypotheses were evaluated. Results showed a significant difference between different nationalities in terms of LD. Also, from the perspective of first language and gender, writings of Arab subjects indicated significantly more LD than Chinese, and the texts produced by women indicated more LD than the texts of men. On the other hand, the two hypotheses related to the LD and age and university education were not confirmed, because these relationships were non-significant. The results and findings of this research can help teachers and examiners in the field of teaching Persian as a second/foreign language obtain a suitable tool for evaluating the lexical richness of written texts and gain insights on how to use lexical richness criteria in the evaluation of learners' texts.

Keywords: Active vocabulary; lexical richness; lexical diversity; Guiraud's index; writing.

Received: 5 November 2022
Received in revised form: 15 January 2023
Accepted: 6 February 2023

¹ PhD Candidate of Teaching Persian to Non-Persian Speakers, Allameh Tabataba'i University, Tehran, Iran.

² Corresponding author, Professor, General Linguistics, Department of Linguistics, Literature and Foreign Languages Faculty, Allameh Tabataba'i University, Tehran, Iran; rofessor in Linguistics, Allameh Tabataba'i University, Tehran, Iran;

Email: sahraei@atu.ac.ir, ORCID: <https://orcid.org/0000-0001-9953-0789>

³ Associate Professor, English Language Education, Department of English Language & Literature, Allameh Tabataba'i University, Tehran, Iran;
ORCID: <https://orcid.org/0000-0001-8555-8481>

1. Introduction

By analyzing the vocabulary in the learners' linguistic productions, a more accurate and desirable assessment of the vocabulary index can be achieved. One way for assessing the lexical knowledge is to measure and evaluate the "lexical richness" of the subjects' writing and speech. The concept of lexical richness is related to the use of words when producing speech or writing. The aim of this study was to analyze the lexical breadth in the writing of non-Iranian Persian language learners, or in other words, in a part of their Active vocabulary. This analysis was based on one of the components of lexical richness, namely Lexical Diversity (LD) (using diverse words and avoiding repetition of words).

Research Question(s)

- 1: Is there a significant difference between Persian learners by nationality in terms of the amount of LD?
- 2: Is there a significant difference between Arabic-speaking and Chinese-speaking Persian learners in terms of the amount of LD?
- 3: Is there a significant difference between male and female Persian learners in terms of LD?
- 4: Is there a significant relationship between LD and the age of Persian learners?
- 5: Is there a significant relationship between LD and the university education of Persian learners?

2. Literature Review

Researchers have used various terms to describe learners' vocabulary. For example, Fan (2000, p. 105) states that vocabulary is generally divided into two categories: receptive/passive and productive/active. According to Schmidt (2000), the terms Active and Passive are alternative terms for Productive and

Receptive. Productive lexical knowledge can be divided into two categories: Free and Controlled. In this way, lexical knowledge generally has three parts: "receptive/passive knowledge", "controlled" and "free"; the first type of productive lexical knowledge, i.e. controlled knowledge, involves the production of words when faced with a task; for example, completing the word "fragrant" in the sentence "the garden was full of fra____ flowers." Free productive knowledge also relates to the use of words at will and discretion, without the need for specific stimuli to produce words (such as free writing) (Laufer, 1998). According to Laufer, it is necessary to distinguish between free and controlled active vocabulary as not all learners who use infrequent vocabulary when forced to do so will also use it.

Regarding second/foreign language learners, it is generally believed that the number of words that learners know in the target language is greater than the number they can actually use (Fan, 2000). Since the "size of the vocabulary" of language learners in speech and writing is one of the most important factors determining their level of proficiency in a second/foreign language, "lexical richness measurement criteria" can help teachers and researchers in evaluating language learners' written and oral productions in this regard. According to Read (2000, p. 200), after completing the writing assignments, we should begin our statistical analyses to measure the efficiency of the learners' vocabulary, which we do with the criteria and components proposed under the concept of lexical richness. In fact, the lexical richness of a written or oral text is the result of the underlying lexical knowledge of the writer or speaker of that text, but through measures of lexical richness, we do not measure the depth of individuals' lexical knowledge, but rather the breadth of their knowledge.

The lexical richness of learners' texts can be described using several components and measures, including Lexical Diversity and Lexical Density. Read (2000, p. 2000) considers Lexical Richness to be an umbrella term that includes components such as "lexical density", "lexical complexity", "lexical diversity" and "low number of errors".

One of the most important components that has been commonly stated as a component of lexical richness is “Lexical Diversity.” This importance is so great that the term “lexical diversity” is sometimes considered equivalent to “lexical richness.” Lexical Diversity (LD) is an indicator of the number of different words in a text. In other words, LD indicates the use of a range of different words and the avoidance of lexical repetitions in learners' texts.

The most basic formula and measure for LD is the type-token ratio (TTR), which is the ratio between the number of different words in a text and its total number of words (number of words multiplied by 100 divided by the number of tokens) (Laufer & Nation, 1995).

Regarding the aforementioned formula, there is a problem that TTR is sensitive to the length of the text; that is, text that contains a large number of items yields lower TTR values, and vice versa; According to Daller et al. (2007, p. 14), TTR has been a subject of debate for almost a century, but it is still used today; for example, some researchers have used this measure with controlled text length. Thus, TTR, while still useful, is length-sensitive, with the longer the text, and the smaller texts will yield smaller numbers. Some studies have even gone further and reported that because TTR falls as the token size increases, it is not an indicator of lexical richness and is an ineffective and unreliable tool for measuring lexical growth (McKee, et al, 2000; Vermeer, 2000; Johansson, 2008; McCarthy & Jarvis, 2007 & 2010). For this reason, in cases where our texts are of different lengths, TTR can become an unstable tool.

Researchers have proposed alternative methods to reduce this effect of text length. Some alternatives calculate the LD of each text based on word frequency lists derived from the speech or writing of native speakers. However, some alternatives have made formulaic and mathematical changes to the simple TTR, one of which is the "Guiraud's Index". This index is also known as the "Root TTR ". The Guiraud's Index uses the following formula to calculate the LD of a text. As can be seen from the formula, the Guiraud's

Index (unlike the TTR) does not rely on a simple count of tokens.

In the present study, since it was based on the use of formulaic methods and did not rely on word frequency lists, the Guiraud's Index was used to examine the LD of written texts.

Regarding the conducted researchs, it should be said that no research was observed in the field of analyzing the active vocabulary of non-Iranian Persian language learners and measuring the lexical richness of the writing of this group of Persian learners, and the few existing studies have the least common aspect (i.e., the Persian language) with the subject of the present study and have been conducted in the field of aphasia and stuttering. In summary, researches on the study of LD in writing has been conducted in two main areas: (1) measuring the accuracy of formulas/metrics and computer tools in calculating LD (such as Jarvis, 2002; Espinosa, 2005; Abbassian and Shiri, 2011); (2) Investigating the relationship between LD and variables such as first language, overall writing quality, level of proficiency in second language, gender, nationality, etc. (e.g. Mozaffari, 2014; Gregori-Signes and Clavel-Arroitia, 2015; Tabandeh Fard, 2016; Yu, 2009). Another aspect of interest is the process of growth and development of learners' Interlanguage in the field of lexical knowledge, in order to understand the path and stages of their lexical development and to test various measures (e.g. Johansson, 2008).

3. Methodology

The corpus of the present study consisted of a collection of written productions of non-Iranian Persian language learners of the Persian language Center (PLC) of Imam Khomeini International University in the advanced/supplementary course final exams, which included the written language productions of 251 Chinese, Lebanese, Syrian, and Iraqi learners in several rounds of final exams. After transcribing and extracting of tokens this corpus contained 36,711 tokens and 19,889 types (on average, 146.259 tokens

and 79.23 types per subject). In order to utilize computer and software facilities and to increase the accuracy and speed of extracting and counting, compared to manual counting, a software called LancsBox (#LancsBox 6.0) was used. After transcribing and normalizing the texts, it was time to prepare the texts for counting the number of tokens and types in order to calculate the LD of each text. At this stage, each text was entered in the form of a separate docx file (i.e., one file for each text), then each file was supplied separately as input to LancsBox to obtain the number of tokens and types in each text.

4. Results

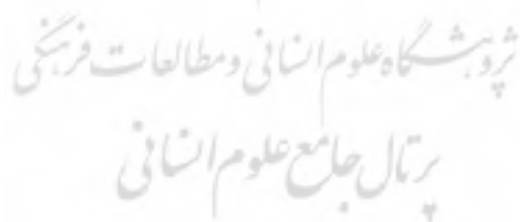
The research findings for the first and second hypotheses showed that nationality and first language can have a significant effect on LD. As a whole, in terms of first language, it was found that Arabic speakers have more LD in writing compared to Chinese speakers. These results can be considered in line with the results of Dewale and Pavlenko (2003) (in the context of bilingualism) who showed that the first language of individuals can affect and make a difference in the lexical richness of their texts, but do not support the results of Yu (2009) who made a comparison between Filipinos and Chinese (no significant difference between the lexical richness of the texts of Filipino and Chinese subjects).

Another finding was that gender can have a significant and significant effect on the level of LD of the subjects. The results obtained in terms of the effect of gender on LD are in line with the results of Singh (2001) and Dewale and Pavlenko (2003), which showed the effect of gender (although weak) on LD. Also, this finding confirms the study of Härnquist et al. (2003), but does not support the study of Yu (2009), which did not find significant relationships between these two variables.

The findings for fourth hypothesis showed that not only is there no positive relationship between age and LD and that increasing age does not increase the LD of texts, but the relationship between these two variables is negative and

insignificant. In fact, when writing in Persian as a second language, increasing the age can be a factor inhibiting or reducing LD. This finding may be in line, albeit in the opposite direction, with the results of Johansson (2008), who states that LD, compared to lexical density, can be a good measure for distinguishing differences between age groups, but it is not consistent with the research of Jarvis (2002), who concluded that with increasing the age, the level of LD of the texts certainly increases.

Regarding the relationship between education and LD, it was also concluded that there is a negative and, of course, insignificant relationship. In fact, with increasing education level, LD does not increase and in some cases it even decreases, which of course is not in line with the results of studies such as Härnquist et al. (2003), who admit that educational background can be a predictor of lexical richness.





پروہشگاہ علوم انسانی و مطالعات فرهنگی
پرتال جامع علوم انسانی



دوماهنامه بین‌المللی

د ۱۷، ش ۲ (پیاپی ۹۲)، تابستان ۱۴۰۵، صص ۳۳-۷۵

مقاله پژوهشی

https://lrr.modares.ac.ir/article_7116.html

تحلیلی بر دایره واژگانی فعال در نوشتار فارسی آموزان غیرایرانی:

تنوع واژگانی بر پایه شاخص گیراد

سید اکبر جلیلی^۱، رضامراد صحرایی^{۲*}، امیر زندمقدم^۳

۱. دانشجوی دکتری آموزش زبان فارسی به غیرفارسی‌زبانان، دانشگاه علامه طباطبائی، تهران، ایران.

۲. استاد گروه زبان‌شناسی دانشگاه علامه طباطبائی، تهران، ایران.

۳. دانشیار گروه زبان و ادبیات انگلیسی و گروه زبان‌شناسی دانشگاه علامه طباطبائی، تهران، ایران.

تاریخ پذیرش: ۱۴۰۱/۱۱/۱۷

تاریخ دریافت: ۱۴۰۱/۰۸/۱۴

چکیده

هدف پژوهش، بررسی دایره واژگانی فعال در متون کتبی فارسی آموزان غیرایرانی بر پایه مؤلفه «تنوع واژگانی» بود. بدین منظور، با استفاده از یکی از سنجه‌های فرمولی تنوع واژگانی با عنوان «شاخص گیراد»، تنوع واژگانی متون کتبی آزمودنی‌های پژوهش حاضر محاسبه گردید و ارتباط آن با متغیرهایی همچون ملیت، جنسیت، سن، زبان اول و تحصیلات دانشگاهی مشخص شد. ابتدا، براساس اصول و قواعد پیاده‌سازی و نرمال‌سازی پیکره و با استفاده از نرم‌افزار لنکس‌باکس، موردواژه‌ها و یکتاواژه‌های متون ۲۵۱ فارسی‌آموز لبنانی، چینی، سوریه‌ای و عراقی استخراج و شمارش شد. سپس، تنوع واژگانی هر متن، براساس شاخص گیراد محاسبه و فرضیه‌های پژوهش ارزیابی شد. طبق یافته‌ها، از نظر تنوع واژگانی تولیدات کتبی، بین ملیت‌های مختلف تفاوت معنی‌دار وجود داشت و این معنی‌داری ناشی از اختلاف بین چینی‌ها با لبنانی‌ها، و لبنانی‌ها با عراقی‌هاست و بین چینی‌ها با عراقی‌ها و سوریه‌ای‌ها و همچنین سوریه‌ای‌ها با لبنانی‌ها و عراقی‌ها تفاوتی وجود ندارد. همچنین، از منظر زبان اول و جنسیت، تنوع واژگانی عربی‌زبانان به‌طور معنی‌داری بیشتر از تنوع چینی‌زبانان بود و متون زن‌ها تنوع واژگانی بیشتری نسبت به متون مردها داشت، اما دو فرضیه مطرح‌شده درخصوص رابطه تنوع واژگانی با سن و تحصیلات مورد تأیید قرار نگرفت، زیرا این روابط، غیرمعنی‌دار (و البته منفی) بودند. نتایج این پژوهش می‌تواند به مدرّسان و آزمونگران حوزه آموزش زبان فارسی کمک کند تا به ابزاری مناسب برای ارزیابی غنای واژگانی و کسب بینش درمورد نحوه به‌کارگیری معیارهای غنای واژگانی در ارزیابی متون زبان‌آموزان دست یابند.

