

به کارگیری نظریه ساختار بلاغی برای بهبود

بازنمایی متن با شبکه‌های عصبی عمیق

عرفانه غروی^۱ و هادی ویسی^{۲*}

دانشکده علوم داده، دانشگاه ویرجینیا، ویرجینیا، آمریکا

دانشکده علوم و فنون نوین، دانشگاه تهران، تهران، ایران

چکیده

یافتن یک بازنمایی معنایی غنی با ابعاد کم برای متون طولانی یکی از چالش‌های اساسی در فعالیت‌های مختلف پردازش زبان طبیعی به شمار می‌رود. این بازنمایی باید اطلاعات معنایی و نحوی متن را در بر گرفته، و همچنین، بر حسب وظیفه مد نظر ارتباط و تشابه متون را در ابعاد کم الگوسازی کند. در این مقاله تلاش بر آن است تا با بهره‌گیری از نظریه ساختار بلاغی و شبکه‌های عصبی عمیق چالش‌های مطرح شده مرتفع شود. نظریه ساختار بلاغی با ارائه یک ساختار سلسله‌مراتبی به توصیف اهمیت عبارات موجود در متن و روابط بین آن‌ها می‌پردازد. در اینجا تأثیر به کارگیری این ساختار درختی بر دو وظیفه بازیابی اطلاعات و تحلیل احساسات بررسی شده است. در وظیفه بازیابی اطلاعات، جهت الگوسازی وابستگی معنایی بین مستندات، یادگیری بازنمایی سند با شبکه‌های عصبی بازگشتی عمیق دوقلو انجام شد؛ به طوری که ذخیره و بازیابی مستندات متنی تسهیل شود. این شبکه از دو زیرشبکه بازگشتی عمیق تشکیل شده است. این شبکه‌های بازگشتی، مبتنی بر ساختار درختی حاصل از تجزیه متن با نظریه ساختار بلاغی هستند. این روش‌شناسی بر روی دو مجموعه داده خبری شامل اخبار بی‌بی‌سی و همچنین، زیرمجموعه‌ای از دادگان رويترز ارزیابی شد. نتایج نشان می‌دهد بازنمایی ارائه شده با این ساختار، کارایی بالاتری از بازنمایی‌های سنتی مبتنی بر سبک کلمه دارد. این رویکرد کارایی را به میزان شش درصد بر روی مجموعه داده بی‌بی‌سی و سه درصد بر روی مجموعه داده رويترز نسبت به بهترین روش کلاسیک بهبود داده است. در وظیفه تحلیل احساسات، نخست، به کمک شبکه عصبی بازگشتی عمیق مبتنی بر درخت ساختار بلاغی به ایجاد بازنمایی و در نهایت دسته‌بندی احساسات نظرات افراد پرداخته، سپس، سایر اطلاعات موجود در درخت جهت بهبود الگو استفاده شد. این اطلاعات شامل آگاهی از اهمیت هر بخش از متن با استفاده از درخت ساختار بلاغی است. با تشخیص بخش‌های مرکزی متن و اعمال سازوکار توجه بر آن در شبکه عمیق بازگشتی بازنمایی غنی‌تری برای متن ایجاد می‌شود. این بازنمایی کارایی الگوی تحلیل احساسات را بر روی دادگان اینترنتی نظرات بینندگان فیلم در مقایسه با روش‌های پایه به میزان ۳ درصد افزایش داده است. نتایج حاصل از این بررسی، بهبود بازنمایی متن با استفاده از شبکه‌های عمیق مبتنی بر نظریه ساختار بلاغی را نشان می‌دهد. بهبود بازنمایی به کمک ساختاردهی متن غیر ساختاریافته بر روی زبان‌های دیگر از جمله زبان فارسی می‌تواند راستی آزمایی قرار شود.

