

A Comparative Analysis of Persian Translations of the Qur'an Based on Automatic Word Alignment

Ahmad Rabieizadeh^{1✉} and Majid Lesani²

1. Deputy for Technology, Noor Computer Research Center of Islamic Sciences, and Director of the Artificial Intelligence Laboratory, Qom, Iran. Email: rabieizadeh@noornet.net
2. Researcher, Artificial Intelligence Laboratory, Research Institute for Digital Islamic and Humanities Studies, Qom, Iran. Email: mlesani@noornet.net

Article Info

Article type:
Research Article

Article history:

Received:

29/08/2025

Received in revised form:

30/10/2025

Accepted:

15/01/2026

Available online:

4/04/2026

Keywords:

Automatic Word
Alignment;
Digital Islamic Studies;
Qur'anic Translation;
Natural Language
Processing.

ABSTRACT

Achieving an accurate understanding of the meanings of the verses of the Holy Qur'an has been a central concern for Muslims, as well as for non-Muslims interested in Islam, since the earliest period of Islamic history. Given that the Qur'an was revealed in Arabic, the most accessible means for individuals unfamiliar with the Arabic language to understand its message is through Qur'anic translations. As a substantial proportion of Muslims are Persian speakers, numerous Persian translations of the Qur'an have been produced and are readily available. Conducting a systematic qualitative analysis of these translations, however, is an extremely time-consuming and resource-intensive task. This study proposes a framework for the automatic alignment of Qur'anic words with their Persian translations using natural language processing (NLP) techniques. Employing this approach, the words of 88 Persian translations of the Holy Qur'an, spanning from the fourth century AH to the present, were automatically aligned with their corresponding Qur'anic words. The results of this research have been incorporated into the Noor Comprehensive Tafsir Software 4. This achievement provides a robust platform for conducting a variety of analytical studies on Qur'anic translations. Three major types of analysis were carried out: (1) assessing the degree of each translator's fidelity to the original Qur'anic text; (2) identifying the extent of variation among translators' lexical choices for individual Qur'anic words; and (3) examining the degree of similarity among the Persian translations.

Cite this article: Rabieizadeh, A. & Lesani, M. (2026). A Comparative Analysis of Persian Translations of the Qur'an Based on Automatic Word Alignment. *Digital Islamic Studies and Humanities*, 2 (3), 43-66. <https://doi.org/10.22034/disah.2026.2094333.1090>



© The Author(s). **Publisher:** Research Center for Digital Islamic Studies and Humanities (RCDISAH).

DOI: <https://doi.org/10.22034/disah.2026.2094333.1090>

تحلیل مقایسه‌ای ترجمه‌های فارسی قرآن مبتنی بر ترازبندی خودکار کلمات

احمد ربیعی‌زاده^۱ و مجید لسانی^۲^۱. نویسنده مسئول، معاون فناوری مرکز تحقیقات کامپیوتری علوم اسلامی و مدیر آزمایشگاه هوش مصنوعی، قم، ایران، رایانامه:

rabieizadeh@noormet.net

^۲. پژوهشگر آزمایشگاه هوش مصنوعی پژوهشگاه علوم اسلامی و انسانی دیجیتال، قم، ایران، رایانامه: mlesani@noormet.net

اطلاعات مقاله

چکیده

نوع مقاله:

مقاله پژوهشی

تاریخ دریافت:

۱۴۰۴/۰۶/۰۷

تاریخ بازنگری:

۱۴۰۴/۰۸/۰۸

تاریخ پذیرش:

۱۴۰۴/۱۰/۲۵

تاریخ انتشار:

۱۴۰۵/۰۱/۱۵

کلیدواژه‌ها:

ترازبندی خودکار واژگان،

علوم اسلامی دیجیتال،

ترجمه قرآن کریم،

پردازش زبان طبیعی.

فهم دقیق معنای آیات قرآن کریم از دوران صدر اسلام تا کنون، از جمله دغدغه‌های اصلی مسلمانان و همچنین غیرمسلمانان کنجکاو به اسلام بوده است. برای این منظور، باتوجه به عربی‌بودن قرآن کریم، نزدیک‌ترین راه برای افراد ناآشنا به زبان عربی، بهره‌گیری از ترجمه‌های قرآن کریم است. بخش زیادی از مسلمانان فارسی‌زبان هستند و به همین دلیل ترجمه‌های زیادی از قرآن کریم به این زبان در دسترس است. تحلیل کیفی این ترجمه‌ها به‌صورت سیستماتیک، فرایند بسیار پرهزینه و زمان‌بری است. در این پژوهش، چارچوبی برای ترازبندی ماشینی کلمات قرآن و ترجمه‌های فارسی آن مبتنی بر روش‌های پردازش زبان طبیعی پیشنهاد شده و با به‌کارگیری این روش، بر روی ۸۸ ترجمه فارسی قرآن کریم از قرن چهارم قمری تا قرن حاضر، کلمات هر یک از ترجمه‌ها به کلمه معادل آن از قرآن کریم، ترازبندی شدند. خروجی این پژوهش در نرم افزار جامع التفاسیر^۴ نور ارائه شد. بدین ترتیب، بستر مناسبی برای انجام تحلیل‌های مختلف بر روی ترجمه‌ها فراهم آمد. سه نوع تحلیل انجام‌شده عبارتند از بررسی میزان وفاداری هر یک از مترجمین به متن قرآن کریم، شناسایی میزان تنوع نظر مترجمین به ازای هر یک از کلمات قرآنی و همچنین بررسی میزان شباهت ترجمه‌ها به یکدیگر.

استناد: ربیعی‌زاده، احمد و لسانی، مجید (۱۴۰۵). تحلیل مقایسه‌ای ترجمه‌های فارسی قرآن مبتنی بر ترازبندی خودکار کلمات.

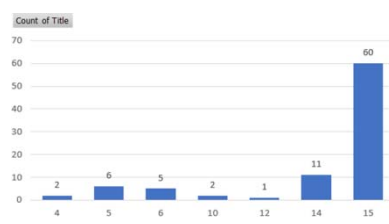
علوم انسانی و اسلامی دیجیتال، ۲ (۳)، ۴۳-۶۶. <https://doi.org/10.22034/disah.2026.2094333.1090>

ناشر: پژوهشگاه علوم اسلامی و انسانی دیجیتال (مرکز تحقیقات کامپیوتری علوم اسلامی نور). © نویسندگان.