واژه‌های کلیدی: دایره واژگانی فعال، غنای واژگانی، تنوع واژگانی، شاخص گیراد، نوشتار.

۱. مقدمه

واژه یکی از ارکان مهم یادگیری و یاددهی زبان دوم/خارجی است. درواقع، می‌توان گفت که «واژه‌ها، آجرهای زبان هستند و در مرکز مهارت‌های زبانی قرار دارند» (Web & Nation, 2017, p. 25). بدون دستور، بسیار کم می‌توان منتقل کرد، [اما] بدون واژه‌ها هیچ چیز قابل انتقال نیست (Wilkins, 1972; as cited in Williams, 2012, p. 1). به همین دلیل، در ارزیابی عملکرد زبانی زبان‌آموزان در آزمون‌های تولیدی مانند مکالمه و نگارش نیز دایره واژگانی یکی از مهم‌ترین و تعیین‌کننده‌ترین شاخص‌هایی است که سنجیده می‌شود. اما چگونه بدانیم که آزمودنی‌های ما از دایره واژگانی مطلوبی برخوردارند؟ به عبارتی دیگر، طبق تعریف مفاهیم عمق واژگانی^۱ (دانش فرد درباره ابعاد یک واژه) و گستردگی واژگانی^۲ (تعداد واژه‌هایی که فرد می‌داند) (Zimmerman, 2014, p. 289)، دایره واژگانی آزمودنی‌ها در تولیدات کتبی و شفاهی، از چه میزانی از گستردگی و عمق برخوردار است یا باید باشد؟ درواقع، میزان دایره واژگانی فعال زبان‌آموزان می‌تواند نشانگر چه چیزی باشد؟ از یک منظر، شاید بتوان به کارگیری تعداد زیادی واژه در تولیدات زبانی را ملاکی برای تسلط و مهارت زبانی یک زبان‌آموز از لحاظ شاخص واژگانی دانست؛ هر چه گفتار و نوشتار زبان‌آموزان دربرگیرنده تعداد زیادی از واژه‌هایی باشد که آموخته‌اند، از سطح زبانی بالاتری برخوردار هستند؛ یعنی کاربرد تعداد زیادی از واژه‌ها را نشانگر تسلط زبان‌آموزان بدانیم. از این‌رو، قطعاً بررسی و تحلیل واژه‌های فعال آنان (یعنی واژه‌های موجود در گفتار و نوشتار) می‌تواند بینش‌های خوبی برای پژوهشگران و دست‌اندرکاران حوزه آموزش زبان ایجاد کند. از سویی دیگر، در آزمون‌های تولیدی، ارزیاب‌ها و آزمونگرها در زمان ارزیابی تولیدات شفاهی و کتبی آزمودنی‌ها به سازوکارهایی نیاز دارند تا بتوانند ارزیابی دقیق‌تری از دانش واژگانی آزمودنی‌ها صورت دهند و ارزیابی خود را به عینیت^۳ نزدیک کنند و از قضاوت‌های شمی اجتناب نمایند. اکنون این پرسش ایجاد می‌شود که ارزیاب‌ها و مصححان در آزمون‌های تولیدی، چه تعریفی از شاخص دایره واژگانی دارند و مبنای قضاوت آنان چیست؟ درواقع،

در ارزیابی عملکرد آزمودنی‌ها در آزمون‌های تولیدی، ارزیابی دانش واژگانی آزمودنی‌ها را باید بر چه اساسی انجام داد و چه ملاک و معیاری می‌توان داشت؟

یکی از راهکارهای ارزیابی دانش واژگانی آزمودنی‌ها، سنجش و ارزیابی «غنای واژگانی»^۵ نوشتار و گفتار آزمودنی‌هاست. مفهوم غنای واژگانی، مرتبط با به‌کارگیری واژه‌ها در زمان تولید گفتار یا نوشتار است. به عبارتی دیگر، سنجش غنای واژگانی متون، همان سنجش گستردگی واژگانی در بافت^۶ است که متمایز از سنجش گستردگی واژگانی به روش بافت‌زدوده در آزمون‌های ویژه دانش واژگانی (آزمون‌های مختص گستردگی یا عمق واژگانی) است. بدیهی است که به‌کارگیری و کاربرد واژه‌ها در گفتار و نوشتار، ریشه در دانش واژگانی افراد دارد. بنابراین، می‌توان گفت که معیارهای غنای واژگانی، به ارزیابی کاربرد واژه‌ها توسط زبان‌آموزان می‌پردازند، اما نباید فراموش کرد که این کاربرد، از دانش واژگانی دریافتی و البته «عمق واژگانی» سرچشمه می‌گیرد. درواقع، غنای واژگانی یک متن کتبی یا شفاهی، نتیجه «دانش واژگانی زیربنایی»^۷ نویسنده یا گوینده آن متن می‌باشد، اما از طریق سنجش‌های غنای واژگانی، به اندازه‌گیری عمق دانش واژگانی افراد پرداخته نمی‌شود، بلکه گستردگی دانش واژگانی آنان مورد سنجش و ارزیابی قرار می‌گیرد.

با توجه به مطالب فوق، می‌توان ضرورت انجام پژوهش حاضر را نیز این‌گونه بیان کرد که با تحلیل واژه‌های موجود در تولیدات زبانی زبان‌آموزان می‌توان به ارزیابی مطلوب‌تر و دقیق‌تری در زمینه شاخص دایره واژگانی دست یافت. از این رو، هدف این پژوهش، تحلیل و واکاوی گستردگی دایره واژگانی در نوشتار فارسی‌آموزان غیرایرانی یا به عبارتی دیگر، در بخشی از «دایره واژگانی فعال» آنان بود. این تحلیل و واکاوی براساس یکی از مؤلفه‌های غنای واژگانی با عنوان «تنوع واژگانی»^۸ (به‌کارگیری واژه‌های متنوع و پرهیز از تکرار واژه‌ها) صورت پذیرفت. از آنجاکه در پیشینه پژوهش‌های صورت‌گرفته در حوزه آموزش زبان فارسی، پژوهشی درخصوص تحلیل تولیدات زبانی فارسی‌آموزان غیرایرانی در زمینه «غنای واژگانی» یا «تنوع واژگانی» انجام نشده است، پژوهش حاضر می‌تواند تا حدودی این خلأ را پر کند. بنابراین، در این پژوهش، تحلیل و واکاوی دایره واژگانی فعال فارسی‌آموزان غیرایرانی در قالب تولیدات کتبی آنان

در آزمون‌های پایان دوره و از طریق محاسبه «تنوع واژگانی» متون با استفاده از سنج‌های با عنوان «شاخص گیراد» صورت پذیرفت. در همین راستا، پرسش - فرضیه‌های زیر که ناظر بر پنج متغیر هستند، مطرح شدند:

- ۱: آیا از لحاظ میزان تنوع واژگانی، تفاوت معنی‌داری بین فارسی‌آموزان برحسب ملیت وجود دارد؟
 - ۲: آیا از لحاظ میزان تنوع واژگانی، تفاوت معنی‌داری بین فارسی‌آموزان عربی‌زبان و چینی‌زبان وجود دارد؟
 - ۳: آیا از لحاظ میزان تنوع واژگانی، تفاوت معنی‌داری بین فارسی‌آموزان مرد و زن وجود دارد؟
 - ۴: آیا رابطه‌ای معنی‌دار بین تنوع واژگانی و سن فارسی‌آموزان وجود دارد؟
 - ۵: آیا رابطه‌ای معنی‌دار بین تنوع واژگانی و تحصیلات فارسی‌آموزان (تحصیلات دانشگاهی) وجود دارد؟
- فرضیه ۱: از لحاظ تنوع واژگانی، تفاوت معنی‌داری بین فارسی‌آموزان برحسب ملیت وجود ندارد.
- فرضیه ۲: از لحاظ تنوع واژگانی، تفاوت معنی‌داری بین عربی‌زبانان و چینی‌زبانان وجود ندارد.
- فرضیه ۳: از لحاظ تنوع واژگانی، تفاوت معنی‌داری بین فارسی‌آموزان مرد و زن وجود ندارد.
- فرضیه ۴: بین تنوع واژگانی و سن رابطه‌ای معنی‌دار وجود ندارد.
- فرضیه ۵: بین تنوع واژگانی و تحصیلات رابطه‌ای معنی‌دار وجود ندارد.

۲. پیشینه پژوهش

در زمینه تحلیل دایره واژگانی فعال فارسی‌آموزان غیرایرانی و سنجش غنای واژگانی نوشتار این دسته از فارسی‌آموزان، پژوهشی مشاهده نشد و محدود پژوهش‌های موجود نیز کم‌ترین وجه مشترک (یعنی زبان فارسی) را با موضوع پژوهش حاضر دارند و در حوزه زبان‌پریشی و لکت زبان صورت پذیرفته‌اند. پژوهش توکلی و همکاران (۱۳۹۵) با موضوع تنوع واژگانی در گفتار کودکان ایرانی پس از کاشت حلزون شنوایی، پژوهش احدی و همکاران (۱۳۹۸) در زمینه بررسی غنای واژگانی کودکان دوزبانه طبیعی (آذری-فارسی) و کم‌شنوای دارای کاشتینه حلزونی، پژوهش گلپور و همکاران (۱۳۸۵) با موضوعی تقریباً مشابه، پژوهش روحی و همکاران (۱۳۹۵) در زمینه سنجش غنای واژگانی در گفتار دوزبانه‌های آذری-فارسی دچار زبان‌پریشی، پژوهش اکبری و قربانی (۱۳۹۴) با هدف مقایسه غنای واژگانی در گفتار آزاد و توصیفی

بزرگسالان کم‌توان ذهنی و ۱۰ کودک طبیعی، و پژوهش حسین‌زاده اشرفی (۱۳۹۱) در زمینه تأثیر بازگویی داستان بر مؤلفه‌هایی همچون غنای واژگانی در کودکان سندرم داون نمونه‌هایی از این دسته از پژوهش‌ها هستند. با توجه به این خلأ پژوهشی در زمینه آموزش زبان فارسی که خود حاکی از ضرورت انجام پژوهش حاضر است، به پژوهش‌های زبان‌های دیگر درخصوص غنای واژگانی در نوشتار زبان‌آموزان اشاره می‌شود.

جارویس^۸ (۲۰۰۲) به منظور اندازه‌گیری تنوع واژگانی از طریق منحنی، میزان دقت پنج فرمول را از نظر توانایی در مدل‌سازی منحنی‌های تی‌تی‌آر (رابطه بین طول متن و یکتاواژه‌ها در یک نمودار) برای متون کتبی زبان‌آموزان و گویشوران بومی انگلیسی زبان سنجید و رابطه تنوع واژگانی و سن، آموزش زبان دوم، بسندگی در زبان دوم، زبان اول، کیفیت نگارش و دانش واژگانی را بررسی کرد؛ آزمودنی‌ها باید یک فیلم صامت را به صورت کتبی توصیف می‌کردند. نتایج نشان داد که دو فرمول از فرمول‌های فوق، منحنی‌های تی‌تی‌آر دقیقی برای بیش از ۹۰٪ از متن‌ها به‌دست‌می‌دهند. متن‌هایی که مدل‌های دقیق برای آن‌ها به‌دست‌آمد، مورد تجزیه و تحلیل بیشتر قرار گرفتند و نتایج، رابطه‌ای روشن بین تنوع واژگانی و میزان آموزش، و رابطه پیچیده‌تری بین تنوع واژگانی و زبان اول، کیفیت نوشتن و دانش واژگانی نشان داد.

اسپینوسا^۹ (۲۰۰۵) ابزار رایانه‌ای P-Lex را برای اندازه‌گیری غنای واژگانی تولیدات نوشتاری انگلیسی‌آموزان (به‌عنوان زبان خارجی) به‌کار گرفت و به این نتیجه رسید که P-Lex قادر به نشان‌دادن تمایز بین آزمودنی‌ها با سطوح بسندگی گوناگون نیست، اما همچنان اعتقاد دارد که استفاده از یک ابزار ارزیابی رایانه‌ای برای مدرّسان و پژوهشگران از اهمیت بالایی برخوردار است، زیرا چنین ابزارهایی یک شاخص عینی ارائه می‌دهند که مبنای قضاوت ذهنی مدرّسان در زمینه ارزیابی دانش واژگانی قرار می‌گیرد.

