

Application of Distributional Semantics in the Qur'an, A Case Study of the Root-word "*Farah*"

Asieh Zouelm¹  & Nosrat Nilsaz^{*2} 

Vol. 16, No. 6, Tome 90
pp. 37-72
January & February
2026

Abstract

Distributional semantics is a neoconstructionist method that focuses on the contextual use of words in real texts to determine the approximate meaning of a word relative to other words. Since one of the goals of applying semantics to the Qur'an is to ascertain the meaning of words according to their practical context, the use of this method in Qur'anic studies becomes important. In this research, in order to introduce the application of distributional semantics in the Qur'an, the descriptive-analytical method is employed to explain the implementation steps, the challenges, and the solutions to overcome them, by examining a case study. The steps are: determining comparable words for the main target word; determining distributional features (surah, equivalent to document, and phrase); linguistic pre-processing of the Qur'anic text and removal of stopwords; forming the context vector and co-occurrence matrix and weighting its elements; and finding and analyzing the semantic similarity of words. The most important challenges of using this method in the Qur'an are the small volume of the Qur'anic text, the lack of suitable software for calculations within the Qur'anic context, and the great difference in the length of surahs in the word-document pattern. Solutions to overcome some of these challenges include: paying closer attention to the basis of the distributional method (distributional hypothesis); avoiding the use of this method for very low-frequency words; and comparing the results obtained from the word-document and word-phrase patterns. For the root-word *Farah*, it is found that the meaning derived from the Qur'anic context (closer to the meaning of pride) is different from the meaning mentioned in standard dictionaries (joy).

Key words: Distributional Semantics, Context Vector, Semantical Similarity, Quran, Farah

Received: 22 September 2022
Received in revised form: 17 April 2023
Accepted: 10 May 2023

¹. Ph.D candidate, Department of Qur'an and Hadith sciences, Faculty of Humanities, Tarbiat Modares University, Tehran, Iran; ORCID ID: <http://orcid.org/0000-0001-8061-5774>

². Corresponding author: Associate Professor, Department of Qur'an and Hadith sciences, Faculty of Humanities, Tarbiat Modares University, Tehran, Iran; Email: nilsaz@modares.ac.ir, ORCID ID: <http://orcid.org/0000-0002-8580-5574>

1. Introduction

Distributional semantics, as one of the neo-constructural methods in semantics, identifies the meaning of words by examining their use in natural texts. This method is based on the distributional hypothesis, which states: *'Word with similar distributional properties have similar meanings.'* Applying this method to the Quranic corpus is helpful in understanding the meaning of Quranic words by examining their usage context, especially in the case of words that are disputed by philologists. Thus, by comparing different words from the perspective of their distribution in the Quran (including co-occurrence with other words or use in different *Surahs*), and finding their distributional similarity, it identifies the scope of the meaning of the disputed word by obtaining the semantic distance between the words. Applying the distributional method in the Quran, due to the characteristics of the Arabic language and the Quranic corpus, requires localizing the method and overcoming its possible challenges. The method of implementing the method in the Quran is demonstrated practically using the case example of the root word "*Farah*", the meaning of which is disputed.

Research Question(s)

This research includes two main questions about the method of implementing distributional semantics in the Quran :

- 1- What are the practical steps for implementing distributional semantics in the Quran, considering its linguistic and textual properties ?
- 2- What are the challenges of applying distributional method in the Quran and how to overcome them ?

After that, in the practical implementation of the method in the Quran, another question is also raised :

- 3- What does the use of distributional semantics reveal about the meaning of the root word 'Farah' in the Quran?

2. Literature Review

Many works have been written in the field of distributional semantics, but we have not found a work in Persian that fully introduces this method and specifically implements it in the Quran. However, the translator of the two books *Machine learning for text* (Aggarwal, 2018) and *Networks* (Newman, 2018) into Persian (Ayub Turkian), added the appendices "Practical Text Mining Training" (Turkian, 2019a) and "Text Network" (Turkian, 2019b), respectively, which also examined the semantical similarity of verses of the Holy Quran and the clustering of nouns in the Quran. These two cases can be considered the closest to the subject of this research.

3. Methodology

The present study can be divided into two main parts. In the first part, which includes the explanation of how to implement distributional semantics in the Quran and the examination of its challenges, the research method is descriptive-analytical and the data collection method is documentary. In the second part, which is the implementation of distributional semantics on a case study in the Quran, first a number of words are selected for comparison with the main word "farah", which, based on the views presented in dictionaries, are in the two semantic domains of 'joy' and 'arrogance'. Then, two models of the distributional method are implemented for comparing words. In the word-phrase model, considering a 15-word context window ($n=15$), a number of phrases are selected as features and words are compared with each other, based on how they co-occur. In the word-document model, surahs are considered as documents (distributional features), and the occurrence of words in surahs is compared. Then, the results of implementing these two patterns of the distribution method on words are compared to obtain the final result.

4. Results

Distributional semantics is applied to the Quranic corpus with the aim of identifying the approximate meaning, i.e., determining the semantic domain of the disputed words. The implementation of this method in the present study includes the following steps:

1. Finding words to compare with the original word based on dictionaries;
2. Determining the distributional features, including documents (Surahs) and phrases that co-occur with the word (by grammatical or non-grammatical relations);
3. Linguistic pre-processing of words, including finding roots and effective forms of each root, disambiguation and merging words;
4. Removing stopwords;
5. Forming the co-occurrence matrix and its weighting (in the word-surah model with the TFIDF rule and in the word-phrase model with the PPMI rule);
6. Calculating the similarity of words by the cosine of the angle between the weighted context vectors;
7. Interpretation of the similarity between words using dictionaries and semantic relations between sentences.

The challenges of applying distributional semantics, especially in the Quran, and some ways to overcome them include:

1. General challenges of the distributional semantics, including the problem of automatic language preprocessing, interpretation of word similarity, and speaker intention recognition.
2. The challenges of applying distributional semantics to the Quran including the small size of the Quranic corpus for statistical calculations, the

lack of appropriate software for effective rooting, disambiguation, and recognizing stopwords, and the large variation in the length of surahs in the word-surah model. Some solutions to overcome the challenges include: paying more attention to the distributional hypothesis as the basis of the method, which, based on the wisdom of the speaker, results in greater credibility in the Quran; and correcting the error caused by the difference in the length of the documents in the word-surah model by comparing the results with the word-phrase model.

Applying the distributional method to the root-word Farah in the Quran shows the difference in the meaning obtained from the Quranic corpus (in the semantic domain of arrogance) compared to the dominant meaning mentioned in dictionaries, especially modern dictionaries (joy); that represents the semantic evolution of this word.



دوماهنامه بین‌المللی

د ۱۵، ش ۶ (پیاپی ۹۰)، بهمن و اسفند ۱۴۰۴، صص ۷۲-۳۷

مقاله پژوهشی

https://lrr.modares.ac.ir/article_7139.html

کار بست معنی‌شناسی توزیعی در قرآن، مطالعه نمونه موردی ریشه - واژه «فرح»

آسیه ذوعلم^۱، نصرت نیل ساز^{۲*}

۱. دانشجوی دکتری رشته علوم قرآن و حدیث، دانشگاه تربیت مدرس، تهران، ایران

۲. دانشیار گروه علوم قرآن و حدیث، دانشگاه تربیت مدرس، تهران، ایران

تاریخ پذیرش: ۱۴۰۲/۲/۲۰

تاریخ دریافت: ۱۴۰۱/۶/۲۱

چکیده

معنی‌شناسی توزیعی، روشی نوساختگراست که برای شناخت معنی، به کاربرد واژگان در متون واقعی توجه دارد و به کارگیری آن، حدود معنایی یک واژه را در مقایسه با دیگر واژگان مشخص می‌کند. از آنجاکه یکی از اهداف معنی‌شناسی قرآن، شناخت واژگان با توجه به بافت کاربردی آن‌هاست، بهره‌گیری از این روش در قرآن اهمیت می‌یابد. در این پژوهش به منظور معرفی چگونگی کار بست معنی‌شناسی توزیعی در قرآن، با روش توصیفی - تحلیلی، گام‌های پیاده‌سازی این روش در قرآن، چالش‌ها و راهکارهایی برای برون‌رفت از آن، با بررسی یک نمونه موردی تشریح شده است. گام‌های اجرای روش در قرآن عبارت‌اند از: تعیین واژگانی برای مقایسه با واژه اصلی، تعیین ویژگی‌های توزیعی (سوره معادل سند، و عبارت)، پیش‌پردازش زبانی متن قرآنی و حذف ایست‌واژه‌ها، تشکیل بردار بافتی و ماتریس هم‌وقوعی و وزن‌دهی عناصر آن، اندازه‌گیری مشابهت معنایی واژگان و تحلیل آن. مهم‌ترین چالش‌های استفاده از روش توزیعی در قرآن، حجم کم پیکره قرآنی، عدم وجود نرم‌افزار مناسب جهت انجام محاسبات و تفاوت زیاد طول سوره‌ها در الگوی واژه - سند است. توجه بیشتر به مبنای روش توزیعی در متن قرآن، عدم استفاده از این روش برای واژگان بسیار کم‌بسامد و مقایسه نتایج به‌دست‌آمده از الگوی واژه - سند و واژه - عبارت، راهکارهایی برای برون‌رفت از برخی چالش‌هاست. اجرای روش توزیعی برای ریشه - واژه فرح در قرآن، تفاوت معنای به‌دست‌آمده از بافت قرآنی (نزدیک به حوزه معنایی تکبر)، نسبت به معنای غالبی ذکر شده در لغت‌نامه‌ها (شادی) را نشان می‌دهد.

واژه‌های کلیدی: معنی‌شناسی توزیعی، بردار بافتی، مشابهت معنایی، قرآن، فرح.

۱. مقدمه

معنی‌شناسی توزیعی^۱ روشی نو ساختگرا در حوزهٔ زبان‌شناسی پیکره‌ای^۲ است؛ که با الهام از برخی مبانی ساختگرایی مانند توجه به روابط هم‌نشینی و جانشینی و بسط آن شکل گرفته، اما در عمل از برخی مبانی آن، مانند تمایز کاربردشناسی از معنی‌شناسی، فاصله گرفته است. این روش به‌عنوان رویکردی کاربردمدار، با استفاده از داده‌های متون واقعی و به‌عبارتی «زبان طبیعی»^۳ به شناخت معنی واژگان با توجه به بافت و نسبت به سایر واژه‌ها می‌پردازد.

یکی از اهداف معنی‌شناسی قرآن، شناخت معنی واژگان قرآنی با استفاده از بافت آن است. احتمال تطور معنایی واژگان از عصر نزول تا ظهور اولین کتاب‌های لغت عربی (حدود دو قرن بعد) و وجود اختلاف نظر میان لغت‌شناسان، ایجاب می‌کند که برای شناخت واژگان قرآن، از روش‌های غیرمستقیم شناخت معنی، مانند رویکردهای تاریخی یا مراجعه به سایر متون استفاده شود. معنی‌شناسی توزیعی یکی از این روش‌هاست؛ که با نگاهی آماری به چگونگی کاربرد واژگان در مقایسه با یکدیگر در متن، حدودی از معنای واژگان به‌دست می‌دهد.

