

A Critical Corpus-based Study of Translation Universals: A Comparison of Translational and Original Persian

Vol. 13, No. 1, Tome 67
pp. 701-745
March & April 2022

Morteza Taghavi¹  & Mohammad Reza Hashemi^{2*} 

Abstract

One of the crucial topics discussed in descriptive translation studies is that of Translation Universals (TUs), which addresses typical, salient features of translational language that make it distinguished from other linguistic variants. Taking into consideration the differences between languages, the key question here is whether the purported universal features (mainly articulated based on examining European languages) exist in non-European, less- or uninvestigated languages. Employing Chesterman's categorization of universals into 'S-universals' and 'T-universals', the present study aimed at examining the latter, less investigated group of universals. A comparable corpus was made of original and translated Persian expository texts to investigate two T-universals, namely simplification and explicitation. In the light of linguistic features of translational Persian obtained, the present study challenges the purported universals as none of the extracted features were in line with the previous studies' prepositions.

Keywords: Translation Universals, Comparable Corpus, Translational Persian, T-universals.

Received: 30 June 2020
Received in revised form: 23 August 2020
Accepted: 5 September 2020

1 Ph.D. in Translation Studies, English Department, Ferdowsi University of Mashhad, Mashhad, Iran; ORCID ID: <https://orcid.org/0000-0002-3243-2137>

2. Corresponding Author: Professor of Applied Linguistics, English Department, Ferdowsi University of Mashhad, Mashhad, Iran; Email: hashemi@um.ac.ir, ORCID ID: <https://orcid.org/0000-0001-9437-131X>

1. Introduction

One of the most significant topics within Descriptive Translation Studies (DTS) is Translation Universals (TUs), first clearly articulated by Mona Baker in her work (1993). TU hypotheses are concerned with typical linguistic features that makes translational language different from other linguistic variants. According to Hansen and Teich (2001), "it is commonly assumed in translation studies that translations are specific kinds of texts that are different not only from their original source language (SL) texts, but also from comparable original texts in the same language as the target language (TL)" (p. 44). Such claims can be examined either manually or by means of corpus-based analytical tools. Corpora have been a reliable popular tool among researchers since the convergence of corpus-based empirical methodology and linguistic studies, including Translation Studies, during the 1990s.

Over the last three decades, many studies have been conducted on theories of TUs and evidence of specific linguistic features of translational language has been provided, but almost all of these studies have been on Western languages, especially English. Chesterman (2004) divides the TUs into "S-universals" and "T-universals." The first category refers to "universal differences between translations and their source texts" (ibid, p. 39) and the second category refers to the differences in the linguistic features of translations (target texts) as compared to non-translated, native TL texts. Although some TUs may usually be investigated within one category, some can be examined from the perspective of both groups.

Some interpret the so-called TUs as an inseparable part of any translational language and are in line with the theories presented in the literature of translation studies. It should be noted, however, that if a linguistic feature is to be considered a "universal," it should be found in translations into all languages, but, in fact, almost all the literature on TUs is devoted to research on Western languages, especially English. Only a few examples can be found (e.g. Xiao & Hu, 2015) that have studied universals in non-European

languages. In addition, research on S-universals outweighs work on T-universals. Hence, such claims as the existence of 'universal' features, in the strict sense of the word 'universal', is highly debatable, unless they are scrutinized in other languages, especially those that are different from English in terms of word order, syntactic structures, stylistic features and the like. The hypothesis of the present study is that due to the differences between Persian and other Indo-European languages, including English, in those aspects (word order, syntactic structures, stylistic features, etc.), the claimed universal features shown in the other languages are not present in Persian, not at least with the same quality.

Investigation of TUs in Persian has also been largely neglected and faces some drawbacks. Some of these drawbacks are due to research methods (such as manual investigation and not benefiting from corpus investigation tools) and some are related to the limitation of data to only literary texts and novels. Focusing on comparison of source with target text(s) (addressing only S-universals using parallel corpora) and neglecting the examination of T-universals (using comparative corpora) is another limitation of such studies on Persian language. The present study intends to study the salient and distinctive linguistic features of translational Persian using corpus methodology and comparing translated texts with other comparable original writings, thus shedding light on the existence of the claimed universal features in Persian. In this regard, the T-universals of simplification and explicitation were selected and specific linguistic features were examined that can signal the presence of the selected universals. As an instance of non-literary writings, a medium-sized corpus consisting of two sets of expository academic and general humanities texts, namely philosophy and texts about literature (such as general informative texts about literature, review and critique of literary works, etc.), was analysed.

2. Background

The search for TUs dates back to the mid-nineties, where this topic led to a surge of interest among researchers, especially since the emergence of corpora as a research tool in Translation Studies. Searching through the existing literature on TUs shows that the research carried out on S-universals outweigh the studies on T-universals. Among others, simplification and explication are the two T-universals investigated in the present study.

A number of studies have been done on simplification as a universal feature in translation at the lexical, syntactic and stylistic level (*e.g.* see Laviosa-Braithwaite, 1996; Malmkjær, 1997; Laviosa, 1998; Cvrček & Chlumská, 2015). Taking a closer look at these studies highlighted some disagreements. For example, regarding mean sentence length, Laviosa (1998) (English), Xiao and Yue (2009) (Chinese), and Ilisei *et al.* (2009) (Spanish) showed that mean sentence length in translational language is significantly higher than original writings. But, contrary to these three studies, Malmkjær (1997) believed that stronger punctuations may result in shorter sentences in translated texts. Also, Xiao (2010) and Xiao and Hu (2015) found that sentences in original Chinese are relatively longer than translated texts, although this difference was not significant.

A number of other studies have further demonstrated evidence for explication or the tendency in translational language to make explicit what has been implicit in the source text, thus making it different from original writings (Blum-Kulka, 1986; Toury, 1991; Baker, 1996; Øverås, 1998; Olohan & Baker, 2000; Xiao, 2010). Although this feature is found in translations at different lexical, syntactic, and textual levels, "there is variation even in these results, which could be explained in terms of the level of language studied, or the genre of the texts" (Mauranen, 2007, p. 39). There is still a long way to go to determine whether explication is a universal feature or not, as most of the data in the literature is based on Western languages, especially English.

Much of the criticism that the topic of TUs has attracted relates to the fact

that most studies have only focused on Western languages and failed to move beyond and scrutinize others, a fact that is also reflected in the small body of literature on Persian. What we know about the possible presence of TUs in translational Persian is mostly based on studies that were carried out on S-universals and are limited in one way or another (*e.g.* Ghamkhah & Khazaei Farid, 2011; Salimi & Askarzadeh Torghabeh, 2015; Vahedi Kia & Ouliaeinia, 2016; Ahangar & Rahnemoon, 2019). In general, these limitations can be classified into seven categories:

- direction being restricted to comparison of source with target text(s) (addressing only S-universals)
- lack of variety in the source language (almost all studies feature English as the source)
- universals (all studies are on the four recurrent features of translation proposed by Baker (1996))
- genre (almost all studies focus on literary texts)
- size (very small-scale studies, mainly on selected parts of one or two books)
- methodology (using manual investigation and not benefiting from corpus investigation tools)
- source of data collection (all data were collected from books and published works, ignoring online translated materials available)

3. Corpus Design and Method

In the present study, a comparative corpus was used which includes two subcorpora: original Persian texts and English-Persian translated texts. Each component consisted of one hundred extracts, each of 3000-word length, taken randomly from books and webpages, thus amounting to 300,000 words for

each subcorpus and 600,000 words on the whole. The current literature on Persian language has failed to move beyond literary texts. Contrary to the predominance of studies on European languages and small-sized corpus-based works on Persian, all confined to literary texts and books, the data for the present research was collected from books and webpages on two non-literary fields in Humanities, philosophy and texts about literature (such as general informative texts about literature, review and critique of literary works, etc.). Finally, both sides of the corpus are comparable in terms of number of samples, size, genre and sampling period.

After collecting each sample, a header was assigned to it. For samples collected from books, this header contains information about the book title and year of publication, and for websites, it includes the title of the text, date of the post and the webpage URL. To normalize the data, we employed Virastyar, a Persian MS-Word add-in. Moreover, for segmentation, tokenization and POS tagging, we utilized tools developed by Mojgan Seraji (2015) for Persian, namely SeTPer (sentence segmenter and tokenizer) and TagPer (POS tagger). In addition, after analyzing different corpus-analyzer tools (namely WordSmith, AntConc, Sketch Engine, and LancsBox), it was found that the best and most adaptable software for analyzing Persian texts is "WordSmith".

Two universal features of simplification and explication were selected for investigation. The presence of universals was identified through a number of features. For simplification, the study used the three signs discussed in Laviosa-Braithwaite (1996) where she concluded that translational language uses lower lexical density, shows less lexical variety, and reports greater mean sentence length. For explication, the higher frequency of connectives and cohesive ties in translated than non-translated language was employed (Olohan & Baker, 2000; Chen, 2006).

4. Results

The four different lexical and syntactic features of translational Persian were

examined in the corpus under investigation, namely lexical density, lexical variety, mean sentence length, and frequency of connectives. First, regarding simplification, it was found that translational Persian in the comparative corpus used in this study has a higher lexical density, although this difference was minor and was not statistically significant. Also, the lexical variety in translational Persian was greater than non-translational texts. In addition, the study of the mean sentence length showed that sentences in original texts are slightly longer than translated texts. Comparing the two subcorpora, the texts "about literature" show higher lexical density and variety (or richness), and the sentences in philosophical texts were longer, which can be interpreted as field (also genre) variations and their idiosyncratic linguistic features. Finally, regarding explicitation, the total number of connectives was higher in the original texts than in translated texts. However, no clear overall tendency was detected in either subcorpus favoring connectives more than the other. Some connectives were more frequent in translations and some in original texts. Further, some connectives followed no trend as they were more frequent in one field but less frequent in the other.

5. Discussion & Conclusion

Moreover, the data and findings provide further support for the controversies over the strong version of TU hypotheses and raise intriguing questions regarding the presence of universal features in translations as none of the results for the four features addressed were in line with previously proposed T-universals. Therefore, the results of this study support the hypothesis that the claimed universal features, due to linguistic differences, are not present in Persian (at least to the same quality). Contrary to many previous studies (such as the detailed investigation of Ilisei *et al.* (2009)), features like lower lexical richness and density, greater mean sentence length and higher frequency of connectives might possibly not be among the most salient, *universal* (at least in its global sense) features indicative of the simplification and explicitation

hypotheses. Therefore, it can be concluded that the findings of this study indicate the specific features of translational (from an English source) and original Persian texts. In general, the present study shows that, in contrast to what might be assumed, simplification and explicitation as so-called translation universals may not be really universal as discussed by Baker (1993) and Eskola (2004), because they are not universally present in all translated texts, at least as far as this research accounts for translational Persian psychology and sociology.

Whereas a number of studies support simplification and explicitation as translation universals, these linguistic features have been challenged by some other studies, especially when language pairs and genres vary and move from the more investigated languages and genres (Western languages, literary texts) to the less investigated ones (non-Western languages, non-literary texts) (Chesterman, 2004; Mauranen, 2007; Xiao & Hu, 2015). It seems that the assumption of the presence of similar linguistic features in all translations needs to be revised. Therefore, it is better to be cautious in presenting such generalizations and to reclassify them under what Eskola (2004) labels *local translation law* rather than *universal translation law*. In fact, it should be noted that some of the theories presented have been formulated using only a pair of specific language pairs or texts, and may not apply to other languages or genres and should therefore be limited and narrowed down. As it was shown in this study, both T-universals examined here were not present in Persian with the same quality as indicated by previous research.

Since it is not possible to proceed with any claim about the presence of universal tendencies in translations without validation, further work needs to be done to establish whether TU hypotheses are supported, at least in their current account, in other, especially unexamined, languages and genres. Although the results of the present study did not support any of the hypotheses presented in the previous studies, this may not be a good reason to dismiss the universals altogether. The authors believe that, instead of abandoning the whole possibility of translations displaying common features, we may find, at

least, new tendencies that are different from those of the previous hypotheses; for example, simplification in translational language may be universally manifested through features other than lower lexical density or less lexical variety. Nevertheless, the present study indicated that the claim of the existence of "universal" features in its absolute sense (in all languages and text types) is unfounded. Much more research should be done on translational Persian and other non-European languages in order to clarify the validity and nature of TUs and the role of language, text type, translator skills and other intervening aspects involved in the manifestation of certain linguistic features in translations.



دوماهنامه بین‌المللی

۱۳، ش ۱ (پیاپی ۶۷) فروردین و اردیبهشت ۱۴۰۱، صص ۷۴۵-۷۰۱

مقاله پژوهشی

<http://dori.net/dor/20.1001.1.23223081.1401.0.0.48.0>

جهانی‌های ترجمه در بوته نقد پیکره‌ای: مقایسه فارسی ترجمه‌ای با فارسی تألیفی

مرتضی تقوی^۱، محمدرضا هاشمی^{۲*}

۱. دانش‌آموخته دکتری مطالعات ترجمه، گروه زبان انگلیسی، دانشگاه فردوسی مشهد، مشهد، ایران.

۲. استاد گروه زبان انگلیسی، دانشگاه فردوسی مشهد، مشهد، ایران.

تاریخ پذیرش: ۱۳۹۹/۰۶/۱۵

تاریخ دریافت: ۱۳۹۹/۰۴/۱۰

چکیده

یکی از مسائل مهم در حوزه مطالعات توصیفی ترجمه مفهوم جهانی‌های ترجمه است که موضوع آن یافتن ویژگی‌های خاص و بارز زبان ترجمه‌ای به مثابه زبانی متفاوت از زبان مبدأ و مقصد است. پرسشی که در اینجا مطرح می‌شود این است که با توجه به تفاوت‌های بین زبان‌ها آیا ویژگی‌های جهانی ادعایی (که عمدتاً حاصل بررسی بر روی زبان‌های اروپایی هستند) در زبان‌های غیراروپایی نیز وجود دارند یا خیر. مطالعه حاضر با استفاده از تقسیم‌بندی چسترمن (۲۰۰۴) از جهانی‌ها به دو نوع «جهانی‌های مبدأ» و «جهانی‌های مقصد»، دسته دوم را موردکاوی قرار می‌دهد. این پژوهش، با استفاده از یک پیکره مقایسه‌ای برساخته از متون توضیحی فارسی تألیفی و ترجمه‌ای، دو جهانی «ساده‌سازی» و «تصریح» را موردبررسی قرار می‌دهد. نتایج این مطالعه، در پرتو یافته‌های مربوط به ویژگی‌های فارسی ترجمه‌ای، وجود ویژگی‌های جهانی ادعایی در فارسی ترجمه‌ای را به‌چالش می‌کشد و نشان می‌دهد که هیچ‌کدام از ویژگی‌های مستخرج از پیکره مقایسه‌ای زبان فارسی با نظریات پیشین مطرح‌شده همخوانی ندارند.