مقدمه

فهم دقیق معنای آیات قرآن کریم از دوران صدر اسلام تا کنون، از جمله دغدغه‌های اصلی مسلمانان و همچنین غیرمسلمانان کنجکاو به اسلام بوده است. برای این منظور، باتوجه به عربی بودن قرآن کریم، نزدیک‌ترین راه برای افراد ناآشنا به زبان عربی، بهره‌گیری از ترجمه‌های قرآن کریم است. البته در ابتدا برخی افراد با برگردان متن قرآن به سایر زبان‌ها مخالف بوده‌اند. (رشید رضا ۱۳۴۴)

بخش زیادی از مسلمانان فارسی زبان هستند و به همین دلیل ترجمه‌های زیادی از قرآن کریم به این زبان در دسترس می‌باشد. در برخی از پژوهش‌ها به سیر ترجمه‌های قرآن به زبان فارسی اشاره شده است (انوار ۱۳۸۲). تحلیل کیفی این ترجمه‌ها به صورت سیستماتیک، فرایند بسیار پرهزینه و زمانبری می‌باشد. هدف این پژوهش، ارائه چارچوبی برای یافتن معادل فارسی هر یک از کلمات قرآن از میان کلمات ترجمه‌های فارسی آن با استفاده از روش‌های پردازش زبان طبیعی است.



شکل ۱ نمودار توزیع زمانی ترجمه‌های فارسی قرآن

در این پژوهش، ۸۸ ترجمه از پایگاه قرآن نور ۱ هدفگذاری شده‌اند. توزیع زمانی این ترجمه‌ها بر اساس قرن قمری وفات مولف در شکل ۱ ارائه شده است.

پیدا کردن هر کلمه از قرآن در میان این ۸۸ ترجمه به صورت دستی، زمان زیادی خواهد گرفت. معادلیابی کلمات قرآن در ترجمه‌ها، باعث صرفه جویی در زمان پژوهشگران و افراد عادی می‌شود. علاوه بر آن، گاهی مشاهده معنای یک واژه در سایر ترجمه‌ها باعث روشن‌تر شدن معنای دقیق یک واژه برای خواننده خواهد شد. همچنین پس از شناسایی معادلهای مختلف یک واژه از قرآن در تمام ترجمه‌ها، داده ارزشمندی مهیا خواهد شد که در صورت تکمیل و تصحیح، علاوه در پردازش‌های لغوی قرآنی، می‌توان از آن در سایر پژوهش‌های دو زبانه ماشینی عربی-فارسی نیز استفاده کرد. در نهایت نیز پس از ترازبندی ماشینی کلمات، بستر مناسبی برای انجام تحلیل‌های مختلف بر روی ترجمه‌ها فراهم خواهد آمد.

پیشینه پژوهش

در مورد ارزیابی ترجمه‌های قرآن دو رویکرد وجود دارد. رویکرد ارزیابی انسانی که در آن به جنبه‌های مختلف کیفیت ترجمه پرداخته می‌شود. برخی از معیارهای ارزیابی ترجمه‌های قرآنی در پژوهش‌ها بررسی شده است (Ereksoussi 2021). ارزیابی انسانی غالباً بدلیلی هزینه بالا به صورت محدود روی تعداد کمی از ترجمه‌ها و برای آیات و سوره‌های معدودی انجام شده است (بمانی نائینی ۱۳۹۲؛ اکبرکراسی و اصغری ۲۰۲۲؛ عامری و حسینی ۲۰۱۶؛ انوشیروانی ۲۰۰۶؛ واعظی ۲۰۰۹). رویکرد دوم، ارزیابی سیستماتیک است که محدودیت‌های ارزیابی انسانی را نخواهد داشت، اما از سوی دیگر ممکن است در حد ارزیابی انسانی، به جنبه‌های مختلف کیفیت ترجمه توجه نداشته باشد و روی ابعاد خاصی تمرکز داشته باشد (Shenassa و Khalvandi 2008). از جمله معیارهای ارزیابی سیستماتیک ترجمه‌ها بررسی وضعیت ترازبندی متن و ترجمه است که به دو گونه دستی و ماشینی قابل انجام است.

غَايِدُونَ: يَا كَاتِبُوا الْمَلِكِ قَبْلَ أَنْ يَكْتُبَ غَايِدُونَ (۵:۴۱)	غَايِدُونَ: وَخُذُوا لَهُ عَدِيدُونَ (۱۳۸:۲)
بار کثیر خوشی شوند ۲۲، بار گردیت ۳۰،	نرسند گان ۲، برسند گان ۳، بروند ۱۰،
ببار گردانید گان ۳۲، ببار گردید ۳۹،	نروند ۱۰، برسند گان ۲۲، برسند گان ۵۵،
بار گردند گان ۴۱، بار گرداند گان ۶۱، بار آید	می برسیم ۶۱، برسند گانیم ۸۲، برسارن ۱۲۲،
۹۶، بار گرداند گان ۱۲۱، بار گردند ۱۲۴، بر سر	برسنش کنند گان ۱۲۹، برنده کنند گان ۱۳۶
شوند گان ۱۳۱، روند گان ۱۳۸، گردانید گان ۱۳۹	
غَايِلًا: وَوَجَدَكَ عَالِيًا فَاغْنِ (۸:۹۳)	بِسْمِ الْكُتُبِ الْكُتُبِ الْكُتُبِ الْكُتُبِ (۱۱۲:۹)
می خواسته ۲۱، رویش ۲۴	برسون کناران ۱، برسند گان ۲، خدای برسان
غَايِدٌ: وَلَا أَلَا غَايِدًا مَاتِدْتُمْ (۴:۱۰۹)	۳، بندگی کنند گان ۳۹، خدای نرسند گان ۵۵،
برسند ۲، برسند ۷۹، برسند ۱۰۲	خدای برسند گان ۸۰، برسند گان خدا ۸۸،
	برسند گان ۹۴، عادت کنند گان ۱۱۳، بندگی
	گان ۱۲۶
غَايِدَاتٍ: تَكْبِكُ عَدِيدًا تَكْبِكُ (۵:۶۶)	بِسْمِ الْكُتُبِ الْكُتُبِ الْكُتُبِ الْكُتُبِ (۱۷:۲۳)
برسون کناران ۱، خدا برسان ۲، برسند گان ۳،	برسون کناران ۱، برسند گان ۲، بند گان ۲۴،
خدای برست ۲۹، برستگار ۳۴، بندگی کنند گان	برسند گان ۵۳، می برسند ۷۶، خدایکناران ۸۶

شکل ۲ نمونه صفحه ای از فرهنگنامه قرآنی یاحقی

یکی از کارهای دستی انجام شده برای معادل یابی کلمات قرآن در ترجمه‌های فارسی، کتاب ارزشمند فرهنگنامه قرآنی یاحقی است که به صورت دستی بر روی ۱۴۲ نسخه خطی از ترجمه‌های کهن قرآن به زبان فارسی انجام شده است (یاحقی ۱۹۹۸). در شکل ۲ نمونه ای از این کتاب را مشاهده می‌کنیم که به ازای هر کلمه از قرآن، معادل‌های موجود در هر یک از کتب ترجمه، شناسایی و ذکر شده است.