روایی‌سنجی «رخ‌نمای بسامد واژگانی»^{۱۰} به‌عنوان سنجه‌ای برای غنای واژگانی متون کتبی انگلیسی‌آموزان ایرانی، هدف پژوهش عباسیان و شیری^{۱۱} (۲۰۱۱) بوده است. آن‌ها به این نتیجه رسیدند که این رخ‌نما (=پروفایل)‌های واژگانی، ابزارهایی حدوداً پایا و روا هستند، اما به اندازه کافی قوی نیستند که بتوانند به‌تنهایی به‌عنوان سنجه‌ای برای ارزیابی غنای واژگانی به‌کار روند.

گرگوری-سیگنس و کلاول-ارویتیا^{۱۲} (2015) نیز تراکم واژگانی و تنوع واژگانی را در تولیدات کتبی دو گروه از دانشجویان سال اول دانشگاه تحلیل کردند و نتایج حاصل را با نتایج یک گروه دیگر مقایسه نمودند. آن‌ها تراکم واژگانی را با استفاده از Textalyser و تنوع واژگانی را با نرم‌افزار RANGE بررسی کردند. نتایج این پژوهش دستیابی به سنج‌ای پایا برای غنای واژگانی برای متون نوشته‌شده توسط آزمودنی‌های مشابه را امکان‌پذیر می‌سازد.

مظفری^{۱۳} (2014) در یک تحقیق پیکره‌ای، رابطه میان پیچیدگی دستوری و غنای واژگانی با کیفیت کلی نگارش انگلیسی‌آموزان ایرانی را بررسی کرد. پیکره‌ای شامل ۳۵ متن (۱۸۰۰۰۰ واژه) بود. یافته‌ها نشان داد که تقریباً بین تمام سنج‌های پیچیدگی دستوری با کیفیت نگارش در تمام سطوح بسندگی رابطه معنی‌داری وجود دارد. همچنین، درمورد تنوع و ظرافت واژگانی، تفاوت‌های معنی‌داری به لحاظ آماری وجود داشت، اما در مورد تراکم واژگانی، تفاوت معنی‌داری مشاهده نشد. تابنده فرد^{۱۴} (2016) نیز به بررسی تأثیر خواندن خبر از طریق برنامه‌های موبایلی بر غنای واژگانی متون نگارش یافته توسط انگلیسی‌آموزان ایرانی پرداخت. یافته‌های پژوهش حاکی از تأثیر مثبت این فعالیت بر دانش واژگانی و مهارت نگارش زبان‌آموزان بود.

ارزیابی توأمان گفتار و نوشتار نیز در پژوهش یو^{۱۵} (2009) مورد توجه قرار گرفت. یو با هدف بررسی این ادعا در مقیاس‌های ارزیابی^{۱۶} آزمون‌های بین‌المللی و همچنین در نظام‌های ارزشیابی خودکار (مانند e-rater) که می‌گوید بین تنوع واژگانی، کیفیت کلی متون کتبی یا شفاهی، و بسندگی زبانی آزمودنی‌ها رابطه‌ای مثبت وجود دارد، یک مطالعه روایی‌سنجی پسینی^{۱۷} با استفاده از نمونه‌ای از داده‌های بایگانی‌شده یک آزمون زبان بین‌المللی صورت داد تا از نظر تجربی بررسی کند که این روابط تا چه حد وجود دارد. اهمیت پژوهش یو در مقایسه با پژوهش‌های پیشین در این است که تنوع واژگانی تولیدات کتبی و شفاهی آزمودنی‌های یکسان را مقایسه کرده است. بنابراین، علاوه بر بررسی تفاوت در تنوع واژگانی بین نوشتار و گفتار، تأثیر موضوعات نگارش بر تنوع واژگانی متون کتبی نیز بررسی شد. در این پژوهش که از شاخص D (Malvern & Richards, 2002 & 1997) استفاده کرده، مشخص شد که D از نظر آماری، همبستگی

مثبت و معنی‌داری با ارزیابی‌های کلی کیفیت^{۱۸} در زمینه عملکرد آزمودنی‌ها در نوشتار و گفتار و همچنین با بسندگی زبانی کلی آزمودنی‌ها دارد. با وجود این، روابط قابل‌توجهی از نظر جنسیت، زبان اول، هدف از شرکت در آزمون و موضوعات نگارش بین زیرگروه‌های نمونه به دست نیامد. موضوعات مختلف نگارش نیز تأثیرات مهمی بر تنوع واژگانی داشتند (به‌ویژه موضوعاتی که داوطلبان با آن‌ها بسیار آشنا بودند)، حتی پس از کنترل توانایی نگارش و بسندگی زبانی کلی، تنوع واژگانی آزمودنی‌ها در نوشتن و صحبت کردن تقریباً در یک سطح بود.

یکی دیگر از جنبه‌های مورد توجه، فرایند رشد و تکامل «زبان میانی»^{۱۹} زبان‌آموزان در حوزه دانش واژگانی بوده تا از این طریق، هم به مسیر و مراحل رشد واژگانی آنان پی ببرند، هم سنجه‌های گوناگون را بیازمایند. یوهانسن^{۲۰} (۲۰۰۸) برحسب سنجه‌های تنوع واژگانی و تراکم واژگانی، بر الگوهای رشدی تمرکز کرد. درواقع، یوهانسن علاوه بر بررسی میزان ارتباط دو سنجۀ تراکم و تنوع واژگانی، به این پرسش پرداخت که آیا این دو سنجه، به ژانر (روایی در برابر توضیحی) و نحوه بیان (نوشتن در مقابل صحبت کردن) حساس هستند یا خیر. نتایج تحلیل داده‌های وی که متشکل از ۳۱۶ متن روایی و توضیحی شفاهی و مکتوب از چهار گروه سنی سوئدی‌زبان بود، نشان داد که هرچند می‌توان تراکم و تنوع واژگانی را برای تفاوت‌های شیوۀ بیان و تفاوت‌های رشدی به‌کار گرفت، یک واکاوی دقیق‌تر در شرایطی که هر دو سنجه درمورد داده‌هایی مشابه استفاده شوند، نشان می‌دهد که آن‌ها را نمی‌توان به جای یکدیگر به‌کار برد و قابل جایگزینی با یکدیگر نیستند. نتیجۀ دیگر یوهانسن (با استناد به تمایز میان این دو سنجه) این بود که یک روند رشدی بارزتر برای تنوع واژگانی، نسبت به تراکم واژگانی، احساس می‌شود و این یعنی تنوع واژگانی، سنجۀ بهتری برای تشخیص تفاوت بین گروه‌های سنی است.

پژوهش کراسلی و مک‌نامارا^{۲۱} (۲۰۰۹) در زمینه تحلیل تفاوت‌های واژگانی مرتبط با انسجام و مدل‌های پیوندگرا^{۲۲} در متون انگلیسی گویشوران بومی و زبان‌دوم‌آموزان صورت گرفت. تحلیل پیکره‌های پژوهش (یعنی متون انگلیسی گویشوران اسپانیایی‌زبانان از پیکره ICLE و پیکره گردآوری‌شده از دانشجویان دانشگاه دولتی می‌سی‌سی‌پی) نشان داد که متن‌های زبان اول و دوم از لحاظ بهره‌گیری نویسنده از

انتخاب‌های واژگانی، از ابعاد مختلفی که با عمق واژگانی، تنوع و ظرافت واژگانی رابطه دارند، متفاوت هستند.

به‌طور خلاصه، پژوهش‌های مربوط به بررسی تنوع واژگانی در نوشتار، در دو زمینه عمده صورت گرفته‌اند: (۱) سنجش میزان دقت فرمول/سنجه‌ها و ابزارهای رایانه‌ای در محاسبه تنوع واژگانی؛ (۲) بررسی رابطه بین تنوع واژگانی و متغیرهایی همچون زبان اول، کیفیت کلی نگارش، سطح بسندگی در زبان دوم، جنسیت، ملیت و در پژوهش حاضر نیز رابطه میان تنوع واژگانی با پنج متغیر مختلف بررسی شده است.

۳. چارچوب نظری

۳-۱. دایره واژگانی فعال

«مدت‌هاست که به وجود ابعاد مختلفی برای دانستن یک واژه و درجات گوناگون این دانش پی برده‌ایم که تمایز دریافتی-تولیدی^{۲۳}، شناخته‌شده‌ترین آن‌هاست» (Laufer & Nation, 1999, p. 34). پژوهشگران از عناوین مختلفی برای توصیف دایره واژگانی زبان‌آموزان استفاده کرده‌اند. به‌عنوان مثال، فن^{۲۴} (2000, p. 105) این‌گونه عنوان می‌دارد که دایره واژگانی عموماً به دو دسته دریافتی/غیرفعال^{۲۵} و تولیدی/فعال^{۲۶} تقسیم می‌شود. ناتینگر (Nattinger, 1988; as cited in Fan, 2000, p. 105) «دانش واژگانی غیرفعال» را با عنوان درک معنای واژه‌ها و ذخیره آن‌ها در حافظه توصیف می‌کند و «دانش واژگانی فعال» را بازیابی واژه‌ها از حافظه با کاربردشان در موقعیت‌های مناسب می‌داند. نیشن (1990, p. 5) این دو مقوله را با واژه یادگیری هم‌نشین می‌سازد: «یادگیری دریافتی» یعنی توانایی تشخیص یک واژه و به یادآوری معنای آن در زمان مواجهه با آن و «یادگیری تولیدی»، به معنی توانایی تشخیص یک واژه و به یادآوردن معنای آن به همراه توانایی نوشتن یا گفتن آن واژه در زمان مناسب است». اشمیت^{۲۷} (2000, P. 4) نیز این دو مقوله واژگانی را با مهارت‌های زبانی پیوند می‌زند: «توانایی درک یک واژه "دانش دریافتی" نامیده می‌شود که عموماً با خواندن و شنیدن در ارتباط است و اگر بتوانیم یک واژه را در صحبت کردن و نگارش تولید کنیم،

این دانش تولیدی محسوب می‌شود». طبق نظر اشمیت اصطلاحات فعال و غیرفعال اصطلاحاتی جایگزین برای دو اصطلاح تولیدی و دریافتی هستند (ibid). زیمرمن (2014, p. 289) نیز با به‌کارگیری اصطلاحات «دانش دریافتی» و «دانش تولیدی» از تبلور این دو دانش در مهارت‌های زبانی می‌گوید: «دانش دریافتی» یعنی تشخیص یک واژه در خواندن یا شنیدن، و «دانش تولیدی» یعنی استفاده از یک واژه در نوشتن یا صحبت کردن». همچنین، وب و نیشن (2017) «دانش تولیدی» یا همان دانش فعال را دانشی می‌دانند که برای به‌کارگیری یک واژه (نوشتن و صحبت کردن) لازم است، مثلاً صورت مکتوب یا ملفوظ، اشتقاق‌ها و تصریفات، تداعی‌ها و باهم‌آیی‌های آن واژه؛ درواقع، دانش مورد نیاز برای درک واژه‌ها در حین خواندن و شنیدن «دانش دریافتی» یا همان «دانش غیرفعال» است که به ما امکان می‌دهد صورت شفاهی، صورت مکتوب، اشتقاق‌ها و تصریفات، معانی، تداعی‌ها، باهم‌آیی‌ها، کارکردهای دستوری یک واژه را تشخیص دهیم و بدانیم که چه زمانی به‌کار می‌روند.