واژه‌های کلیدی: جهانی ترجمه، پیکره مقایسه‌ای، فارسی ترجمه‌ای، جهانی‌های زبان مقصد.

۱. مقدمه

مبحث «جهانی‌های ترجمه»^۱ و نظریات مرتبط با آن یکی از مهم‌ترین موضوعات در مطالعات توصیفی ترجمه است. این بحث را اولین بار به شکل مستقل مونا بیکر در مقاله مهم خود (1993) مطرح کرد. موضوع بحث، بررسی ویژگی(های) بارز و خاص زبان ترجمه به مثابۀ زبانی متفاوت و متمایز و مستقل از سایر انواع زبانی است. از نظر هانسن و تایک «تعبیری رایج در مطالعات ترجمه وجود دارد که ترجمه‌ها نوع خاصی از متون هستند که نه تنها با متون مبدأشان متفاوت هستند، بلکه با متون تألیفی مشابه خودشان در همان زبان نیز تفاوت‌هایی دارند» (2001, p.44). چنین ادعاهایی را می‌توان به روش‌های مختلف اعم از دستی و ماشینی، بررسی کرد. با گسترش روش تحقیق تجربی در مطالعات زبانی (ازجمله مطالعات ترجمه) از دهۀ ۱۹۹۰ به بعد «پیکره» به مثابۀ ابزاری ماشینی و قابل اعتماد و پیوسته مورد استفاده محققان قرار گرفته است.

در طول سه دهۀ اخیر، مطالعات زیادی بر روی نظریات جهانی‌های ترجمه انجام شده و شواهدی از ویژگی‌های زبانی خاص ترجمه‌ها ارائه شده است، ولی تقریباً همه مطالعات مزبور بر روی زبان‌های غربی، به خصوص انگلیسی بوده است. چسترمن (2004) جهانی‌های ترجمه را به دو دسته «جهانی‌های مبدأ»^۲ و «جهانی‌های مقصد»^۳ تقسیم می‌کند. دسته اول به «تفاوت‌های جهان‌شمول بین ترجمه‌ها و متون مبدأشان» برمی‌گردد (ibid: 9) و دسته دوم با تفاوت ویژگی‌های زبانی متون ترجمه‌شده و متون تألیفی در یک موضوع در زبان مقصد اشاره می‌کند. منظور چسترمن از جهانی‌های مبدأ این است که متن مبدأ بالقوه کوتاه‌تر از متن ترجمه‌شده است؛ به عبارت دیگر، گرایش غالب در ترجمه‌ها این است که بلندتر از متن مبدأشان باشند. مترجمان با استفاده از دو راهکار، پاک‌سازی^۴ و تصریح^۵ متنی بلندتر از متن مبدأ می‌آفرینند. اما منظور چسترمن از جهانی‌های مقصد این است که متون مقصد بالقوه ویژگی‌های زبانی مشترکی دارند که از آن‌ها با تعبیری مانند ساده‌سازی^۶، متعارف‌سازی^۷ یا عادی‌سازی^۸ و الگوی واژگانی غیرمعمول یاد می‌شود. گرچه برخی از جهانی‌ها ممکن است در یک گروه قرار گیرند، اما برخی را می‌توان از منظر هر دو گروه نیز بررسی کرد. برای بررسی گروه اول

جهانی‌های مبدأ) نیازمند یک پیکره موازی^۹ از متون مبدأ و مقصد هستیم و برای بررسی جهانی‌های مقصد باید یک پیکره مقایسه‌ای^{۱۰} از متون ترجمه‌ای و متون تألیفی مشابه در همان زبان داشته باشیم.

برخی وجود این گونه ویژگی‌ها در زبان ترجمه‌ای را بدیهی می‌دانند و نظریات ارائه‌شده در این باره در پیشینه مطالعات ترجمه را پذیرفته‌اند. اما باید توجه داشت که اگر یک ویژگی زبانی «جهانی ترجمه» به‌شمار می‌آید، باید در نمونه‌های ترجمه به تمام زبان‌ها قابل‌مشاهده باشد، اما درواقع تقریباً تمام پیشینه موجود درباره نظریات جهانی‌های ترجمه به پژوهش‌ها بر روی زبان‌های غربی، به‌ویژه انگلیسی، اختصاص دارد و مطالعه زبان ترجمه‌ای در سایر جفت‌زبان‌ها کم‌تر بررسی شده است. تنها می‌توان نمونه‌های اندکی (Xiao & Hu, 2015) را یافت که به پژوهش بر روی زبان‌های غیراروپایی پرداخته‌اند. به‌علاوه، تعداد مطالعات بر روی جهانی‌های مبدأ بیشتر از جهانی‌های مقصد است. حال این پرسش مطرح می‌شود که آیا ویژگی‌های ادعایی که عمدتاً برخاسته از بررسی موارد در برخی از زبان‌های هندواروپایی و به‌ویژه انگلیسی هستند، با توجه به تفاوت‌های واژگانی، نحوی و سبکی بین زبان‌های زیرمجموعه زبان هندواروپایی، در زبان فارسی نیز یافت می‌شوند؟ ادعاها درمورد «جهانی‌های» ترجمه محل تأمل است و تنها هنگامی می‌توان به صحت وجود آن‌ها پی برد که حضورشان در زبان‌های دیگر (خصوصاً آن‌ها که ذاتاً با انگلیسی تفاوت دارند) نیز بررسی شوند. فرضیه پژوهش حاضر این است که با توجه به تفاوت‌های زبان فارسی با دیگر زبان‌های هندواروپایی و ازجمله با زبان انگلیسی، از جنبه‌هایی چون ترتیب چینش کلمات، ساختارهای نحوی، ویژگی‌های سبکی و ...، ویژگی‌های همگانی ادعا شده در این زبان (حداقل به همان کیفیت گفته شده) وجود ندارند.

بررسی جهانی‌های ترجمه در زبان فارسی نیز تا حدود زیادی مغفول مانده و با ضعف‌هایی مواجه است. بخشی از این ضعف‌ها به روش تحقیق (مانند بررسی‌های دستی و عدم بهره‌مندی از ابزارهای تحلیل پیکره) و بخشی نیز به محدودیت داده‌ها به متون ادبی و رمان‌ها برمی‌گردد. توجه بیش از حد پژوهش‌ها به جهانی‌های مبدأ (با استفاده از پیکره‌های موازی) و غفلت از بررسی جهانی‌های مقصد (با استفاده از پیکره‌های مقایسه‌ای) از طریق مقایسه متون ترجمه‌ای و تألیفی فارسی یکی از نکات ضعف در این گونه مطالعات در زبان فارسی است. پژوهش

حاضر قصد دارد تا با استفاده از روش‌شناسی پیکره‌ای و با مقایسه متون ترجمه‌شده با سایر متون مشابه تألیفی در زبان فارسی به بررسی ویژگی‌های زبانی برجسته و متمایزکننده فارسی ترجمه‌ای بپردازد و از این رهگذر پاسخی برای بود یا نبود ویژگی‌های جهانی مقصد در زبان فارسی ارائه دهد. در این راستا، دو جهانی ساده‌سازی و تصریح انتخاب و برای هرکدام نشانه‌های زبانی مشخصی تعریف شد. داده‌ها برای توسعه پیکره مقایسه‌ای و به‌مثابه نمونه‌ای از متون غیرادبی، از ژانر متون توضیحی و متون دو حوزه مختلف یعنی متون فلسفی و متونی درباره ادبیات (مانند معرفی ادبیات، بررسی و نقد آثار ادبی و ...) گردآوری شد.

۲. پیشینه و چارچوب نظری

۲-۱. جهانی‌های ترجمه: در جست‌وجوی جهانی‌های زبان مقصد

با تغییر نگاه از رویکردهای تجویزی به توصیفی و پررنگ شدن مطالعات ترجمه توصیفی در دهه ۱۹۸۰، توجه به زبان مقصد بیش‌ازپیش اهمیت یافت. همچنین، استفاده از زبان‌شناسی پیکره‌ای در مطالعات ترجمه در اوایل دهه ۱۹۹۰، مطالعات توصیفی «محصول‌محور»^{۱۱} (Toury, 1995) را به پیش راند و روش‌شناسی نظام‌مند و داده‌های تجربی قابل‌اعتمادی برای توصیف ماهیت ترجمه‌ها به‌دست آمد.

در مقایسه ترجمه‌ها با متون اصلی و همچنین سایر متون تألیفی در جامعه مقصد، برخی ویژگی‌ها تنها در متون ترجمه به چشم می‌خورند و به متون مزبور شخصیتی مستقل و متفاوت می‌بخشند و سبب تمایز آن‌ها با سایر انواع متون می‌شوند. ویلیام فراولی (1984, p. 168) برای نخستین‌بار گفت تقابل زبان مبدأ و مقصد در فرایند ترجمه به پدید آمدن «زبان سوم»^{۱۲} منجر می‌شود، یعنی زبانی که طی ترجمه شکل می‌گیرد زبانی منحصربه‌فرد است. در عین حال بیکر (1993) بود که پیشنهاد کرد «ایده جهانی‌های ترجمه زبانی در کانون توجهات مطالعات ترجمه قرار گیرد» (Mauranen & Kujamäki, 2004, p. 1).

تلاش‌هایی نیز برای دسته‌بندی ویژگی‌های مشترک ترجمه‌ها انجام شده که چسترمن یکی از برجسته‌ترین افراد است. او در بحث از مسیر تاریخی حرکت به ورای جزءنگری و تعمیم‌سازی درباره ترجمه‌ها، سه مسیر در پیشینه ترجمه شامل (۱) مسیر سنتی تجویزی، (۲)

مسیر سنتی انتقادی و تقلیل‌دهندگی، و ۳) مسیر یافتن جهانی‌ها با استفاده از پیکره‌ها را ذکر می‌کند (2004, p.33). وی در مسیر سوم یعنی تعمیم‌سازی به یافتن ویژگی‌های خاص ترجمه‌ها می‌پردازد و جهانی‌های ترجمه را به دو دسته کلی تقسیم می‌کند:

۱) جهانی‌های مبدأ: تفاوت‌هایی که از مقایسه همه ترجمه‌ها با متون مبدأ به دست می‌آیند؛ مانند طولانی‌شدن^{۱۳}، پاک‌سازی (مانند استفاده از همایندهای رایج‌تر)، تصریح و دو قانون پیشنهادی توری یعنی استانداردسازی^{۱۴} و مداخله^{۱۵}.

۲) جهانی‌های مقصد: تفاوت‌هایی که از مقایسه ترجمه‌ها با سایر متون غیرترجمه‌ای (تألیفی) مشابه در زبان مقصد به دست می‌آیند و در همه ترجمه‌ها دیده می‌شوند؛ مانند ساده‌سازی (شامل استفاده از واژگان با تنوع (تنوع دایره واژگانی) کمتر، به‌کارگیری نسبت اقلام واژگان محتوایی^{۱۶} به کارکردی^{۱۷} (تراکم واژگانی) پایین‌تر و استفاده بیشتر از واژگان با بسامد بالا (Laviosa-Braithwaite, 1996))، متعارف‌سازی، الگوی واژگانی غیرمعمول و کم‌رنگ کردن ویژگی‌ها و موارد خاص زبان مقصد.

بررسی دقیق و گسترده جهانی‌ها مستلزم استفاده از رویکرد پیکره‌بنیاد است. در این راستا، برای بررسی جهانی‌های مبدأ، باید «پیکره موازی» متشکل از متون مبدأ و مقصد تشکیل داد و برای بررسی جهانی‌های مقصد، «پیکره مقایسه‌ای» متشکل از متون ترجمه‌ای و تألیفی تدوین کرد. اگرچه چسترمن برخی از جهانی‌ها را به شکل طبیعی و براساس پژوهش‌های انجام‌شده در یک گروه قرار می‌دهد، اما به دلیل هم‌پوشانی‌ها، برخی از آن‌ها در قالب گروه دیگر نیز بررسی شده‌اند، مانند تصریح (Xiao, 2010, p. 9) و ساده‌سازی. در پواقع، اینکه این نظریات از چه منظری بررسی شود به هدف و داده‌های پژوهش وابسته است.

گرچه تعداد مطالعات بر روی جهانی‌های مبدأ بیشتر از مقصد است (مانند Blum-kulka, Molés- Xiao & Hu, 2015; Mauraenen, 2007; Olohan & Baker, 2000; 1986 Cases, 2019)، اما پژوهشگرانی نیز بوده‌اند که با مقایسه متون ترجمه‌شده و تألیفی، ویژگی‌های متمایز زبان ترجمه‌ای را با روش‌های گوناگون استخراج کرده‌اند. دو ویژگی جهان‌شمول که در پژوهش حاضر بر آن‌ها تمرکز شده عبارت‌اند از ساده‌سازی و تصریح. ساده‌سازی به «تمایل برای ساده کردن زبان در ترجمه‌ها» اشاره دارد (Baker, 1996, pp.)

181-182) بدین معنا که زبان استفاده‌شده در ترجمه‌ها از نظر واژگانی، نحوی و حتی سبکی از زبان نوشته‌های تألیفی ساده‌تر است (Blum-kulka & Levenston, 1983).