در برخی از مقالات نیز از ترازبندی خودکار برای ارزیابی ترجمه‌های ماشینی آیات قرآن استفاده شده است و به بررسی کیفی آنها پرداخته شده است. (Al-Sukhni و همکاران ۲۰۱۶)

روش پیشنهادی برای ترازبندی ماشینی کلمات

در این پژوهش از تکنیک ترازبندی خودکار متن و ترجمه استفاده شده است. ترازبندی کلمات متون، مسئله شناخته‌شده‌ای است که پژوهش‌های مختلفی بر روی آن انجام پذیرفته است. راهکارهای حل این مسئله به دو روش دستی و ماشینی تقسیم می‌شوند. در آغاز این پژوهش، از روش‌های عمومی ترازبندی دوزبانه استفاده شد که نتایج آن‌ها در مورد ترجمه‌های قرآن، دقت مطلوب را نداشت. به همین دلیل، تصمیم به طراحی یک الگوریتم جدید گرفته شد. دقت بالای موردنیاز برای متن قرآن کریم، در برابر متون عمومی، قابل مقایسه نیست. یکی از الگوریتم‌های بررسی‌شده Simalign (Sabet و همکاران ۲۰۲۰) بود که تا حدودی نتایج قابل قبولی داشت ولی با دقت موردنظر بسیار فاصله داشت.

روش انجام ترازبندی در این پژوهش، الگوریتم مفصلی است که با استفاده از اطلاعات لغت‌نامه‌های مختلف و تطابق احتمالی و اطلاعات صرفی و نحوی در چند فاز صورت می‌پذیرد که شمای کلی آن در شکل ۳ ارائه شده است.



شکل ۳ فرایند ترازبندی

فرایند اجمالی در این الگوریتم به قرار زیر است:

در فاز اول برای هر آیه ابتدا آن را به کلمات و زیر بخش‌های معنایی مناسب تقسیم می‌کنیم. در فاز دوم به جمع‌آوری معادل‌های فارسی هر یک از کلمات آیه مورد نظر با استفاده از لغت‌نامه‌ها می‌پردازیم و یک لغت‌نامه جامع و مختص به کلمات آن آیه ایجاد می‌کنیم. در فاز سوم، ابتدا کلمات هر ترجمه را پیراسته کرده و به کلمات مناسب تقسیم می‌کنیم. برای هر کدام از کلمات ترجمه با استفاده از لغت‌نامه جامع آیه، همه کاندیداهای احتمالی را به دست می‌آوریم. در مورد کاندیداهای هر لغت یک امتیازبندی و مرتب‌سازی انجام می‌دهیم تا کاندیداهای مناسب‌تر در اولویت قرار گیرند؛ زیرا در قسمت‌های بعدی الگوریتم از کاندیدای اول هر لغت برای ارزیابی سایر کاندیداهای لغات استفاده خواهد شد. سپس برای هر کلمه ترجمه، با در نظر گرفتن معیارهایی از قبیل معنای مناسب‌تر، پوشش حداکثری و رعایت ترتیب، مناسب‌ترین کاندیدا را بر می‌گزینیم. پس از این با روش‌هایی مثل جستجو در سطح ریشه، کلمات فاقد معادل از متن آیه یا ترجمه فارسی را تا حد ممکن در میان موارد باقیمانده معادل‌یابی می‌کنیم. در فاز چهارم، تمام فرایند مذکور فاز سوم را با استفاده از معادل‌های پیداشده دوباره انجام می‌دهیم تا از ترازبندی ترجمه‌های سابق برای ترجمه‌های دیگر استفاده شود. علاوه بر این در آیات طولانی در اجرای دوباره از حدود تقریبی بخش‌های آیه که در اجرای مرحله اول به دست می‌آید برای محدود کردن دامنه کاندیداها استفاده می‌شود. پس از انجام ترازبندی نوبت به گردآوری ترجمه‌های مختلف یک لغت و سایر تحلیل‌ها بروی ترجمه‌های مختلف با استفاده از ترازبندی خواهد رسید. در ادامه به توضیح تفصیلی هر یک از این مراحل می‌پردازیم.

مرحله اول: قطعه‌بندی آیات

در این مرحله، برای اینکه بتوان به درستی معادل هر قسمت از آیه را در ترجمه‌ها استخراج کرد، لازم است آیه را به کلمات و اجزای معنایی مناسب قطعه‌بندی کرد. رویکرد انتخاب‌شده برای قطعه‌بندی، فقط مبتنی بر فاصله یا بر اساس جزئیات صرفی نیست. اگر چه در نهایت پس از اتمام

morphological tokens

word tokens

multi-word tokens

در سطر بالا اجزای صرفی هر جزء با یک رنگ، در سطر وسط اجزای معنایی مورد استفاده در معادل‌یابی و در سطر پایین، حالت نهایی پس از ادغام و گسترش این اجزای معنایی برای ارائه نهایی، نشان داده شده است. به طور مثال «لا» و «یاتی» تز حیث صرفی، دو کلمه هستند که با فاصله از هم جدا شده‌اند؛ ولی از نظر معنایی، معادل یک کلمه فارسی به معنای «نمی‌آید» هستند. برای به دست آوردن اجزای معنایی از اطلاعات صرفی مختلف و قواعد متعددی استفاده شده است؛ مثلاً در اعداد یا «ال» آغازین یا «ما» زائده و «لا» نهی و ضمائر فاعلی، اجزاء ادغامی شوند. در حدود چهل قاعدهٔ مختلف صرفی مشابه، در این مرحله استفاده شده است. لازم به ذکر است برخی عبارات، در اکثر ترجمه‌ها به صورت یک واحد معنایی ترجمه شده‌اند؛ مانند «اولوا الالباب» که بیشتر به صورت تک کلمه «خردمندان» به جای «صاحبان خرد» ترجمه شده یا «الذین آمنوا» که بیشتر به صورت تک کلمه «مومنان» به جای «کسانی که ایمان آورده‌اند» ترجمه شده یا «اصحاب النار» که بیشتر به صورت تک کلمه «جهنمیان» به جای «اهل جهنم» ترجمه شده است؛ که این قبیل عبارات نیز در هم ادغام خواهند شد. برای بهبود الگوریتم شناسایی این قبیل موارد اهمیت بالایی دارد؛ به همین دلیل، حدود دویست مورد از موارد پرتکرار شناسایی شده و به صورت یک جزء معنایی در نظر گرفته می‌شوند؛ هرچند که دو یا چند کلمه‌ی جدا شده و با فاصله باشند.