دانش واژگانی تولیدی را می‌توان به دو دسته آزاد و کنترل‌شده تقسیم کرد. در این حالت، دانش واژگانی به‌طور کلی، سه بخش دارد: «دانش دریافتی/غیرفعال»، «کنترل‌شده» و «آزاد»؛ نوع اول از دانش واژگانی تولیدی یعنی دانش کنترل‌شده، شامل تولید واژه در زمان مواجهه با یک تکلیف است؛ مثلاً کامل کردن واژه *fragrant* در جمله «the garden was full of fra___ flowers». دانش تولیدی آزاد نیز مربوط به کاربرد واژه‌ها از روی اراده و اختیار و بدون نیاز به محرک‌هایی برای تولید واژه‌هایی معین می‌شود (مثل نگارش آزاد) (Laufer, 1998). به اعتقاد لاوفر، تمایز قائل شدن میان دایره واژگانی فعال آزاد و کنترل‌شده ضرورت دارد، زیرا تمام زبان‌آموزانی که واژه‌های کم‌بسامد (کم‌تر رایج) را در زمانی که مجبور هستند، به‌کار می‌برند، زمانی که انتخاب واژه‌ها با خودشان باشد، آن واژه‌ها را به‌کار نخواهند برد. به‌هرحال، آنچه که از تعاریف و دسته‌بندی‌ها برمی‌آید این است که «بیشتر انگاره‌های دانش واژگانی بین دایره واژگانی غیرفعال/دریافتی و فعال/تولیدی تمایز قائل می‌شوند» (Laufer, & Paribakht, 1998, p. 368). توانایی کاربرد یک واژه را چه دانش بدانیم چه کنترل، پژوهشگران عموماً توافق دارند که درک واژه به‌طور خودکار، پیش‌بینی‌کننده کاربرد صحیح واژه نیست، و دانش غیرفعال/دریافتی معمولاً از دانش فعال پیشی می‌گیرد (ibid, p. 369). درخصوص زبان‌آموزان نیز باور کلی این است که واژه‌هایی که زبان‌آموزان از زبان

مقصد می‌دانند، بیشتر از آن تعدادی است که واقعاً می‌توانند به‌کار ببرند (Fan, 2000). پژوهشگران همواره به سنجش و اندازه‌گیری دو مقوله دایره واژگانی دریافتی و تولیدی و ارتباط این دو نوع دانش واژگانی با یکدیگر نیز توجه کرده‌اند. در بخش بعدی، به مسائل مرتبط با نحوه آزمودن و اندازه‌گیری دانش واژگانی تولیدی/فعال از طریق مفهوم غنای واژگانی پرداخته می‌شود.

۳-۲. غنای واژگانی

بدیهی است که آزمون‌های تولیدی زبان‌آموزان را نمی‌توان بدون در نظر گرفتن شاخص‌ها و معیارهایی همچون واژه و دستور ارزیابی کرد. درواقع، یکی از ارکان ارزیابی تولیدات نوشتاری و گفتاری زبان‌آموزان، شاخص و معیار دایره واژگانی است که در پژوهش‌های مربوط به سنجش دایره واژگانی تولیدی زبان‌آموزان، با اصطلاح «غنای واژگانی» معرفی می‌شود. از آنجاکه «اندازه دایره واژگانی»^{۲۸} زبان‌آموزان در گفتار و نوشتار، یکی از مهم‌ترین عوامل تعیین‌کننده میزان تسلط آن‌ها بر یک زبان دوم/خارجی است، «معیارهای سنجش غنای واژگانی» می‌توانند به مدرّسان و پژوهشگران در ارزیابی تولیدات کتبی و شفاهی زبان‌آموزان در این زمینه کمک کنند. به اعتقاد رید^{۲۹} (2000, p. 200) «پس از اتمام تکالیف نگارش باید تحلیل‌های آماری خود را برای سنجش کارآمدی دایره واژگانی زبان‌آموزان آغاز کنیم که این کار را با معیارها و مؤلفه‌های مطرح‌شده ذیل مفهوم غنای واژگانی انجام می‌دهیم». طبق تعریف لاورفر و نیشن (Laufer & Nation, 1995, p. 307) این معیارها برای کمی‌سازی میزان به‌کارگیری متنوع و زیاد واژه‌ها توسط یک نویسنده [یا گوینده] به‌کار می‌روند. پس می‌توان گفت که غنای واژگانی زبان‌آموزان را می‌توان مرتبط با به‌کارگیری واژه‌ها در زمان تولید گفتار یا نوشتار دانست، اما بدیهی است که این به‌کارگیری و کاربرد واژه‌ها، در دانش واژگانی افراد ریشه دارد. بنابراین، معیارهای غنای واژگانی، به ارزیابی کاربرد واژه‌ها توسط زبان‌آموزان می‌پردازند، اما نباید فراموش کرد که این کاربرد، از دانش واژگانی دریافتی یا همان عمق واژگانی زبان‌آموزان سرچشمه می‌گیرد. درواقع، غنای واژگانی یک متن کتبی یا شفاهی، نتیجه دانش واژگانی زیربنایی نویسنده یا گوینده آن متن است، اما از طریق سنجش‌های

غناي واژگاني، به اندازه‌گیری عمق دانش واژگاني افراد نمی‌پردازیم، بلکه گستردگی دانش آنان را می‌سنجیم. غناي واژگاني زبان‌آموزان را می‌توان با استفاده از چندین مؤلفه و سنجه توصیف کرد که تنوع واژگاني و تراکم واژگاني از آن جمله هستند. رید (Read, 2000, p. 200) غناي واژگاني را اصطلاحی کلی و فراگیر می‌داند که شامل مؤلفه‌هایی همچون «تراکم واژگاني»، «پیچیدگی واژگاني»، «تنوع واژگاني» و «تعداد کم خطاها» است. به‌طور کلی، مؤلفه‌ها و سنجه‌های غناي واژگاني مربوط به گستردگی دانش واژگاني هستند و متفاوت از مؤلفه‌های عمق دانش واژگاني (مانند باهم‌آیی و...) می‌باشند. در ادامه، به معرفی یکی از این مؤلفه‌ها با عنوان «تنوع واژگاني» پرداخته می‌شود، زیرا در پژوهش حاضر، تحلیل نوشتار فارسی‌آموزان غیرایرانی براساس مؤلفه «تنوع واژگاني» است.

۳-۱-۲. تنوع واژگاني

یکی از مهم‌ترین مؤلفه‌هایی که به‌عنوان مؤلفه رایج غناي واژگاني بیان شده، «تنوع واژگاني» است. این اهمیت تا جایی است که گاهی اصطلاح «تنوع واژگاني» معادل «غناي واژگاني» محسوب شده است. «تنوع واژگاني» شاخصی برای تعداد واژه‌های گوناگون در یک متن است. به عبارتی دیگر، «تنوع واژگاني»، به‌کارگیری طیفی از واژه‌های گوناگون و پرهیز از تکرارهای واژگاني در متون زبان‌آموزان را نشان می‌دهد. به‌طور سنتی، تنوع واژگاني یک متن، از تقسیم کردن تعداد واژه‌های گوناگون (یکتاواژه^{۲۰}) بر تعداد کل واژه‌های متن (موردواژه^{۲۱}) حاصل می‌شود. بنابراین، می‌توان گفت که ابتدایی‌ترین فرمول و سنجه‌ای که برای تنوع واژگاني مطرح شده، همان نسبت یکتاواژه-موردواژه (=تایپ-توکن) (موسوم به تی‌تی‌آر)^{۲۲} است؛ یعنی نسبت میان واژه‌های گوناگون در متن و تعداد کل واژه‌های آن (تعداد واژه‌ها ضربدر ۱۰۰ تقسیم بر تعداد موردواژه‌ها) (Laufer & Nation, 1995):

$$LV = \frac{\text{number of types} \times 100}{\text{number of tokens}}$$

در فرمول ذکرشده در کتاب رید (2000, p. 203) همین فرمول با استفاده از واژه قاموسی (یعنی

(lexeme) صورت‌بندی شده است:

$$LV = \frac{\text{number of different lexemes in the text}}{\text{the total number of lexemes in the text}}$$

با نگاهی به عملکرد فرمول فوق به آسانی می‌توان دریافت که هرچه یک متن از واژه‌های متنوع‌تری برخوردار باشد، تنوع واژگانی بالاتری هم خواهد داشت. پس برای اینکه زبان‌آموزان، متنی متنوع به‌لحاظ واژگانی داشته باشند، باید تا حد امکان از تکرار واژه‌ها اجتناب کنند و تنوع را با به‌کارگیری واژه‌های گوناگون نشان دهند. اما درخصوص فرمول مذکور که با عنوان تی‌تی‌آر نیز شناخته می‌شود، این مشکل وجود دارد که تی‌تی‌آر، به طول متن حساس است؛ یعنی «متنی که حاوی تعداد زیادی موردواژه است، مقادیر کم‌تری از تی‌تی‌آر را به دست می‌دهد و بالعکس؛ دلیل این اتفاق هم این است که تعداد موردواژه‌ها یا یکتاواژه‌ها می‌تواند بی‌نهایت افزایش یابد، زیرا برای تولید یک «واژه واژگانی» جدید، نویسنده یا گوینده مجبور است از چندین «واژه نقشی» استفاده مجدد کند و در نتیجه، معمولاً مقدار تی‌تی‌آر متعلق به متن طولانی‌تر، کم‌تر از متن کوتاه‌تر می‌شود» (Johansson, 2008, p. 62). از این‌رو، پژوهشگران همواره کوشیده‌اند که این نقصان یعنی حساسیت به حجم نمونه را در تی‌تی‌آر برطرف کنند. به گفته دالر و همکاران (Daller, et al., 2007, p. 14) «تی‌تی‌آر به مدت تقریباً یک قرن محل بحث بوده، اما امروزه نیز از آن استفاده می‌شود؛ مثلاً برخی از پژوهشگران از این سنج به طول متن کنترل‌شده استفاده کرده‌اند»، بنابراین، تی‌تی‌آر، با آن که همچنان کاربرد دارد، به طول حساس است و هرچه متن طولانی‌تر باشد، عدد به‌دست‌آمده کوچک‌تر می‌شود. برخی پژوهش‌ها حتی فراتر از این رفته‌اند و گزارش کرده‌اند که چون تی‌تی‌آر با افزایش اندازه موردواژه‌ها سقوط می‌کند، نشانگر غنای واژگانی نیست و برای سنجش رشد واژگانی، ابزاری ناکارآمد و فاقد پایایی است (McCarthy, Johansson, 2008; Vermeer, 2000; McKee, et al, 2000). ورمیر (2000, p. 77) تأکید می‌کند از آنجاکه تغییرات موضوعی می‌تواند به شدت بر نتایج سنج‌های واژگانی تأثیر بگذارد، باید متون مربوط به موضوعات مشابه را مقایسه کنیم و از تکالیف کنترل‌شده استفاده نماییم. به همین دلیل است که مک‌کارتی و جارویس (2007) معتقدند که بیش از ۶۰ سال است که تلاش‌های زیادی برای دستیابی به یک شاخص پایا برای تنوع واژگانی صورت می‌گیرد.

در نتیجه، بسیاری از محققان و پژوهشگران تلاش کرده‌اند که جایگزین‌هایی برای تی‌تی‌آر ساده (همان تقسیم تعداد یکتاواژه‌ها بر مجموع موردواژه‌ها) بیابند. یکی از نتایج این تلاش‌ها، سنجه‌ای با عنوان «شاخص گیراد» است که در بخش بعدی بیان می‌شود.

۳-۲-۲. «شاخص گیراد»^{۳۳}

در بخش پیشین بیان شد که تی‌تی‌آر ساده، از طول متن تأثیر می‌پذیرد. به همین دلیل، در مواردی که متون ما طول متفاوتی دارند، تی‌تی‌آر می‌تواند به ابزاری ناپایا تبدیل شود. پژوهشگران، برای کاهش این تأثیرپذیری از طول متن، روش‌هایی جایگزین پیشنهاد کرده‌اند. برخی از جایگزین‌ها، تنوع واژگانی هر متن را براساس فهرست‌های بسامدی برآمده از گفتار یا نوشتار گویشوران بومی محاسبه می‌کنند. برای مثال، لاور و نیشن (۱۹۹۵) در پژوهش خود پس از بررسی مشکلات و معایب روش‌های محاسبه مؤلفه‌های غنای واژگانی (همچون تراکم واژگانی، تنوع واژگانی با فرمول تی‌تی‌آر ساده، اصالت واژگانی و...)، یک رخ‌نما (=پروفایل) را معرفی کرده‌اند. این رخ‌نما که با عنوان «رخ‌نمای بسامد واژگانی» لاور و نیشن شناخته می‌شود و برای اجتناب از نواقص سنجه‌های پیشین برای اندازه‌گیری غنای واژگانی زبان‌آموزان ارائه شده است، مبتنی بر سه فهرست واژگانی است. «تی‌تی‌آر پیشرفته»^{۳۴} و «گیراد پیشرفته»^{۳۵} نیز نمونه‌هایی دیگر از سنجه‌های مبتنی بر فهرست‌های واژگانی هستند که توسط دالر^{۳۶} و همکاران (۲۰۰۳) معرفی شده‌اند. اما برخی جایگزین‌ها تغییراتی به لحاظ فرمولی و ریاضی در تی‌تی‌آر ساده ایجاد کرده‌اند که یکی از آن‌ها «شاخص گیراد» است. این شاخص که با عنوان «جذر تی‌تی‌آر»^{۳۷} نیز شناخته می‌شود، همچون تی‌تی‌آر وابسته به فهرست‌های بسامدی نیست (برعکس رخ‌نمای لاور و نیشن). شاخص گیراد از فرمول زیر برای محاسبه تنوع واژگانی متن استفاده می‌کند. همان‌طور که در فرمول مشاهده می‌شود، شاخص گیراد (برعکس تی‌تی‌آر)، یک شمارش ساده موردواژه‌ها را در دستور کار خود قرار نمی‌دهد:

$$G = \text{types} / \sqrt{\text{token}}$$

«تی‌تی‌آر اصلاح‌شده»^{۳۸} فرمول دیگری است که توسط کُروِل^{۳۹} مطرح شد. هر دو فرمول کُروِل و گیراد فرض می‌کنند که ریزش، متناسب با جذر شمارش موردواژه‌هاست و اساساً به همان اندازه می‌رسند که تی‌تی‌آر ضرب در $\sqrt{N}$ (گیراد) یا $\sqrt{N}/2$ (کُروِل) می‌شود (Malvern, et al., 2004):

$$CTTR = \frac{\text{Types}}{\sqrt{2 \times \text{tokens}}}$$

Guiraud's Root TTR	Carroll's Corrected TTR
$RTTR = \frac{V}{\sqrt{N}}$ $= \frac{\sqrt{N}}{N} V$ $= \sqrt{N} \frac{V}{N} = \sqrt{N} \times TTR \quad (2.2)$	$CTTR = \frac{V}{\sqrt{2N}}$ $= \frac{\sqrt{N}}{N\sqrt{2}} V$ $= \sqrt{\frac{N}{2}} \frac{V}{N} = \sqrt{\frac{N}{2}} \times TTR \quad (2.3)$