تحقیقات متعددی درمورد ساده‌سازی به‌مثابه یک ویژگی جهان‌شمول در ترجمه در سطح واژگانی، نحوی و سبکی انجام شده است. لایوزا - برایثویت (1996) ساده‌سازی واژگانی را با استفاده از یک پیکره مقایسه‌ای متون تألیفی و ترجمه‌ای زبان انگلیسی بررسی کرد و چهار الگوی اصلی کاربرد واژگانی در ترجمه‌ها را نشان داد که می‌توانند شواهدی بر وجود ساده‌سازی در متون ترجمه‌ای باشند: ۱) پایین بودن تقریبی نسبت واژگان محتوایی به واژگان کارکردی در متون ترجمه‌ای نسبت به متون تألیفی، ۲) میزان بالای استفاده از واژگان پربسامد نسبت به کم‌بسامد، ۳) تکرار بیشتر پربسامدترین واژگان، و ۴) تنوع کمتر در واژگان پربسامد. همچنین، لایوزا (1998) در تحقیقی نشان داده که تراکم واژگانی در متون ترجمه‌ای نثر روایی به شکل معناداری پایین‌تر است. علاوه‌براین، میزان واژه‌های پربسامد در متون ترجمه‌ای به شکلی معنادار بیشتر از بخش مشابه در متون تألیفی پیکره است؛ و این به معنای تکرار بیشتر، تنوع کمتر و سادگی بیشتر است. سیوروچک و کلامسکا (2015) نیز پژوهشی بر روی تنوع دایره واژگانی^{۱۸} کمتر به‌مثابه تبلور ساده‌سازی انجام دادند. مطالعه مزبور که از پیکره مقایسه‌ای متشکل از متون ترجمه‌ای و غیرترجمه‌ای ادبیات داستانی و ادبیات حرفه‌ای به زبان چک استفاده می‌کرد، درنهایت، به همان نتیجه لایوزا - برایثویت (1996) دست یافت که در متون ترجمه‌شده به زبان چک در مقایسه با متون غیرترجمه‌ای در همان زبان از واژگان متفاوت نسبتاً کم‌تری استفاده می‌شود.

علاوه‌بر مطالعات درمورد ساده‌سازی واژگانی در ترجمه، ساده‌سازی نحوی و سبکی در ترجمه نیز موردتوجه محققان بوده است. اولوهان و بیکر (2000) با مطالعه روی ساده‌سازی نحوی در انگلیسی ترجمه‌ای نشان دادند که بسامد حرف ربط *that* در متون ترجمه‌ای بیشتر است و این حرف در متون ترجمه‌ای نسبت به متون تألیفی با ساختارهای نحوی ساده‌تری استفاده شده است. درمورد ساده‌سازی سبکی، ملمکیار (1997) دریافت که علائم سجاوندی در ترجمه‌ها شکلی قوی‌تری به خود می‌گیرند؛ مثلاً ویرگول معمولاً با نقطه‌ویرگول یا نقطه، و نقطه‌ویرگول معمولاً با نقطه جایگزین می‌شوند. درنتیجه جملات بلند و پیچیده در متون مبدأ به

جملات کوتاه‌تر و با پیچیدگی کمتر تبدیل می‌شوند که به کاهش پیچیدگی ساختار و سهولت در خوانش نیز منجر می‌شود.

دقت در یافته‌های مطالعاتی که ذکر شد تضاد و ناهمخوانی‌هایی را نشان می‌دهد. مثلاً درخصوص معیار طول جمله، لایوزا (۱۹۹۸) (زبان انگلیسی)، ژیانو و یو (۲۰۰۹) (زبان چینی) و ایلوسی و همکاران (۲۰۰۹) (زبان اسپانیایی) نشان داده‌اند که میانگین طول جمله در زبان ترجمه‌ای به شکل معناداری بیشتر از متون تألیفی است. اما برخلاف این سه مطالعه، ملمکیار (۱۹۹۷) معتقد بود استفاده از صورت‌های قوی‌تر علائم سجاوندی در ترجمه‌ها احتمالاً به جملات کوتاه‌تر منجر می‌شود. همچنین، ژیانو (۲۰۱۰) و ژیانو و هو (۲۰۱۵) نشان دادند که در زبان چینی جملات در متون تألیفی اندکی طولانی‌تر هستند، گرچه این تفاوت معنادار نبود. کورپاس - پاستور و همکاران (۲۰۰۸) نیز با اتخاذ یک رویکرد خاص در پردازش زبان طبیعی^{۱۹} نتیجه گرفتند که پیکره مقایسه‌ای زبان اسپانیایی از نظر پایین‌تر بودن تراکم واژگانی در متون ترجمه ساده‌سازی را تأیید می‌کرد، اما معیارهای طول جمله و استفاده جملات ساده در مقابل پیچیده به نحوی نبود که حضور این ویژگی جهانی را نشان دهد. لذا، ساده‌سازی به مثابه یک ویژگی جهان‌شمول ترجمه، در همه زبان‌ها امری قطعی و پذیرفته‌شده نیست، چرا که برخی آن را تأیید و برخی دیگر، مانند مائورانن (۲۰۰۰)، ژانتونین (۲۰۰۱) با آن مخالفت کرده‌اند.

تصریح یا آشکارسازی اطلاعات نهفته و قابل استنباط از متن مبدأ و ذکر آن در متن مقصد یکی دیگر از جهانی‌های مقصد است. آنچه را امروزه به عنوان تصریح می‌شناسیم، اولین بار بلوم - کولکا مطرح کرد. وی با بررسی متون منفرد دریافت که مترجمان تمایل دارند حروف ربط را که در متن مبدأ به شکل پنهان و بدون نمود ظاهری استفاده شده‌اند، به صورت آشکار و صریح استفاده کنند. بلوم - کولکا تأکید می‌کند که این راهکار جهانی «به‌طور ذاتی در دل فرایند میانجیگری زبانی قرار دارد» (۱۹۸۶، p.21).

تحقیقات متعددی درمورد تصریح در ترجمه انجام شده است. توری (۱۹۹۱، cited in Baker, 1993, p. 244) تصریح را ویژگی مشترک تمام رخدادهای زبانی میانجی (ازجمله تعامل به زبان خارجی) می‌شمارد، اما سؤالاتی را نیز مطرح می‌کند ازجمله اینکه آیا در سطح و شیوه تصریح مثلاً زبان‌آموزان و مترجمان، یا مترجمان حرفه‌ای و غیرحرفه‌ای یا حتی ترجمه

کتبی و شفاهی تفاوتی وجود دارد یا خیر. بیکر (1996) پیشنهادهایی درمورد نموده‌های احتمالی تصریح در ترجمه‌ها و نحوه بررسی آن‌ها ارائه می‌دهد و عنوان می‌کند که استفاده بیشتر حروف ربط در متون ترجمه نسبت به متون تألیفی از جمله موارد مشترک در یافته‌های تحقیقاتی درباره تصریح است. اولوهان و بیکر (2000) نیز با مقایسه نتایج بررسی دو پیکره انگلیسی ترجمه‌ای^{۲۰} و پیکره ملی بریتانیا^{۲۱} (متون تألیفی) نشان دادند که بسامد حرف ربط *that* و افعال خبری *say* و *tell* در متون انگلیسی ترجمه‌ای بسیار بیشتر است و ربط اجزای جمله بدون استفاده از حروف ربط آشکار (پیوندهای صفر^{۲۲})، در متون تألیفی بیشتر از متون ترجمه‌ای است. این نتایج شواهدی قوی از تصریح نحوی در انگلیسی ترجمه‌ای است که برخلاف «اضافه کردن ارادی توضیحات اضافی برای پر کردن خلأ اطلاعاتی بین خوانندگان متن مبدأ و مقصد، پدیده‌ای ناخودآگاه، غیرمستقیم و نهفته در ذات فرایند ترجمه است» (Laviosa, 2002, p. 68).

پیم (2005) با بررسی مفصل تاریخچه و انواع تصریح، این نظریه را از زاویه چارچوب مدیریت ریسک بررسی می‌کند و می‌گوید هرگاه ریسک ترجمه (به دلیل فاکتورهایی مثل مشتری، هنجارها و اخلاقیات) بالا باشد، مترجم احتمالاً از هر فرصتی برای کاهش ریسک استفاده می‌کند (گرچه اجباری در کار نیست) و تصریح یکی از راه‌های اجتناب از ریسک است. ژیانو و یو (2009)، ژیانو (2010) و ژیانو و هو (2015) نیز نشان داده‌اند که بسامد حروف ربط در زبان چینی ترجمه‌ای بیشتر از زبان تألیفی است.

اگرچه این ویژگی در سطوح مختلف واژگانی، نحوی و متنی در ترجمه‌ها یافت شده است، اما «حتی در این نتایج نیز تفاوت‌هایی دیده می‌شود که می‌توان علت آن را به زبان مورد مطالعه یا ژانر متن ربط داد» (Mauranen, 2007, p. 39). برای اینکه دریابیم تصریح یک جهانی ترجمه است یا خیر هنوز راه زیادی در پیش است، چراکه بیشتر داده‌های موجود در پیشینه مطالعات براساس زبان‌های غربی و به‌خصوص زبان انگلیسی است.

۲.۲. نقد جهانی‌های ترجمه

علاوه بر نقدهایی که درمورد دو جهانی مقصد مدنظر این پژوهش مطرح شد، در طول دو دهه گذشته، برخی صاحب‌نظران هم ایراداتی را بر بیان ادعاهای کلی درباره ترجمه‌ها وارد

دانسته‌اند. سه مورد از این ایرادات به ترتیب زیر است:

۱) زبان‌ها و فرایندهای زبانی متنوع هستند و امکان تعمیم به همه زبان‌ها وجود ندارد (Tymoczko, 1998).

۲) مشخص نیست که آیا ویژگی‌های موردنظر حاصل فرایند ناخودآگاه در ذهن مترجم و در دل فرایند ترجمه هستند (و در این صورت می‌توانند ویژگی ذاتی ترجمه‌ها خوانده شوند) یا ناشی از تصمیم آگاهانه مترجم براساس شرایط اجتماعی - ذهنی - فیزیکی (Malmkjær, 2007).

۳) ترجمه‌ها به شدت وابسته به زبان‌هایی هستند که از/ به آن‌ها ترجمه می‌شوند و حتی اگر چند جفت زبان را بررسی کنیم و الگوهایی را بیابیم، این الگوها در نهایت الگوهای همان چند جفت زبان هستند و نه بیشتر؛ تنها زمانی می‌توان از ویژگی‌هایی مانند تصریح یا ساده‌سازی صحبت کرد که به شکل دقیق تعریف شوند و شیوه بررسی آن‌ها مشخص شود (House, 2008).

پژوهش‌های بعدی بر روی خصوصیات شناختی و روانی موجود در جهانی‌ها (مانند Zaslavskiy, 2019) به دغدغه ملمس‌تر (2007) پاسخ داد و مشخص شد که برخی از نکات به صورت ناخودآگاه در عملکرد مترجمان به چشم می‌خورد. همچنین، در پاسخ به سایر نقدهای منفی درمورد جهانی‌ها، می‌توان برای بررسی جهانی‌های ترجمه برخی علل و مزایا را ذکر کرد. اولین مزیت مربوط به بحث روش تحقیق است. پژوهش بر روی نظریات جهانی‌های ترجمه سبب شده تا مطالعات ترجمه از نظر پژوهش تجربی پیشرفت کند، ادعاها دقیق‌تر و ارائه نظریات و بررسی آن‌ها عمیق‌تر شود و همچنین پژوهشگران از پیکره‌های بزرگ و مزیت‌های آن‌ها بهره‌مند شوند. چسترمن (2010, p.45) نیز معتقد است این حوزه پژوهشی ما را ترغیب به پرسیدن سؤال‌های «چرایی» کرده و با حرکت از توصیف به سمت توضیح به تقویت جنبه تجربی مطالعات ترجمه کمک کرده است. در همین راستا، توری (2004) نیز تأکید دارد که بررسی ویژگی‌های مشترک ترجمه‌ها به دلیل توان توضیحی این‌گونه پژوهش‌ها درباره ماهیت ترجمه‌ها، ارزشمند است.

اهمیت تمایز بین جهانی‌ها با قوانین^{۲۳} و هنجارها نکته مهم دیگری است که باید به آن توجه

کرد. بیکر با پذیرش این تمایز معتقد است ویژگی‌های جهانی ترجمه‌ها محصول «محدودیت‌ها و سازوکارهایی هستند که در ذات فرایند ترجمه قرار دارند و همین مسئله جهانی بودنشان را نشان می‌دهد»، اما هنجارها «ویژگی‌های ترجمه‌ای هستند که پیوسته در نوع خاصی از ترجمه‌ها و بافت فرهنگی - اجتماعی و تاریخی مشاهده می‌شوند» (1993, p.242). همان‌طور که اسکولا (2004) درباره تمایز این دو اشاره می‌کند، هنجارها عمده‌تاً ماهیتی تجویزی دارند و فرهنگ‌بند هستند، درحالی که جهانی‌ها یا قوانین، توصیفی و پیش‌بینی‌کننده هستند؛ لذا نباید برای ارجاع به ویژگی‌های رایج در ترجمه‌ها آن‌ها را به جای یکدیگر به کار ببریم. در همین راستاست که اسکولا (*ibid*)، بین «قوانین بومی ترجمه» و «قوانین جهانی ترجمه» تمایز قائل می‌شود. اینکه آیا نظریات مطرح‌شده «زبان‌بند» یا «جهانی» هستند را می‌توان با تحقیق روی زبان‌هایی که کم‌تر موردبررسی قرار گرفته‌اند، مشخص کرد؛ هدفی که موردنظر این تحقیق نیز هست.