برای هر آیه با استفاده از ۲۵ مجموعه لغتنامه، تمام معانی کلمات هر آیه گردآوری می‌شود. این لغت‌نامه‌ها شامل چند دسته‌اند:

- برخی از این داده‌ها از سایت قاموس نور^۱ استخراج شده است که قبلاً به روش‌های دستی و بعضاً ماشینی از لابلای لغت‌نامه‌ها گردآوری شده بودند.
- همچنین از فرهنگ‌های عربی به فارسی قرآنی مثل «لغت نامه تفسیری استاد بهرام‌پور» و «فرهنگ‌نامه قرآنی یاحقی» و برخی دیگر از فرهنگ‌های برگرفته از تفسیرهای فارسی استفاده شد.
- علاوه بر این از لغت‌نامه‌های خاص محتوای اسلامی کهن مثل صحیفه سجادیه و نهج البلاغه نیز استفاده شده است.
- برای گسترش کاندیداهای احتمالی از ترکیب لغت‌نامه‌ها با زبان‌های دیگر نیز استفاده شد؛ به عنوان نمونه، معادل‌های انگلیسی کلمات قرآن با فرهنگ انگلیسی فارسی ادغام شد و در نتیجه معادل‌هایی برای واژگان عربی قرآن بدست آمد. همین روش، برای ترکیب لغت‌نامه عربی به فارسی و فارسی به فارسی نیز به کار برده شد. البته به علت اصیل نبودن، امتیاز کمتری به این مجموعه در الگوریتم داده شده است.
- همچنین از برخی لغت‌نامه‌های عمومی عربی فارسی و نیز از سرویس ترجمه گوگل^۲ برای ترجمه لغات و عبارات استفاده شده است.
- از بردارهای جانمایی کلمات^۳، برای معادلیابی استفاده شد؛ در این‌گونه روش‌ها هر لغت، تبدیل به بردار می‌شود به طوری که معانی نزدیک به هم بردارهای نزدیک به هم دارند و در حالت چندزبانه، لغات هم معنا در زبانهای مختلف به نزدیک به هم خواهند بود. مبتنی بر این مدل‌های برداری، نزدیکترین کلمات ترجمه با هر کلمه از آیه شناسایی می‌شود. در این پژوهش از مدل زبانی BERT (Devlin و همکاران ۲۰۱۹) به صورت چندزبانه استفاده می‌کنیم. البته این تکنیک در ابتدا پیش از افزایش مجموعه لغت‌نامه‌ها باعث افزایش دقت ترازبندی بود؛ ولی در نهایت پس از افزودن مجموعه‌های متعدد و بهبود الگوریتم، بر اساس آزمایش‌های انجام‌شده، به دلیل تقریبی بودن معادلها، تصمیم بر حذف این مجموعه از لغت‌نامه‌ها گرفته شد.

1. Qamus.inoor.ir
 2. Translate.google.com
 3. word embedding

- با تحلیل کلمات پوشش داده نشده، حدود ۱۳ هزار جفت، داده‌های دستی نیز شناسایی شد که پس از تایید متخصصین زبان، به این دایره لغات افزوده شدند.
- در مورد حروف و افعال ربطی نیز که تعداد آنها محدود است به خاطر مناسب نبودن برخی لغت نامه‌ها در این مورد و ایجاد و تکثیر خطا ما تنها از معادل‌های متنوع تهیه شده دستی استفاده کرده‌ایم. علاوه بر دایره لغات فوق، با استفاده از داده‌های صرفی و نحوی و قواعد مشخص به گسترش معانی جمع آوری شده می‌پردازیم. در ادامه به ذکر چند نمونه از این قواعد می‌پردازیم:
- معادل‌یابی برای یک کلمه جمع با داشتن گونه مفرد آن در لغت‌نامه‌ها؛ برای مثال با توجه به اینکه کلمه «مُتَّقِینَ»، جمع سالم کلمه «مُتَّقِی» است، با اضافه کردن پسوند «ها»، معادل‌های جدیدی برای آن از کلمه «مُتَّقِی» ساخته می‌شود.
- تولید معادل برای کلمات حاوی حروف تاکید از قبیل لام و نون؛ برای مثال معنای جدیدی برای «لَیَقُولَنَّ» از معانی «یَقُولُ» را به بانک معادل‌ها، اضافه می‌کنیم.
- توجه به نقش نحوی کلمه؛ برای مثال با توجه به نقش تمییزی در واژه‌ی «عِلْمًا» معادل «از نظر علم» را از روی کلمه ساده‌ی «العِلْم» تولید می‌شود. یا با توجه به حالیه بودن واو، معادل احتمالی «در حالی که» را اضافه می‌کنیم.
- لازم به ذکر است در این مرحله فقط معانی جدید اضافه نمی‌شود بلکه ممکن است معنای برخی حروف و افعال ربطی را با توجه به نوع و نقش آنها محدودتر کنیم. مثلاً با توجه به نافی بودن کلمه «ما» در یک آیه خاص، معنای احتمالی موصولی را برای آن حذف می‌کنیم.
- در نهایت به هر کدام از لغت‌نامه‌ها و معادل‌های بدست آمده با توجه به کیفیت آن منبع امتیاز خاصی تعلق گرفته است تا در مراحل بعدی الگوریتم لحاظ شود. همچنین چنانچه یک معادل در لغت‌نامه‌های مختلف تکرار شود، به همان نسبت، با امتیاز بیشتری ثبت خواهد شد.

مرحله سوم: معادل‌یابی در هر یک از ترجمه‌ها

همانطور که ذکر شد مرحله سوم مرحله اصلی معادل‌یابی است که به طور خاص برای هر ترجمه از یک آیه خاص اجرا می‌شود و شامل زیرمراحل مختلفی است که به ترتیب به بیان آن‌ها می‌پردازیم:

- پیرایش ترجمه و قطعه‌بندی آن به اجزای لغوی مناسب
- پیدا کردن همه‌ی کاندیداهای احتمالی به ازای هر لغت از ترجمه
- مرتب‌سازی کاندیداهای هر لغت ترجمه بر اساس معیارهای مشخص
- یافتن معادل نهایی برای هر لغت با شروع از مناسب‌ترین لغت‌ها
- پیدا کردن معادل‌های باقیمانده‌ی احتمالی فاقد کاندید

پیرایش ترجمه و قطعه‌بندی آن به اجزای لغوی مناسب

در مورد ترجمه‌ها برای هر ترجمه ابتدا به پیرایش ترجمه می‌پردازیم. با توجه به تنوع حروف مماثل فارسی از قبیل کاف و یاء، حروف را نرمال می‌کنیم و مواردی مثل عبارات تشریف (مثل سلام الله علیه) یا آیه‌ها را در میان ترجمه‌ها حذف می‌کنیم. سپس کلمات را با توجه به فاصله‌ها یا در مواردی حتی ریزتر تفکیک کرده یا ادغام می‌کنیم تا تطابق در مراحل بعد راحت‌تر صورت گیرد. به عنوان نمونه:

- ضمائر، جدا در نظر گرفته می‌شوند. مثلاً عبارت «می‌گویدش» را به «می‌گوید» و «ش» تبدیل می‌کنیم تا برای هر کدام از معادل‌های عربی این دو واژه که اجزای معنایی مختلفی هستند راحت‌تر معادل مناسب خودش پیدا شود.

- افعال مرکب به صورت یک کلمه واحد در نظر گرفته می‌شوند. مثلاً افعالی مانند «می‌زند» یک کلمه خواهند بود یا افعال ربطی یا اصطلاحاً افعال «هم‌گرد» به عنوان یک کلمه خواهند بود؛ مثلاً دو کلمه‌ی «عزیز می‌کند» را با اینکه با فاصله جدا شده‌اند یک واحد می‌شوند تا فرایند معادل‌یابی، دقیق‌تر انجام شود.