پژوهش حاضر به‌منظور استفاده از معنی‌شناسی توزیعی در قرآن با هدف شناخت معنای واژگان با استفاده از بافت، با روش توصیفی - تحلیلی، در درجهٔ اول به دو سؤال اصلی پاسخ می‌دهد: گام‌های عملیاتی معنی‌شناسی توزیعی در قرآن، با توجه به ویژگی‌های زبانی و متنی آن چیست؟ در به‌کارگیری این روش در قرآن چه چالش‌هایی وجود دارد و راه‌های برون‌رفت از آن‌ها چیست؟ با توجه به هدف پژوهش، پیاده‌سازی معنی‌شناسی توزیعی در قرآن، با ذکر نمونهٔ موردی ریشه - واژه «فرح» به‌طور دستی و غیررایانه‌ای صورت می‌گیرد. بنابراین سؤال دیگری که این پژوهش به صورت ضمنی درصدد پاسخ به آن است، چیست معنای ریشه - واژه فرح در قرآن است که درمورد آن میان لغت‌شناسان اختلاف نظر وجود دارد. فرضیهٔ مرتبط با سؤال اخیر این است که با استفاده از معنی‌شناسی توزیعی حوزهٔ معنایی فرح در قرآن، به آنچه مفسران و پژوهشگران قرآنی بیان کرده‌اند (معنایی نزدیک به تکبر) نزدیک‌تر است؛ تا آنچه لغت‌شناسان بدون لحاظ کاربردهای قرآن ذکر کرده‌اند (شادی). جمع‌آوری اطلاعات پژوهش کتابخانه‌ای است.

۲. پیشینه پژوهش

در زبان فارسی آثار بسیاری با موضوع معنی‌شناسی با رویکردها و روش‌های مختلف عرضه شده است. اگرچه پژوهش‌های حوزه معنی‌شناسی توزیعی در خارج از ایران، با نگرش زبان‌شناسانه، دامنه‌ای گسترده دارد، اما به دلیل نیاز این روش به ابزارهای ریاضیاتی و رایانه‌ای، در ایران به صورت محدودی با این نگاه به آن پرداخته شده است. برای نمونه می‌توان به بخشی از کتاب *نظریه‌های معنی‌شناسی واژگانی*^۱ از دیرک گیرترس (2010)، ترجمه کوروش صفوی (۱۳۹۳) اشاره کرد که ضمن پرداختن به سیر تاریخی معنی‌شناسی، در میان روش‌های نو ساختگرا به «تحلیل توزیع پیکره‌ای» می‌پردازد. از جمله برخی آثار این حوزه پژوهشی که از سوی متخصصان رایانه ارائه شده نیز می‌توان به کتاب‌های *متن‌کاوی یادگیری ماشین*^۲ (آگاروال، ۲۰۱۸، ترجمه ترکیان، ۱۳۹۸) و *شبکه‌ها*^۳ (نیومن، ۲۰۱۸، ترجمه ترکیان، ۱۳۹۸) اشاره کرد. مترجم این دو اثر ضمن ترجمه، به ترتیب به آن‌ها پیوسته‌های «آموزش عملی متن‌کاوی» (ترکیان، ۱۳۹۸ الف) و «شبکه متن» (ترکیان، ۱۳۹۸ ب) را افزوده، و ضمن آن به بررسی مشابهت معنایی آیات قرآن کریم و نیز خوشه‌بندی اسم‌ها در قرآن پرداخته است. اما اثری به زبان فارسی که روش توزیعی را به صورت کامل و با غلبه نگرش زبان‌شناسانه معرفی کند یافت نشد.

در حوزه پژوهش‌های قرآنی نیز، آثار فراوانی به روش‌های متعدد معنی‌شناسی در قرآن پرداخته‌اند. پژوهش‌های قرآنی سال‌های اخیر با هدف شناخت معنی واژگان قرآنی با استفاده از بافت، عمدتاً با نام معنی‌شناسی ساختگرا، یا ساختاری رواج یافته است. از نخستین آثار در معرفی این روش، پایان‌نامه *معناشناسی تدبر در قرآن کریم* است؛ که ضمن تقسیم هم‌نشینان واژه به انواع مکملی، اشتدادی، تقابلی و توزیعی، جانشینان واژه‌ها را براساس هم‌نشینان مکملی شناسایی و به توصیف معنای واژه براساس آن‌ها می‌پردازد (سلمان‌نژاد، ۱۳۹۱). این روش در آثار دیگری نیز از جمله پایان‌نامه *معناشناسی کلمه در قرآن کریم* (شفیع‌زاده، ۱۳۹۱) و رساله *معناشناسی کرم در قرآن* (قاسم‌احمد، ۱۳۹۶) به کار رفته است و در پژوهش‌های مرتبط با شناخت معنا مانند بررسی ترجمه نیز، به نوعی استفاده از آن مشاهده می‌شود (گلی ملک‌آبادی و همکاران، ۱۳۹۶). در برخی پژوهش‌ها تصریح شده که این روش، بومی‌سازی روش‌های معنی‌شناسی در علم زبان‌شناسی است. به نظر می‌رسد روش رایج موسوم به معنی‌شناسی ساختگرا، از جهاتی برگرفته از معنی‌شناسی توزیعی است. اما نقدهایی به آن وارد است که بررسی آن مجال گسترده‌تری را می‌طلبد. اما به طور خلاصه می‌توان

به چند نکته اشاره کرد. اولاً چگونگی استناد کلیت این روش و نیز جزئیات گام‌های اجرایی آن به منابع اصلی مشخص نشده است. به عبارتی، بومی‌سازی یک روش، مستلزم شناخت اصل آن و سپس ایجاد تغییراتی متناسب با فضای مورد استفاده است؛ که در معرفی این روش، مشاهده نمی‌شود. همین امر سبب شده است که اجرای روش موسوم به ساختگرا، با نقدهایی منطقی روبه‌رو باشد. به‌علاوه مبنای روش نام‌برده، یعنی توجه به زبان کاربردی برای دریافت معنی، بر خلاف نام آن، اساساً با ساختگرایی در مفهوم اصیل تفاوت دارد.

نوآوری این پژوهش، پس از تبیین گام‌های معنی‌شناسی توزیعی با غلبه نگرش زبان‌شناسانه، معرفی چگونگی به‌کارگیری آن - به عنوان روشی برای شناخت معنی از بافت - در قرآن، معرفی چالش‌های این کاربرست و راه‌های برون‌رفت از آن، با ذکر نمونه موردی ریشه - واژه فرح است. گفتنی است ریشه - واژه فرح پیش از این با روش موسوم به معنی‌شناسی ساختگرا در پایان‌نامه معنی‌شناسی فرح در قرآن و روایات (ذوعلم، ۱۳۹۷) و مقاله «معنی‌شناسی فرح در قرآن کریم؛ پاسخی به شبهه نكوهش شادی در قرآن» (ذوعلم و دیگران، ۱۳۹۹) بررسی شده است، اما فارغ از اعتبار نتایج پژوهش نام‌برده، در پژوهش حاضر، استفاده از روشی مستند به منابع اصلی و با رعایت گام‌های اصولی معنی‌شناسی توزیعی، در شناخت معنای این واژه مورد نظر است.

۳. مفاهیم نظری: معنی‌شناسی توزیعی و گام‌های اجرای آن

معنی‌شناسی توزیعی رویکردی کاربردمدار به زبان است که با توجه به بافت کاربردی، به شناخت معنی واژه در مقایسه با سایر واژه‌ها می‌پردازد (گیررتس، ۲۰۱۰، ترجمه صفوی، ۱۳۹۸، صص. ۳۴۵ و ۳۴۶). نمونه‌هایی از کاربردهای این روش شامل شناخت معنای واژگان مورد اختلاف، ابهام‌زدایی از واژگان چندمعنا، و مقایسه معنای یک واژه در دوره‌های زمانی مختلف می‌شود (Boleda, 2020, p.217-229). ایده بنیادین در پس معنی‌شناسی توزیعی به زبان، فرضیه توزیعی^۷ است که با توجه به روش توزیعی مطرح شده از سوی هریس^۸ در دهه ۱۹۵۰ به‌وجود آمده است. این فرضیه بیان می‌کند: «واژگان با ویژگی‌های توزیعی مشابه، معانی مشابه دارند» (Sahlgren, 2006, p. 21). بر این اساس اگر دو واژه، ویژگی^۹ توزیعی مشابه در متن داشته باشند، می‌توان گفت به یک دسته واژگانی تعلق دارند (Harris, 1954, p. 156). بعدها دو نوع ویژگی توزیعی برای واژگان، شامل متن (سند^{۱۰})هایی که واژه

در آن وقوع یافته و هم‌نشینان یک واژه معرفی شد. بر پایه ایده هریس در هر نوع، تعدادی ویژگی توزیعی، مشخص شده، مشابهت واژگان براساس هم‌وقوعی^{۱۱} هر واژه با این ویژگی‌ها سنجیده می‌شود. بر این اساس، یک مدل واژه - فضا، یا مدلی محاسباتی از معنی شکل می‌گیرد؛ که ویژگی‌های توزیعی، با ابعاد فضا بازنمایانده می‌شود. در این مدل، شباهت‌های میان واژگان، با مشاهده شواهد تجربی کاربرد واقعی زبان، استخراج می‌شود (Sahlgren, 2006, p. 21). هر چه تعداد ویژگی‌های توزیعی (ابعاد فضای واژه‌ای) بیشتر باشد، پیچیدگی‌های بیشتری از معنای واژه‌ها بازنمایانده می‌شود (Schutze, 1993, p. 896).

در این قسمت گام‌های اجرایی معنی‌شناسی توزیعی معرفی می‌شود.

۱-۳. تشکیل پیکره

اولین گام روش توزیعی تشکیل پیکره به معنای مجموعه متونی با ویژگی‌های موقعیتی مشخص است. ویژگی‌های موقعیتی شامل مشارکت‌کنندگان (گوینده، نویسنده، مخاطبان) و روابط میان آن‌ها، شیوه ارتباط (نوشته، سخنرانی و ...)، شرایط تولید (وجود یا عدم برنامه‌ریزی، ویراستاری و ...)، وضعیت (فضا و زمان)، اهداف ارتباط و موضوع می‌شود (کرافورد و سزومی، ۲۰۱۶، ترجمه صفوی، ۱۳۹۹، صص. ۲۴ و ۲۵). بسته به هدف پژوهش و براساس میزان پرسش‌ها، پیکره‌ها عمومی یا تخصصی هستند (کرافورد و سزومی، ۲۰۱۶، ترجمه صفوی، ۱۳۹۹، صص. ۹۳ و ۹۴). تعداد واژگان برخی پیکره‌های عمومی، از ۱.۹ میلیارد تا ۵۰ میلیون واژه گزارش شده است (کرافورد و سزومی، ۲۰۱۶، ترجمه صفوی، ۱۳۹۹، ص. ۷۲). اما پیکره‌های تخصصی می‌تواند تعداد واژگانی از درجه میلیون واژه داشته باشد.^{۱۲}

۲-۳. تعیین ویژگی‌های توزیعی واژه‌ها

دو دسته ویژگی توزیعی واژه، سندها و واژگان هم‌وقوع با یک واژه هستند؛ که دو الگوی مستقل «واژه - سند» و «واژه - واژه» به دست می‌دهند. دسته اول ویژگی‌ها، علاوه بر سندهایی مانند مقاله می‌تواند به بافت‌های کوتاه‌تری مانند بخشی از یک مقاله یا یک بند در متن تعمیم یابد (Turney, 2010, p. 148). دسته دوم ویژگی‌ها نیز می‌تواند از واژه، به عبارت^{۱۳} (مانند ریشه یا عبارت‌های چندواژه‌ای) گسترش یافته، الگوی واژه - عبارت را تشکیل دهد (Sahlgren, 2006, p. 31) و با لحاظ رابطه دستوری با واژه

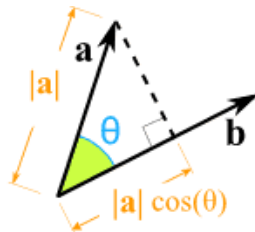
و یا بدون لحاظ این رابطه تعیین شود (Lenci, 2018, p. 158). اگر رابطه دستوری در این الگو در نظر گرفته نشود، قرار گرفتن در محدوده جمله، یا در زنجیره‌ای با اندازه مشخص از واژه‌ها و به عبارتی «پنجره بافتی»^{۱۴} اهمیت دارد (Schutze, 1992, pp. 787, 793)؛^{۱۵} که البته هیچ اصل نظری برای مشخص کردن بهترین اندازه پنجره وجود ندارد (Lenci, 2018, p. 153). در الگوی واژه - سند، مشابهت واژگان براساس هم‌نشینی آن‌ها در سندهای مختلف بررسی می‌شود. اما در الگوی واژه - عبارت، که ویژگی‌های توزیعی، هم‌نشینی واژه هستند، مشابهت جانشینان واژه با آن سنجیده می‌شود.