نکته دیگر در بررسی پیشینه جهانی‌های ترجمه این است که بخش عظیمی از این پیشینه دچار نوعی اروپامحوری (غرب‌محوری) است و آنچه از ویژگی‌های جهانی می‌دانیم هم عمدتاً مستخرج از مطالعات روی زبان‌های اروپایی یا کشورهای غربی، به‌ویژه انگلیسی، است، درحالی که زبان‌های دیگر به جز تعداد معدودی (مثلاً Xiao & Hu, 2015 بر روی زبان چینی) عمدتاً مورد مطالعه قرار نگرفته‌اند. اگر قرار است ویژگی‌های پیشنهادی برای متون ترجمه‌ای را واقعاً ویژگی‌هایی «جهانی» بشماریم، باید شواهد زبان‌های غیراروپایی را نیز ارائه کرد، به‌خصوص زبان‌هایی که ساختاری متفاوت دارند تا مشخص شود که آیا جهانی‌های ترجمه در ترجمه به زبان‌های غیراروپایی نیز به همان شکل و به همان کیفیت مشاهده می‌شوند یا خیر.

۲-۳. جهانی‌های مقصد در فارسی ترجمه‌ای

آنچه درباره وجود احتمالی جهانی‌های ترجمه در فارسی ترجمه‌ای می‌دانیم غالباً به پژوهش‌هایی برمی‌گردد که بر روی جهانی‌های مبدأ انجام شده است و هرکدام دچار محدودیت‌ها و ضعف‌هایی هستند (مانند غمخواه و خزاعی‌فرید، ۱۳۹۰؛ Salimi & Ahangar & Vahedi Kia & Ouliaeinia, 2016؛ Askarzadeh Torghabeh, 2015؛ Rahnemoon, 2019). به‌طور کلی، این محدودیت‌ها را می‌توان به هفت دسته طبقه‌بندی کرد:

- ۱) نوع جهانی‌ها: محدود بودن به مقایسه متن/متون مبدأ با مقصد (تمرکز بر جهانی‌های مبدأ)؛
- ۲) نبود تنوع در زبان مبدأ ترجمه‌ها (تقریباً در تمام مطالعات انگلیسی به‌عنوان زبان مبدأ است)؛
- ۳) نبود تنوع در نظریات موردبررسی (تنها نظریات چهارگانه بیکر بررسی شده‌اند)؛
- ۴) ژانر (محدودیت تقریباً تمام مطالعات به ژانر ادبی)؛
- ۵) حجم پیکره (مطالعات با حجم داده کم، معمولاً بر روی بخش‌های انتخابی از یک یا دو کتاب)؛
- ۶) روش‌شناسی (عمدتاً بررسی دستی و بهره نبردن از ابزارهای تحلیل پیکره)؛
- ۷) منبع جمع‌آوری داده (همه مطالعات بر روی کتاب‌ها و آثار چاپی و غفلت از ترجمه‌های موجود در سطح اینترنت و مطالب ترجمه‌شده برخط).

در بررسی پیشینه، به جز دو پژوهش (Alibabae & Salehi, Esteki et al., 2010; 2012)، هیچ مطالعه دیگری یافت نشد که به بررسی جهانی‌های مقصد با استفاده از پیکره مقایسه‌ای متون ترجمه‌ای و تألیفی فارسی پرداخته باشد. استکی و همکاران (2010) به‌منظور صحت‌سنجی جهانی ساده‌سازی پیکره‌ای مقایسه‌ای از متون اقتصادی تدوین کردند و سه ویژگی و مشخصه ساده‌سازی را موردبررسی قرار دادند و به این نتیجه رسیدند که جملات متون ترجمه‌ای طولانی‌تر از جملات متون تألیفی‌اند، اما برخلاف انتظار، نسبت استفاده از واژه‌های متفاوت به کل واژگان (تنوع دایره واژگانی) و تراکم واژگانی نیز در متون تألیفی کمتر از متون ترجمه‌ای بود. در مطالعه‌ای دیگر، علی‌بابایی و صالحی (2012) متوجه شدند که دو معیار تراکم واژگانی و تنوع دایره واژگانی، در متون ترجمه‌شده سیاسی بیشتر از متون تألیفی سیاسی است و طول جملات در زبان ترجمه‌ای کمتر از زبان تألیفی است. به عبارت دیگر، هیچ‌کدام از مؤلفه‌هایی که نشانه ساده‌سازی در زبان ترجمه باشد در این مطالعه وجود نداشت. در کل مطالعات انجام‌شده درباره فارسی ترجمه‌ای در هر دو گروه از جهانی‌های مبدأ و مقصد، تنها داده‌های سه پژوهش از متون غیرادبی استخراج شده است: مطالعه استکی و همکاران (2010) بر روی متون اقتصادی، علی‌بابایی و صالحی (2012) بر روی متون سیاسی و آهنگر و رهنمون (2019) بر روی متون پزشکی. این مسئله حکایت از غفلت شدید نسبت به متون غیرادبی در پژوهش‌های جهانی‌های ترجمه در زبان فارسی و فقر پژوهشی در این زمینه دارد، درحالی که جهانی‌های ترجمه مربوط به ترجمه همه ژانرها هستند، نه فقط متون ادبی.

۳. پیکره و روش تحقیق

روش تحقیق پژوهش حاضر دارای سه بخش است: طراحی پیکره، نشانه‌گذاری و ابزارهای مورد استفاده برای ساخت پیکره، و جهانی‌های منتخب.

۳-۱. طراحی پیکره

در پژوهش حاضر از پیکره مقایسه‌ای، شامل دو زیرپیکره متون ترجمه‌ای و تألیفی استفاده شد. عموماً متون دو سمت پیکره باید از جنبه‌های مختلف (موضوع، حجم، تاریخ تولید، شیوه جمع‌آوری، نحوه پردازش و ...) مشابه یکدیگر باشند تا بتوان نتایج را مقایسه کرد. پیکره مقایسه‌ای پژوهش حاضر شامل دو زیرپیکره است: متون فارسی تألیفی و متون ترجمه‌شده از انگلیسی. هر سمت پیکره شامل ۱۰۰ نمونه سه هزار کلمه‌ای است که از کتاب‌ها و صفحه‌های مختلف اینترنتی گردآوری شده‌اند. بنابراین، هر زیرپیکره شامل ۳۰۰ هزار کلمه و حجم کل پیکره مقایسه‌ای مشتمل بر ۶۰۰ هزار کلمه است. برای اجتناب از گردآوری سویه‌دار نمونه‌ها از یک منبع به صورت نامتعادل و از یک بخش خاص، هر متن به شکل تخمینی به سه قسمت فرضی تقسیم شد (اول، وسط و آخر) و از هر قسمت به شکل تصادفی حدود ۱۰۰۰ کلمه استخراج شد (با رعایت این قاعده که آخرین جمله همیشه باید کامل باشد). قطعاً برخی نمونه‌ها کمی بیشتر یا کم‌تر از سه هزار کلمه هستند. به علاوه، هر جا حجم متنی کم بود و امکان نمونه‌گیری تصادفی وجود نداشت (به خصوص در صفحات اینترنتی)، پس از ترکیب چند متن، نمونه‌برداری انجام شد. همچنین، تمام تیتراهای متون و عنوان‌های بخش‌های مختلف حذف شد، چراکه عنوان‌ها ویژگی‌های متنی خاص خود را دارند و نباید با آن‌ها مانند متن عادی رفتار کرد.

بیشتر مطالعات پیکره‌ای و غیرپیکره‌ای قبلی در زبان فارسی بر روی متون ادبی انجام شده است (نک. بخش ۲.۳)، داده‌های مورد نیاز برای پژوهش حاضر، به مثابه نمونه‌ای از متون غیرادبی، از کتاب‌ها و صفحات اینترنتی مربوط به دو حوزه در علوم انسانی گردآوری شد: متون فلسفی و متونی که درباره ادبیات نوشته شده‌اند (مانند پیشینه و معرفی آثار ادبی، نقد ادبی، مکاتب ادبی و ...). در باب چرایی انتخاب متون ترجمه‌شده از انگلیسی، براساس آمار سایت «خانه کتاب» (مستخرج ۱۳۹۸/۱۰/۲۷)، در طول دهه گذشته ۲۲/۴ درصد از آثار چاپ‌شده در ایران ترجمه بوده‌اند و اکثر آن‌ها از زبان انگلیسی به فارسی ترجمه شده‌اند. در

سایت‌های فارسی زبان نیز بیشتر مطالب از زبان انگلیسی ترجمه می‌شوند.

هدف اولیه این پژوهش، جمع‌آوری داده از متون تولیدشده در دهه اخیر بود، اما در عمل با کمبود کتاب ترجمه‌شده در این دو حوزه مواجه شدیم، خصوصاً نسخه‌هایی که به شکل دیجیتالی در دسترس باشند. نشر آثار ترجمه‌ای در این دو حوزه نسبتاً کم است و بسیاری از آن‌ها نیز به شکل دیجیتالی در دسترس نیستند. اگرچه ممکن بود از یک موضوع خاص میزان قابل‌توجهی از داده داشت (مثلاً تاریخ ادبیات یا آثاری از دو سه فیلسوف نامدار) و بتوان همان بازه زمانی یک دهه اخیر را انتخاب کرد، اما در این صورت تنوع در داده‌ها به شدت پایین می‌آمد و امکان استفاده از طیف متنوعی از موضوعات در پیکره فراهم نمی‌شد و پیکره، معرف واقعی متون نمی‌بود. لذا، بازه انتخاب داده‌ها به ۲۵ سال اخیر افزایش پیدا کرد تا بتوان متون دیجیتالی بیشتری را در اختیار داشت. درنهایت، هر دو سوی پیکره مقایسه‌ای که در این پژوهش مورد استفاده قرار گرفت از نظر تعداد نمونه‌ها، حجم داده، ژانر و زمان تولید داده‌ها کاملاً مشابه و قابل‌قیاس هستند. جدول ۱ حاوی اطلاعات پیکره تدوین شده است:

جدول ۱: اطلاعات پیکره مقایسه‌ای فارسی تألیفی و ترجمه‌ای

Table 1: Details of the comparable corpus

فارسی ترجمه‌ای		فارسی تألیفی		نوع متن
تعداد نمونه	تعداد کلمه	تعداد نمونه	تعداد کلمه	
۵۰	۱۵۱,۹۵۳	۵۰	۱۵۰,۸۲۰	درباره ادبیات
۵۰	۱۵۱,۴۸۹	۵۰	۱۵۰,۹۹۲	فلسفی
۱۰۰	۳۰۳,۴۴۲	۱۰۰	۳۰۱,۸۱۲	کل

در نمونه‌برداری از کتاب‌ها با سه چالش اصلی مواجه بودیم. اولاً، اکثر کتاب‌هایی که موضوعی غیر از داستان و رمان دارند، خصوصاً درمورد دو حوزه انتخاب‌شده، معمولاً به شکل دیجیتالی در دسترس نیستند. ثانیاً، نسخه‌هایی هم که به صورت دیجیتالی در دسترس‌اند به شکل فایل‌های اسکن‌شده هستند، به این معنی که نمی‌توان متن داخل آن‌ها را در محل دیگری

کپی کرد و کیفیت پایین بسیاری از آن‌ها مانع تشخیص و خوانش برای نرم‌افزار نویسه‌خوان^{۲۴} است. ثالثاً، دشواری یافتن نرم‌افزار نویسه‌خوان توانمند و دقیق. به این منظور، نویسه‌خوان‌های مختلفی به‌صورت آزمایشی موردبررسی قرار گرفتند و درنهایت دو نرم‌افزار نویسه‌خوان «متن‌یار» و ابزار نویسه‌خوان گوگل درایو برای اهداف این تحقیق مناسب تشخیص داده شدند. البته همین دو ابزار نیز محدودیت‌های خود را داشتند. مثلاً، متن‌یار درونداد متنی محدودی را قبول می‌کند یا نویسه‌خوان گوگل در هنگام دریافت متن فارسی، ساختار پاراگراف متن را رعایت نمی‌کند و بعضاً انتهای یک پاراگراف دقیقاً در شروع پاراگراف بعدی قرار می‌گیرد که نیاز به اصلاح دستی توسط کاربر دارد. علاوه بر این‌ها، در هر دو مواردی از خطا و تشخیص اشتباه واژگان و نویسه‌ها به‌چشم می‌خورد. اما با مقایسه این دو، درنهایت نویسه‌خوان گوگل درایو به‌دلیل ارائه نتایج دقیق‌تر برای پژوهش حاضر برگزیده شد. تصویرهایی که از بخش‌های مختلف هر کتاب گرفته شده بود به این ابزار نویسه‌خوان داده شد و خروجی آن‌ها در کنار یکدیگر یک نمونه سه‌هزار کلمه‌ای را تشکیل می‌داد. همچنین با وجود وقت‌گیر و مشکل بودن، نمونه‌های نهایی با منبع اصلی به دقت مقایسه و تعارض‌های احتمالی (تشخیص اشتباه واژه‌ها، نویسه‌ها، انتهای پاراگراف‌ها، علائم سجاوندی و ...) با بررسی دستی تصحیح شد.

۳-۲. نشانه‌گذاری پیکره و ابزارهای مورد استفاده

پس از جمع‌آوری هر نمونه یک سرصفحه به آن اضافه شد. برای نمونه‌های گردآوری‌شده از کتاب‌ها این سرصفحه شامل اطلاعات اسم کتاب و سال نشر آن است و برای صفحه‌های اینترنتی نیز شامل عنوان متن، تاریخ مطلب و نشانی صفحه (یو.آر.آل^{۲۵}) است. اگر در جمع‌آوری نمونه‌ها، تعداد کلمات یک صفحه اینترنتی اندک می‌بود و نمونه سه‌هزار کلمه‌ای با ترکیب چند صفحه اینترنتی آماده می‌شد، سرصفحه آن نمونه، اطلاعات همه صفحه‌ها را نشان می‌داد. برای یکسان‌سازی^{۲۶}، از افزونه ویراستیار در واژه‌پرداز ورد استفاده شد که کاربردهای مختلف آن عبارت‌اند از: غلطیابی املائی، تکمیل خودکار کلمات، اصلاح علائم نگارشی، استانداردسازی و اصلاح نویسه‌های فارسی، پیش‌پردازش املائی متن، تبدیل پی‌نگلیش، تبدیل اعداد و تبدیل تاریخ. به‌علاوه، لیستی از موارد اختلافی و متغیر از به‌کارگیری نیم‌فاصله، صورت‌های نوشتاری

واژگان، وندها و ... که احتمال می‌دادیم ویراستیار آن‌ها را پشتیبانی نکند، تهیه شد و هر دو زیرپیکره نیز با لیست مقایسه و اصلاحات موردنظر پیاده‌ش تا یکسان‌سازی متن تاحد ممکن رعایت شود و با بردن داده‌ها به سمت صورت‌های استاندارد نوشتاری، تنوع نگارشی و در نتیجه خطا در نتایج تا حد زیادی از بین بروید (برای اطلاعات بیشتر درباره چالش‌های نویسه‌ها و املاک واژگان فارسی نک. سراجی، ۲۰۱۵؛ بی‌جن‌خان و همکاران، ۲۰۱۰).