پیدا کردن همه‌ی کاندیداهای احتمالی به ازای هر لغت از ترجمه

در این مرحله سعی می‌کنیم با پیمایش بروی هر کدام از لغات ترجمه کاندیداهای احتمالی قابل تطابق را از اجزای معنایی قرآن در آیه با استفاده از لغت‌نامه‌ای که برای آیه ساخته بودیم بدست آوریم. به عنوان نمونه همانطور که در شکل ۵ مشاهده می‌کنید، برای واژه خدا، سه معادل احتمالی «الله» در آیه مورد نظر شناسایی شده است.



دا و پیغمبر او رضا می‌دادند و می‌گفتند خدا ما را بس است، زود باشد که خدا از کرم خویش به ما عطا کند و نیر

شکل ۵ نمونه ای از یک کلمه با کاندیداهای متعدد ترازبندی

برای بدست آوردن کاندیداهای بیشتر، این کار در چند مرحله صورت می‌گیرد:

- در مرحله اول معادلهای دقیق را که در لغت‌نامه‌ی ما دقیقاً موجودند انتخاب می‌کنیم.
- معادل‌های تقریبی با توجه به پیراسته سازی پسوندها و پیشوندها و ریشه‌یابی کلمات شناسایی می‌شوند.
- معادل‌های تقریبی را با استفاده از شباهت ظاهری کاراکتری کلمه فارسی با یکی از معادل‌ها و در صورت کم‌بودن فاصله لونشتاین^۱ به دست می‌آوریم.
- معادل‌های اجزای چند بخشی را نیز در صورت نبود معادل، به صورت دقیق یا تقریبی بدست می‌آوریم.
- در پایان از مرادفات فارسی برای پیدا کردن معادل‌های بیشتر بهره می‌بریم.

مرتب‌سازی کاندیداهای هر لغت

در این مرحله، با توجه به معیارهای مختلف به مرتب سازی کاندیدها به ازای هر لغت از ترجمه فارسی می‌پردازیم. این معیارها عبارتند از: توجه به نوع و کیفیت لغتنامه حاوی آن کاندیدا، تعداد انطباق‌های مشابه، نوع انطباق که دقیق یا تقریبی بوده است و نیز از همه مهم‌تر فاصله کلمه کاندیدا با معادل‌های کلمات مجاور، به عنوان مثال، طبق شکل ۵، اگر برای لغت فارسی «خدا»، سه کاندیدا در جمله عربی وجود داشته باشد، کاندیدایی که فاصله کمتری با معادل واژه‌ی مجاور آن (یعنی معادل کلمه «کرم» که «فضله» بوده است) دارد، امتیاز بیشتری خواهد داشت. در نهایت کاندیداهای هر لغت بر اساس امتیاز مرتب خواهند شد.

1. Levenshtein distance

یافتن معادل نهایی برای هر لغت

پس از مرتب‌سازی کاندیداهای هر لغت، تمام لغات را بر اساس مناسب‌بودن برای آغاز معادل‌یابی نهایی مرتب می‌کنیم و برای هر لغت فارسی یک معادل یکتا در میان لغات معادل عربی انتخاب می‌کنیم. کاندیدایی مناسب‌تر است که قطعیت بیشتری در انطباق بتوانیم به آن نسبت دهیم؛ مثلاً تعداد کاندیداهای متفاوت آن کمتر باشد و کاندید اول آن نسبت به سایرین (اگر وجود داشته باشند) امتیاز بالاتری داشته باشد. کلمات مبهم از قبیل ضمائر و افعال ربطی که بار معنایی پایینی دارند و ممکن است به راحتی با بخش‌های مختلف جمله انطباق یابند، اولویت پایین‌تری برای شروع تطابق دارند و در انتهای فرایند معادل‌یابی قرار می‌گیرند. لذا با شروع از مناسب‌ترین لغت معادل‌یابی نشده، سعی می‌کنیم بهترین کاندیدا را برای آن پیدا کنیم. برای پیدا کردن بهترین کاندیدا، غیر از امتیازات قبلی کاندیداها، میزان پوشش و ترازبندی حاصل، بعد از انتخاب این کاندیدا در انتخاب آن موثر خواهد بود. به عبارت دیگر، به ترتیب، زنجیره کاندیداهایی انتخاب خواهند شد که در نهایت، باعث ترازبندی تعداد بیشتری از لغات فارسی و عربی معادل خواهند شد. ولی مسئله اینجاست که ما هنوز برای لغات دیگر معادل‌یابی نکرده‌ایم! پس برای ساده‌سازی مسأله، فرض می‌کنیم که ما برای همه لغات فارسی معادل‌یابی نشده، اولین کاندیدا را به عنوان معادل انتخاب خواهیم کرد که کاندیدای عربی با بیشترین امتیاز داده شده در مرحله قبل خواهد بود.

علاوه بر میزان پوشش، مسئله‌ی دیگری که در انتخاب کاندیدای مناسب اهمیت دارد رعایت حداکثری ترتیب کلمات است، به عبارت دیگر، کاندیدایی که با انتخاب آن، توافق بیشتری بین ترتیب معادل‌های آیه با کلمات ترجمه باشد، مقدم خواهد بود. برای به دست آوردن میزان مرتب بودن ترازبندی، از الگوریتم «طولانی‌ترین زیررشته صعودی»^۱ استفاده می‌کنیم. نمونه این کار در شکل ۶، ارائه شده است. رشته ورودی، دنباله‌ای از اندیس لغت عربی پیدا شده به ازای لغات فارسی ترجمه است. در نمونه ارائه‌شده، دو زیررشته صعودی با طول ۴ شناسایی شده است که در صورت تغییر ترتیب در رشته ورودی ممکن بود طول زیررشته صعودی شناسایی شده بیشتر یا کمتر باشد. لازم به ذکر است که در این پژوهش، از این الگوریتم به صورت وزن‌دهی شده استفاده شده

1. Longest Increasing Subsequence (LIS)

است به نحوی که برای کلمات طولانی‌تر امتیاز بیشتری در نظر گرفته خواهد شد. بدین منظور، کافی است که کلمات طولانی را به تعداد حروف آنها در رشته ورودی الگوریتم مذکور تکرار کنیم.

Input Sequence	6	9	8	2	3	5	1	4	7
LIS 1	2	3	4	7					
LIS 2	2	3	5	7					

شکل ۶ نمونه دو خروجی الگوریتم بزرگترین زیر رشته صعودی برای رشته اعداد ورودی

پیدا کردن معادل‌های باقیمانده احتمالی

پس از معادلیابی لغات هر ترجمه برخی لغات عربی یا فارسی هستند که معادل خاصی در طرف مقابل، برای آنها پیدا نشده است در مرحله بعد سعی می‌کنیم با استفاده از سایر روش‌ها حتی الامکان معادل ممکن را برای آنها پیدا کنیم. در این رویه سعی می‌کنیم با سخت‌گیری کمتر با استفاده از مواردی مثل شباهت ریشه‌ای یا انطباق‌های تقریبی با معیارهایی کلی‌تر از مرحله سابق، لغات عربی و فارسی بدون معادل را به هم نسبت بدهیم. البته با حفظ اینکه لغات در جای مناسب خود قرار بگیرند و ترتیب تقریبی ترجمه حفظ شود.