تصویر ۱. برابری شاخص گیراد و تی‌تی‌آر اصلاح‌شده (Malvern et al., 2004: 26)

Figure 1. Equality of Guiraud index and modified TTR (Malvern et al., 2004: 26)

از آنجاکه در پژوهش حاضر، بنا بر استفاده از روش‌های فرمولی و غیرمتکی بر فهرست‌های بسامدی بود، از شاخص گیراد برای بررسی تنوع واژگانی تولیدات کتبی استفاده شد. درخصوص دلیل عدم بهره‌گیری از نرم‌افزارهایی همچون کومتریکس نیز باید گفت که علاوه بر این که دسترسی به این نرم‌افزارها برای پژوهشگران به‌ویژه پژوهشگران ایرانی دشوار یا غیرممکن است، با زبان فارسی نیز سازگار نیستند و در تشخیص کاراکترهای خط فارسی و محاسبه تنوع واژگانی متون فارسی کارایی مطلوبی ندارند. همچنین، «شاخص گیراد» نتایج مطلوب‌تری در مقایسه با تی‌تی‌آر ساده به دست می‌دهد. جدول ۱ مقایسه کارآمدی «شاخص گیراد» و تی‌تی‌آر ساده در زمینه برآورد میزان تنوع واژگانی متون و تفاوت نتایج محاسبات هر دو را به‌خوبی نشان می‌دهد.

جدول ۱. مقایسه نتایج محاسبه تنوع واژگانی با تی تی آر و گیراد

Table 1. Comparison of the results of calculating the LD (TTR and Guiraud)

گیراد	تی تی آر	یکتاواژه‌ها	موردواژه‌ها	
۵/۷۸	۰/۵۴	۶۲	۱۱۵	زبان آموز ۱
۶/۱۵	۰/۶۰	۶۳	۱۰۵	زبان آموز ۲
۶/۷۴	۰/۶۰	۷۵	۱۲۴	زبان آموز ۳
۶/۹۴	۰/۴۵	۱۰۸	۲۴۲	زبان آموز ۴
۸/۶۹	۰/۵۸	۱۳۱	۲۲۷	زبان آموزه

طبق این جدول، براساس تی تی آر ساده، تمایز مطلوبی بین زبان آموزی که متنی کوتاه و زبان آموزی که متنی طولانی تر تولید کرده ایجاد نمی شود و حتی به یکی از زبان آموزان قوی (زبان آموز ۴) که هم متن طولانی تری داشته و هم یکتاواژه های بیشتری نسبت به زبان آموزان ردیف های اول تا سوم داشته، عدد ۰/۴۵ اختصاص یافته (کمتر از زبان آموزان ردیف های یک تا سه)؛ همچنین، مشاهده می شود که تی تی آر ساده نتوانسته تمایزی بین زبان آموز دوم و سوم یا حتی پنجم قائل شود و نمره تنوع واژگانی اختصاص یافته به هر کدام از این سه زبان آموز، تقریباً یکسان است (۰/۶۰؛ ۰/۶۰؛ ۰/۵۸). این مسئله یادآور سخن دالر و همکاران (2007, p. 15) است که می گویند «با این معیار، گویندگانی که متن های طولانی تری تولید می کنند به نحوی نظام مند، نمره کمتری می گیرند و کسانی که متن کوتاه تری تولید می کنند (اغلب، نشان دهنده سطح بسندگی پایین تر) نمرات بالاتری می گیرند». این در حالی است که با نگاهی به نتایج مندرج در ستون مربوط به «شاخص گیراد» درمی یابیم که تمایز میان زبان آموزان ضعیف تا قوی به خوبی رعایت شده و نمره تنوع واژگانی زبان آموزان از ۴/۷۳ آغاز شده و با روندی صعودی افزایش یافته است.

۴. روش شناسی

۴-۱. پیکره پژوهش

پیکره پژوهش حاضر را مجموعه ای از تولیدات کتبی فارسی آموزان غیرایرانی مرکز آرفای دانشگاه

بین‌المللی امام خمینی^(ه) در آزمون‌های پایان‌دوره دوره پیشرفته/تکمیلی تشکیل داد که مشتمل بر تولیدات زبانی کتبی ۲۵۱ فارسی‌آموز چینی، لبنانی، سوریه‌ای و عراقی در چندین دوره از آزمون‌های پایان‌دوره بود. این پیکره، پس از پیاده‌سازی و استخراج موردواژه‌ها طبق موارد بیان‌شده در بخش‌های ۴-۴ و ۴-۵، شامل ۳۶۷۱۱ موردواژه و ۱۹۸۸۹ یکتاواژه بود (به‌طور میانگین، ۱۴۶/۲۵۹ موردواژه و ۷۹/۲۳ یکتاواژه به‌ازای هر آزمودنی). در جدول ۲، توصیفی آماری از تعداد آزمودنی‌ها (برحسب ملیت، جنسیت، سن و تحصیلات) قابل مشاهده است.

جدول ۲. توصیف آزمودنی‌ها
Table 2. Description of subjects

درصد	تعداد		
۳۱.۵٪	۷۹	چینی	ملیت
۴۱.۸٪	۱۰۵	لبنانی	
۹.۲٪	۲۳	عراقی	
۱۷.۵٪	۴۴	سوری	
۴۷٪	۱۱۸	۱۶ تا ۴۰	سن
۳۵.۹٪	۹۰	۲۵ تا ۲۱	
۱۷.۱٪	۴۳	بالای ۲۵	
۴۰.۶٪	۱۰۲	زن	جنسیت
۵۹.۴٪	۱۴۹	مرد	
۷۷.۳٪	۱۹۴	دیپلم یا پایین‌تر	تحصیلات
۱۷.۵٪	۴۴	کارشناسی	
۵.۲٪	۱۳	ارشد/ دکتری	

۲-۴. آزمون‌های نگارش مرکز آزفا

آزمون‌های یادشده، آزمون‌های پایان‌دوره دوره‌های پیشرفته/تکمیلی فارسی‌آموزی هستند که هر سال

حداقل دو بار، در پایان هر نیمسال تحصیلی معمولاً ۱۶ هفته‌ای برگزار می‌شوند. این آزمون‌ها به‌منظور ارزشیابی توانایی زبانی فارسی‌آموزان در زمینه چهار مهارت زبانی شنیدن، خواندن، مکالمه، و نگارش برگزار می‌شوند. در پژوهش حاضر، تمرکز ما بر آزمون‌های مربوط به مهارت تولیدی نوشتن بوده است. روند کلی و رایج در آزمون‌های نگارش بدین نحو است که به آزمودنی‌ها یک موضوع نگارش ارائه می‌شود و از آن‌ها خواسته می‌شود که در زمان مشخص‌شده، متنی را در حداقل ۱۵ و حداکثر ۲۰ خط (تقریباً ۲۵۰ واژه) درباره آن موضوع بنویسند. در ادامه، نمونه‌ای از یک موضوع نگارش آمده است.

یک متن درباره‌ی موضوع زیر در ۲۰-۱۵ خط بنویسید. هر خط باید ۱۲-۱۰ واژه داشته باشد.
امروزه مشاغل/شغل‌های خانگی رواج زیادی در جامعه پیدا کرده است. وضعیت مشاغل خانگی در کشور شما چگونه است؟ به نظر شما داشتن شغل خانگی چه مزایا و معایبی دارد؟

۳-۴. لنکس باکس

از آنجاکه محاسبه «تنوع واژگانی» متون بر پایه شاخص گیراد انجام شد، ضرورت داشت که تعداد دقیق موردواژه‌ها و یکتاواژه‌های هر متن از پیکره پژوهش به صورت جداگانه استخراج شود. به منظور بهره‌گیری از امکانات رایانه‌ای و نرم‌افزاری و برای افزایش دقت و سرعت استخراج و شمارش موارد یادشده، در مقایسه با شمارش دستی، از یک نرم‌افزار بهره برده شد. به‌طور خلاصه، برای شمارش موردواژه‌ها و یکتاواژه‌ها، موارد زیر در دسترس بودند:

۱) نرم‌افزار لنکس باکس (LancsBox 6.0)

۲) نرم‌افزار ورداسمیت

۳) نرم‌افزار آنت کانک

۴) ابزار اینترنتی کومتریکس

برای شمارش موردواژه‌ها و یکتاواژه‌ها، هر چهار ابزار مذکور آزمایش شد تا ابزاری انتخاب شود که بیشترین سازگاری را با زبان فارسی داشته باشد و علاوه بر سرعت در پردازش و محاسبه، کار با آن

آسان‌تر باشد. از این‌رو، ۱۰ متن از متون کتبی فارسی‌آموزان به‌عنوان نمونه‌هایی از داده‌های پژوهش، به‌عنوان درونداد به ابزارهای مذکور عرضه شد تا عملکرد آن‌ها در زمینه شمارش موردواژه‌ها و یکتاواژه‌ها، براساس سه شاخص «سازگاری با زبان فارسی» «سرعت پردازش» و «کاربرد آسان» (اصطلاحاً، کاربرد دوستی^{۴۱}) بررسی شود. نتایج حاصل از بررسی و آزمایش این نرم‌افزارها و ابزار اینترنتی کومتریکس نشان داد که کومتریکس هیچ‌گونه سازگاری و تناسبی با زبان فارسی ندارد و از بین سه نرم‌افزار دیگر، با توجه به نتایج بررسی براساس سه معیار گفته‌شده، نرم‌افزار لَنکس‌باکس برای شمارش موردواژه‌ها و یکتاواژه‌ها انتخاب شد.

۴-۴. پیاده‌سازی و نرمال‌سازی

به‌منظور محاسبه تنوع واژگانی متون هر آزمودنی، اقدام به پیاده‌سازی متون در واژه‌پرداز «مایکروسافت ورد» و ذخیره‌سازی در قالب فایل‌های docx و نرمال‌سازی^{۴۲} آن‌ها شد. در ادامه، نحوه پیاده‌سازی و نرمال‌سازی متون بیان می‌شود. روش پیاده‌سازی به این صورت بود که متون، به صورت دقیق و بدون کم و کاست، توسط یک نفر (یا با ابزار اینترنتی «تایپ صوتی گوگل»^{۴۳} یا به روش تایپ دستی) پیاده‌سازی می‌شد، سپس بازخوانی و تطبیق متن روی برگه با متن تایپ‌شده، و کنترل/بازبینی نهایی توسط پژوهشگران صورت می‌پذیرفت. دلیل این تنوع در روش تایپ (تایپ دستی و تایپ صوتی) قطعی گاهگاهی اینترنت یا کاهش سرعت آن در زمان تایپ صوتی بود که تایپیست را به استفاده از تایپ دستی وادار می‌کرد. سپس نرمال‌سازی صورت گرفت. به فرایندی که طی آن، متن به یک فرمت استاندارد تبدیل می‌شود، نرمال‌سازی می‌گویند؛ در جریان این تبدیل، متون مخدوش و نامرتب به متونی ویراسته به لحاظ نگارشی و ویرایشی بدل می‌شوند (میرزائی، ۱۳۹۶، ص. ۲۲):

- چسباندن علائم نگارشی به واژه پیش از خود و فاصله گرفتن از واژه بعدی
- اتصال اجزای یک واژه به یکدیگر و در صورت لزوم، به‌کارگیری نیم‌فاصله بین اجزا
- جداسازی واژه‌ها از یکدیگر با یک فاصله کامل و حذف فاصله‌های اضافی بین واژه‌ها و سطرها در