برای تقطیع جملات و واحدسازی و همچنین برچسب‌زنی نقش دستوری کلمات از دو ابزار طراحی شده توسط سراجی (۲۰۱۵) استفاده شد؛ ابزار *SeTPer* (تقطیع‌کننده جملات و واحدسازی) و ابزار *TagPer* (برچسب‌زنی نقش دستوری کلمات). ابزار *SeTPer* برای زبان فارسی بهینه شده است و قدرت تشخیص علائم سجاوندی فارسی از قبیل ویرگول و گیومه را نیز دارد (۸۸-۹۰). ابزار برچسب‌زنی نقش دستوری کلمات *TagPer* نیز بر مبنای برچسب‌زنی آماری هان‌پوز^{۳۷} برای فارسی بسط داده شده و با پیکره فارسی دانشگاه اوسلا آموزش داده شده است (این پیکره نسخه بهبودیافته پیکره بی‌جن‌خان محسوب می‌شود) (ibid: ۹۱-۹۶). پس از آموزش و ارزیابی این برچسب‌زنی توسط سازنده، دقت آن در تشخیص نقش دستوری واژگان مختلف، ۹۷/۴۶ درصد اعلام شد.

چون پژوهش حاضر بر روی زبان فارسی است، ابزار تحلیل پیکره آن هم باید فارسی را پشتیبانی کند و با آن همخوانی داشته باشد. ابزارهای مختلفی برای تحلیل پیکره و داده‌کاوی آن در دنیا تدوین شده‌اند که امکانات مختلفی را در قالب یک نرم‌افزار در اختیار پژوهشگر قرار می‌دهند. برخی از برجسته‌ترین آن‌ها عبارت‌اند از: ورداسمیث^{۳۸}، انتکانک^{۳۹}، اسکچ‌انجین^{۴۰} و لنکس‌باکس^{۴۱}. در بررسی نرم‌افزارهای مزبور مشخص شد همه‌ی آن‌ها به‌جز نرم‌افزار ورداسمیث یا با زبان فارسی سازگاری نداشتند یا امکانات محدودتری نسبت به این نرم‌افزار ارائه می‌دادند. بنابراین، بهترین و جامع‌ترین نرم‌افزار تحلیل پیکره برای زبان فارسی، نرم‌افزار «ورداسمیث» است.

۳-۳. جهانی‌های منتخب

برای پژوهش حاضر، دوجاهانی زبان مقصد یعنی ساده‌سازی و تصریح، برای بررسی انتخاب

شدند. بنا به نظر بیکر ساده‌سازی «گرایش به ساده کردن زبان ترجمه‌ها» است (1996, pp.181-182). این ویژگی ممکن است در سطوح واژگانی، نحوی و سبکی خود را نشان دهد (Laviosa-Braithwaite, 1996; Blum-kulka & Levenston, 1983). تصریح را نیز بلوم - کولکا (1986) اولین بار مطرح کرد. به عقیده وی ترجمه‌ها از متون غیرترجمه‌ای نظیر خود صریح‌تر هستند (چه در مقایسه با متن مبدأ و چه در مقایسه با متن تألیفی مشابه در زبان مقصد). اگرچه ادعای بلوم - کولکا براساس حروف ربط در ترجمه‌ها و متون مبدأشان است، اما مطالعات دیگری نیز وجود دارند (مانند Chen, 2006, cited in Xiao, 2010) که متون ترجمه‌ای و تألیفی را بررسی کرده‌اند.

در این مطالعه، برای بررسی وجود هر یک از جهانی‌ها در زبان ترجمه‌ای، ویژگی‌های زبانی خاصی در نظر گرفته شد. برای بررسی ساده‌سازی، به مطالعه لایووزا - برایثویت (1996) مراجعه کردیم و سه ویژگی این جهانی در زبان ترجمه‌ای را انتخاب کردیم. وی در رساله دکتری‌اش یک پیکره مقایسه‌ای متشکل از متون ترجمه‌ای و تألیفی انگلیسی طراحی کرده و با تحلیل آن به این نتیجه رسیده است که زبان ترجمه‌ای از تراکم واژگانی کم‌تری بهره می‌برد، تنوع دایره واژگانی کم‌تری از خود نشان می‌دهد و میانگین طول جمله بیشتری دارد. این ویژگی‌ها در مطالعات بسیاری بر روی زبان‌های غربی و حتی زبان چینی تأیید شده است. مثلاً ژیانو و یو (2009) نشان داده‌اند که متون ترجمه‌ای داستانی چینی به شکل معناداری میانگین طول جمله بیشتری نسبت به متون مشابه تألیفی در همان زبان دارند.

برای تصریح نیز معیار بیشتر بودن بسامد حروف ربط در متون ترجمه‌ای در مقایسه با متون تألیفی انتخاب شد (Chen, 2006; Olohan & Baker, 2000). در زبان فارسی دو نوع غالب از حروف ربط داریم: ساده (یک-کلمه‌ای) و مرکب (دو یا چندکلمه‌ای). برای این مطالعه پانزده مورد از پرکاربردترین حروف ربط به شرح جدول ۲ انتخاب شد. استخراج بسامد حروف ربط انتخابی در جدول ۲ به بررسی این نظریه کمک می‌کند که آیا تعداد حروف ربط در متون ترجمه‌ای بالاتر از تألیفی است یا خیر.

جدول ۲: حروف ربط منتخب

Table 2: Fifteen commonly used cohesive ties in Persian

۱۵ حرف ربط پرکاربرد در زبان فارسی				
و	که	اما	ولی	اگر
چون	زیرا	بنابراین	سپس	پس
بلکه	اگرچه	برای اینکه	لیکن	هرچند

۴. نتایج و بحث

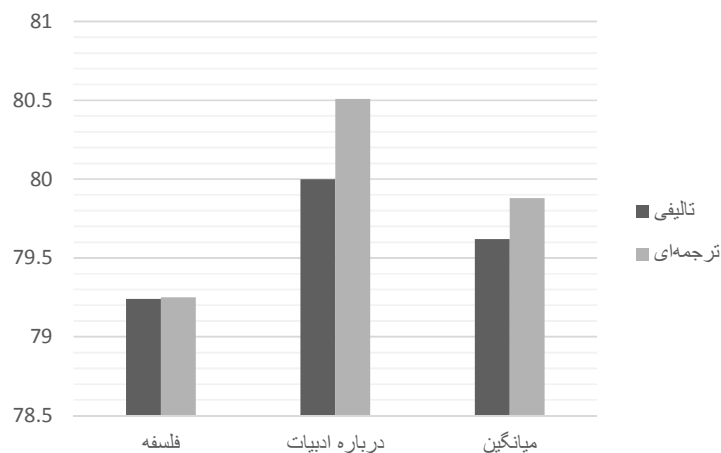
در این بخش به ارائه نتایج تحلیل چهار ویژگی واژگانی و نحوی در پیکره مقایسه‌ای فارسی تألیفی و ترجمه‌ای می‌پردازیم: تراکم واژگانی، تنوع دایره واژگانی، میانگین طول جمله و بسامد حروف ربط.

۴.۱. تراکم واژگانی و تنوع دایره واژگانی

یکی از نظریات مطرح درمورد جهانی ساده‌سازی این است که ترجمه‌ها تراکم واژگانی و تنوع دایره واژگانی کم‌تری دارند. تراکم واژگانی، با تعریف استابز (1996) عبارت است از نسبت اقلام واژگان (محتوایی) به کل واژگان. به این منظور، نقش دستوری کلمات هر دو سوی پیکره مقایسه‌ای مورد استفاده در این پژوهش با استفاده از ابزار برچسب‌زن تعیین و بسامد واژگان محتوایی و کارکردی به دست آمد.

رویکرد دیگر در تحلیل واژگانی نسبت واژگان متفاوت به کل واژگان است که تعداد واژگان منحصر به فرد نسبت به کل واژگان را شامل می‌شود. مهم‌ترین ضعف این معیار آن است که حساسیت بالایی به اندازه متن و پیکره دارد و لذا تنها هنگامی باید استفاده شود که هر دو (زیر)پیکره اندازه یکسانی داشته باشند. همچنین، زمانی که متن از اندازه خاصی بیشتر می‌شود، افزایش در واژگان متفاوت کندتر می‌شود، نکته‌ای که به مقادیر پایین‌تر این نسبت در پیکره‌های بزرگ‌تر منجر می‌شود (Cvrček & Chlumská, 2015). به همین دلیل اسکات (2004) نسخه بهبود یافته این معیار را با عنوان نسبت «استاندارد شده» واژگان متفاوت به کل واژگان پیشنهاد داد. در نسخه جدید، میانگین همان نسبت قدیمی ولی به ازای هر n واژه (معمولاً تکه‌های ۱۰۰۰ کلمه‌ای متوالی) محاسبه می‌شود. به طور کلی می‌توان نتیجه گرفت که «نسبت واژگان متفاوت به کل واژگان» تنوع دایره واژگانی را نشان می‌دهد و «تراکم واژگانی» بار و تراکم اطلاعاتی متن را.

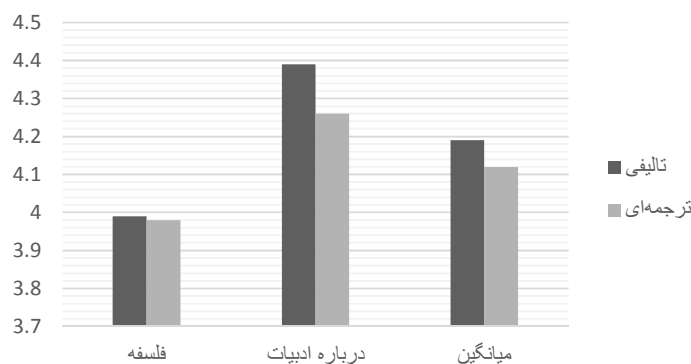
در بررسی نتایج تحلیل پیکره، ابتدا به سراغ مقایسه تراکم واژگانی متون فارسی تألیفی و فارسی ترجمه‌ای می‌رویم. در بررسی پیشینه به مطالعات مختلفی برمی‌خوریم که تراکم واژگانی کم‌تری را برای زبان ترجمه‌ای عنوان کرده‌اند. مثلاً، لایوزا (1998) با بررسی زبان انگلیسی پی برد که تراکم واژگانی در متون ترجمه‌ای نثر روایی انگلیسی (۵۲/۸۷ درصد) به شکل معناداری پایین‌تر از متون تألیفی مشابه (۵۴/۹۵ درصد) است. این مسئله در زبان‌های غیراروپایی نیز نشان داده شده است. برای نمونه، ژیاو و یو (2009) با بررسی متون ادبیات داستانی تألیفی و ترجمه‌ای در زبان چینی به این نتیجه رسیده‌اند که تراکم واژگانی در متون ترجمه‌ای (۵۸/۶۹ درصد) کم‌تر از متون چینی تألیفی (۶۳/۱۹ درصد) بوده است. ژیاو (2010) هم همین نتایج را گزارش کرده است (۶۱/۵۹ درصد در چینی ترجمه‌ای در مقابل ۶۶/۹۳ درصد در چینی تألیفی). سؤال اصلی این است که آیا این مشخصه در فارسی ترجمه‌ای و تألیفی نیز وجود دارد یا خیر. پیکره برچسب‌خورده بررسی و تعداد واژه‌های محتوایی و دستوری از یکدیگر تفکیک و درنهایت تراکم واژگانی هر دو سوی پیکره محاسبه شد. شکل ۱ تراکم واژگانی فارسی ترجمه‌ای و تألیفی در هر دو ژانر موردبررسی و نمرات میانگین آن‌ها را نشان می‌دهد:



شکل ۱: تراکم واژگانی در فارسی تألیفی و ترجمه‌ای

Figure 1: Lexical density in original vs. translational Persian

همان‌طور که از شکل ۱ پیداست، برخلاف یافته‌های مطالعات قبلی، تراکم واژگانی در زیرپیکره متون فارسی ترجمه‌ای بیشتر از تألیفی است، اما این تفاوت از نظر آماری معنادار نیست ($t = 1.08095$ for 198 d.f., $p = .140517$). در مقایسه متون دو حوزه نیز زیرپیکره متون درباره ادبیات تراکم واژگانی (بار اطلاعاتی) بیشتری نسبت به متون فلسفی از خود نشان داد. به علاوه، بیشتر بودن تراکم واژگانی متون ترجمه‌ای نسبت به تألیفی در داده‌های هر دو حوزه مشاهده شد، گرچه این تفاوت از نظر آماری در هیچ‌کدام از دو حوزه معنادار نبود. نتایج به دست آمده برخلاف انتظار محققان بود، چون معمولاً در فضای آموزشی و آکادمیک فرض بر این است که متون ترجمه‌ای بار اطلاعاتی و محتوای کم‌تری دارند و ساده‌تر هستند. همچنین، این نتایج برخلاف یافته‌های لایوزا (۱۹۹۸) در مورد انگلیسی ترجمه‌ای و ژائو و هو (۲۰۱۵) و ژائو (۲۰۱۰) در مورد زبان چینی است. نتایج تحقیقات استکی و همکاران (۲۰۱۰) (متون اقتصادی) و علی‌بابایی و صالحی (۲۰۱۲) (متون سیاسی) نیز با نتایج پژوهش حاضر درباره درست نبودن گزاره رایج در پیشینه مطالعات در مورد فارسی ترجمه‌ای، هم‌راستاست. لایوزا (۱۹۹۸) شیوه دیگری را برای سنجش تراکم واژگانی یا بار اطلاعات متنی ارائه می‌دهد. در این شیوه، نسبت واژگان محتوایی به واژگان کارکردی^{۳۲} محاسبه می‌شود. در مطالعه حاضر، حتی با استفاده از این شیوه (شکل ۲) نیز نتیجه مشابهی به دست آمد.