مرحله چهارم: اجرای مجدد معادلیابی با استفاده از نتایج ترازبندی سایر ترجمه‌ها

پس از آنکه معادلیابی به‌ازای یک ترجمه انجام شد، از نتایج آن می‌توان برای معادلیابی بهتر سایر ترجمه‌ها نیز استفاده کرد؛ لذا این معادل‌های جدید را (که حتی ممکن است از روش‌های تقریبی بدست آمده باشند) به لیست واژه‌نامه‌های اصلی اضافه کرده و الگوریتم معادلیابی را دوباره برای آیات قبلی تکرار می‌کنیم؛ چرا که یک تغییر جزئی در بانک معادل‌ها، ممکن است باعث بهبود بسیاری از نتایج شود.

يٰۤاَيُّهَا الَّذِيْنَ ءَامَنُوْا لَا تَجْلُوْا سَعْتِرَ اللَّهِ وَلَا السَّعِرَ الْحَرَامَ وَلَا الْهَدْيَ وَلَا الْقَلَائِدَ وَلَا
ءَاثِيْنَ الْبَيْتِ الْحَرَامِ يَنْتَعُوْنَ فَضْلًا مِّنْ رَبِّهِمْ وَرِضْوَانًا وَإِذَا حَلَلْتُمْ فَاصْطَادُوا وَلَا يَجْرِمَنَّكُمْ
شَكَاكُ قَوْمٍ أَن صَدُّوْكُمْ عَنِ الْمَسْجِدِ الْحَرَامِ أَن تَعْتَدُوا وَيَعَاوَنُوا عَلَى الْإِثْرِ وَالْعُقُوْبِ وَلَا
تَعَاوَنُوا عَلَى الْإِثْرِ وَالْعُدُوْبِ وَأَتَّقُوا اللَّهَ إِنَّ اللَّهَ شَدِيْدُ الْعِقَابِ ﴿٥﴾

شکل ۷ نمونه بخش‌بندی آیه بر اساس علائم وقف

[illegible]

در ادامه، نمونه‌ای از خروجی نهایی ترازبندی ماشینی کلمات قرآن به ازای تعدادی از ترجمه‌های فارسی در شکل ۸ قابل مشاهده است. در هر ترجمه، لغات فارسی معادل با لغت عربی با رنگ آن لغت مشخص شده است. لغات سیاه نمایانگر کلماتی هستند که معادلی برای آنها پیدا نشده است. در نهایت، خروجی این پژوهش در نسخه ۴ نرم افزار جامع التفاسیر نور^۱ ارائه شد.

1. <https://www.noorsoft.org/fa/software/View/74155>

تحلیل ترجمه‌ها مبتنی بر خروجی نهایی ترازبندی

پس از اجرای الگوریتم ترازبندی خودکار برای تمامی ترجمه‌های همه آیات قرآن کریم، نوبت به تحلیل خروجی نهایی می‌باشد. از جنبه‌های مختلفی می‌توان به تحلیل کلان ترجمه‌های قرآنی با استفاده از خروجی ترازبندی کلمات پرداخت. در این پژوهش، ۳ نمونه از تحلیل‌های کاربردی و با اهمیت ارائه شده است که عبارتند از بررسی میزان وفاداری هر یک از مترجمین به متن اصلی قرآن، تحلیل میزان تنوع نظر مترجمین به ازای هر یک از کلمات قرآنی و بررسی میزان شباهت ترجمه‌ها به یکدیگر. در ادامه به توضیح هر یک از این موارد می‌پردازیم.

الف. بررسی میزان وفاداری هر یک از مترجمین به متن اصلی قرآن

یکی از گونه‌های دسته بندی ترجمه‌ها، تقسیم آنها به ترجمه کلمه‌به‌کلمه یا تحت‌اللفظی و ترجمه آزاد است (معرفت ۱۳۸۲) (Catford 1965). با توجه به اینکه در ترجمه، امکان جایگزینی صددرصدی یک متن از یک زبان به زبان دیگر وجود ندارد، بهترین روش ارزیابی کیفیت یک ترجمه، بررسی میزان همسانی واژگانی و ساختاری است (Sultana 1968) با توجه به این نکته، ترازبندی ماشینی کلمات ترجمه‌ها و متن قرآن می‌تواند معیار مناسبی برای سنجش میزان وفاداری هر یک از این ترجمه‌ها به متن قرآن کریم یا به عبارت دیگر میزان تحت‌اللفظی بودن ترجمه باشد. نرخ میزان تحت‌اللفظی بودن یک جفت آیه-ترجمه، با نسبت تعداد کلمات ترازبندی‌شده به تعداد کل کلمات تعریف می‌شود و از دو نقطه نظر متن اصلی و متن ترجمه، قابل محاسبه است. برای مثال، به ازای آیه و ترجمه ارائه شده شکل ۹، نرخ ترازبندی بر مبنای متن اصلی آیه ۱۱/۷ و بر مبنای متن ترجمه ۱۲/۷ می‌باشد.

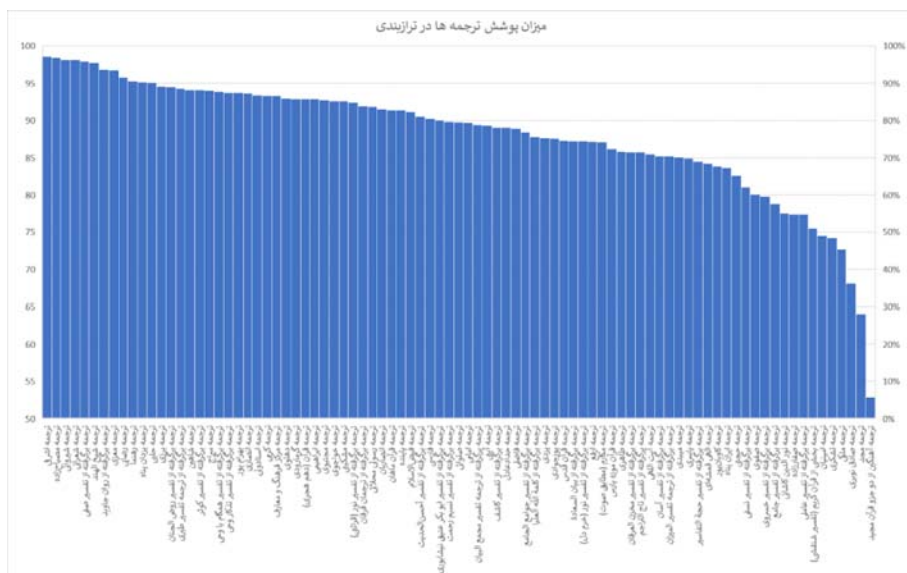


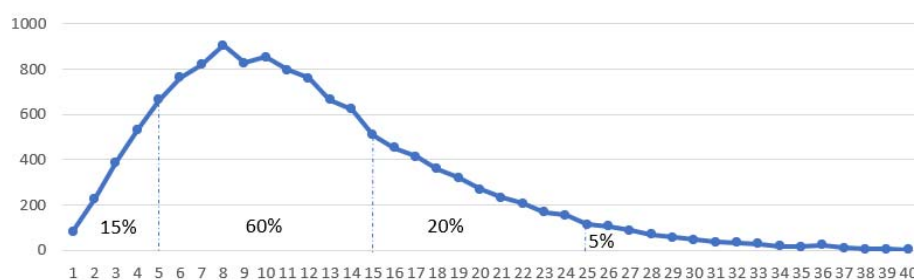
شکل ۹ نمونه یک آیه و ترجمه ترازبندی‌شده در سطح کلمه