طول متن

- یکدست‌سازی فونت‌ها (مثلاً نباید ی و ک عربی هم‌زمان با نوع فارسی آن‌ها در متن باشد). یکی از موارد نیازمند توجه در تایپ صوتی، قطعه‌بندی^{۹۹} است. «به فرایند تشخیص مرز کلمات، عبارات، جملات، بندها و اساساً هر نوع واحد معنی‌داری در متن، قطعه‌بندی می‌گویند» (ibid). درواقع، گاهی مرزبندی واژه‌ها و پاره‌گفته‌ها، با نوع خوانش پاره‌گفته ارتباط زیادی دارد و بازخوانی و بررسی مجدد، به تعیین دقیق مرز و تعداد واژه‌ها کمک می‌کند (مزیتی که ابزار رایانه‌ای تایپ صوتی از آن بی‌نصیب است). بنابراین، با توجه به این مسائل و همچنین، با عنایت به برخی موارد مشاهده‌شده در پیکره متون، این ضرورت وجود داشت که تصمیماتی اتخاذ شود و رویه‌ای ثابت در پیش گرفته شود. ازاین‌رو، پس از پیاده‌سازی و نرمال‌سازی، اقدام به کنترل/بازبینی نهایی و تطبیق متون آزمودنی‌ها با متون پیاده‌سازی‌شده، و آماده‌سازی پیکره نهایی براساس موارد زیر شد:

۱- قاعده کلی، ثبت پاره‌گفته‌های معنی‌دار بود. بنابراین، اگر پاره‌گفته‌ای قابل فهم و معنی‌دار و درعین‌حال، حاوی خطاهای دستوری بود، به همان صورت در پیکره حفظ شد. نمونه‌هایی از پاره‌گفته‌های معنی‌دار، اما حاوی خطا:

- این بحران‌ها روی انسان‌ها تأثیر می‌گذارد و باید جلوگیری جنگ‌ها باشد.
- امروزه مردم برای درآمد افزایش کار می‌کند. چطور برای بیماری‌ها کاهش می‌دهد. چون نیاز وقت زیاد برای استراحت می‌کند.

و نمونه‌هایی از پاره‌گفته‌های نامفهوم و بی‌معنی که از پیکره حذف شدند:

- آن‌ها باید به کدام طور با چیزهای طبیعت بلد باشید.
- انسان‌ها باید با ترسناک و حیوانات با کدام طور معالجه کنند.
- دوست دارم روی عادت و تقلید معرفی می‌کنیم.
- ۲- پاره‌گفته‌های غیرمرتبط با موضوع اصلی نگارش حذف نشدند؛ یعنی مهارت تولیدی نوشتن به مهارت ادراکی خواندن و درک موضوع نگارش گره زده نشد.

- ۳- خطاهای املایی و رسم الخطی کم‌اهمیت و قابل اغماض نادیده گرفته شد و همان صورت صحیح واژه‌ها در پیکره ثبت گردید. نمونه: فوتبال (فوتبال)؛ جنک (جنگ)؛ داست (داشت)؛ اشنا (آشنا)؛ آسب (آسیب)؛ انترنت (اینترنت)؛ تنک (تنگ)؛ گوچک (کوچک)؛
- ۴- از لحاظ به‌کارگیری گونه گفتاری/غیررسمی یا نوشتاری/رسمی واژه‌ها، هر صورتی که آزمودنی به‌کار برده بود، همان در پیکره ثبت شد، هرچند در عمده متون، غلبه با گونه نوشتاری/غیررسمی بود.
- ۵- اگر آزمودنی اعداد را به صورت حروفی نوشته بود، همان صورت حروفی در پیکره ثبت شد و در غیر این صورت، شکل رقمی/عددی حفظ شد.

۵.۴. موردواژه‌یابی^{۴۵}

پس از پیاده‌سازی و نرمال‌سازی متون، نوبت به آماده‌سازی متون برای شمارش تعداد موردواژه‌ها و یکتاواژه‌ها به منظور محاسبه تنوع واژگانی هر متن رسید. در این مرحله، هر متن، در قالب یک فایل docx جداگانه (یعنی به صورت هر متن، یک فایل) درآمد، سپس هر فایل به صورت جداگانه، به عنوان درونداد به لنکس‌باکس عرضه شد تا تعداد موردواژه‌ها و یکتاواژه‌های هر متن به دست آید. خوشبختانه، لنکس‌باکس علاوه بر سرعت بالا در انجام دستورات، سهولت در کاربرد و برخورداری از رابط کاربری بسیار مطلوب، سازگاری با زبان فارسی، سازگاری مطلوبی نیز با فایل‌های متنی docx دارد و پژوهشگر نیازی به تغییر این فایل‌ها و تبدیل آن‌ها به فرمت «متن ساده»^{۴۶} ندارد. در این مرحله نیز باید تصمیماتی در خصوص نحوه پالایش و آماده‌سازی متون، پیش از عرضه به لنکس‌باکس و شمارش موردواژه‌ها و یکتاواژه‌ها اتخاذ می‌شد تا به یک وحدت رویه در مورد نحوه آماده‌سازی متون برسیم. از این رو، عمدتاً دستورالعمل موردواژه‌یابی (ذکرشده در میرزائی، ۱۳۹۶، ص. ۲۴) مبنای کار قرار گرفت که براساس آن، هر واژه نحوی، یک موردواژه محسوب می‌شود؛ واژه نحوی می‌تواند ساده، اشتقاقی، ترکیبی یا اشتقاقی-ترکیبی باشد (مثال: مهر، مهربان، مهری، مهرجو، نامهربان و...). در شمارش موردواژه‌ها اجزای واژه‌های مرکب از هم جدا نمی‌شوند (ibid). علاوه بر آن، وندهای تصریفی نیز همراه با پایه با هم یک موردواژه می‌سازند (مانند

وندهای تصریفی شمار، جنس و حالت در اسم یا ویژگی‌های نمود، زمان، جهت، وجه در فعل و مانند آن)، اما اجزای «واژه‌های واجی»^{۴۷} باید در شمارش موردواژه‌ها از هم جدا شوند؛ به بیان دیگر، برای موردواژه‌یابی، لازم است واژه‌بست‌ها^{۴۸} از پایه جدا شوند و کسره اضافه بسته به رویکردی که در ماهیت آن اتخاذ می‌شود یا بنا به قرارداد، در شمارش وارد شود یا نشود (ibid).

۵. واکاوی داده‌ها

۱-۵. آمار توصیفی متغیرهای اصلی پژوهش

از آنجاکه برای سنجش تنوع واژگانی متون فارسی‌آموزان، یکی از فرمول‌های موسوم به «شاخص گیراد» انتخاب شد، آزمون فرضیه‌ها نیز در پژوهش حاضر با این شاخص صورت پذیرفت. «شاخص گیراد» با بهره‌گیری از دو مقوله «موردواژه» یا همان توکن و «یکتاواژه» یا همان تایپ، تنوع واژگانی متون را اندازه‌گیری می‌کند. بنابراین، دو متغیر موردواژه و یکتاواژه از متغیرهای اصلی و تأثیرگذار این پژوهش محسوب می‌شوند که خود، متغیر سومی را با عنوان شاخص «تنوع واژگانی» به دست می‌دهند. در جدول ۳ اطلاعات - توصیفی مربوط به متغیرهای اصلی پژوهش ذکر شده است.

جدول ۳. آمار توصیفی متغیرها

Table 3. Descriptive statistics of variables

متغیر	میانگین	انحراف معیار	کمینه	بیشینه
موردواژه‌ها	۱۴۶/۲۵۹	۴۱/۲۴۲	۵۶	۲۶۸
یکتاواژه‌ها	۷۹/۲۳۹	۲۳/۳۵۹	۲۸	۱۴۱
تنوع واژگانی	۶/۵۱۲	۱/۲۱	۳/۶۷	۹/۹۲

۲-۵. آزمون کولموگروف - اسمیرنوف جهت بررسی نرمال بودن توزیع متغیرها

به منظور بررسی نرمال بودن توزیع متغیرها، از آزمون کولموگروف-اسمیرنوف استفاده شد. این آزمون جهت بررسی ادعای مطرح شده در مورد توزیع داده‌های یک متغیر کمی مورد استفاده قرار می‌گیرد. فرض‌های آماری مربوط به توزیع نرمال به صورت زیر مطرح می‌شود:

H₀: متغیرها دارای توزیع نرمال هستند.

H₁: متغیرها دارای توزیع نرمال نیستند.

اگر داده‌ها دارای توزیع نرمال بودند، می‌توان از آزمون‌های پارامتریک نظیر آزمون میانگین یک جامعه، دو جامعه یا چند جامعه استفاده کرد. اگر داده‌ها دارای توزیع نرمال نبودند از آزمون‌های ناپارامتریک استفاده می‌شود. با توجه به جدول ۴ و به دلیل اینکه سطح معنی‌داری متغیر بزرگ‌تر از ۰.۰۵ است، فرض صفر تأیید و ادعای نرمال بودن توزیع این متغیر پذیرفته شد.

جدول ۴. آزمون نرمال بودن توزیع متغیرهای تحقیق

Table 4. The normality test of the distribution of variables

وضعیت	سطح معنی‌داری	آماره آزمون	متغیر
نرمال	۰.۹۸۴	۰.۴۵۹	تنوع واژگانی

۳-۵. آزمون فرضیه‌ها

فرضیه اول پژوهش این بود که «از لحاظ تنوع واژگانی، تفاوت معنی‌داری بین فارسی‌آموزان برحسب ملیت وجود ندارد». با توجه به نتایج آزمون تحلیل واریانس (جدول‌های ۵ و ۶) ملاحظه می‌شود که آزمون معنی‌دار بوده است ($\text{sig} = ۰/۰۰۱ < ۰/۰۵$)، بدین معنی که بین ملیت‌های مختلف، از نظر تنوع واژگانی، تفاوت معنی‌دار وجود دارد (رد فرضیه).

جدول ۵. مقایسه تنوع واژگانی برحسب ملیت

Table 5. Comparison of LD according to nationality

منبع تغییرات	مجموع مربعات	درجه آزادی	میانگین مربعات	آماره f	sig
بین گروهی	۲۳/۷۸۰	۳	۷/۹۲۷	۵/۷۱۸	۰/۰۰۱
درون گروهی	۳۴۲/۳۸۴	۲۴۷	۱/۳۸۶		
خطا	۳۶۶/۱۶۵	۲۵۰			

جدول ۶. میانگین و انحراف معیار تنوع واژگانی برحسب ملیت

Table 6. Mean and SD of LD according to the nationality

ملیت	میانگین	انحراف معیار	
چینی	۶/۲۴۰	۱/۵۵۱	۶/۱±۲۴۰/۵۵۱
لبنانی	۶/۸۵۲	۰/۹۸۸	۶/۰±۸۵۲/۹۸۹
عراقی	۶/۰۱۷	۰/۸۹۵	۶/۰±۰۱۷/۸۹۵
سوری	۶/۴۴۶	۰/۹۰۷	۶/۰±۴۴۶/۹۰۷

با توجه به معنی داری آزمون تحلیل واریانس، برای بررسی محل اختلاف، از آزمون تعقیبی توکی استفاده شد. طبق نتایج این آزمون که در جدول ۷ آمده، معنی داری تحلیل واریانس ناشی از اختلاف بین ملیت های چینی با لبنانی، و لبنانی با عراقی بود، در حالی که بین چینی ها با عراقی ها و سوری ها و همچنین، سوری ها با لبنانی ها و عراقی ها تفاوتی از نظر تنوع واژگانی وجود نداشت.

جدول ۷. سطح معنی داری آزمون توکی برای مقایسه عملکرد ملیت ها

Table 7. The significance level of Tukey's test for comparing the performance of nationalities

ملیت	چینی	لبنانی	عراقی
چینی	-	-	-
لبنانی	۰/۰۰۳	-	-
عراقی	۰/۸۵۳	۰/۰۱۲	-
سوری	۰/۷۸۹	۰/۲۲۴	۰/۴۸۹

فرضیه دوم ناظر بر تفاوت عملکرد آزمودنی‌ها براساس زبان اول بود و بدین صورت پیکربندی شده بود که «از لحاظ تنوع واژگانی، تفاوت معنی‌داری بین عربی‌زبانان و چینی‌زبانان وجود ندارد». برای راستی‌آزمایی این فرضیه از «آزمون تی مستقل» استفاده شد که براساس نتایج (جدول‌های ۸ و ۹) ملاحظه می‌شود که آزمون معنی‌دار بوده است ($\text{sig} = 0.04 < 0.05$)، بدین معنی که بین چینی‌زبانان و عربی‌زبانان از نظر تنوع واژگانی، تفاوت معنی‌داری وجود دارد (رد فرضیه) و درواقع، عربی‌زبانان نسبت به چینی‌زبانان، از لحاظ تنوع واژگانی برتری چشمگیری دارند.