۳-۳. تشکیل بردار بافتی^{۱۶} و واژه و ماتریس هم‌وقوعی^{۱۷}

«بردار»^{۱۸} در تعریف پاره‌خطی جهت‌دار است که در فضای n بُعدی، به صورت مجموعه‌ای از مؤلفه‌ها، نقطه‌ای از فضا را تعریف می‌کند. بردار بافتی واژه، به صورت مجموعه‌ای از بسامدهای هم‌وقوعی واژه با ویژگی‌های توزیعی تعریف می‌شود. هر مؤلفه این بردار در الگوی واژه - سند، تعداد وقوع واژه در یک سند؛ و در الگوی واژه - عبارت، بسامد هم‌وقوعی واژه با یک عبارت است (Sahlgren, 2006, pp. 29, 30). مجموعه مؤلفه‌های بردارهای بافتی واژه‌ها در متن، به صورت جدولی با محور عمودی واژه‌ها و محور افقی ویژگی‌ها، «ماتریس هم‌وقوعی» را تشکیل می‌دهد.

۴-۳. محاسبه مشابهت^{۱۹} واژگان

مشابهت معنایی واژه‌ها با مقایسه بردارهای بافتی آن‌ها تعیین می‌شود؛ که کاربردی‌ترین معیار آن، تعیین کوسینوس زاویه میان بردارهاست (Schutze, 1992, p. 787). هرچه این زاویه کمتر باشد، کوسینوس آن نزدیک‌تر به یک، و مشابهت دو بردار بیشتر است و اگر دو بردار عمود (بدون مؤلفه مشترک) باشند، این مقدار صفر است.



شکل ۱: ضرب داخلی دو بردار

Figure 1: Inner product of two vectors

کوسینوس زاویه دو بردار، از تقسیم ضرب داخلی آن‌ها (جمع حاصل ضرب مؤلفه‌های متناظر)، بر ضرب اندازه‌های آن‌ها^{۲۰} تعریف می‌شود.

$$a = (a_1, a_2, \dots)$$

$$b = (b_1, b_2, \dots)$$

$$a \cdot b = a_1 b_1 + a_2 b_2 + \dots$$

$$\cos \theta = \frac{a \cdot b}{|a| \cdot |b|} = \frac{a_1 b_1 + a_2 b_2 + \dots}{\sqrt{(a_1^2 + a_2^2 + \dots)(b_1^2 + b_2^2 + \dots)}}$$

و مشابهت بردارهای a و b از این رابطه به دست می‌آید:

$$\text{sim}(a, b) = \cos \theta$$

۵-۳. گام‌های اصلاحی اجرای روش توزیعی

برای بهبود عملکرد روش توزیعی از جهت سهولت اجرا و نتیجه‌گیری مؤثرتر، اقداماتی اصلاحی صورت می‌گیرد که عبارت‌اند از:

۵-۳-۱. پیش‌پردازش^{۲۱} زبانی

این عمل به منظور به کارگیری مؤثرتر عناصر ماتریس هم‌وقوعی، بر پایه اطلاعات زبانی بر روی واژگان صورت می‌گیرد (Sahlgren, 2006, p. 71, 72). نمونه‌هایی از این پیش‌پردازش، ریشه‌یابی^{۲۲} واژگان (Ibid, p. 12)، برچسب زدن پاره گفتار^{۲۳} (به معنی تعیین مقوله‌های واژگان) و ابهام‌زدایی^{۲۴} از واژگان چندمعناست.

۵-۳-۲. کاهش بُعد ابتدایی ماتریس هم‌وقوعی

ماتریس هم‌وقوعی در شکل اولیه، با ابعاد بسیار زیاد، تمام واژگان را بر حسب تمام ویژگی‌ها بیان می‌کند. اما بسیاری از ویژگی‌ها، در تعیین مشابهت واژگان تأثیر اندکی دارند. کاهش بُعد، بازنمایی اطلاعات مفید در فضایی با ابعاد کوچک‌تر است؛ که در ابتدا با اطلاعات زبانی صورت می‌گیرد. حذف واژه‌های متعلق به مقوله‌های دستوری بسته،^{۲۵} و یا دارای خواص آماری نامطلوب (مانند بسیار پرسامد یا کم‌پسامد)^{۲۶} از آن جمله است. هم‌وقوعی واژگان بسیار پرسامد - که معمولاً همان واژه‌های متعلق به

مقوله‌های دستوری بسته هستند - با واژه مورد بررسی، اطلاعات معنایی مفیدی دربر ندارد و از این رو ایست‌واژه^{۳۷} نامیده می‌شوند (Sahlgren, 2006, pp. 38, 39).

۳-۵-۳. وزندهی^{۳۸} عناصر ماتریس هموقوعی

میزان ارزش عناصر مختلف ماتریس هموقوعی، برای شناخت معنی، مساوی نیست. از این رو با انجام عملیاتی روی بسامدهای هموقوعی اولیه (بسامد خام^{۳۹})، هر یک از آن‌ها ارزش‌گذاری می‌شود (Turney, 2010, p.156).

در الگوی واژه - سند، هرچه طول سند شامل واژه کوتاه‌تر باشد، و نیز هرچه واژه در سندهای کم‌تری وقوع یافته باشد، هموقوعی سند و واژه اهمیت بیشتری برای یافتن مشابهت با دیگر واژگان دارد. بر همین اساس نوعی از وزندهی (خانواده TFIDF) بر روی ماتریس پیاده می‌شود. در نمونه‌ای از این وزندهی، بسامد خام هموقوعی واژه i با سند j (TF_{ij}) در دو ضریب IDF_i (تابعی از معکوس تعداد سندهایی دارای عبارت i)^{۴۰} و S_j ^{۴۱} (معکوس تعداد واژگان سند j) ضرب می‌شود و هر عنصر از ماتریس $f_{ij}=TF_{ij}.IDF_i.S_j$ به دست می‌آید (Sahlgren, 2006, pp. 31, 32).

در الگوی واژه - عبارت، اگر بسامد عبارت هموقوع با واژه مورد نظر، در کل متن زیاد نباشد، هموقوعی آن با واژه مهم‌تر است. برای وزندهی، احتمال هموقوع شدن مورد انتظار واژه و عبارت (p_{exp}) ^{۴۲} (براساس بسامد هر یک از آن‌ها در کل متن)، با بسامد هموقوعی مشاهده شده این دو در متن (p_{obs}) ^{۴۳}، مقایسه می‌شود. اگر واژه و عبارتی از لحاظ آماری مستقل باشند، هموقوع شدن آن‌ها کاملاً تصادفی است و انتظار می‌رود احتمال هموقوعی آن‌ها، حاصل ضرب احتمال وقوع هریک در متن باشد (بسامد مورد انتظار). وزندهی باید به گونه‌ای باشد که اگر بسامد هموقوعی مشاهده شده، برابر با بسامد مورد انتظار باشد، این عنصر تأثیری در تعیین معنی واژه نداشته، عنصر متناظر در ماتریس صفر شود (Turney, 2010, p.158). مدل‌های مختلف وزندهی ماتریس هموقوعی بر همین اساس معرفی شده؛ که در نمونه‌ای از آن، موسوم به قاعده PPMI، هر عنصر از ماتریس هموقوعی این‌گونه تعریف می‌شود:

$$\begin{cases} \log_2 \left(\frac{p_{obs}}{p_{exp}} \right) & \text{اگر } p_{obs} > p_{exp} \\ 0 & \text{اگر } p_{obs} \leq p_{exp} \end{cases}$$

۴-۵-۳. کاهش بُعد پیشرفته در ماتریس هم‌قوعی

ماتریس هم‌قوعی به‌طور طبیعی عناصر بسیاری با مقدار صفر دارد؛ بعضی ویژگی‌ها اطلاعات تکراری دارند؛ و ویژگی‌های توزیعی پنهانی وجود دارند که با لحاظ ویژگی‌های صرفاً ظاهری، از دست می‌روند (Lenci, 2018, p.156). بنابراین از روش‌های پیشرفته‌ای برای کاهش بُعد، مانند تجزیه مقدار تکین^{۳۵} (SVD) استفاده می‌شود؛ که با عملیاتی ریاضی، داده‌های کم‌اهمیت ماتریس را حذف و ویژگی‌های مؤثر پنهان را بازمی‌نمایاند (Lenci, 2018, p.156).

۶-۳. ترسیم نقشه معنایی^{۳۶}

نقشه معنایی، شامل تصویر واژگان مورد مقایسه در یک صفحه است، که با نمایش دوری و نزدیکی واژگان نسبت به هم می‌تواند تا حد خوبی، حوزه‌های معنایی واژگان را از یکدیگر متمایز سازد. هر چه دو واژه در این نقشه به هم نزدیک‌تر (همسایه^{۳۷}) باشند، مشابهت معنایی بیشتری دارند.

۷-۳. تفسیر مشابهت بردارهای بافتی واژگان

نتیجه اصلی روش توزیعی، مدل کردن شباهت‌های معنایی است. اما «مشابهت معنایی»، خود یک مفهوم مبهم است؛ که شامل انواع روابط واژگانی مانند شمول معنایی، تقابل، جزءواژگی، روابط مکانی و دیگر روابط دسته‌بندی نشده^{۳۸} می‌شود (Lenci, 2018, p.161). شناخت این رابطه بخش مهمی از معنی‌شناسی توزیعی است. در این میان رابطه «تقابل» یکی از چالش‌های مهم است؛ که تمایل به وقوع در ویژگی‌هایی مشابه «هم‌معنایی» دارد. یکی از راه‌های تمیز آن از هم‌معنایی، استفاده از مدل‌های ترکیبی است که از اطلاعات منابع واژگانی نیز بهره می‌گیرد (Lenci, 2018, p.162).

۴. به‌کارگیری معنی‌شناسی توزیعی در قرآن، بررسی نمونه موردی ریشه -

واژه «فرح»

آغاز به‌کارگیری روش توزیعی، تشکیل پیکره با جمع‌آوری متون طبیعی دارای ویژگی‌های بافتی و موقعیتی مشترک است. ویژگی‌های موقعیتی مشترک در قرآن شامل این موارد است: مشارکت‌کنندگان و رابطه میان آن‌ها (گوینده خداوند، و مخاطب پیامبر و مردم)، شیوه ارتباط (وحی)، شرایط تولید

(انتقال کامل پیام و تنظیم متن با ترتیب توقیفی)، هدف ارتباط (هدایت انسان‌ها) و وضعیت ارتباط (بازه زمانی و مکان مشخص). بر این اساس متن قرآن می‌تواند به عنوان پیکره‌ای طبیعی با رویکرد توزیعی مورد بررسی قرار گیرد.