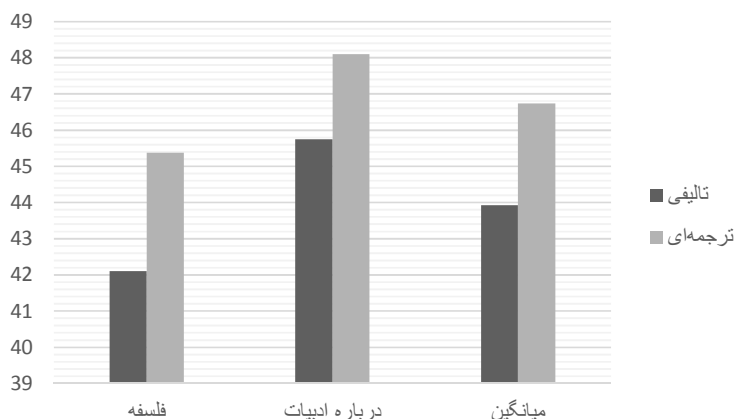
شکل ۲: نسبت واژگان محتوایی به کارکردی در فارسی تألیفی و ترجمه‌ای

Figure 2: Lexical-to-function word ratios in original and translational Persian

شکل ۲ نشان می‌دهد که در زیرپیکره فارسی تألیفی نسبت واژگان محتوایی/کارکردی بیشتر از فارسی ترجمه‌ای است ($4/19$ در مقابل $4/12$ ، تفاوت میانگین: $0/07$)، اما این تفاوت معنادار نبود ($t = 1.14672$ for 198 d.f., $p = .126441$). همچنین غلبه متون تألیفی بر ترجمه‌ای در داده‌های متنی هر دو حوزه مشاهده شد، گرچه میزان این نسبت و تفاوت تألیفی - ترجمه‌ای در متون درباره ادبیات بیشتر بود. نتیجه به دست آمده با نتایج برخی از پژوهش‌های قبلی، مانند نظریه لویوزا (1998) (نسبت واژگان محتوایی/دستوری در انگلیسی ترجمه‌ای پایین‌تر از انگلیسی تألیفی است) و مشاهدات ژیاو و یو (2009) و ژیاو (2010) در زبان چینی همخوانی نداشت. به طور کلی، «بسامد واژگان محتوایی» و «نسبت واژگان محتوایی به کارکردی» بار اطلاعاتی و تراکم واژگانی را نشان می‌دهند.

رویکرد دوم در تحلیل واژگانی، محاسبه نسبت استاندارد شده تعداد واژگان متفاوت به کل واژگان است. همان‌طور که در بالا اشاره شد، برای جلوگیری از هرگونه تأثیرگذاری حجم پیکره بر خروجی تنوع دایره واژگانی، در مطالعه حاضر از معیار نسبت استاندارد شده تعداد واژگان متفاوت به کل واژگان استفاده شد. همانگونه که شکل ۳ نشان می‌دهد، سه نکته مهم در تحلیل تنوع دایره واژگانی دو زیرپیکره مشخص شد. اول اینکه فارسی ترجمه‌ای برخلاف انتظار نمره میانگین بیشتری را در تنوع دایره واژگانی نسبت به فارسی تألیفی نشان داد ($46/73$ در مقابل $43/93$) و این تفاوت میانگین‌ها ($2/8$) از نظر آماری معنادار است ($t = -5.17141$ for 198 d.f., $p < .00001$). نکته دوم این است که در متون هر دو حوزه، تنوع دایره واژگانی فارسی ترجمه‌ای بیشتر از فارسی تألیفی است، اگرچه این تفاوت در متون فلسفی ملموس‌تر است. یکی از دلایل آن می‌تواند این باشد که نهضت ترجمه در سال‌های اخیر توانسته مفاهیم و اصطلاحات متنوع‌تر و بدیع‌تری را وارد پیکره و قاموس زبان فارسی تخصصی کند؛ درحالی که این مفاهیم و اصطلاحات هنوز به فارسی تألیفی راه نیافته است و لذا فارسی تألیفی فاقد این تنوع واژگانی است. نکته سوم این است که با تحلیل پیکره تدوین شده مشخص شد متون «درباره ادبیات» به طور کلی، از تنوع دایره واژگانی بیشتری (هم در زیرپیکره ترجمه‌ای و هم تألیفی) در مقایسه با متون فلسفی برخوردارند. وجود این ویژگی می‌تواند ناشی از این باشد که حوزه دانشی مزبور (به خصوص در متونی مانند نقد ادبی) در سال‌های اخیر با جهش‌های قابل ملاحظه مفهومی - اصطلاحی مواجه شده است، درحالی که شیب تحول مفاهیم و اصطلاحات

در حوزه‌ای مانند فلسفه یا حتی دین یا هنر ملایم‌تر بوده است.



شکل ۳: تنوع دایره واژگانی استاندارد شده در فارسی تألیفی و ترجمه‌ای

Figure 3: Standardized TTR in original and translational Persian

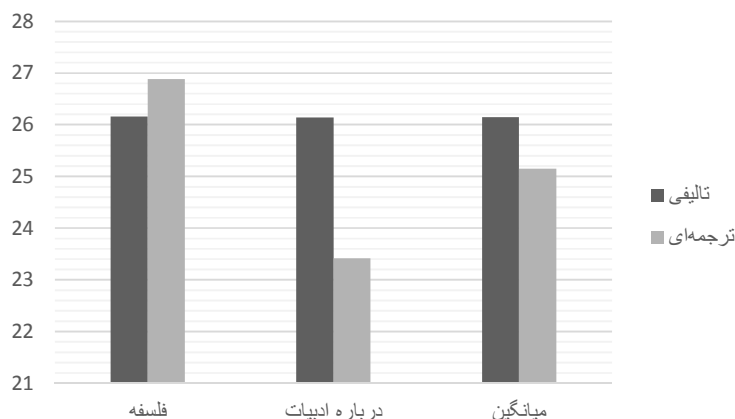
از این یافته‌ها می‌توان نتیجه گرفت که تنوع دایره واژگانی در فارسی ترجمه‌ای نسبت به فارسی تألیفی بیشتر است. این یافته برخلاف نتایج برخی از مطالعات پیشین در زبان‌های دیگر است که در آن‌ها ترجمه‌ها تنوع دایره واژگانی کم‌تری نسبت به متون تألیفی دارند، از جمله مطالعه لایووزا - برایثویت (1996) بر روی زبان انگلیسی و سوروچک و کلامسکا (2015) در زبان چک.

به‌طور کلی، در پیکره مورد بررسی، فارسی ترجمه‌ای تراکم واژگانی کمتر (تفاوت غیرمعنادار) و تنوع دایره واژگانی بیشتری نسبت به فارسی تألیفی از خود نشان داد. این دو ویژگی با یافته‌های پیشین در تضاد است، یافته‌هایی دال بر تراکم واژگانی پایین‌تر و تنوع دایره واژگانی کمتر زبان ترجمه‌ای نسبت به متون تألیفی. به‌علاوه، نتایج حاصل از پیکره مقایسه‌ای این پژوهش در مورد تنوع دایره واژگانی، یافته‌های مطالعات استکی و همکاران (2010) و علی‌بابایی و صالحی (2012) مبنی بر متنوع‌تر بودن واژگان در فارسی ترجمه‌ای را تأیید می‌کند. در مورد تراکم واژگانی نیز نتیجه‌ای مشابه با دو مطالعه ذکر شده یعنی بالاتر بودن میزان

واژگان محتوایی در زبان فارسی ترجمه‌ای به‌دست آمد، گرچه تفاوت مقادیر کل در پژوهش کنونی از نظر آماری معنادار نبود. از طرفی، محاسبهٔ «نسبت واژگان محتوایی به واژگان کارکردی» گرچه روند متفاوتی (بالا تر بودن نسبت در متون تألیفی) را نشان داد، اما باز هم تفاوت‌ها از نظر آماری معنادار نبود. در نهایت، می‌توان نتیجه گرفت که بودن تراکم واژگانی و تنوع دایرهٔ واژگانی در متون ترجمه‌ای نشانهٔ حضور جهانی ساده‌سازی در زبان فارسی نیست و بنابراین، نظریهٔ ساده‌سازی به‌شکل مطرح‌شده نمی‌تواند یک ویژگی «جهانی» تمام زبان‌های مقصد باشد.

۲-۴. میانگین طول جمله

برخی پژوهشگران میانگین طول جمله را یکی از جهانی‌های زبان مقصد در ترجمه می‌شمارند، بدین معنا که طول جملات متون ترجمه‌شده نسبت به متون تألیفی مشابه بیشتر است. در این مطالعه با اندازه‌گیری و مقایسهٔ میانگین طول جمله در دو زیرپیکرهٔ متون تألیفی و ترجمه‌ای فارسی به بررسی این فرضیه پرداخته‌ایم. پیشینهٔ مطالعات نشان می‌دهد که پژوهشگران مختلف نظرات متفاوت و بعضاً متعارضی دربارهٔ معیار میانگین طول جمله به‌عنوان نشانه‌ای از ساده‌سازی داشته‌اند (نک. بخش ۲-۱). نتایج مطالعهٔ این معیار در زبان فارسی که در شکل ۴ نشان داده شده است با انتظارات ما همخوانی نداشت. در فضای دانشگاهی، به‌خصوص کلاس‌های تربیت مترجم در ایران، فرضی وجود دارد که به‌طور کلی جملات ترجمه‌شده از جملات متن اصلی و حتی متون تألیفی فارسی طولانی‌ترند؛ فرضی که نتایج پژوهش حاضر در تضاد با آن است. تحلیل زیرپیکرهٔ متون فارسی ترجمه‌ای نشان داد که میانگین طول جمله در این متون کوتاه‌تر از طول جمله در زیرپیکرهٔ متون فارسی تألیفی است، اگرچه تفاوت میانگین آن‌ها معنادار نیست ($t = 0.14525$ for 198 d.f., $p = .44233$).



شکل ۴: میانگین طول جمله در فارسی تألیفی و ترجمه‌ای

Figure 4: Mean sentence length in original and translational Persian

شکل ۴ نشان می‌دهد که میانگین طول جمله در متون ترجمه‌ای حوزه فلسفه بیشتر از متون تألیفی است، درحالی که در متون «درباره ادبیات»، این روند برعکس است. به علاوه، اختلاف میانگین طول جمله بین متون تألیفی و ترجمه‌ای در حوزه ادبیات، بیشتر از متون فلسفی است. به طور کلی، تحلیل این معیار نشان می‌دهد که میانگین طول جمله بیشتر به ژانر و موضوع متون وابسته است تا ترجمه‌ای یا تألیفی بودن آن.

فرضیه میانگین طولانی‌تر بودن جمله در زبان ترجمه‌ای به عنوان بخشی از جهانی ساده‌سازی توسط پژوهشگران مختلفی از جمله لایوزا - برایثویت (1996) (زبان انگلیسی)، ژیاو و یو (2009) (زبان چینی) و اسفندیاری و همکاران (2012) و استکی (2010) (زبان فارسی) مطرح شده است. نتایج حاصل از پژوهش حاضر این فرضیه را تأیید نمی‌کند، اما نتیجه این مطالعه با این استدلال ملمکیر (1997) همخوانی دارد که کاربرد شکل قوی‌تر علائم سجاوندی در ترجمه‌ها احتمالاً موجب ظهور جملات کوتاه‌تر می‌شود. این نتیجه همچنین با یافته‌های کورپاس - پاستور و همکاران (2008) در زبان اسپانیایی هم‌راستاست. علی‌بابایی و صالحی (2012) نیز با پیکره‌ای کوچک‌تر نتیجه مشابهی گرفته‌اند که کاربرد جملات کوتاه‌تر در فارسی ترجمه‌ای بیشتر از فارسی تألیفی است.

۳-۴. بسامد حروف ربط به‌عنوان نشانه‌ای بر تصریح

همان‌طور که قبلاً اشاره شد، حروف ربط یکی از موضوع‌های مهم در پژوهش‌های جهانی‌های ترجمه هستند و معیار بسامد بالاتر حروف ربط در متون ترجمه‌ای به‌عنوان نمودی از جهانی تصریح اولین بار توسط بلوم - کولکا (1986) مطرح شد. برای صحت‌سنجی این موضوع، در پژوهش حاضر، ۱۵ مورد از پرکاربردترین حروف ربط در فارسی انتخاب و بسامد هریک در پیکره مقایسه‌ای زبان ترجمه‌ای و تألیفی فارسی با استفاده از نرم‌افزار ورداسمیث استخراج شد (جدول ۳).