جدول ۸. مقایسه تنوع واژگانی برحسب زبان اول

Table 8. Comparison of LD according to the first language

خطای معیار تفاوت	میانگین تفاوت	سطح معنی‌داری آزمون t	درجه آزادی	آماره t	سطح معنی‌داری آزمون لون	آماره f آزمون لون	
۰/۱۶۲	۰/۳۹۵	۰/۰۱۶	۲۴۹	۲/۴۳۱	۰/۰۰۰	۲۷/۲۱۶	برابری واریانس
۰/۱۹۰	۰/۳۹۵	۰/۰۴	۱۰۸/۶۴۲	۲/۰۸۰			عدم برابری واریانس

جدول ۹. میانگین و انحراف معیار تنوع واژگانی برحسب زبان اول

Table 9. Mean and SD of LD according to the first language

ملیت	میانگین	انحراف معیار
چینی	۶/۲۴	۱/۵۵۱
عرب	۶/۶۳۶	۰/۹۹۷

برای راستی‌آزمایی فرضیه سوم مبنی بر این که «از لحاظ تنوع واژگانی، تفاوت معنی‌داری میان فارسی‌آموزان مرد و زن وجود ندارد»، نتایج آزمون تی مستقل (جدول‌های ۱۰ و ۱۱) نشان داد که آزمون معنی‌دار بوده ($\text{sig} = 0.01 < 0.05$)؛ بدین معنی که بین زنان و مردان از نظر تنوع واژگانی تفاوت معنی‌داری

وجود دارد (رد فرضیه).

جدول ۱۰. مقایسه تنوع واژگانی برحسب جنسیت

Table 10. Comparison of LD according to gender

خطای معیار تفاوت	میانگین تفاوت	سطح معنی‌داری آزمون t	درجه آزادی	آماره t	سطح معنی‌داری آزمون لون	آماره f آزمون لون	
۰/۱۴۸	۰/۷۷۰	۰/۰۰۱	۲۴۹	۵/۲۰۴	۰/۰۵۳	۳/۷۶۸	برابری واریانس
۰/۱۵۱	۰/۷۷۰	۰/۰۰۱	۲۰۰/۷۵۲	۲/۰۸۰			عدم برابری واریانس

جدول ۱۱. میانگین و انحراف معیار تنوع واژگانی برحسب جنسیت

Table 11. Mean and SD of LD according to gender

ملیت	میانگین	انحراف معیار	
زن	۶/۹۶۹	۱/۲۲۷	۶/۱±۹۶۹/۲۲۷
مرد	۶/۱۹۹	۱/۰۹۷	۶/۱±۱۹۹/۰۹۷

فرضیه چهارم ناظر بر رابطه میان تنوع واژگانی و سن بود: «میان تنوع واژگانی و سن رابطه‌ای معنی‌دار وجود ندارد». با توجه به اینکه فرضیه از نوع رابطه‌ای است، برای راستی‌آزمایی آن از «ضریب همبستگی» استفاده شد. همچنین با توجه به کمی بودن متغیر سن، برای بررسی رابطه میان سن و تنوع واژگانی، از «ضریب همبستگی پیرسون» استفاده شد. آزمون فرضیه چهارم نشان داد که میان متغیر سن و تنوع واژگانی متون، رابطه‌ای معنی‌دار وجود ندارد (تأیید فرضیه چهارم).

جدول ۱۲. نتایج بررسی رابطه تنوع واژگانی با متغیر سن

Table 12. The results of examining the relationship between LD and age

نتیجه آزمون	سن	متغیر
		متغیر
رد فرضیه	-۰/۰۸۸ (۰/۱۶۶)	تنوع واژگانی

رابطه تحصیلات آزمودنی‌ها (بدون/با تحصیلات دانشگاهی) با تنوع واژگانی در فرضیه پنجم بیان شد: «بین تنوع واژگانی و تحصیلات رابطه‌ای معنی‌دار وجود ندارد». با توجه به کیفی و رتبه‌ای بودن متغیر تحصیلات، برای بررسی رابطه آن و تنوع واژگانی از «ضریب همبستگی رتبه‌ای اسپیرمن» استفاده شده است. در جدول ۱۳ ملاحظه می‌شود که ضریب همبستگی بین سطح تحصیلات و تنوع واژگانی تولیدات کتبی، غیرمعنی‌دار ($\text{sig}=0/۸۲۹>0/۰۵$) بوده (تأیید فرضیه) و مشخص شد که بین تحصیلات آزمودنی‌ها و تنوع واژگانی متون آنان رابطه‌ای معنی‌دار وجود ندارد.

جدول ۱۳. نتایج بررسی رابطه تنوع واژگانی با تحصیلات

Table 13. The results of examining the relationship between LD and education

نتیجه آزمون	سطح تحصیلات	متغیر
		متغیر
تأیید فرضیه	-۰/۰۱۴ (۰/۸۲۹)	تنوع واژگانی

۶. نتیجه

یکی از زمینه‌های قدیمی تحقیق و توسعه در ارزیابی دایره واژگانی زبان‌آموزان، تخمین اندازه دانش واژگانی است که اغلب به‌عنوان «گسترده‌ی دانش واژگانی» شناخته می‌شود (Read, 2007, p. 107). برای تخمین گسترده‌ی دانش واژگانی زبان‌آموزان می‌توان از روش‌های گوناگونی بهره برد که یکی از آن‌ها

استفاده از سنج‌ها و شاخص‌های کمی و ریاضی‌گونه است. در همین راستا، به بررسی دایره واژگانی فعال و سنجش «تنوع واژگانی» تولیدات زبانی کتبی فارسی‌آموزان غیرایرانی پرداخته شد. برای سنجش این مؤلفه، با استفاده از «شاخص گیراد» و با بهره‌گیری از نرم‌افزار لنکس باکس، به شمارش موردواژه‌ها و یکتاواژه‌های متون کتبی ۲۵۱ فارسی‌آموز پرداخته شد. متناسب با هدف پژوهش، پنج پرسش - فرضیه مطرح شد و با استفاده از آزمون‌های آماری گوناگون، فرضیه‌ها راستی‌آزمایی شدند.

یافته‌های پژوهش در واکاوی فرضیه نخست و دوم نشان داد که ملّیت و زبان اول می‌توانند تأثیر چشمگیری بر تنوع واژگانی داشته باشند. به‌طور کلی، از لحاظ زبان اول مشخص شد که عربی‌زبانان در مقایسه با چینی‌زبانان، از تنوع واژگانی بیشتری در نوشتار برخوردارند. این نتایج می‌توانند همسو با نتایج دیوکه و پاولنکو^۹ (2003) (در بافت دوزبانگی) تلقی شوند که نشان داد زبان اول افراد می‌تواند روی غنای واژگانی متون آنان تأثیرگذار باشد و تفاوت ایجاد کند، اما مؤید نتایج یو (2009) که مقایسه‌ای بین فلیپینی‌ها و چینی‌ها انجام داده بود، نیست (عدم وجود تفاوت معنی‌دار بین غنای واژگانی متون آزمودنی‌های فلیپینی و چینی). به‌طور کلی، رابطه زبان اول و تنوع واژگانی، رابطه‌ای پیچیده است. به‌عنوان مثال، جارویس (2002) که متون کتبی ۱۴۰ نوجوان فنلاندی و ۷۰ نوجوان فنلاندی-سوئدی که در فنلاند زندگی می‌کردند و ۶۶ نوجوان بومی انگلیسی‌زبان آمریکایی (همگی ۱۰ ساله تا ۱۵ ساله) را بررسی کرد، به این نتیجه رسید که زبان اول نمی‌تواند پیش‌بینی‌کننده خوبی برای تنوع واژگانی در نوشتار زبان‌آموزان باشد. به گفته جارویس (2002, p. 75)، گویشوران بومی، نسبت به زبان‌آموزان، از سطوح بالاتری از تنوع واژگانی برخوردارند، در حالی که توصیف اثرات بالقوه زبان اول بر تنوع واژگانی - در صورت وجود - و احتمالاً ترکیب آن با عواملی همچون سن، بسندگی در زبان دوم، و قرابت زبان اول و دوم بسیار دشوار است. جارویس اذعان داشته است که بررسی رابطه متغیرهایی مانند زبان اول، سن، میزان آموزش و... در پژوهش‌های گوناگونی صورت گرفته، اما به دلیل عدم اطمینان از پایایی شاخص‌ها و سنج‌های ارزیابی تنوع واژگانی، نتایج این پژوهش‌ها مورد اطمینان قطعی نیست. بنابراین، هم درخصوص ملّیت و زبان اول که در بالا گفته شد و هم در مورد دیگر متغیرهای پژوهش که در ادامه مطرح خواهند شد، باید به این نکته مطرح‌شده توسط جارویس دقت داشت که

همسویی یا عدم همسویی نتایج پژوهش‌ها با یکدیگر در حوزه تنوع واژگانی، می‌تواند تا حدودی ناشی از ماهیت شاخص‌ها و سنجه‌های به‌کارگرفته‌شده در هر پژوهش باشد.

درخصوص تأثیر زبان اول ذکر این نکته ضروری است که هرچند آزمودنی‌های پژوهش، تنها از دو زبان اول برخوردار بودند (چینی و عربی)، از سویی دیگر، از چهار ملیت بودند بین عربی‌زبانان نیز تفاوت‌هایی برحسب ملیت مشاهده می‌شود. بنابراین، می‌توان گفت که زبان اول و ملیت می‌تواند دو متغیر جداگانه باشند و لزوماً افرادی که از ملیت‌های گوناگون هستند، اما از زبان اول مشترکی برخوردارند، عملکرد یکسانی در زمینه غنای واژگانی نخواهند داشت.

یافته دیگر این بود که جنسیت می‌تواند تأثیر چشمگیر و معنی‌داری بر میزان تنوع واژگانی آزمودنی‌ها داشته باشد. نتایج به‌دست‌آمده از نظر تأثیر جنسیت بر تنوع واژگانی، همسو با نتایج پژوهش سینگ^{۶۰} (2001) و پژوهش دیوکه و پاولنکو (2003) است که تأثیر جنسیت (هرچند ضعیف) بر تنوع واژگانی را نشان داده بود. همچنین، این یافته تأییدکننده پژوهش هارنکوئیست^{۶۱} و همکاران (2003) است، اما مؤید پژوهش یو (2009) نیست که به روابط قابل‌توجهی بین این دو متغیر دست نیافته بود.

یافته‌های حاصل از راستی‌آزمایی فرضیه چهارم نشان داد که نه‌تنها بین سن و تنوع واژگانی رابطه مثبتی وجود ندارد و بالا رفتن سن باعث افزایش تنوع واژگانی تولیدات نمی‌شود، بلکه رابطه بین این دو متغیر از نوع منفی و غیرمعنی‌دار است؛ درواقع، در زمان نوشتن به زبان فارسی به‌عنوان زبان دوم، افزایش سن می‌تواند یک عامل بازدارنده یا کاهنده تنوع واژگانی باشد. این یافته می‌تواند به‌نحوی، هرچند با جهتی معکوس، همسو با نتایج یوهانسن (2008) باشد که می‌گوید تنوع واژگانی، نسبت به تراکم واژگانی، می‌تواند سنجۀ خوبی برای تشخیص تفاوت بین گروه‌های سنی باشد، اما با پژوهش جارویس (2002) که به این نتیجه رسید که با افزایش سن، قطعاً میزان تنوع واژگانی آزمودنی‌ها نیز افزایش می‌یابد، همخوانی ندارد. درمورد رابطه تحصیلات و تنوع واژگانی متون نیز این نتیجه حاصل شد که رابطه‌ای منفی و البته غیرمعنی‌دار وجود دارد؛ درواقع، با افزایش سطح تحصیلات، تنوع واژگانی افزایش نمی‌یابد و در مواردی دچار کاهش نیز می‌گردد که البته این نتیجه همسو با نتایج پژوهش‌هایی همچون پژوهش هارنکوئیست و