شناخت معنی واژه با استفاده از اطلاعات توزیع واژه در پیکره قرآنی، عملاً برای واژه‌های مورد اختلاف معنایی سودمند است. به علاوه، از آنجاکه نتیجه معنی‌شناسی توزیعی، شناخت حدودی معنی واژگان در مقایسه با یکدیگر است و تفاوت‌های ظریف معنایی در این روش کمتر مورد توجه قرار می‌گیرد، برای شناخت واژگانی کاربرد دارد که در مورد معنای آن‌ها اختلاف نظر کلی وجود دارد و به عبارتی حوزه معنایی آن‌ها مورد اختلاف است. ویژگی‌های متن قرآن و نیز زبان عربی سبب می‌شود در کاربست قرآنی معنی‌شناسی توزیعی نکاتی مورد توجه قرار گیرد. در این بخش گام‌های پیاده‌سازی روش توزیعی در قرآن، به ترتیب اجرا، با نمونه موردی ریشه - واژه «فرح» بررسی می‌شود؛ که در قرآن ۲۲ بار به دو شکل فعل ثلاثی مجرد (۱۶ بار) و صفت مشبیه فَرِحَ (شش بار) به صورت مفرد یا جمع به کار رفته است (عبدالباقی، ۱۳۹۳، صص. ۶۶۲، ۶۶۳).

۱-۴. انتخاب واژگانی برای سنجش مشابهت با واژه اصلی

از آنجا که در روش غیر رایانه‌ای، عملاً بررسی مشابهت تمام واژگان قرآنی با یک واژه امکان‌پذیر نیست، تعداد محدودی واژه به این منظور انتخاب می‌شود. برای واژه‌ای که در مورد معنی آن اختلاف وجود دارد، معانی ذکرشده در لغت‌نامه‌ها و تفاسیر، و واژگان مشمول حوزه معنایی آن‌ها، می‌تواند برای مقایسه با واژه اصلی گزینش شوند.

لغت‌شناسان برای واژه فَرِحَ، دو معنای اصلی بیان کرده‌اند: ۱ - سرور (ابن درید، ۱۹۸۸، ج ۱، ص. ۵۱۸؛ ازهری، ۱۴۲۱ق، ج ۵، ص. ۱۵؛ جوهری، ۱۳۷۶ق، ج ۱، ص. ۳۹۰؛ آذرنوش، ۱۳۸۳، ص. ۵۰۰)؛ ۲ - تکبر و واژگان حوزه معنایی آن (ابوعبیده، ۱۳۸۱ق، ج ۲، ص. ۱۱۱؛ ابن قتیبه، ۱۴۱۱ق، ص. ۲۸۶؛ جوهری، ۱۳۷۶ق، طبرسی، ۱۳۷۲، ج ۷، ص. ۴۱۶). بنابراین واژگان حوزه‌های معنایی سرور (مانند استبشار، ضحک، رضایت) و تکبر (مانند مرح، فخر، اختیال، علو، بغی) و نیز واژگان و اصطلاحات نزدیک به این حوزه‌ها، برای مقایسه با فرح مورد توجه هستند. شایان ذکر است با دقت در ارتباط میان معنای ذکرشده و نوع لغت‌نامه، به نظر می‌رسد لغت‌شناسان، و به ویژه در دوره معاصر بیشتر به معنای

سرور برای فرح تأکید دارند و معنای تکبر غالباً از سوی مفسران یا لغت‌شناسانی بیان شده است، که کاربردهای قرآنی را مبنای شناخت معنا دانسته‌اند.

۲-۴. تعیین نوع ویژگی توزیعی

در الگوی واژه - سند، چگونگی مطابقت بخش‌های قرآن با سندها اهمیت دارد. در پیکره‌های زبانی مورد بررسی در روش توزیعی، سندها متن‌هایی جداگانه دارای ماهیتی نسبتاً مستقل‌اند؛ به‌عنوان نمونه مقالات یک نشریه (نک: Dumaïs et.al, 1988). در قرآن اگرچه برخی، واحد «آیه» را به‌عنوان سند برگزیده‌اند (ترکیان، ۱۳۹۸ الف، ص. ۶۲۳)، اما در بسیاری موارد آیه‌ها، وابسته به هم و حتی گاهی مکمل یکدیگر برای تشکیل جمله‌اند.^{۳۹} با توجه به قرارگیری توقیفی آیات در واحدهای مستقل سوره، سوره به‌عنوان سند، گزینه مناسب‌تری به نظر می‌رسد.

علاوه بر مفهوم اصطلاحی سوره که بخش‌های متمایز آیات (عمدتاً همراه با بسمله) است، برخی پژوهشگران «رکوعات قرآنی» را با عنوان «سوره‌وار»، برای گونه‌ای از دسته‌بندی آیات با واحدهایی مستقل مشابه سوره، از زمان پیامبر (ص) معرفی می‌کنند. در این دسته‌بندی مجموعه آیات قرآن به ۵۵۵ رکوع تقسیم می‌شود که در برخی نسخه‌های قدیمی با حرف «ع» از هم جدا شده‌اند (لسانی فشارکی و مرادی زنجانی، ۱۳۹۵، صص. ۸۸ - ۹۶). هرچند با دقت در محدوده رکوعات قرآنی (لسانی فشارکی و مرادی زنجانی، ۱۳۹۵، صص. ۲۲۸ - ۲۳۲) به نظر می‌رسد هدف از این دسته‌بندی، سهولت در قرائت و حفظ آیات، و نه دسته‌بندی موضوعی است، اما به‌هرحال سبب تشکیل سندهایی بسیار کوچک‌تر از سوره و پررنگ شدن نقش هم‌وقوعی دو واژه نزدیک در متن به‌ویژه برای سوره‌های طولانی می‌شود. این روش را می‌توان معادل تعمیم سندها به بخش‌ها یا بندهایی از یک مقاله در این الگو (Turney, 2010, p. 148) دانست.

در الگوی واژه - عبارت، دو نوع هم‌وقوعی واژگان می‌تواند مورد توجه باشد: قرارگیری واژه‌ها در یک پنجره بافتی، یا رابطه نحوی میان واژگان. در نوع اخیر می‌توان علاوه بر روابط نحوی تعریف شده (مانند رابطه «فاعل و فعل»)، رابطه‌های دیگری را نیز، براساس ارتباط معنایی جملات (مانند شرط و جواب شرط) در نظر گرفت. شناخت این نوع هم‌وقوعی به دلیل لزوم توجه به معنا دشوارتر از نوع اول است. در این پژوهش برای معنی‌شناسی ریشه - واژه فرح در قرآن، الگوی واژه - عبارت (بدون لحاظ

روابط دستوری) و الگوی واژه - سند (معادل سوره) مورد توجه قرار می‌گیرد.

۳-۴. پیش‌پردازش زبانی متن

در ریشه‌یابی واژگان قرآنی برای قرار گرفتن در ماتریس هم‌وقوعی، ریشه‌محور بودن ساخت واژگان در زبان عربی نقش مهمی دارد. وجود باب‌ها و صیغه‌های متعدد از یک ریشه، معانی متفاوتی ایجاد می‌کند که گاهی تنها برخی از آن‌ها مورد نظر است. به عنوان مثال برای بررسی مشابهت واژگان حوزه معنایی سرور با فرح، مصدر «استبشار» به معنای «سرور به سبب بشارت» (عسکری، ۱۴۰۰ق، ص. ۲۶۰) و واژگان مرتبط با این مصدر مورد توجه قرار می‌گیرد و سایر باب‌های ریشه بشر به دلیل ارتباط اندک معنایی بررسی نمی‌شود. از ریشه کبر نیز، مصدرهای «تکبر» و «استکبار»، با اندکی اختلاف معنایی به معنای «خودبزرگ‌بینی» (نک: ایزوتسو، ۱۳۷۸، ص. ۲۸۹)، مورد نظر است و از سایر واژگان این ریشه صرف‌نظر می‌شود. در ویژگی‌های توزیعی فرح در ماتریس واژه - عبارت نیز، از میان واژه‌های ریشه فضل، اسم مصدر «فَضْل» به دلیل بار معنایی ویژه‌ای که آن را از صیغه‌های دیگر این ریشه متمایز می‌کند، انتخاب می‌شود. در توجه به یک صیغه ویژه از یک ریشه، برخی معتقدند ریشه‌یابی اسم‌ها در قرآن تنها تبدیل جمع به مفرد است و شامل یافتن حروف اصلی واژه نمی‌شود (ترکیان، ۳۹۸ الف، ص. ۶۱۹ و ۶۲۷). با این حال در بسیاری موارد، کاربرد صیغه‌های فعلی و اسمی یک ریشه، تفاوت عمده‌ای ندارد.^{۴۰} به‌طور کلی بسته به مورد پژوهش، توجه به کل واژگان ریشه و یا لحاظ واژه یا واژه‌هایی از یک ریشه، می‌تواند مقبول باشد.

در ابهام‌زدایی، واژگانی با ظاهر مشابه اما با معانی مختلف^{۴۱} از هم متمایز؛ و به تناسب نیاز تنها برخی در ماتریس هم‌وقوعی بررسی می‌شوند. از واژه‌های مورد مقایسه با فرح، ریشه «سرر» در قرآن دارای دو معنای «پنهان کردن» و «شادی» است؛ که تنها معنای دوم، با فرح مقایسه می‌شود. ریشه «بَغَى» نیز در شکل ثلاثی مجرد در اصل به معنای طلب و در بسیاری کاربردها «طلب تجاوز از میانه‌روی» است؛ که به معنای دیگری مانند ظلم، زنا، حسد و تکبر گسترش یافته (فراهیدی، ۱۴۰۹ق، ج ۴، ص. ۴۵۳؛ ازهری، ۱۴۲۱ق، ج ۸، ص. ۱۷۹؛ راغب، ۱۴۱۲ق، ج ۱، صص. ۱۳۶ و ۱۳۷) و به علاوه گاهی به معنای «خواستن و جست‌وجو کردن»^{۴۲} (متعدی) و گاه به معنای «چیزی را به صورتی خواستن»^{۴۳} (دومفعولی) به کار رفته است. برخی از این معانی در معنی‌شناسی فرح حذف می‌شوند. در میان

ویژگی‌های توزیعی فرح در الگوی واژه - عبارت نیز، واژه «سیئه» تنها به معنای «اتفاق ناگوار» و نه به معنای «عمل بد» مورد توجه است.

علاوه بر ریشه‌یابی و ابهام‌زدایی، ادغام واژه‌های هم‌معنا یا قریب‌المعنا نیز از جمله موارد پیش‌پردازش زبانی پیکره است. در ماتریس هم‌وقوعی فرح، می‌توان واژه‌های جهنم و عذاب، و نیز واژه‌های اعراض، تَوَلَّى، ادبار و تصعیر خَد (به معنی رویگردانی) را ادغام کرد.

۴-۴. کاهش بُعد ابتدایی ماتریس هم‌وقوعی

کاهش بُعد ابتدایی، شامل حذف ایست‌واژه‌هاست که به دلیل هم‌وقوعی با بسیاری واژگان، هم‌نشینی آن‌ها، اطلاعات چندانی برای شناخت معنا در بر ندارد. واژه‌های متعلق به مقوله‌های دستوری بسته، مانند حروف، ضمایر و افعال ناقصه از این جمله اند. در برخی بررسی‌های زبان‌شناسانه قرآن، تعدادی از این ایست‌واژه‌ها (با عنوان «واژه‌های توقف») ذکر شده که برخی ترکیب دو یا چند ایست‌واژه‌اند (ترکیان، ۱۳۹۸، ص. ۶۲۳).

جدول ۱: تعدادی از ایست‌واژه‌ها در قرآن

Table 1: some of the stopwords in the Qur'an

نمونه واژه‌های توقف						
الا	إِنه	علینا	فلا	کل	لنا	منها
الذی	أم	علیه	فهم	کله	له	منهم
الذین	أولئک	علیهم	فی	کلهم	له	منهما
إذ	أیها	عند	فیه	کما	لهم	هم
إذا	بما	عندنا	ق	کنتم	لهم	هما
إن	به	عنده	قال	لا	ما	هو
إنما	یهم	عندی	قالوا	لک	مما	هی
إنتا	بین	عنه	قل	لک	من	والذی
إینه	ثم	عنها	قول	لکم	منا	والذین
	علی	عنهم	قیل	لم	منه	وإن
					ذلك	

۵-۴. تشکیل ماتریس هم‌وقوعی و وزن‌دهی عناصر آن

تشکیل ماتریس هم‌وقوعی، در دو الگوی ذکرشده، پس از پیش‌پردازش زبانی و حذف ایست‌واژه‌ها، با شمارش هم‌وقوعی‌ها صورت می‌گیرد. در ماتریس واژه - عبارت برای ریشه فرح، پنجره هم‌وقوعی ۱۵ واژه‌ای پیشنهاد می‌شود. پس از آن وزن‌دهی عناصر این دو ماتریس مورد نظر است.