جدول ۳: بسامد حروف ربط در فارسی تألیفی و ترجمه‌ای

Table 3: Occurrences of connectives in original and translational Persian

حروف ربط	فلسفه (تألیف)			فلسفه (ترجمه)			درباره ادبیات (تألیف)			درباره ادبیات (ترجمه)		
	بسامد	واژگان	درصد از کل	بسامد	واژگان	درصد از کل	بسامد	واژگان	درصد از کل	بسامد	واژگان	درصد از کل
و	۷,۶۵۷	۵/۱	۰/۹۳	۵,۸۴۱	۳/۷۹	۰/۹۴	۸,۷۹۷	۵/۵۸	۰/۹۸	۶,۳۲۲	۳/۸۱	۰/۹۰
که	۴,۱۵۱	۲/۷۷	۰/۹۶	۵,۰۴۶	۳/۲۷	۰/۹۷	۴,۰۶۱	۲/۵۸	۰/۹۸	۴,۹۲۹	۲/۹۷	۰/۹۷
اما	۴۶۲	۰/۳۱	۰/۹۲	۴۸۵	۰/۳۱	۰/۸۸	۴۵۲	۰/۲۹	۰/۹۰	۶۳۴	۰/۳۸	۰/۹۴
ولی	۱۹۸	۰/۱۳	۰/۸۶	۱۶۱	۰/۱	۰/۷۰	۱۴۸	۰/۰۹	۰/۸۰	۱۱۴	۰/۰۷	۰/۷۴
اگر	۴۱۸	۰/۲۸	۰/۸۸	۳۴۰	۰/۲۲	۰/۸۷	۱۶۳	۰/۱	۰/۸۹	۲۵۳	۰/۱۵	۰/۹۱
چون	۱۴۷	۰/۱	۰/۸۶	۱۰۸	۰/۰۷	۰/۸۳	۱۷۴	۰/۱۱	۰/۸۴	۱۹۷	۰/۱۲	۰/۸۸
زیرا	۹۱	۰/۰۶	۰/۸۶	۱۸۳	۰/۱۲	۰/۸۳	۶۱	۰/۰۴	۰/۸۳	۴۷	۰/۰۳	۰/۶۷
بنابراین	۱۲۶	۰/۰۸	۰/۸۳	۸۸	۰/۰۶	۰/۸۲	۴۳	۰/۰۳	۰/۷۱	۴۵	۰/۰۳	۰/۷۳
سپس	۴۹	۰/۰۳	۰/۷۶	۳۸	۰/۰۲	۰/۸۶	۴۷	۰/۰۳	۰/۷۹	۳۴	۰/۰۲	۰/۷۱
پس	۲۰۷	۰/۱۴	۰/۸۶	۱۷۵	۰/۱۱	۰/۸۱	۲۶۴	۰/۱۷	۰/۹۲	۲۲۵	۰/۱۴	۰/۹۲
بلکه	۲۰۸	۰/۱۴	۰/۹	۲۲۷	۰/۱۵	۰/۹۱	۱۲۹	۰/۰۸	۰/۸۵	۱۴۹	۰/۰۹	۰/۸۸
اگرچه	۲۳	۰/۰۲	۰/۶۲	۲۹	۰/۰۲	۰/۸۱	۳۱	۰/۰۲	۰/۷۸	۲۶	۰/۰۲	۰/۷۳
برای اینکه	۱۵	۰/۰۱	۰/۶۶	۱۵	۰/۰۱	۰/۷۱	۱۰	۰/۰۱	۰/۷۵	۱۸	۰/۰۱	۰/۶۲

حروف ربط	فلسفه (تألیف)			فلسفه (ترجمه)			درباره ادبیات (تألیف)			درباره ادبیات (ترجمه)		
	پراکنش	درصد از کل	بسامد	پراکنش	درصد از کل	بسامد	پراکنش	درصد از کل	بسامد	پراکنش	درصد از کل	بسامد
لیکن	۱۳	۰/۰۱	۰/۲۰	۳۰	۰/۰۲	۰/۷۶	۴	۰/۰۰	۰/۳۰	۳	۰/۰۰	۰/۳۰
هرچند	۳۹	۰/۰۳	۰/۶۷	۲۹	۰/۰۲	۰/۴۷	۱۸	۰/۰۱	۰/۷۳	۳۴	۰/۰۲	۰/۶۷
کل	۱۳۸۰۴	۹/۲۱٪	—	۱۲۷۹۵	۸/۲۹٪	—	۱۴۴۰۲	۹/۱۴٪	—	۱۳۰۳۰	۷/۸۶٪	—

داده‌های جدول ۳ و مقادیر کل نشان می‌دهد در متون هر دو حوزه، بسامد حروف ربط در فارسی تألیفی بیشتر از فارسی ترجمه‌ای است (۱۳۸۰۴ در برابر ۱۲۷۹۵ در متون فلسفی و ۱۴۴۰۲ در برابر ۱۳۰۳۰ در متون درباره ادبیات). نکته شایان توجه دیگر این است که مقایسه خروجی تحلیل هر حرف ربط در متون ترجمه‌ای و تألیفی به صورت منفرد هیچ روند و الگوی مشخصی را نشان نمی‌دهد که گواه توجه بیشتر به استفاده از حروف ربط در یکی از زیرپیکره‌ها باشد. درحقیقت، برخی حروف ربط در فارسی ترجمه‌ای بسامد بالاتری دارند («که»، «اما» و «بلکه»)، برخی دیگر در فارسی تألیفی («و»، «ولی»، «سپس» و «پس») و چند مورد هم بسته به نوع متن متغیرند («اگر»، «چون»، «زیرا»، «بنابراین»، «اگرچه»، «برای اینکه»، «لیکن» و «هرچند»). این بی‌نظمی حتی در حروف ربط پرتکرار نیز قابل مشاهده است. مثلاً حروف ربط «و»، «ولی» و «پس» در زیرپیکره تألیفی و حروف ربط «که»، «اما» و «بلکه» در زیرپیکره ترجمه‌ای بسامد بیشتری دارند.

یافته‌های این بخش نشان می‌دهد که فرضیه مطرح شده در برخی از مطالعات پیشین (مانند Xiao, 2010؛ Laviosa-Braithwaite, 1996؛ Øverås, 1998 در زبان نروژی، و در زبان چینی) مبنی بر اینکه تصریح در استفاده بیشتر از حروف ربط در ترجمه یک ویژگی همگانی است، حداقل در زبان فارسی فاقد اعتبار است، زیرا یک ویژگی وقتی به معنای اخص خود «جهانی» محسوب می‌شود که در همه زبان‌ها خود را نشان دهد. نتیجه مطالعه پورتینن (2004) نیز نشان می‌دهد که فنلاندی تألیفی و ترجمه‌ای از هیچ الگوی کلی مشخصی در به کارگیری بیشتر حروف ربط تبعیت نمی‌کنند.

به طور کلی، همان‌طور که اسکولا (2004) و مائورانن (2007) نیز اشاره کرده‌اند، به نظر

می‌رسد این نظریه نیز مانند برخی ویژگی‌های زبانی ژانربنیاد یا زبان‌بند باشد و با تغییر هر یک از این دو عامل این احتمال وجود دارد که ویژگی مزبور به جای «ویژگی جهانی ترجمه» یک «ویژگی بومی ترجمه» باشد (نک. بخش ۲-۰۲). مسئله دیگر، نوع گفتمان به‌کار رفته در این متون است. مثلاً، گفتمان متون فلسفی استدلالی است و لذا حروف ربط بیشترین تواتر را دارند، اما مثلاً در حوزه پزشکی که نوع گفتمان نه استدلالی و نه ترغیبی، بلکه صرفاً اطلاعاتی است، ساخت‌های واژگانی از تواتر و بسامد بیشتر برخوردار است. به‌طور طبیعی، در حوزه زبان تخصصی پزشکی (چه تألیفی و چه ترجمه‌ای) بسامد بالای حروف ربط به‌عنوان راهبرد تصریح‌نمی‌تواند محلی از اعراب داشته باشد. بنابراین، با توجه به نتایج این بخش و نوع گفتمان مورد استفاده در متون مختلف، نمی‌توان گفت که در پیکره مقایسه‌ای زبان فارسی بسامد زیاد یا اندک حروف ربط که در جهانی تصریح گفته شده حائز اهمیت است یا تبلور و ظهور آن معنادار است. به عبارت دیگر، با توجه به آنچه درباره حروف ربط به‌عنوان نمودی از جهانی تصریح گفته شده است، به نظر نمی‌رسد انتخاب بسامد حروف ربط به‌عنوان نشانه‌ای از جهانی تصریح در پیکره زبان فارسی یا زبان‌هایی که نتیجه‌ای مغایر به‌دست داده‌اند، پایایی و اعتبار چندانی داشته باشد.

۵. نتیجه

هدف این مطالعه بررسی ویژگی‌های زبان‌شناختی خاص فارسی ترجمه‌ای با استفاده از پیکره مقایسه‌ای فارسی ترجمه‌ای و تألیفی بود تا بتوان میزان و کیفیت ساده‌سازی و تصریح به‌مثابه جهانی‌های زبان مقصد در یک زبان غیر اروپایی (در اینجا فارسی) را بررسی کرد. متون دو حوزه علوم انسانی (متون فلسفی و درباره ادبیات) در ژانر توضیحی و چهار ویژگی پیشنهادی درمورد زبان ترجمه‌ای که می‌توانند دال بر وجود دوج جهانی ساده‌سازی و تصریح باشند، انتخاب شد؛ چهار ویژگی عبارت بودند از بسامد پایین تراکم واژگانی، تنوع کمتر واژگانی، جملات طولانی‌تر (همه مربوط به ساده‌سازی) و بسامد بالاتر حروف ربط (مربوط به تصریح). اولاً، در بررسی جهانی ساده‌سازی مشخص شد که فارسی ترجمه‌ای در پیکره مقایسه‌ای مورد استفاده در این پژوهش تراکم واژگانی بیشتری دارد، گرچه این تفاوت بسیار اندک بود و از نظر آماری معنادار محسوب نمی‌شد. همچنین، تنوع دایره واژگانی در فارسی ترجمه‌ای

بیشتر از فارسی تألیفی بود. به علاوه، بررسی معیار میانگین طول جمله نشان داد که جملات در فارسی تألیفی اندکی طولانی‌تر از فارسی ترجمه‌ای هستند. در دو معیار غنا و تنوع دایره واژگانی، متون «درباره ادبیات» و در معیار طول جمله، متون فلسفی مقادیر بیشتری از خود نشان دادند که این خود می‌تواند نمایانگر تنوع و تفاوت در متون حوزه‌ها و ژانرهای مختلف و بیانگر ویژگی‌های خاص هر یک باشد. در نهایت، در بررسی جهانی تصریح، بسامد حروف ربط در فارسی تألیفی بیشتر از متون ترجمه‌شده بود، اما بررسی کلی بسامد حروف ربط مختلف هیچ گرایش مشخصی در هیچ‌کدام از زیرپیکره‌ها مبنی بر استفاده بیشتر از حروف ربط مشاهده نشد؛ برخی حروف ربط در زیرپیکره فارسی ترجمه‌ای بیشتر بودند و برخی دیگر در داده‌های فارسی تألیفی و گروهی نیز هیچ روند مشخصی را دنبال نمی‌کردند؛ بدین معنا که در یک نوع متن در داده‌های ترجمه‌ای بیشتر بودند و در نوع متن دیگر کمتر بودند.

به علاوه، نتایج این پژوهش جهان‌شمول بودن نظریه جهانی‌های ترجمه را به چالش می‌کشد (نک. بخش ۲-۱) و سؤالاتی را درباره وجود ویژگی‌های مشترک و همگانی در همه ترجمه‌ها مطرح می‌کند، چراکه نتایج درمورد چهار ویژگی موردنظر با هیچ یک از نظریات ارائه‌شده در باب جهانی‌های زبان مقصد همخوانی ندارد. بنابراین، نتایج این تحقیق مؤید این فرضیه است که ویژگی‌های همگانی موردادعا، به دلیل تفاوت‌های زبانی، در زبان فارسی (حداقل به همان کیفیت گفته‌شده) صادق نیست. برخلاف بسیاری از مطالعات پیشین (مانند مطالعه جامع Ilisei et al., 2009)، ویژگی‌هایی مانند پایین‌تر بودن تراکم و تنوع دایره واژگانی، بیشتر بودن میانگین طول جمله و بسامد بیشتر حروف ربط در زبان ترجمه‌ای نمی‌توانند جزو برجسته‌ترین ویژگی‌های جهانی (حداقل در مفهوم مطلق آن) باشند. حتی درمورد محاسبه تراکم واژگانی با استفاده از شیوه «نسبت واژگان محتوایی به کارکردی» که نسبت بیشتری را در متون تألیفی نشان داد و با الگوی ارائه‌شده در نظریه ساده‌سازی همخوانی داشت، تفاوت مقادیر کل در دو زیرپیکره معنادار نبود. لذا، یافته‌های این پژوهش احتمالاً بیانگر ویژگی‌های خاص زبان فارسی ترجمه‌ای (ترجمه‌شده از انگلیسی) و فارسی تألیفی است. به طور کلی، می‌توان گفت مطالعه حاضر، برخلاف تصور، نشان می‌دهد نظریات ساده‌سازی و تصریح که به عنوان جهانی‌های مقصد در ترجمه مطرح شده‌اند واقعاً «جهانی» نیستند چون در تمام متون ترجمه‌ای در زبان‌های مختلف، و حداقل در متون فارسی ترجمه‌ای فلسفی و درباره ادبیات، مصداق ندارند.

اگرچه نتایج تعدادی از مطالعات مؤید نظریات ساده‌سازی و تصریح است (نک. بخش ۲)، اما نتایج برخی مطالعات دیگر این نظریات را به‌چالش می‌کشند، خصوصاً وقتی زبان یا ژانر متون از مواردی که بیشتر بررسی شده‌اند (زبان‌های اروپایی و متون ادبی) به‌سمت موارد کم‌تر بررسی‌شده (زبان‌های غیراروپایی و متون غیرادبی) تغییر می‌کند (Chesterman, 2004; Xiao & Hu, 2015; Mauranen, 2007). به‌نظر می‌رسد فرض اشتراک همه ترجمه‌ها در بروز و نمود جهانی‌های ترجمه نیازمند اصلاح است؛ لذا بهتر است در ارائه چنین تعمیم‌ها یا گرایش‌های همگانی جانب احتیاط را رعایت کرد و آن‌ها را، به تعبیر اسکولا (2004)، «قوانین» بومی ترجمه» بنامیم نه «قوانین جهانی ترجمه». درحقیقت، باید توجه داشت که برخی از نظریات مطرح‌شده تنها با استفاده از جفت زبان‌ها یا متون خاص ارائه شده‌اند و این احتمال وجود دارد که درمور برخی زبان‌ها یا ژانرهای دیگر صدق نکنند و لذا باید محدود شوند، همان‌گونه که نمونه آن را در این مطالعه دیدیم و نشان داده شد که هر دو جهانی مقصد موردبررسی، در زبان فارسی با کیفیت بیان‌شده نمود نداشتند.