$$SAR_{B_1} = \frac{\sum_{n=1}^N \frac{|ast_n|}{|st_n|}}{N} \quad TAR_{B_1} = \frac{\sum_{m=1}^M \frac{|att_m|}{|tt_m|}}{M}$$

$$AR = \frac{SAR + TAR}{2}$$

نمودار وضعیت کتاب‌های ترجمه از لحاظ نرخ میزان ترازبندی در شکل زیر ارائه شده است.





شکل ۱۲ نمودار فراوانی میزان تنوع کلمات قرآن

بر اساس نمودار ارائه شده در شکل ۱۲، در مجموع ۱۷۶۶ کلمه از قرآن، یعنی حدود ۱۴ درصد، زیر ۵ معادل فارسی داشته‌اند به عبارت دیگر، مترجمین در ترجمه آنها اتفاق نظر بالایی داشته‌اند. این کلمات غالباً شامل گونه‌های خاصی بوده‌اند از قبیل بعضی حروف اضافه مشخص (مثل او و لو به ترتیب به معنای یا و اگر)، برخی اسامی خاص مثل اسامی اشخاص و جای‌ها مانند (بابل، عرفات، موسی، طالوت، نوح، زکریا، عیسی)، اعداد (مثل ثلاثین، سته، سبع) و در نهایت، برخی اسامی و افعال عربی که در فارسی دارای معادل‌های مشخص و محدودی هستند (مثل أصابع، یدی، تصوموا به ترتیب معنای انگشتان، دودست، روزه بگیرید).

در مقابل، برای برخی از کلمات، تنوع بالایی توسط مترجمین ارائه شده و معادل‌های متعددی تا ۴۰ نوع ترجمه، توسط مترجمین بیان شده است. برای بیان مفهوم دقیق این کلمات، بعضاً مفاهیم تفسیری نیز در ترجمه دخالت داده شده است. مانند کلمه ماعون (به معنای زکات، اسباب، اثاث، وسایل، رخت خانه، مایحتاج، نیکوکاری و ...) یا مانند کلمه صریم (به معنای خاکستر، شب تاریک، شب سیاه، تل ریگ، باغ سوخته، زراعت برچیده و ...) یا مانند اسامی مبهمی که ذاتاً مبهم هستند و در جمله معنای دقیقشان مشخص می‌شود (مانند ضمائر و اسامی موصول).

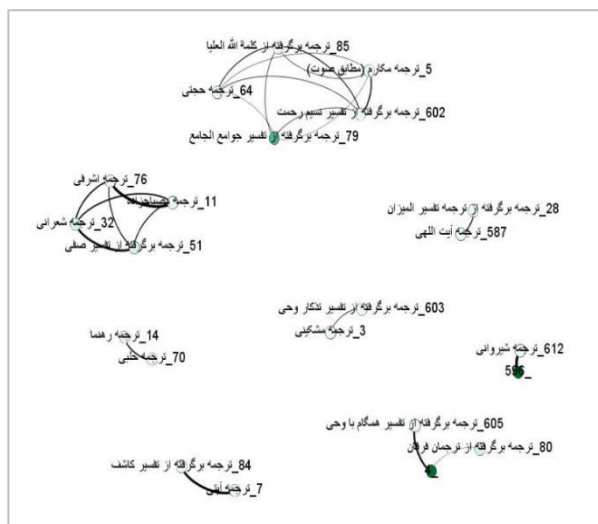
ج. بررسی میزان شباهت ترجمه‌ها به یکدیگر

میزان شباهت بالای دو ترجمه قرآن، می‌تواند نشانه‌ای قوی از احتمال تاثیرپذیری ترجمه متاخر از ترجمه متقدم باشد. (Svensson 2019) تعداد کلمات مشترک معادل بیشتر در دو ترجمه به ازای کلمات یک آیه واحد، نشانگر شباهت بیشتر آن دو ترجمه به یکدیگر است و اگر این

مقدار به ازای تمام آیات قرآن در معادل‌های دو کتاب ترجمه محاسبه شود، می‌توان میزان شباهت کلی دو کتاب را نسبت به یکدیگر شناسایی کرد. در این مطالعه، برای نرمال کردن امتیاز شباهت، تعداد اشتراک به دست آمده به کل کلمات دو ترجمه تقسیم می‌شود. همچنین مشاهده شد که کلماتی مثل حروف ربط که از ارزش کمی برخوردار هستند، بسیار مشترک واقع می‌شوند. لذا برای کاهش تاثیر اینگونه کلمات که غالباً طول کوتاهی دارند، به جای شمارش تعداد کلمات مشترک، از مجموع طول آنها در فرمول نهایی استفاده شد. فرمول نهایی شباهت دو کتاب ترجمه قرآن، در ادامه ارائه شده است:

$$Sim_{B_1B_2} = \frac{\sum_{n=1}^N Len(B_1T_n \cap B_2T_n)}{2 \times \sum_{n=1}^N Len(B_1T_n \cup B_2T_n)}$$

طبق فرمول فوق، میزان شباهت هر یک از کتاب‌های ترجمه با سایر کتب محاسبه شد و گراف مشابهت کل کتب با یکدیگر رسم شد. برای شناسایی کتاب‌های مشابه، یال‌های با نرخ تشابه کمتر از ۰.۵ حذف شدند. گراف نهایی در شکل ۱۳ زیر ارائه شده است.



شکل ۱۳ گراف شباهت کتب ترجمه قرآن

همانطور که در تصویر قابل مشاهده است، برخی از ترجمه‌ها بسیار مشابه یکدیگر می‌باشند و بر اساس تقدم زمانی می‌توان قدیمی‌ترین ترجمه اصلی را به ازای هر زیرگراف شناسایی کرد.