۱-۵-۴. ماتریس واژه - سوره

در وزن‌دهی TFIDF در ماتریس واژه - سند (در قرآن واژه - سوره)، پس از به‌دست آوردن بسامد هم‌وقوعی واژه در سوره‌ها، دو ضریب اعمال می‌شود: ۱ - واژه IDF، که نشان می‌دهد که هرچه واژه یا عبارتی در سوره‌های کم‌تری تکرار شده باشد از نظر زبان‌شناسی اهمیت بیشتری دارد؛^{۴۴} ۲ - سوره S، که معکوس طول سوره^{۴۵} است و نشان می‌دهد هم‌وقوعی دو واژه در سوره کوتاه‌تر، بیانگر مشابهت بیشتر آن‌هاست.

درمجموع عنصر متناظر با یک سوره و یک واژه در ماتریس هم‌وقوعی، چنین به‌دست می‌آید:

$$TFIDF = \frac{\text{تعداد وقوع واژه در سوره}}{\text{طول سوره}} \times \log_2 \left(\frac{\text{تعداد کل سوره‌ها}}{\text{تعداد سوره‌هایی که واژه در آن آمده}} \right)$$

برای نمونه ابتدا ماتریس هم‌وقوعی واژه - سوره فرح، برای مقایسه با واژه‌های فخر، اختیال، مرح و سرور، بدون اعمال ضریب نشان داده می‌شود. در زیر نام هر سوره، تعداد واژگان آن (فرهنگ نامۀ علوم قرآنی، ۱۳۹۴، صص. ۲۶۲۱ - ۲۷۳۳ و ۲۷۵۱)؛ و در مقابل نام واژه‌ها تعداد سوره‌هایی که این واژه‌ها در آن‌ها وقوع یافته (نک: عبدالباقی، ۱۳۹۳) درج شده است.

جدول ۲: ماتریس واژه - سوره فرح با بسامدهای خام

Table 2: Word-Sura matrix of Farah with raw frequencies

	انشق اق ۱۰۹	انسان ۲۴۰	حدید ۵۴۴	شوری ۸۶۶	غافر ۱۱۹۹	لقمان ۵۴۲	روم ۸۱۹	قص ص ۱۴۴۱	نمل ۱۱۴۹	مؤمنون ۱۸۴۰	اسراء ۱۵۳۱	رعد ۸۵۵	هود ۱۷۱۵	یونس ۱۸۳۱	توبه ۴۰۹۸	انعام ۳۸۵۰	نساء ۳۷۴۵	آل عمران ۳۴۸۰	بقره ۶۲۲۱
فرح (۱۳)	۰	۰	۱	۱	۲	۰	۳	۲	۱	۱	۰	۲	۱	۲	۲	۱	۰	۳	۰
فخر (۴)	۰	۰	۲	۰	۰	۱	۰	۰	۰	۰	۰	۰	۱	۰	۰	۰	۱	۰	۰
اختیال (۳)	۰	۰	۱	۰	۰	۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱	۰	۰
مرح (۳)	۰	۰	۰	۰	۱	۱	۰	۰	۰	۰	۱	۰	۰	۰	۰	۰	۰	۰	۰

انشقاق اقی ۱۰۹	انسان ۲۴۰	حدید ۵۴۴	شوری ۸۶۶	غافر ۱۱۹۹	لقمان ۵۴۲	روم ۸۱۹	قص ص ۱۴۴۱	نمل ۱۱۴۹	مؤمنون ۱۸۴۰	اسراء ۱۵۳۱	رعد ۸۵۵	هود ۱۷۱۵	یونس ۱۸۳۱	توبه ۴۰۹۸	انعام ۳۸۵۰	نساء ۳۷۴۵	آل عمران ۳۴۸۰	بقره ۶۲۲۱
۲	۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱	۰	۰	۱	۱
سور (۵)																		

ماتریس وزن‌دهی شده چنین به‌دست می‌آید (در زیر نام هر سوره، طول آن^۶ و در مقابل نام واژه‌های مورد

مقایسه با فرح، عامل دوم در فرمول وزن‌دهی، $\log_2 \left(\frac{114}{\text{تعداد سوره‌هایی که واژه در آن آمده}} \right)$ ذکر شده است):

جدول ۳: ماتریس واژه - سوره فرح وزن‌دهی‌شده

Table 3: Weighted Word- Sura matrix of Farah

انشقاق ۲	انسان ۴	حدید ۱	شوری ۱۶	غافر ۲۲	لقمان ۱۹	روم ۱۹	قص ۲۶	نمل ۲۴	مؤمنون ۲۴	اسراء ۱۶	رعد ۱۶	هود ۲۲	یونس ۲۴	توبه ۲۴	انعام ۲۲	نساء ۲۴	آل عمران ۱۶	بقره ۱۶
فرح (۳۰)	۰	۰	۳۰۱	۱۰۹۴	۲۸۲	۶۰۷	۲۰۲۸	۱۰۴۸	۰۰۹۱	۰	۳۸۸	۰۰۹۷	۱۸۲	۰۰۸۳	۰۰۴۴	۰	۱۰۴۵	۰
فخر (۴۸)	۰	۰	۹۰۶	۰	۰	۰	۰	۰	۰	۰	۰	۱۰۵	۰	۰	۰	۰۰۷۰	۰	۰
اختیال (۵۰)	۰	۰	۵۰۲	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰۰۷۵	۰	۰
مرح (۵۰)	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱۰۸۶	۰	۰	۰	۰	۰	۰	۰	۰
سورور (۴۵)	۴۵	۱۱۰۲۵	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰۰۷۰	۰۰۳۹

۴-۵-۲. ماتریس واژه - عبارت

در این پژوهش الگوی پنجره بافتی (بدون لحاظ رابطه دستوری) برای هم‌وقوعی واژگان مورد نظر است؛ که پس از حذف ایست‌واژه‌ها محاسبه می‌شود. برای وزن‌دهی ماتریس واژه - عبارت، بسامد مشاهده شده هم‌وقوعی دو واژه (p_{obs}) با بسامد مورد انتظار برای هم‌وقوعی (p_{exp}) آن‌ها مقایسه می‌شود.

در پیکره‌ای دارای N واژه، که واژه الف در آن m بار تکرار شده است، احتمال وقوع واژه الف در پنجره بافتی n واژه‌ای، با فرض بزرگ بودن مقدار N ، برابر است با $\frac{n \times m}{N}$ ؛ که اگر دو واژه الف و ب، با بسامد وقوع m و m' در متن، از لحاظ آماری مستقل باشند، انتظار می‌رود احتمال وقوع هر دو آن‌ها در یک پنجره بافتی n واژه‌ای (p_{exp})، حاصل ضرب احتمال وقوع هر یک در آن باشد. احتمال مشاهده شده هم‌وقوعی دو واژه الف و ب (p_{obs}) نیز، از تقسیم بسامد هم‌وقوعی مشاهده شده (f)، بر تعداد پنجره‌های بافتی به دست می‌آید ($p_{obs} = \frac{n \times f}{N}$)؛ که در محاسبه هم‌وقوعی دو واژه، در صورت نادیده گرفتن جهت هم‌وقوعی، یک پنجره از طرف راست و یک پنجره از طرف چپ واژه اصلی لحاظ می‌شود. براساس قاعده PPMI، عناصر ماتریس هم‌وقوعی (برای دو واژه با بسامد وقوع m و m' در متن و بسامد هم‌وقوعی f در پنجره بافتی) چنین به دست می‌آید:

$$\begin{cases} \log_2 \left(\frac{p_{obs}}{p_{exp}} \right) = \log_2 \left(\frac{f \times N}{n \times m \times m'} \right) & \text{اگر } p_{obs} > p_{exp} \\ 0 & \text{اگر } p_{obs} \leq p_{exp} \end{cases}$$

برای پیاده‌سازی در قرآن، به منظور حذف ایست‌واژه‌ها فرض می‌شود نسبت آن‌ها به کل واژه‌ها در بخش‌های مختلف پیکره قرآنی تقریباً مساوی است. بنابراین نسبت «تعداد کل واژگان» به «تعداد واژگان باقی‌مانده پس از حذف ایست‌واژه‌ها»، برای تعدادی از سوره‌ها محاسبه، و به کل قرآن تعمیم داده می‌شود. این نسبت به‌طور تقریبی حدود ۶۷ درصد به دست می‌آید^۷ و بنابراین اندازه پیکره قرآنی با حذف ایست‌واژه‌ها، شامل ۵۲۱۳۰ واژه از تعداد کل ۷۷۸۰۷ واژه (فرهنگ نامه علوم قرآنی، ۱۳۹۴، ص. ۱۷) است. وزن‌دهی عناصر ماتریس هم‌وقوعی براساس قاعده PPMI، برای نمونه با احتساب پنجره

هم‌وقوعی ۱۵ واژه‌ای (n=۱۵) در قرآن، چنین به‌دست می‌آید:

$$\begin{cases} \log_2 \left(\frac{f \times 3475}{m \times m'} \right) & \text{اگر } f \times 3475 > m \times m' \\ 0 & \text{اگر } f \times 3475 \leq m \times m' \end{cases}$$

به‌کارگیری دقیق معنی‌شناسی توزیعی مستلزم درنظر گرفتن تمام هم‌نشینان واژگان مورد مقایسه است. هرچه ویژگی‌های توزیعی بیشتری در ماتریس لحاظ شود، طول بردار بافتی واژه مورد مقایسه افزایش، و طبق محاسبات، مشابهت دو واژه کمتر به‌دست می‌آید. واژه‌های مورد مقایسه با فرح، براساس اشتراک با آن در هم‌نشینی با واژه‌ها و عبارت‌هایی (ویژگی‌های توزیعی)، مورد مقایسه قرار می‌گیرند. تعدادی از عبارت‌های چندواژه‌ای که به‌عنوان ویژگی توزیعی فرح استفاده شده است عبارت‌اند از: «فی الارض»، «الله لایحب» و «اذاقه (چشاندن) رحمه».

نمونه کوچکی از ماتریس هم‌وقوعی واژه - عبارت برای مقایسه فرح، فخر، اختیال، مَرَح و سرور در زیر آمده است. به‌دلیل محدودیت اندازه‌گیری بدون نرم‌افزار، تنها تعداد کمی ویژگی توزیعی بیان شده است. در مقابل هر واژه میزان تکرار آن در کل قرآن ذکر شده است.