همکاران (۲۰۰۳) نیست که اذعان می‌دارند پیشینه تحصیلی می‌تواند پیش‌بینی‌کننده غنای واژگانی باشد. یافته‌های پژوهش حاضر می‌تواند از چند جهت حائز اهمیت باشد: (۱) می‌تواند این مسیر را پیش روی مدرّسان و دست‌اندرکاران آموزش زبان فارسی به غیرفارسی‌زبانان بگشاید که با استفاده از راهکارهایی مناسب، حجم انبوه واژه‌های غیرفعال/دریافتی فارسی‌آموزان را به واژه‌هایی فعال/تولیدی تبدیل کنند. یکی از این راهکارها مطابق با نظر میلتن^۲ (2009, p. 117) این است که زبان‌آموزان واژه‌ها را به صورت تولیدی بیاموزند و کاربردشان را تمرین کنند تا سهم بیشتری از واژه‌های آنان هم تولیدی و هم دریافتی باشد. (۲) توجه مؤلفان منابع آموزش زبان فارسی را به این نکته جلب کند که کتاب‌ها و منابع آموزشی خود به‌ویژه منابع مربوط به مهارت نگارش را به‌نحوی تهیه و تدوین کنند که علاوه بر ارائه واژه‌های گوناگون و متنوع، راهبردها، تمرینات و تکالیف مطلوبی را برای تبدیل این واژه‌ها به واژه‌های فعال در نوشتار معرفی کنند تا فارسی‌آموزان ترغیب شوند که با اطمینان از درک صحیح مفاهیم واژه‌ها، آن‌ها را در نوشتار خود به‌کار برند. (۳) مدرّسان و آزمونگرها را به حرکت به سمت روش‌های کمی و عینی در زمینه ارزیابی دانش واژگانی فارسی‌آموزان ترغیب کند تا مسیر دستیابی به قضاوت‌های غیرشمی، رایانه‌ای/خودکار و همچنین، جلوگیری ارزیابی‌های انتزاعی^۳ و خطاها و سوگیری‌های ارزیابی هموارتر شود. در پایان، موارد زیر به‌عنوان پیشنهادهایی برای پژوهش‌های آتی مطرح می‌شوند:

- مقایسه واژه‌های فعال و غیرفعال و بررسی رابطه آن‌ها
- واکاوی تأثیر بافت و محیط یادگیری زبان فارسی (به‌عنوان زبان دوم یا خارجی) بر میزان رشد دایره واژگانی فعال
- بررسی تأثیر موضوعات نگارش بر دایره واژگانی فعال و غنای واژگانی

۷. پی‌نوشت‌ها

1. vocabulary depth
2. vocabulary breadth
3. objectivity
4. lexical richness
5. context
6. underlying vocabulary knowledge
7. lexical diversity
8. Jarvis
9. Espinosa
10. Lexical Frequency Profile (LFP)
11. Abbasian & Shiri
12. Gregori-Signes & Clavel-Arroitia
13. Mozaffari
14. Tabandeh Fard
15. Yu
16. rating scales
17. posteriori validation study
18. overall quality ratings
19. interlanguage
20. Johansson
21. Crossley & McNamara
22. cohesion and connectionist models
23. peceptive-productive
24. Fan
25. receptive/passive
26. productive/active
27. Schmitt
28. vocabulary size
29. Read
30. type
31. token
32. Type/Token Ratio (TTR)
33. Guiraud's Index
34. Advanced TTR (A_{TTR})
35. Advanced Guiraud or Guiraud Advanced (A_G)
36. Daller
37. Root TTR (RTTR)
38. Corrected Type/Token Ratio (CTTR)
39. Carroll
۴۰. فقط یک نفر ۱۶ ساله و دو نفر ۱۷ ساله
41. user-friendliness
42. normalization
43. Google voice typing
44. segmentaion
45. tokenization
46. plain text
47. phonological words
48. clitics
49. Dewaele & Pavlenko
50. Singh
51. Härnquist
52. Milton
53. Subjective

۸. منابع

- احدی، ح.، رضایی، ح.، نیلی پور، ر.، و جاراللهی، ف. (۱۳۹۸). مقایسه میانگین طول گفته، غنای واژگانی و آگاهی‌های فرازبانی نحوی و واژگانی در کودکان دوزبانه طبیعی و کم‌شنوا (آذری-فارسی). *فصلنامه کودکان استثنایی*، ۴، ۱۰۵-۱۱۸.
- اکبری، ا.، و قربانی، ع. (۱۳۹۴). مشخصه‌های واژگانی در گفتار بزرگسالان با کم‌توانی ذهنی. *فصلنامه روان‌شناسی افراد استثنایی*، ۲۰، ۹۵-۱۰۸.
- توکلی، م.، جلیله‌وند، ن.، کمالی، م.، مدرسی، ی.، و مقصدی زرنندی، م. (۱۳۹۵). بررسی تنوع واژگانی و پیچیدگی نحوی گفتار کودکان ۹-۸ ساله پس از کاشت حلزون شنوایی. *مجله علوم پیراپزشکی و توان‌بخشی مشهد*، ۱، ۲۰-۲۹.
- حسین‌زاده اشرفی، س. (۱۳۹۱). بررسی تأثیر بازگویی داستان بر بهبود شاخص‌های زبانی میانگین طول گفته، سرعت گفتار و غنای واژگانی در کودکان دارای سندرم داون ۷-۶ ساله عقلی. پایان‌نامه‌ی کارشناسی ارشد. دانشگاه علوم بهزیستی و توان‌بخشی.
- روحی، ا.، عزبفتری، ب.، و عشایری، ح. (۱۳۹۵). بررسی غنای واژگانی و راونی زبان اول و دوم در دوزبانه‌های آذری-فارسی مبتلا به زبان‌پریشی. *جستارهای زبانی*، ۶ (۳۴)، ۳۷۱-۳۸۹.
- گلپور، ل.، نیلی پور، ر.، و روشن، ب. (۱۳۸۵). تحلیل مقایسه‌ای برخی ساختارهای صرفی-نحوی گفتار کودکان کم‌شنوای شدید-عمیق در حال آموزش با کودکان عادی ۵-۴ ساله فارسی‌زبان. *شنوایی‌شناسی-دانشکده توان‌بخشی دانشگاه علوم پزشکی تهران*، ۲، ۲۳-۲۹.
- میرزائی، آ. (۱۳۹۶). *آشنایی با زبان‌شناسی پیکره‌ای*. انتشارات دانشگاه علامه طباطبائی.

References

- Abbasian, G. R., & Shiri Parizad, M. (2011). Validation of lexical frequency profiles as a measure of lexical richness in written discourse. *Journal of Technology & Education*, 5(2), 125-131.

- Ahadi, H., Rezaee, H., Nilipour, R., & Jarollahi, F. (2019). The relationship between mean length of utterance (MLU), lexical richness and syntactical and lexical metalinguistic awareness in bilingual (Azari-Persian) normal and hearing impaired. *Journal of Exceptional Children*, 19(4), 105-118. [In Persian]
- Akbari, E., & Ghorbani, A. (2015). Lexical features in speech of mentally retarded adults. *Psychology of Exceptional Individuals*, 5(20), 95-108. [In Persian]
- Crossley, S. A., & McNamara, D. S. (2009). Computational assessment of lexical differences in L1 and L2 writing. *Journal of Second Language Writing*, 18, 119-135.
- Daller, H., van Hout, R., & Treffers-Daller, J. (2003). Lexical richness in spontaneous speech of bilinguals. *Applied Linguistics*, 24, 197-222.
- Dewaele, J.-M., & Pavlenko, A. (2003). Productivity and lexical diversity in native and non-native speech: A study of cross-cultural effects. In V. Cook (Ed), *Effects of the second language on the first* (pp. 120-141). Multilingual Matters.
- Espinosa, S. M. (2005). Can P_Lex accurately measure lexical richness in the written production of young learners of EFL? *Portada Linguarum*, 4, 7-21.
- Fan, M. (2000). How big is the gap and how to narrow It? An investigation into the active and passive vocabulary knowledge of L2 learners. *RELC Journal*, 31(2), 105-119.
- Golpour, L., Nilipour, R., & Roshan, B. (2006). A comparison between morphological and syntactic features of 4 to 5 years old in education severe to profound hearing impaired and normal children. *Bimonthly Audiology*, 15(2), 23-29. [In Persian]
- Gregori-Signe, C., & Clavel-Arroitia, B. (2015). Analysing lexical density and lexical diversity in university students' written discourse. *Procedia- Social and Behavioral Sciences* 198, 546-556.

- Härnquist, K., Christianson, U., Ridings, D., & Tingsell, J.-G. (2003). Vocabulary in interviews as related to respondent characteristics. *Computers and the Humanities*, 37, 179–204.
- Hosseinzade Ashrafi, S. (2012). *Studying the effects of stories retelling on speech indices such as the mean length of utterances, diversity of utterances and speed of speech in children with the down syndrome at the mental age of 6-7*. Master's Thesis. University of Social Welfare and Rehabilitation Speech therapy. [In Persian]
- Jarvis, S. (2002). Short texts, best fitting curves, and new measures of lexical diversity. *Language Testing*, 19, 57–84.
- Johansson, V. (2008). Lexical diversity and lexical density in speech and writing: A developmental perspective. *Lund University, Dept. of Linguistics and Phonetics Working Papers*, 53, 61-79.
- Laufer, B. (1998). The development of passive and active vocabulary in a second language: Same or different? *Applied Linguistics*, 19(2), 255–271.
- Laufer, B., & Paribakht, T. S. (1998). The relationship between passive and active vocabularies: Effects of language learning context. *Language Learning*, 48 (3), 365-391.
- Laufer, B., & Nation, P. (1995). Vocabulary size and use: Lexical richness in L2 written production. *Applied Linguistics*, 16, 307–322.
- Laufer, B., & Nation, P. (1999). A vocabulary-size test of controlled productive ability. *Language Testing*, 16(1), 33–51.
- Malvern, D., & Richards, B. (1997). A new measure of lexical diversity. In A. Ryan and A. Wray (Eds), *Evolving models of language*. Papers from the Annual Meeting of the British Association of Applied Linguists held at the University of Wales, Swansea, September 1996 (pp. 58–71). Multilingual Matters.

- Malvern, D., & Richards, B. (2002). Investigating accommodation in language proficiency interviews using a new measure of lexical diversity. *Language Testing* 19, 85-104.
- Malvern, D., Richards, B., Chipere, N., & Durán, P. (2004). *Lexical diversity and language development: Quantification and assessment*. Palgrave Macmillan.
- McCarthy, P. M., & Jarvis, S. (2007). VOCED: A theoretical and empirical evaluation. *Language Testing*, 24, 459–488. doi:10.1177/0265532207080767
- McCarthy, P. M., & Jarvis, S. (2010). MTLT, vocd-D, and HD-D: A validation study of sophisticated approaches to lexical diversity assessment. *Behavior Research Methods*, 42(2), 381-392. doi:10.3758/BRM.42.2.381
- McKee, G., Malvern, D., & Richards, B. (2000). Measuring vocabulary diversity using dedicated software. *Literary and Linguistic Computing*, 15(3), 323-337. doi:10.1093/lc/15.3.323
- Milton, J. (2009). *Measuring second language vocabulary acquisition*. Multilingual Matters Publications.
- Mirzaei, A. (2017). *An introduction to corpus linguistics*. Allameh Tabataba'i University Press. [In Persian]
- Mozaffari, A. (2014). *Corpus-based analysis of grammatical complexity and lexical richness in second language writing performance*. Master's Thesis. English Language Department, University of Kashan.
- Nation, I. S. P. (1990). *Teaching and learning vocabulary*. Heinle & Heinle.
- Read, J. (2000). *Assessing vocabulary*. Cambridge University Press.
- Read, J. (2007). Second language vocabulary assessment: Current practices and new directions. *International Journal of English Studies*, 7(2), 105-125.

- Roohi, E., Azabdaftari, B., & Ashayeri, H. (2016). Exploring L1 and L2 lexical richness and speech fluency in aphasic Azari and Persian bilinguals. *Language Related Research*, 7(5), 371-389. [In Persian]
- Schmitt, N. (2000). *Vocabulary in language teaching*. Cambridge University Press
- Singh, S. (2001). A pilot study on gender differences in conversational speech on lexical richness measures. *Literary and Linguistic Computing*, 16, 251-264.
- Tabandeh Fard, N. (2016). *The effect of reading news through mobile applications on lexical richness in EFL writing*. Master's Thesis, Department of Linguistics and Foreign Languages, Payame Noor University.
- Tavakoli, M., Jalilevand, N., Kamali, M., Modarresi, Y., & Motasaddi Zarandy, M. (2015). Measuring lexical diversity and syntactic complexity after cochlear implant in 8-9 years age children's. *Journal of Paramedical Science and Rehabilitation*, 5(1), 20-29. doi: 10.22038/jpsr.2016.6382. [In Persian]
- Vermeer, A. (2000). Coming to grips with lexical richness in spontaneous speech data. *Language Testing*, 17(1), 65-83. doi:10.1177/026553220001700103
- Webb, S., & Nation, P. (2017). *How vocabulary is learned*. Oxford University Press.
- Williams, J. (2012). *An alternative approach to measuring second language productive vocabulary size: a validation study of the capture-recapture methodology*. Master's Thesis. Department of Education, Concordia University.
- Yu, G. (2009). Lexical diversity in writing and speaking task performances. *Applied Linguistics*, 31, 236-259. doi:10.1093/applin/amp024
- Zimmerman, C. B. (2014). Teaching and learning vocabulary for second language learners. In M. Celce-Murcia, et al. (Eds), *Teaching English as a second or foreign language* (4th ed) (pp. 288-302). Boston, MA: National Geographic Learning.