واژگان کلیدی: بازنمایی متن، نظریه ساختار بلاغی، شبکه‌های عصبی عمیق، شبکه‌های دوقلو، سازوکار توجه

Improve Text Representation Using RST-based Deep Neural Networks

Erfaneh Gharavi and Hadi Veisi

School of Data Science, University of Virginia, VA, USA

Data and Signal Processing Lab, Faculty of New Sciences and Technologies, University of Tehran, Tehran, Iran

Abstract

Finding a highly informative, low-dimensional representation for texts, specifically long texts, is one of the main challenges for natural language processing (NLP) tasks. For texts longer than sentences or a paragraph, finding a good representation beyond the bag-of-words model without losing word order is still a challenge. This representation should capture the semantic and syntactic information of the text while

* Corresponding author

* نویسنده عهده‌دار مکاتبات

retaining relevance for large-scale similarity search and accurate text classification. We propose the utilization of Rhetorical Structure Theory (RST) to consider the text structure in the representation. RST is a theory of text organization that describes relations that hold between parts of text and creates a tree-structure format for the text document. RST can model the importance and relationship between sentences or phrases as well. Rhetorical relations or discourse relations are paratactic (coordinate) or hypotactic (subordinate) relations that hold across two or more text spans. These relations are applied recursively in a text, until all units in that text are constituents in an RST relation. RST establishes two different types of units: Nuclei and satellites. Nuclei are considered as the most important parts of text and contains basic information whereas satellites contribute to the nuclei and are secondary and contains additional information about nucleus.

In this paper, we examine the effect of using this structure on two different NLP tasks. In information retrieval, to embed document relevance in distributed representation, we use a Siamese neural network to jointly learn document representations. Our Siamese network consists of two sub-networks of recursive neural networks (RNN) built over the RST tree. It means that two chunks (i.e., edu) of the text are merged according to their relation in the RST tree. For this task, we use a subset of Reuters's news corpus (includes eight topics) and BBC news dataset (includes five topics). In the implementations, DPLP parser is used to parse RST trees. The results show that our approach outperforms conventional text representations like TF-IDF, LDA, LSA and word vector averaging. The proposed representation beats the best conventional method by %6 and %3 in precision at k retrieved documents on BBC and Reuters datasets, respectively. As another task, in the sentiment analysis, first, we use an RST-based recursive neural network to represent movie reviews and classify the polarity of people's opinions (positive and negative). Then, we propose to use the nucleus-satellite information of a node in the RST-tree to build an attention mechanism by deep RNN to generate better discourse representations. We test the effectiveness of our approach on sentiment analysis task, and we prove that considering the importance of the text span improves sentiment analysis performance by %3 on the internet movie review database in comparison with the baseline standard RNN and 2% improvement in comparison with the attention-based RNN. In this paper, we improve the text representation by the RST-based deep neural network. This approach can be further evaluated on the other languages to show the effectiveness of using the semantic information embedded in the RST format of the text.

Keywords: Document Embedding, Semantic Representation, Rhetorical Structure Theory, Deep Neural network, Attention Mechanism

حاصل از این روش‌ها در وظایف مختلف یادگیری ماشین منجر به حصول نتایج عالی شده‌است.

روش‌های بازنمایی کنونی به ترکیب بردار کلمات تا رسیدن به بازنمایی جمله و پاراگراف پرداخته‌اند. فرای آن برای متون طولانی‌تر، بازنمایی‌ها به روش‌هایی مانند سید کلمه محدود می‌شود. این بازنمایی‌ها دارای ابعاد بالایی بوده که هزینه ذخیره و بازیابی را بسیار بالا می‌برد. علاوه بر آن، مشکل این نوع بازنمایی‌ها عدم در نظر گرفتن ترتیب و رابطه بین واحدهای متنی است. در صورتی که بخش‌های مختلف متن به صورت تصادفی در کنار یکدیگر قرار نگرفته‌اند و بین آن‌ها ارتباط باقاعده‌ای وجود دارد. این قاعده می‌تواند با نظریه ساختار بلاغی متن^۵ (RST) الگو شود. نظریه ساختار بلاغی چارچوبی را فراهم می‌کند که در آن متن به واحدهای پایه‌ای تقسیم می‌شود. این واحدها واحد پایه‌ای سخن^۶ (edu) نام دارند و ارتباط این واحدها با روابط ازپیش تعریف شده، تعیین می‌شود. نظریه RST چارچوبی فراهم می‌کند که گزاره‌های ارتباطی^۷