پژوهش حاضر از چند جنبه با اکثر کارهای مشابه پیشین در باب ماهیت زبان‌شناختی فارسی ترجمه‌ای تفاوت دارد که عبارت‌اند از: تمرکز پژوهشگران بر روی جهانی‌های مقصد و به‌کارگیری پیکره مقایسه‌ای، فاصله گرفتن از بررسی ژانرهای ادبی و رفتن به‌سمت ژانرهای غیرادبی و کم‌تر بررسی‌شده، تنوع‌بخشی به داده‌ها (جمع‌آوری داده‌ها هم از کتاب‌ها و هم از وب‌سایت‌ها و صفحات آنلاین) و بهبود روش تحقیق و به‌کارگیری ابزارهای تخصصی پردازش و تحلیل پیکره. همچنین، باید اشاره کرد که متن مبدأ داده‌های ترجمه‌ای مورداستفاده زبان انگلیسی (به‌علت فراوانی زیاد و غالب بودن) بود و ترجمه از سایر زبان‌ها موردبررسی قرار نگرفت. به‌علاوه، پژوهش حاضر به‌دلیل استفاده از پیکره‌ای با اندازه و حجم متوسط محدودیت‌هایی داشت که می‌تواند در مطالعات آتی موردتوجه بیشتر قرار گیرد و برطرف شود. برای صحت‌سنجی ادعاها درباره وجود گرایش‌هایی همگانی به اسم جهانی‌های ترجمه در زبان ترجمه‌ای هنوز نیازمند انجام پژوهش‌های بیشتر به‌خصوص در زبان‌ها و ژانرهای کم‌تر بررسی‌شده هستیم. گرچه نتایج پژوهش حاضر از هیچ‌کدام از فرضیات ارائه‌شده در مطالعات پیشین پشتیبانی نکرد، اما این شاید دلیل خوبی برای کنار گذاشتن تمام و کمال جهانی‌های ترجمه نباشد. به عقیده نویسندگان می‌توان به‌جای کنار گذاشتن کامل احتمال وجود ویژگی‌های

زبانی مشترک در ترجمه‌ها، حداقل به گرایش‌های جدید و متفاوت درمورد ترجمه‌ها ارائه شده بود برسیم؛ مثلاً ساده‌سازی در زبان ترجمه‌ای ممکن است از طریق ویژگی‌هایی غیر از بسامد پایین تراکم و تنوع دایره واژگانی بروز پیدا کند یا جهانی‌های دیگری در کار باشند که باید مورد واکاوی قرار گیرند. باوجوداین، پژوهش حاضر مشخص کرد که ادعای وجود ویژگی‌های «جهانی» در مفهوم کامل و مطلق آن (در همه زبان‌ها و انواع متون) فاقد بنیان استوار است. باید پژوهش‌های بسیار بیشتری بر روی فارسی ترجمه‌ای و سایر زبان‌های غیروپایی صورت گیرد تا عیار و کیفیت جهانی‌های ترجمه و نقش زبان، نوع متن، مهارت مترجم و سایر مؤلفه‌های دخیل در بروز ویژگی‌های زبان‌شناختی ترجمه‌ها بیش‌ازپیش روشن شود.

۶. پی‌نوشت‌ها

1. Translation Universals (TUs)
2. S-universals
3. T-universals
4. sanitization
5. explicitation
6. simplification
7. conventionalization
8. normalization
9. parallel corpus
10. comparable corpus
11. product-oriented DTS
12. third code
13. lengthening
14. growing standardization
15. interference
16. content words
17. function words
18. lexical variety
19. Natural Language Processing
20. Translational English Corpus (TEC)
21. British National Corpus (BNC)
22. zero connectives
23. laws
24. OCR
25. URL

26. normalization
27. HunPoS
28. WordSmith
29. AntConc
30. Sketch Engine
31. LancsBox
32. function words

۷. منابع

- غمخواه، ا. و خزاعی‌فرید، ع. (۱۳۹۰). مقایسه کمی سبک متن اصلی و متن ترجمه‌شده (مورد پژوهی سه ترجمه از پیامبر و یک ترجمه از غریزگی). *مطالعات زبان و ترجمه*، ۴۴(۳)، ۱۱۸-۱۰۳.
- Ahangar, A. & Rahnemoon, S. N. (2019). The Level of explicitation of reference in the translation of medical texts from English into Persian: A case study on Basic Histology. *Lingua*, 228, 102704. <https://doi.org/10.1016/j.lingua.2019.06.005>
- Alibabae, A. & Salehi, Z. (2012). A comparative study of Persian translated and non-translated political texts: Focus on simplification hypothesis. *Journal of Language, Culture, and Translation*. 1(3): 1-13.
- Baker, M. (1993). Corpus linguistics and translation studies: Implications and applications. In Baker, M., Francis, G. and Tognini-Bonelli, E. (eds) *Text and technology: In honor of John Sinclair*. Amsterdam: John Benjamins. 233-250.
- Baker, M. (1995). Corpora in translation studies: An Overview and some suggestions for future research. *Target*. 7(2): 223-243. <https://doi.org/10.1075/target.7.2.03bak>
- Baker, M. (1996). Corpus-based translation studies: The challenges that lie ahead. In Somers, H. (ed.) *Terminology, LSP and Translation Studies in Language Engineering: In Honor of Juan C. Sager*. Amsterdam/Philadelphia: John Benjamins. 175-186.
- Blum-Kulka, S. (1986). Shifts of cohesion and coherence in translation. In House, J. and Blum-Kulka, S. (eds) *Interlingual and intercultural communication: Discourse*

and cognition in translation and second language acquisition studies and applications. Tübingen: Gunter Narr. 17–35.

- Blum-Kulka, S. & Levenston, E. A. (1983). Universals of lexical simplification. In Færch, C. and Kasper, G. (eds) *Strategies in Interlanguage Communication*. London: Longman. 119–139.
- Chen, J. W. (2006). *Explication Through the Use of Connectives in Translated Chinese: A Corpusbased Study*. PhD Dissertation. University of Manchester.
- Chesterman, A. (2004). Beyond the particular. In Mauranen, A. and Kujamäki, P. (eds) *Translation Universals: Do They Exist?* Amsterdam/Philadelphia: John Benjamins. 33–49.
- Chesterman, A. (2010). Why study translation universals?. *Acta Translatologica Helsingiensia (ATH)*. 1. 38–48.
- Corpas Pastor, G., Mitkov, R., Afzal, N. & Pekar, V. (2008). Translation universals: do they exist? A corpus-based NLP study of convergence and simplification. In *8th AMTA conference*. 75–81.
- Cvrček, V. & Chlumská, L. (2015). Simplification in translated Czech: A new approach to type-token ratio. *Russian Linguistics*. 39(3) 309–325. <https://doi.org/10.1007/s11185-015-9151-8>
- Esfandiari, M. R., Mahadi, T. S. T., Jamshid, M. & Rahimi, F. (2012). Explication and simplification in translation of poetic and prosaic genre from Persian into English: the case of Sadi's *Gulistan*. *Elixir Ling. & Trans*. 44. 7130–7133.
- Eskola, S. (2004). Untypical frequencies in translated language: A corpus-based study on a literary corpus of translated and non-translated Finnish. In Mauranen, A. and Kujamäki, P. (eds) *Translation Universals: Do They Exist?* Amsterdam/Philadelphia: John Benjamins. 83–100.
- Esteki, A., Vahid Dastjerdi, H. & Pirnajmuddin, H. (2010). *Testing Simplification Hypothesis: A Corpus-Based Study of Simplification Features in Persian Translated and Untranslated Economic Texts*. M.A. Thesis, University of Isfahan.

- Frawley, W. (1984). Prolegomenon to a theory of translation. In Frawley, W. (ed.) *Translation, Literary, Linguistic and Philosophical Perspectives*. London: Associated University Press. 159-175.
- Ghamkhah, A. & Khazaei Farid, Ali. (2011). A Quantitative Comparison of the Source Text and the Target Text: Three Translations of Jebran's *The Prophet* and a Translation of Al-e-Ahmad's *Gharb Zadegi*. *Journal of Language and Translation Studies*, 44(3). 103-118. [In Persian].
- Hansen, S. & Teich, E. (2001). Multi-layer analysis of translation corpora: Methodological issues and practical implications. In *Proceedings of EUROLAN 2001 workshop on multi-layer corpus-based analysis*. University Alexandru Ioan Cuza of Iasi: Iasi. 44-55.
- House, J. (2008). Beyond intervention: Universals in translation?. *Trans-kom*. 1(1). 6-19.
- Jantunen, J. (2001). Synonymity and lexical simplification in translations: A corpus-based approach. *Across Languages and Cultures*. 2(1). 97-112.
- Ilisei, I., Inkpen, D., Corpas Pastor, G. & Mitkov, R. (2009). Towards simplification: A supervised learning approach. *Proceedings of Machine Translation*. 25. 21-22.
- Laviosa-Braithwaite, S. (1996). *The English Comparable Corpus (ECC): A Resource and a Methodology for the Empirical Study of Translation*. PhD Dissertation, University of Manchester.
- Laviosa, S. (1998). Core patterns of lexical use in a comparable corpus of English narrative prose. *Meta*. 43(4). 557-570. <https://doi.org/10.7202/003425ar>
- Laviosa, S. (2002). *Corpus-based Translation Studies: Theory, Findings, Applications*. Amsterdam: Rodopi.
- Malmkjær, K. (1997). Punctuation in Hans Christian Andersen's stories and their translations into English. In Poyatos, F. (ed.) *Nonverbal Communication and Translation: New Perspectives and Challenges in Literature, Interpretation and the Media*. Amsterdam/Philadelphia: John Benjamins. 151-162.

- Malmkjær, K. (2007). Norms and nature in translation studies. In Anderman, G. M. and Rogers, M. (eds) *Incorporating corpora: The linguist and the translator*. Clevedon: Multilingual Matters. 49-59.
- Mauranen, A. (2000). Strange strings in translated language: A study on corpora. In Meave Olohan (ed.) *Intercultural Faculties. Research Models in Translation Studies 1: Textual and Cognitive Aspects*. Manchester: St. Jerome Publishing. 119-141.
- Mauranen, A. (2004). Corpora, universals and interference, In Mauranen, A. and Kujamäki, P. (eds) *Translation universals: Do they exist?* John Benjamins Publishing Co. 65-82.
- Mauranen, A. & Kujamäki, P. (2004). *Translation Universals: Do they exist?* John Benjamins Publishing Co.
- Mauranen, A. (2007). Universal tendencies in translation. In Anderman, G. M. and Rogers, M. (eds) *Incorporating corpora: The linguist and the translator*. Clevedon: Multilingual Matters. 32-48.
- Molés-Cases, T. (2019). Why typology matters: A Corpus-Based study of explicitation and implicitation of manner-of-motion in narrative texts. *Perspectives*. 27(6). 890–907. <https://doi.org/10.1080/0907676X.2019.1580754>
- Olohan, M. & Baker, M. (2000) Reporting that in translated English. Evidence for subconscious processes of explicitation?. *Across Languages and Cultures*. 1(2). 141–158. <https://doi.org/10.1556/Acr.1.2000.2.1>
- Øverås, L. (1998). In search of the third code: An investigation of norms in literary translation. *Meta*. 43(4). 557–570. <https://doi.org/10.7202/003775ar>
- Puurtinen, T. (2004). Explicitation of clausal relations: A corpus-based analysis of clause connectives in translated and non-translated Finnish children's literature. In Mauranen, A. and Kujamäki, P. (eds) *Translation universals: Do they exist?* John Benjamins Publishing Co. 165-176.
- Pym, A. (2005). Explaining explicitation. In Károly, K. and Ágota, F. (eds) *New Trends in Translation Studies*. Budapest: Akadémiai Kiadó. 29-43.

- Salimi, A. & Askarzadeh Torghabeh, R. (2015). Cohesion shifts and explication in English texts and their Persian translations: A case study of three novels. *Iranian EFL Journal*. 11(3). 331 – 348.
- Scott, M. (2004). *WordSmith Tools version 4*. Oxford: Oxford University Press.
- Seraji, M. (2015). *Morphosyntactic Corpora and Tools for Persian*. PhD Dissertation, Uppsala University: Studia Linguistica Upsaliensia 16.
- Stubbs, M. (1996). *Text and Corpus Analysis. Computer-assisted Studies of Language and Culture*. London: Blackwell.
- Toury, G. (1991). Experimentation in Translation Studies: Achievements, Prospects and some pitfalls. In Tirkkonen-Condit, S. (ed.) *Empirical Research in Translation and Intercultural Studies*. Tübingen: Gunter Narr. 45-66.
- Toury, G. (1995). *Descriptive Translation Studies and Beyond*. Benjamins.
- Toury, G. (2004). Probabilistic explanations in translation studies. Welcome as they are, would they qualify as universals?. In Mauranen, A. and Kujamäki, P. (eds) *Translation Universals: Do They Exist?* Amsterdam/Philadelphia: John Benjamins. 15–32.
- Tymoczko, M. (1998). Computerized corpora and the future of translation studies. *Meta*. 43(4). 652–660. <https://doi.org/10.7202/004515ar>
- Vahedi Kia, M. & Ouliaeinia, H. (2016). Explication across literary genres: Evidence of a strategic device?. *Translation & Interpreting*. 8(2). 82-95.
- Xiao, R. (2010). How different is translated chinese from nativechinese?: A Corpus-based Study of Translation Universals. *International Journal of Corpus Linguistics*. 15(1). Pp: 5–35. <https://doi.org/10.1075/ijcl.15.1.01xia>
- Xiao, R. & Yue, M. (2009). Using corpora in translation studies: The state of the art. In Baker, p. (ed.) *Contemporary Corpus Linguistics*. London: Continuum. Pp: 237–262.
- Xiao, R. & Hu, X. (2015). *Corpus-based studies of translational Chinese in English-*

Chinese translation. Springer.

- Zasiakin, S. (2019). Investigating cognitive and psycholinguistic features of translation universals. *PSYCHOLINGUISTICS*. 26(2). 114–134.
<https://doi.org/10.31470/2309-1797-2019-26-2-114-134>