کارهای پیش رو

- کارهای مختلفی برای تکمیل موتور ترازبندی قابل تصوّر است، از جمله:
- شناسایی و تلاش برای کاهش خطاها و تمرکز بر روی لغات پوشش نیافته و بدون معادل
 - استفاده از داده‌های ترازبندی تایید شده، برای آموزش مدل‌های بزرگ زبانی برای بدست آوردن یک ماشین دقیق ترازبندی خودکار قرآن.
 - اجرای فرایند ترازبندی برای ترجمه‌های قرآن به زبانهای دیگر، به طوری که کاربر بتواند معادل هر لغت قرآن را در آن ترجمه‌ها ببیند. برای این مقصود می‌توان از مدل‌های زبانی چند زبانه نیز بهره برد.
 - انجام ترازبندی برای کلمات احادیث و ترجمه‌ها آنها

فهرست منابع

- اکبرکرکاسی، مانده، و جواد اصغری. ۲۰۲۲. «ارزیابی کیفیت ترجمه قرآن کریم براساس راهبرد 'جبران' وینی و داربلنه؛ مورد پژوهی ترجمه مکارم شیرازی، الهی قمشه‌ای و صادقی تهرانی از سوره یوسف.» آموزش زبان، ادبیات و زبان شناسی ۱۰ (۵): ۴۲-۶۷.
- انوار، سیدامیرمحمود 1382. and. «نگاهی به سیر ترجمه و تفسیر قرآن کریم به فارسی.» دانشکده ادبیات و علوم انسانی دانشگاه تهران (نشریات قدیمی)، ش ۱۷۲: ۳۵-۵۴.
- انوشیروانی، علیرضا. ۲۰۰۶. «مطالعه تطبیقی ترجمه‌های انگلیسی سوره حمد.» پژوهش ادبیات معاصر جهان ۱۱ (۳۱).
- بمانی نائینی، معصومه 1392. and. «نقد و بررسی چهار ترجمه انگلیسی از سوره مبارکه الاسراء.» فقه و تاریخ تمدن، ش ۳۸: ۲۹-۵۲.
- رشید رضا، محمد. ۱۳۴۴. ترجمه القرآن و مافیها من المفسد و منافاة الاسلام.
- عامری، خدیجه، و زینب السادات حسینی. ۲۰۱۶. «بررسی تطبیقی ترجمه‌های انگلیسی قرآن آبربی، شاکر و قرایی از آیات منتخب سوره رعد.» پژوهشنامه تفسیر و زبان قرآن ۴ (۲): ۷-۲۴.
- معرفت، محمد هادی. ۱۳۸۲. تاریخ قرآن. سازمان مطالعه و تدوین کتب علوم انسانی دانشگاه‌ها (سمت) - تهران - ایران.
- واعظی، محمود. ۲۰۰۹. نقد و ارزیابی ترجمه‌های انگلیسی قرآن کریم از سوره مبارکه الانسان.
- یاحقی، محمد جواد. ۱۹۹۸. فرهنگنامه قرآنی. آستان قدس رضوی. بنیاد پژوهشهای اسلامی - مشهد مقدس - ایران.
- Al-Sukhni, Emad, Mohammed N Al-Kabi و Izzat M Alsmadi. 2016. «An automatic evaluation for online machine translation: Holy Quran case study.» International Journal of Advanced Computer Science and Applications 7 (6): 118-23.
- Catford, John Cunnison. 1965. «A linguistic theory of translation.» Oxford University Press.
- Devlin, Jacob, Ming-Wei Chang, Kenton Lee و Kristina Toutanova. 2019. «BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding».

- Ereksoussi, Zubaidah M Kheir. 2021. Towards a Theory on Qur'an Translation: Re-Thinking Equivalence.
- Sabet, Masoud Jalili, Philipp Dufter, François Yvon و Hinrich Schütze. 2020. «SimAlign: High quality word alignments without parallel training data using static and contextualized embeddings.» arXiv preprint arXiv:2004.08728.
- Shenassa, Mohammad Ebrahim و Mohammad Javad Khalvandi. 2008. «Evaluation of different English translations of Holy Koran in scope of verb process type.» 2008 3rd International Conference on Information and Communication Technologies: From Theory to Applications, 1-4.
- Sultana, N. 1968. The translation of literary texts: a paradox in theory. Language and style. ج ۲۰،
- Svensson, Jonas. 2019. «Computing Qur'ans: a suggestion for a digital humanities approach to the question of interrelations between English Qur'an translations.» Islam and Christian-Muslim Relations 30 (2): 211-29.

References

- Akbarkarkasi, M., & Asghari, J. (2022). Evaluating the quality of Qur'an translation based on Vinay and Darbelnet's "compensation" strategy: A case study of Makarem Shirazi's, Elahi Ghomshei's, and Sadeghi Tehrani's translations of Surah Yusuf. *Journal of Language Teaching, Literature and Linguistics*, 10(5), 42–67. [in persian]
- Al-Sukhni, E., Al-Kabi, M. N., & Alsmadi, I. M. (2016). An automatic evaluation for online machine translation: Holy Quran case study. *International Journal of Advanced Computer Science and Applications*, 7(6), 118–123.
- Anoushiravani, A. (2006). A comparative study of English translations of Surah al-Fatihah. *Research in Contemporary World Literature*, 11(31). [in persian]
- Anvar, S. A. M. (2003). A survey of the history of Persian translations and commentaries of the Holy Qur'an. *Faculty of Literature and Humanities, University of Tehran (Old Series)*, 172, 35–54. [in persian]
- Bamani-Naeini, M. (2013). A critical review of four English translations of Surah al-Isra'. *Journal of Jurisprudence and the History of Civilization*, 38, 29–52. [in persian]
- Catford, J. C. (1965). *A linguistic theory of translation*. Oxford University Press.
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding.
- Ereksoussi, Z. M. K. (2021). *Towards a theory on Qur'an translation: Re-thinking equivalence*.
- Ma'refat, M. H. (2003). *History of the Qur'an*. SAMT Publications. [in persian]
- Rashid Rida, M. (1965). *Tarjamat al-Qur'an wa ma fiha min al-mafasid wa munafat al-Islam*. [in persian]
- Sabet, M. J., Dufter, P., Yvon, F., & Schütze, H. (2020). SimAlign: High quality word alignments without parallel training data using static and contextualized embeddings. *arXiv Preprint*, arXiv:2004.08728.
- Shenassa, M. E., & Khalvandi, M. J. (2008). Evaluation of different English translations of the Holy Qur'an in the scope of verb process type. In *2008 3rd International Conference on Information and Communication Technologies: From Theory to Applications* (pp. 1–4).
- Sultana, N. (1968). The translation of literary texts: A paradox in theory. *Language and Style*, 20.

- Svensson, J. (2019). Computing Qur'ans: A suggestion for a digital humanities approach to the question of interrelations between English Qur'an translations. *Islam and Christian-Muslim Relations*, 30(2), 211–229.
- Vaezi, M. (2009). *A critique and evaluation of English translations of the Holy Qur'an: Surah al-Insan*. [in persian]
- Yahaghi, M. J. (1998). *Encyclopedia of the Qur'an*. Astan Quds Razavi, Foundation for Islamic Research. [in persian]
- Ameri, K., & Hosseini, Z. S. (2016). A comparative study of Arberry's, Shakir's, and Qarai's English translations of selected verses from Surah al-Ra'd. *Journal of Qur'anic Exegesis and Language*, 4(2), 7–24. [in persian]