جدول ۴: ماتریس واژه - عبارت فرح، بسامد خام

Table 4: Word- Term matrix of Farah with raw frequencies

لوزن (۹)	شُرک (۱۶۱)	استیشار (۸)	سینه (۱۴)	روگردانی (۳۴)	هزو (۳۴)	فساد (۵۰)	بغی (۲۴)	اختیال (۵)	فخر (۵)	یا (۳)	مَرَح (۳)	عذاب و جهنم	فی الارض (۵) لایحب (۲۲)	فضل	چشاندن (اذاقه)	فرح (۲۲)
۰	۶	۱	۴	۲	۲	۲	۱	۱	۲	۱	۱	۴	۲	۲	۳	۳

فخر (۵)	۱	۱	۳	۱	۲	۱	۰	۳	۰	۰	۱	۲	۰	۰	۰	۰
اختیال (۳)	۰	۱	۳	۱	۱	۱	۰	۳	۰	۰	۰	۲	۰	۰	۰	۰
مرح (۴)	۰	۰	۱	۳	۱		۱	۱	۰	۰	۰	۱	۰	۰	۱	۰
سرور (۶)	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱	۰	۰	۱

اعمال قاعده PPMI برای وزندهی به ماتریس هموقوعی، ماتریس زیر را نتیجه می‌دهد:

جدول ۵: ماتریس واژه - عبارت فرح وزندهی‌شده

Table 5: Weighted Word- Term matrix of Farah

چشاندن (اناقه)	رحمت	فضل	الله (۵) لایحب	فی الارض	عذاب و جهنم	مرح	کبر	فخر	اختیال	بغی	فساد	هزو	رویگردانی	سینه	استبشار	شرک	لون
فرح	۶.۳	۲.۵	۳.۸	۱.۴	۰.۵	۵.۷	۱.۵	۶.۰	۵.۷	۲.۷	۲.۷	۳.۲	۱.۸	۵.۵	۴.۳	۲.۶	۰
فخر	۶.۴	۳.۰	۶.۶	۲.۰	۱.۶	۷.۹	۰		۹.۴	۰	۰	۱	۳.۳	۰	۰	۰	۰
اختیال	۰	۳.۸	۷.۳	۲.۷	۱.۴	۸.۶	۰	۹.۴		۰	۰	۰	۴.۱	۰	۰	۰	۰

لون	تُرک	استبشار	سینه	روگردانی	هزو	فساد	بغی	انقیال	فخر	بکا	فرح	عذاب و جهنم	فی الارض	الله (ه) لا یجبر	فضل	رحمت	چشاندن (اذاقه)
۰	۲۸	۰	۰	۳۱	۰	۰	۰	۸۶	۷۹	۴۳	۱۴	۱۴	۴۳	۵۷	۰	۰	مرح
۶	۰	۰	۵۴	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	سرور

۴-۶. محاسبه مشابهت واژگان و تفسیر آن

پس از وزندهی به عناصر ماتریس هموقوعی، مشابهت واژگان براساس مشابهت بردارهای وزندهی شده بررسی می‌شود. به‌عنوان یک نمونه مشابهت واژه‌های فخر و فرح در الگوی واژه - سوره، با محاسبه کوسینوس زاویه بین بردارهای آن‌ها به‌دست می‌آید.

$$\cos(\text{فخر و فرح}) = \frac{\text{ضرب داخلی بردارهای فخر و فرح}}{\text{اندازه بردار فخر} \times \text{اندازه بردار فرح}}$$

$$\text{مؤلفه } m \text{ بردار بافتی فخر} \times \text{مؤلفه } n \text{ بردار بافتی فرح} = \sum_n \text{ضرب داخلی بردارهای فخر و فرح}$$

$$\text{اندازه بردار بافتی یک واژه} = \sqrt{\sum (\text{هر مؤلفه بردار بافتی واژه})^2}$$

$$\text{ضرب داخلی بردارهای فخر و فرح} = 0/97 \times 1/5 + 3/1 \times 9/6 = 31/22$$

$$\text{اندازه بردار بافتی فرح} = 9/44$$

$$\text{اندازه بردار بافتی فخر} = 10/86$$

$$\rightarrow \cos(\text{فخر و فرح}) = 0/30$$

این روش برای محاسبه مشابهت واژگان با فرح در الگوی واژه سوره و نیز وزندهی PPMI در الگوی واژه - عبارت، با تعداد بیشتری واژه اجرا شده و نتیجه در جدول ۶ آمده است:

جدول ۶: مشابهت فرح با واژگان دیگر

Table 6: Similarity of Farah with other words

واژه‌های مورد مقایسه	مشابهت از طریق الگوی واژه - سوره	مشابهت از طریق الگوی واژه - عبارت (وزن‌دهی PPMI)
فخر و فرح	۰.۳۰	0.7۳
اختیال و فرح	۰.۲۳	0.61
مرح و فرح	۰.۱۴	0.57
کبر و فرح	۰.۱۷	۰.۲۸
بغی و فرح	۰.۳۰	۰.۴۳
علو و فرح	۰.۱۵	۰.۲۰
کفر و فرح	۰.۲۶	۰.۷۹
ظلم و فرح	۰.۳۴	۰.۳۹
شرک و فرح	۰.۴۳	۰.۵۷
سرور و فرح	۰.۰۰	۰.۲۷
استبشار و فرح	۰.۱۲	۰.۱۴
رضا و فرح	۰.۰۴	۰.۲۹
ضحک و فرح	۰.۰۳	۰.۲۵
غرور و فرح	۰.۱۷	۰.۲۲

علاوه بر مشابهت فرح با دیگر واژه‌ها، مشابهت سایر واژگان با یکدیگر نیز قابل محاسبه است؛ که بر این اساس، نرم‌افزارهایی قادر به ترسیم نقشه معنایی (در هر الگو به صورت جداگانه) هستند. فاصله واژه‌ها در نقشه معنایی، حوزه معنایی آن‌ها را مشخص می‌کند؛ که از گام‌های ابتدایی شناخت معنی است. به عبارتی معنی‌شناسی توزیعی، در جایگاه یک روش آماری برای شناخت معنا، تنها حدود معنایی یک واژه را در مقایسه با دیگر واژگان مشخص می‌کند و به دست آوردن معنای دقیق آن، مستلزم استفاده از روش‌های دیگر است. با مقایسه شباهت فرح با واژگان مختلف، اجمالاً و در نگاهی کلی نتیجه می‌شود که این واژه، در مقایسه با حوزه معنایی سرور، با واژگان حوزه معنایی کبر مشابهت بیشتری دارد.^{۴۸} تفاوت نتایج استفاده از الگوهای مختلف محاسبه مشابهت معنا در روش توزیعی، که مرتبط با چالش‌های استفاده از این روش در قرآن است، در ادامه بررسی می‌شود.

bike	1			
car	0.16	1		
dog	0	0	1	
lion	0	0	0.17	1
	bike	car	dog	lion

۵. چالش‌های استفاده از معنی‌شناسی توزیعی در قرآن و راهکارهای برون رفت از آن

به‌کارگیری معنی‌شناسی توزیعی دارای چالش‌هایی است؛ که برخی به ماهیت روش توزیعی و برخی به ویژگی‌های پیکره مورد بررسی مانند زبان، حجم پیکره و ... مربوط می‌شود.

۵-۱. محدودیت‌های معنی‌شناسی توزیعی

یکی از مشکلات استفاده از روش توزیعی، پیش‌پردازش زبانی پیکره است که باید به‌صورت دستی و بدون استفاده از نرم‌افزار انجام گیرد. ابهام‌زدایی و لحاظ روابط دستوری در الگوی واژه - عبارت از این جمله است.

از دیگر چالش‌های کلی معنی‌شناسی توزیعی، تفسیر مشابهت میان واژگان است. به‌کارگیری روش توزیعی، تنها حوزه معنایی یک واژه را - آن هم به صورت تقریبی و نه قطعی - مشخص می‌کند و برای شناخت دقیق معنا کافی نیست. شناخت معنی، نیازمند ترکیب با روش‌های دیگر است. به‌عنوان مثال تشخیص تقابل معنایی دو واژه، با استفاده از منابع لغتی؛ و تشخیص دیگر روابط معنایی واژگان، ضمن استفاده از روش‌های دیگری مانند توجه به روابط معنایی میان جملات میسر است (رک: فائز و همکاران، ۱۳۹۸، صص. ۱۰۵ - ۱۱۱).^{۴۹}

محدودیت دیگر روش توزیعی با وجود همه توانایی‌ها برای شناخت معنی، این است که این روش ادعا نمی‌کند که می‌تواند دقیقاً به مراد گوینده کلام پی برد و اثرات ارجاعی کلام که معنای مورد نظر گوینده است، در آن مدل‌سازی نمی‌شود (Boleda, 2020, p.224).^{۵۰}

۵-۲. بررسی چالش‌های به‌کارگیری روش توزیعی در قرآن

برخی از چالش‌های به‌کارگیری معنی‌شناسی توزیعی در قرآن به هم‌خوانی شرایط اجرای معنی‌شناسی

توزیعی با ویژگی‌های متن قرآن؛ و برخی به عدم وجود ابزارهای لازم برای اجرای روش بازمی‌گردد. از مشکلات مربوط به همخوانی شرایط اجرای روش با ویژگی‌های متن، حجم کم متن قرآنی (حدود ۷۸۰۰۰ واژه) نسبت به پیکره‌های معمول برای پیاده‌سازی زبان‌شناسی پیکره‌ای، به‌عنوان روشی آماری و نیازمند به تعداد زیاد داده است. در غیر این صورت داده‌های توزیعی به تصادفی بودن نزدیک می‌شود. با این حال می‌توان در سطحی کوچک و با حذف داده‌های بسیار کم‌بسامد به‌نوعی آن را اجرا کرد. برای پیاده‌سازی این روش در قرآن، حداقل بسامد ۱۰ بار تکرار برای واژه اصلی، و سه بار تکرار برای واژگان مورد مقایسه با آن پیشنهاد می‌شود. در هر صورت، زیاد بودن بسامد واژگان در قرآن (تا حدی که آن‌ها را به ایست‌واژه تبدیل نکند) و نیز بسامد هم‌وقوعی واژگان، به نتایج واقعی‌تری برای شناخت معنی می‌انجامد. در نمونه مورد بررسی (ریشه - واژه فرح) نیز، نتایج به‌دست آمده تنها تا اندازه‌ای حدود معنایی فرح را بیان می‌کند و به‌دلیل زیاد نبودن بسامد آن در قرآن، نمی‌توان دقت زیادی انتظار داشت. مشابهت به‌دست آمده برای واژه فرح و واژگان کم‌بسامد، دقت کم‌تری نیز دارد. نکته شایان توجه این است که آنچه در به‌کارگیری معنی‌شناسی توزیعی در قرآن، در درجه اول اهمیت دارد، استفاده از «مبنای روش توزیعی»، یا همان فرضیه توزیعی است که با بیان «مشابهت معنایی واژگان دارای توزیع مشابه»، می‌تواند در قرآن نیز مورد توجه قرار گیرد. از آنجا که مبنای این فرضیه زبانی، سخن‌گویی حکیمانه انسان است و حکمت‌گوینده، سبب اعتبار داده‌های توزیع واژگان می‌شود، حکیمانه‌ترین سخن‌ها، سخن پروردگار، دارای داده‌های معتبرتری برای سنجش آماری است.

مشکل مهم دیگر در اجرای روش توزیعی در قرآن مربوط به ابزار به‌کارگیری روش است. در زبان عربی و به‌ویژه در قرآن با وجود نرم‌افزارهایی برای یافتن ریشه واژه و بسامد آن در سوره‌ها، نرم‌افزارهای کارآمدی برای برخی محاسبات (مانند حذف ایست‌واژه‌ها) موجود نیست (ترکیان، ۱۳۹۸ الف، ص. ۶۴۶). جداسازی واژگان باب‌های مختلف یک ریشه، و تشخیص معانی متفاوت یک واژه نیز در قرآن عملاً با نرم‌افزار امکان‌پذیر نیست و بنابراین لازم است بسیاری از محاسبات به‌صورت دستی و غیر خودکار انجام شود. براساس این کاستی در الگوی واژه - عبارت با محاسبات دستی، عملاً امکان به‌شمار آوردن تمام هم‌وقوعی‌ها وجود ندارد و بنابراین نتایج از دقت لازم برخوردار نیست. به‌عنوان نمونه در پژوهش حاضر، تنها برخی هم‌نشینان فرح مورد توجه قرار گرفته و هم‌نشینان اختصاصی سایر واژه‌های مورد مقایسه با فرح به حساب نیامده است. بنابراین طول بردار بافتی همه واژگان مورد مقایسه، کم‌تر از میزان واقعی، و در پی آن مشابهت واژه فرح با واژگان دیگر عمدتاً بیشتر از مقدار واقعی به‌دست می‌آید.

اما در الگوی واژه - سوره، همهٔ سندها (سوره‌ها) لحاظ می‌شود و از این رو نتایج، مشابهت واژه‌ها را در الگوی واژه - سوره کمتر از الگوی واژه - عبارت به دست می‌دهد.

علاوه بر این، در برخی گام‌های اجرای روش توزیعی، مانند کاهش بُعد با روش SVD نیز آشنایی با نرم‌افزارهای محاسباتی لازم است و یکی از علل عدم ورود زبان‌شناسان و مشابه آن پژوهشگران قرآنی به حوزهٔ معنی‌شناسی توزیعی، عدم آشنایی با این ابزارهاست.

مشکل دیگر مرتبط با شرایط اجرای روش توزیعی، که در الگوی واژه - سند، با لحاظ سوره (در معنای اصطلاحی آن) به عنوان سند نمود می‌یابد، تفاوت بسیار زیاد طول سوره‌هاست. با وجود لحاظ تفاوت طول سوره در وزندهی عناصر ماتریس هم‌وقوعی، عملاً نقش قرارگیری واژه‌ها در سوره‌های طولانی، به سبب عدم توجه به دوری و نزدیکی آن‌ها به هم، نادیده گرفته شده، به ویژه میزان مشابهت به‌دست‌آمده برای واژه‌هایی که در سوره‌های کوتاه به‌کار رفته‌اند، نسبت به دیگر واژه‌ها دورتر از واقعیت نشان داده می‌شود. در نمونهٔ مورد بررسی، به دلیل عدم کاربرد واژهٔ فرح در سوره‌های کوتاه، مشابهت فرح با واژگانی که در سوره‌های کوتاه کاربرد دارند، کمتر از میزان واقعی به‌دست می‌آید. به عنوان نمونه مشابهت فرح با سرور (که دو مورد از شش وقوع آن، در سورهٔ انشقاق و یک مورد در سورهٔ انسان است) صفر شده است.

یک راهکار برای حل این مشکل، استفاده از رکوعات قرآنی به عنوان سند است؛ که طول آن‌ها تفاوت اندکی با هم دارد. اما علاوه بر این که لزوماً این دسته‌بندی، واحدهای مستقل دارای موضوع واحد را نشان نمی‌دهد، نرم‌افزارهای موجود در تعیین بسامد وقوع واژه‌ها در رکوعات ناتوانند و این امر محاسبات را دشوار می‌کند. راهکار پیشنهادی برون‌رفت از این چالش، تصحیح خطای مربوط به تفاوت طول سوره‌ها، ضمن مقایسه با الگوی واژه - عبارت، پس از انجام محاسبهٔ شباهت میان واژگان است. با توجه به این که مشابهت در الگوی واژه - عبارت بیش از مقدار واقعی آن به دست می‌آید، می‌توان گفت مقدار واقعی مشابهت، بین دو مقدار به‌دست آمده در الگوی واژه - سوره و الگوی واژه - عبارت است.

۶. نتیجه

۱- معنی‌شناسی توزیعی در قرآن، با هدف شناخت تقریبی معنا و به عبارتی تعیین حوزهٔ معنایی واژگان مورد اختلاف، با رویکردی آماری صورت می‌گیرد. پیاده‌سازی این روش بر روی متن قرآن به منظور

- شناخت معنی یک واژه، شامل گام‌های زیر است:

 ۱. یافتن واژگانی برای مقایسه با واژه اصلی؛
 ۲. تعیین ویژگی‌های توزیعی واژه‌ها، شامل سندها (سوره‌ها با معنای اصطلاحی یا رکوعات قرآنی) و عبارات هم‌وقوع با واژه (با لحاظ رابطه دستوری یا وقوع در پنجره بافتی)؛
 ۳. پیش‌پردازش زبانی متن قرآنی شامل ریشه‌یابی و یافتن صیغه‌های مهم هر ریشه برای حضور در ماتریس هم‌وقوعی، ابهام‌زدایی و ادغام واژگان قریب‌المعنا؛
 ۴. حذف ایست‌واژه‌ها؛
 ۵. تشکیل ماتریس هم‌وقوعی و وزندهی آن (در الگوی واژه - سوره با قاعده TFIDF و در الگوی واژه - عبارت با قاعده PPMI)؛
 ۶. تشکیل بردار بافتی وزندهی شده و محاسبه مشابهت واژگان با کوسینوس زاویه آن‌ها؛
 ۷. تفسیر مشابهت میان واژگان با استفاده از لغت‌نامه‌ها و روابط معنایی میان جملات.

- ۲- چالش‌های به‌کارگیری معنی‌شناسی توزیعی به‌ویژه در قرآن و برخی راه‌های برون‌رفت از آن عبارت است از:

 ۱. چالش‌های کلی مربوط به معنی‌شناسی توزیعی؛ شامل عدم سهولت پیش‌پردازش زبانی به‌صورت خودکار، تفسیر مشابهت میان واژگان، محدودیت در تشخیص مراد متکلم.
 ۲. چالش‌های به‌کارگیری معنی‌شناسی توزیعی در قرآن شامل حجم کم پیکره قرآنی برای محاسبات آماری، عدم وجود نرم‌افزار مناسب جهت ریشه‌یابی مؤثر، ابهام‌زدایی و محاسبه ایست‌واژه‌ها، و تفاوت بسیار زیاد طول سوره‌ها در الگوی واژه - سوره است. توجه بیشتر به مبنای روش توزیعی (فرضیه توزیعی) که بر پایه حکمت گوینده، اعتبار بیشتر در داده‌های قرآنی را نتیجه می‌دهد، و تصحیح خطای ناشی از تفاوت طول سندها در الگوی واژه - سوره با مقایسه نتایج با الگوی واژه - عبارت؛ راهکارهای پیشنهادی برون‌رفت از برخی چالش‌هاست.
 - ۳- علاوه بر پاسخ به سؤال‌های اصلی پژوهش، اجرای روش توزیعی برای ریشه - واژه فرح در قرآن، تفاوت معنای به‌دست آمده از بافت قرآنی (نزدیک به حوزه معنایی تکبر)، نسبت به معنای غالبی ذکرشده در لغت‌نامه‌ها، به‌ویژه دیدگاه لغت‌شناسان غیرقرآنی و به طور اخص، لغت‌نامه‌های امروزی (شادی) را نشان می‌دهد. تفاوت این دو معنا، با توجه به مبنای شباهت زبان قرآن به زبان عربی عصر نزول، تطور معنایی این واژه را نتیجه می‌دهد. شناخت سیر تحول در زمانی معنای این واژه، با

به‌کارگیری روش توزیعی در متون دوره‌های مختلف پس از قرآن، مانند روایات، امکان‌پذیر است و می‌تواند موضوع پژوهش‌های دیگر باشد.^{۹۱}

۷. پی‌نوشت‌ها

1. distributional semantics
2. corpus linguistics
3. natural language
4. *Theories of Lexical Semantics*
5. *Machine learning for text*
6. *Networks*
7. Distributional Hypothesis
8. Zellig Harris
9. feature
10. document
11. co - occurrence

۱۲. نمونه‌ای از پیکره‌های تخصصی، پیکره میشیگان شامل مقالات سطح بالایی دانشجویی: Michigan Corpus of Upper Level Student Papers (MICUSP) است (کرافورد و سزومی، ۲۰۱۶، ترجمه صفوی، ۱۳۹۹، ص. ۳۷)؛ که ۸۳۰ مقاله با ۲.۶ میلیون واژه دارد (اقتباس از <https://eresources.eli.lsa.umich.edu/micusp> - papers/ - academic - written - of - corpus).

13. term

14. context window

۱۵. این انتخاب می‌تواند به‌جهت هم‌نشینی با واژه (راست یا چپ) نیز مقید شود (Lenci, 2018, p.153) و نزدیکی و دوری دو واژه در پنجره بافتی هم مورد توجه باشد (Sahlgren, 2006, p.41؛ ترکیان، ۱۳۹۸، ص. ۴۲۴).

16. context vector

17. co - occurrence matrix

18. vector

19. similarity

۲۰. اندازه یک بردار براساس قضیه فیثاغورث، جذر مجموع مجذور مؤلفه‌های آن است:

$$|a| = \sqrt{(a_1^2 + a_2^2 + \dots)}$$

21. pre - processing

22. lemmatization

23. part - of - speech (POS) tagging

24. Disambiguation

۲۵. مقوله‌های دستوری بسته، به‌ندرت عضو جدید می‌پذیرند؛ مانند ضمائر و حروف ربط و اما در مقوله‌های باز، جایگزینی، تغییر و پذیرش اجزای جدید رخ می‌دهد، مانند اسم، صفت، فعل و ... (Sahlgren, 2006, p. 38).
۲۶. واژگان بسیار کم‌بسامد معمولاً به‌طور تصادفی با واژه اصلی هم‌وقوع شده‌اند.

27. stop words

28. weighting

29. raw frequency

30. Term Frequency

۳۱. Inverse Document Frequency. این تابع به‌صورت لگاریتم نسبت کل اسناد (D)، به سندی که عبارت i

در آن به‌کار رفته، تعریف می‌شود: $IDF_i = \log_2 \left(\frac{D}{DF_i} \right)$. در بسیاری موارد که «نسبت» اهمیت دارد، به دلیل گستره بسیار وسیع مقادیر، از تابع لگاریتم استفاده می‌شود؛ که با در نظر گرفتن مرتبه (توان)، نسبت‌ها را به مقیاس مناسبی درمی‌آورد ($x^z = y$ معادل است با $\log_x y = z$).

32. scaling

33. expected probability

34. observed probability

35. Singular Value Decomposition

36. semantic map

37. neighbor

۳۸. به عنوان نمونه «ماشین و راننده» دارای مفاهیم مرتبط با هم اند، اما در این دسته‌بندی نمی‌گنجند.

۳۹. به‌عنوان نمونه: *إِنَّ شَجَرَةَ الزُّقُومِ، طَعَامُ الْأَثِيمِ* (دخان: ۴۳ و ۴۴).

۴۰. به‌عنوان مثال عبارت‌های «الذین آمنوا» و «مؤمنون» ذیل ریشه «امن» قابل بررسی است.

۴۱. این واژگان دارای رابطه «چندمعنایی» یا «هم‌آوا - هم‌نویسی» هستند. «چندمعنایی» (polysemy) ارتباط معانی مختلف واژه و «هم‌آوا - هم‌نویسی» یا هم‌نامی (homonymy)، چند معنا داشتن واژه بدون ارتباط معنایی را نشان می‌دهد (صفوی، ۱۳۹۹، صص. ۱۱۰ و ۱۱۱؛ افراشی، ۱۳۹۵، صص. ۱۲۹ و ۱۳۰).

۴۲. *أَفَحُكْمَ الْجَاهِلِيَّةِ يَبْغُونَ...* (مائده: ۵۰)

۴۳. *الَّذِينَ يَصُدُونَ عَنْ سَبِيلِ اللَّهِ وَيَبْغُونَهَا عِوَجًا...* (اعراف: ۴۵)

۴۴. با توجه به فرمول، مقدار این عبارت چنین به دست می‌آید: $IDF_{\text{واژه}} = \log_2 \left(\frac{D}{DF_{\text{واژه}}} \right)$

$\log_2 \left(\frac{\text{تعداد کل سوره‌ها}}{\text{تعداد سوره‌هایی که واژه در آن آمده}} \right)$ ؛ که «تعداد کل سوره‌ها» بسته به لحاظ سوره یا سوره‌وار به‌معنای رکوعات قرآنی، ۱۱۴ یا ۵۵۵ است.

۴۵. برای محاسبه طول سوره، تعداد واژگان سوره‌ها، یا نسبت تعداد واژگان سوره‌ها به تعداد واژگان یک

سوره می‌تواند محاسبه شود.

۴۶. برای محاسبه طول سوره، نسبت تعداد واژگان آن با تعداد واژگان سوره حدید، ۵۴۴ واژه (فرهنگ نامه علوم

قرآنی، ۱۳۹۴، ص. ۲۶۷۶) مقایسه شده است ($\frac{\text{تعداد واژگان سوره}}{544} = \text{طول سوره}$).

۴۷. تعداد واژه‌های سوره‌های تغابن، مجادله، قمر و تحریم به ترتیب ۲۴۱، ۴۷۳، ۳۴۲ و ۲۴۰ ذکر شده است (فرهنگ نامه علوم قرآنی، ۱۳۹۴، صص. ۲۶۷۳ - ۲۶۸۵). پس از حذف ایست‌واژه‌ها (با محاسبه دستی)، تعداد واژگان این سوره‌ها به ترتیب ۱۶۱، ۲۸۶، ۲۴۰ و ۱۶۶ به دست آمد. متوسط نسبت تعداد واژگان این سوره‌ها پس از حذف ایست‌واژه‌ها به تعداد کل واژگان سوره، حدود ۶۷ درصد است.

۴۸. ممکن است تصور شود برخی مقادیر به دست آمده، نشانگر رابطه معنایی قابل توجه و معناداری میان واژگان نیست. اما توجه به نتایج برخی پژوهش‌های این حوزه، معنادار بودن این مقدار مشابهت معنایی را نشان می‌دهد. در پژوهشی که واژه‌های dog، car، bike، lion و dog، lion از لحاظ توزیعی بررسی شده، مشابهت میان واژگان چنین به دست آمده است (Lenci, 2018, 156):

bike	1			
car	0.16	1		
dog	0	0	1	
lion	0	0	0.17	1
	bike	car	dog	lion

که مشابهت میان واژه‌های bike و car، و نیز واژه‌های dog و lion، را کم‌تر از ۰.۲ نشان می‌دهد؛ درحالی که آن‌ها به یک دسته مفهومی از واژگان تعلق دارند.

۴۹. برای مثال در معنی‌شناسی فرح، ارتباط میان جملات دربردارنده فرح و فخر در آیه‌های ۲۲ و ۲۳ سوره حدید، شمول معنایی یا هم‌معنایی این واژه‌ها را نتیجه می‌دهد (ذوعلم و همکاران، ۱۳۹۹، ص. ۱۳).

۵۰. این محدودیت معنی‌شناسی توزیعی، در کاربست قرآنی، معادل با عدم شناخت مراد متکلم با این روش است. بر این اساس با توجه به برخی تعاریف ظاهر و باطن قرآن، معنی‌شناسی توزیعی صرفاً به شناخت معنی ظاهری واژگان قرآنی - البته با روشی اصولی و نظام‌مند - می‌پردازد و ادعایی درمورد کشف لایه‌های باطنی آن ندارد. ۵۱. در پایان نامه معنی‌شناسی فرح در قرآن و روایات تا اندازه‌ای به این تحول پرداخته شده است (نک: ذوعلم، ۱۳۹۷).

۸. منابع

- آذرنوش، آ. (۱۳۸۳). فرهنگ معاصر عربی - فارسی. تهران: نی.
- ابن درید، م. (۱۹۸۸). جمهرة اللغة. بیروت: دار العلم للملایین.

- ابن قتیبه، ع. (۱۴۱۱ق). تفسیر غریب القرآن. بیروت: دار و مکتبه الهلال.
- ابو عبیده، م. (۱۳۸۱ق). مجاز القرآن. قاهره: مکتبه الخانجی.
- ازهری، م. (۱۴۲۱ق). تهذیب اللغة. بیروت: دار احیاء التراث العربی.
- افراشی، آ. (۱۳۹۵). مبانی معنا شناسی شناختی. تهران: پژوهشگاه علوم اسلامی و مطالعات فرهنگی.
- ایزوتسو، ت. (۱۹۶۶). مفاهیم اخلاقی - دینی در قرآن مجید. ترجمه ف. بدره ای (۱۳۷۸). تهران: فرزانه.
- ترکیان، ا. (۱۳۹۸ الف). آموزش عملی متن کاوی، پیوست الف متن کاوی یادگیری ماشین. چ. آگاروال، ترجمه ا. ترکیان (صص ۶۱۷ - ۶۵۹). چاپ دوم، تهران: نیاز دانش.
- ترکیان، ا. (۱۳۹۸ ب). شبکه متن، پیوست شبکه ها. م. نیومن، ترجمه ا. ترکیان (صص ۴۳۳ - ۴۵۰). تهران: نیاز دانش.
- جوهری، ا. (۱۳۷۶ق). الصحاح (تاج اللغة و صحاح العربیة). بیروت: دارالعلم للملایین.
- ذوعلم، آ. (۱۳۹۷). معنا شناسی فرح در قرآن و روایات. پایان نامه کارشناسی ارشد. دانشکده الهیات دانشگاه تهران.
- ذوعلم، آ.، فائز، ق. و خوش منش، ا. (۱۳۹۹). معنا شناسی فرح در قرآن کریم؛ پاسخی به شبهه نکوهش شادی در قرآن. تحقیقات علوم قرآن و حدیث، ۴۷ (۳)، ۱ - ۳۴.
- راغب اصفهانی، ح. (۱۴۱۲ق). مفردات ألفاظ القرآن. بیروت: دارالقلم.
- صفوی، ک. (۱۳۹۹). درآمدی بر معنی شناسی. چاپ ششم، تهران: پژوهشکده فرهنگ و هنر اسلامی.
- طبرسی، ف. (۱۳۷۲). مجمع البیان فی تفسیر القرآن. تهران: ناصر خسرو.
- عبد الباقي، م. (۱۳۹۳). المعجم المفهرس لالفاظ القرآن الکریم. قم: نوید اسلام.
- عسکری، ح. (۱۴۰۰ق). الفروق فی اللغة. بیروت: دار الافاق الجدیدة.
- فائز، ق.، خوش منش، ا. و ذوعلم، آ. (۱۳۹۸). بهره گیری از سیاق متنی در معنا شناسی ساختگرای قرآن کریم. پژوهش های قرآن و حدیث، ۵۲ (۱)، ۹۵ - ۱۱۵.
- فراهیدی، خ. (۱۴۰۹ق). کتاب العین. قم: هجرت.
- فرهنگ نامه علوم قرآنی، دفتر تبلیغات اسلامی، قم: پژوهشگاه علوم و فرهنگ اسلامی، ۱۳۹۴.
- کرافورد، و. و سزومی، ا. (۲۰۱۶). زبان شناسی پیکره ای در عمل. ترجمه م. نوبخت (۱۳۹۹). تهران: سمت.
- گلی ملک آبادی، ف.، خاقانی اصفهانی، م. و شکرانی، ر. (۱۳۹۶). نقدی معنا شناختی بر ترجمه های فارسی

واژه «دون» در قرآن کریم. جستارهای زبانی، ۳۶ (۱)، ۲۰۷ - ۲۳۰.

- گیرتس، د. (۲۰۱۰). نظریه‌های معنی‌شناسی واژگانی. ترجمه ک. صفوی (۱۳۹۸). تهران: نشر علمی.
- لسانی فشارکی، م. و مرادی زنجانی، ح. (۱۳۹۵). سوره‌شناسی. چاپ دوم. قم: نصایح.

References

- Afrashi. A. (2016). *Introducing Cognitive Semantics*. Institute for Humanities and Cultural Studies [In Persian].
- Askari. H. (1979). *Alfuroogh Fillughha*. Dar Al'afagh Aljadida [In Arabic].
- Azarnoosh, A. (2004). *A Dictionary of Modern Written Arabic*. Nei [In Persian].
- Azhari. (2000). *Tahzib Allughha*. Dar Al'ihya Altorath Al'arabiy [In Arabic].
- Abu - ubaida. M. (1963). *Majaz alQuran*. Maktaba Alkhanaji [In Arabic].
- Boleda, G. (2020). *Distributional Semantics and Linguistic Theory. Annual Review of Linguistics*, 6, 213– 234.
- Crawford. W., Csomay. E. (2016). *Doing Corpus Linguistics*, Samt [In Persian].
- *Dictionary of Quranic Sciences*. (2015). Islamic Sciences and Culture Academy [In Persian].
- Dumais, S., Furnas, G., Landauer, T. (1988). Using latent semantic analysis to improve access to textual information, In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, New York (pp. 281- 285).
- Faez. Gh., Khoshmanesh. A., Zouelm. A. (2019). The Use of Textual Context in Structural Semantics of the Holy Qur'an, *Quranic Researches and Tradition*, 52(1), 95 - 115 [In Persian].
- Farahidi. Kh. (1988). *Ketab Al'Ain*. Hejrat [In Arabic].
- Geeraerts. D. (2010). *Theories of Lexical Semantics*, Nashre Elmi [In Persian].
- Goli. F, Khaqani. M, Shokrani R. (2017). A Semantic Criticism of the Persian Translation of the Arabic Term "Doon" in the Holy Quran, *Language Relate Research*, 8 (1), 207 - 230 [In Persian].
- Harris, Z. (1954). *Distributional Structure*. *WORD*, 10(2 - 3), 146 - 162.
- Ibn Duraid. M. (1988). *Jumhora allughha*. Dar Al'ilm Almullain [In Arabic].

- Ibn Qutaiba. 'A. (1990). *Tafsir Qarib alQuran*. Dar va Maktaba Alhilar [In Arabic].
- Jawhari. I. (1956). *Alsehah*. Dar Al'ilm Almullain [In Arabic].
- Izutsu. T. (1966). *The structure of meaning in ethico -religious concepts in Quran*. Farzan [In Persian].
- Lenci, A. (2018). Distributional Models of Word Meaning. *Annual Review of Linguistics*, 4, 151–171.
- Lesani Fesharaki. M., Moradi. Zanjani. H. (2016). *Surah Shenasi*. Nasayeh [In Persian].
- Raghieb Isfahani. H. (1991). *Mufradat Alfaz Alquran*. Dar Alghalam [In Arabic].
- Safavi. K. (2020). *An Introduction to Semantics*. Sourey - e - Mehr [In Persian].
- Sahlgren, M. (2006) *The Word - Space Model*, Doctoral Dissertation, Stockholm University, Sweden.
- Schutze, H. (1992). dimensions of meaning, In *ACM/IEEE conference on Supercomputing*, (pp. 787 - 796).
- Schutze, H. (1993). Word space. In *Proceedings of the 1993 Conference on Advances in Neural Information Processing Systems*, NIPS'93 (pp. 895–902).
- Tabrasi. F. (1993). *Majma' Albayan fi Tafsir Alquran*. Naserkhosro [In Arabic].
- Turkian, A. (2019a). Learning Practical text mining, *Machine learning for text*. Aggarwal. C. (pp. 617– 659), Niaze Danesh [In Persian].
- Turkian, A. (2019b). Text Network. *Networks*. Newman.M. (pp. 433– 450), Niaze Danesh [In Persian].
- Turney, P. (2010). From Frequency to Meaning: Vector Space Models of Semantics, *Journal of Artificial Intelligence Research*, 37, 141 - 188.
- Zouelm. A. (2018). *Semantical Study of Farah in Quran and Tradition*. Master thesis, Faculty of Theology and Islamic Studies, Tehran University [In Persian].
- Zouelm. A., Faez. Gh., Khoshmanesh. A. (2020). Semantical Study of Farah in the holy Qur'an, A Response to the Doubt of Disapproval of Joy in the Qur'an. *Journal of Researches of Quran and Hadith Sciences*. 17(3), 1 - 34 [In Persian].