۱- مقدمه

ایجاد یک بازنمایی کم‌بعد اما غنی برای متون طولانی یکی از چالش‌های اصلی پردازش زبان طبیعی به‌شمار می‌رود. اگرچه پژوهش‌های اخیر نشان داده است که الگوهای عمیق می‌توانند بازنمایی فشرده، اما دقیقی از متن ارائه کنند، اما همچنان دست‌یابی به یک بازنمایی با طول ثابت که ترتیب کلمات را در متون طولانی حفظ کند، چالش به‌شمار می‌آید. سوچر انواع مختلفی از شبکه‌های عصبی مانند شبکه عصبی بازگشتی^۱ [1] و شبکه عصبی بازگشتی ماتریس-بردار [2] را برای ترکیب بردار کلمات^۲ [3] به‌کار برده‌است. همچنین، شبکه عصبی حافظه کوتاه‌مدت^۳ [4] و شبکه عصبی پیچشی^۴ [5] برای بهبود بازنمایی استفاده شده‌اند. علاوه بر آن، روش بردار پاراگراف [6]، هر پاراگراف را به‌عنوان یک بردار در نظر گرفته و یک بازنمایی در سطح پاراگراف ایجاد می‌کند. بازنمایی

¹ Recursive Neural Network

² Word Embeddings

³ Long Short Term Memory (LSTM)

⁴ Convolutional Neural Network

⁵ Rhetorical Structure Theory (rst)

⁶ Elementary Discourse Unit (edu)

⁷ relational propositions

نامشخص که تنها در فرآیند تفسیر متن بدست می‌آیند، بررسی شود. از آنجاکه انسجام^۱ متن بستگی به این گزاره‌های ارتباطی دارد، RST در مطالعه انسجام متن مفید است. نظریه ساختار بلاغی، متن را با یک ساختار درختی سلسله‌مراتبی الگو می‌کند. اهمیت *edu* در درخت *rst* با عمق گره [7]، نوع رابطه [8] و همچنین، نوع سلسله‌مراتب مشخص می‌شود. سلسله‌مراتب هر *edu* هسته^۲ یا پیرو^۳ بودن آن را مشخص می‌کند. در این دسته‌بندی هسته دارای اطلاعات مرکزی‌تر نسبت به پیرو است. و پیرو اطلاعات پیش‌زمینه و پشتیبان را برای هسته فراهم می‌کند [9].

هدف از این مقاله، ارزیابی تأثیر به‌کارگیری توأمان نظریه ساختار بلاغی و شبکه‌های عصبی عمیق بازگشتی در اغنای بازنمایی متن است. داده‌های غیرساختاریافته با نظریه ساختار بلاغی به‌صورت درخت دودویی ساختاردهی می‌شوند. شبکه‌های عصبی بازگشتی بهترین گزینه برای الگوسازی دادگان درختی هستند. این شبکه‌ها دو بردار ورودی یا برگ‌ها را با هم ترکیب کرده و بردار گره والد با ابعاد یکسان را تولید می‌کنند. این فرآیند تکرارشونده تا ریشه درخت ادامه می‌یابد. بردار نهایی تولیدشده در شبکه عصبی بازگشتی در ریشه ساختار درختی *rst* بردار بازنمایی متن خواهد بود. ارزیابی بازنمایی ارائه‌شده دو وظیفه بازبینی اطلاعات و تحلیل احساسات در زبان انگلیسی بررسی شد. در وظیفه بازبینی اطلاعات با به‌کارگیری شبکه‌های دوقلوی بناشده بر شبکه‌های عصبی بازگشتی مبتنی بر درخت دودویی *rst*، روشی جهت بازنمایی متون طولانی ارائه‌شده‌است. این روش به یادگیری بازنمایی متون به‌طور توأمان مبادرت می‌ورزد. با مقایسه این روش با سایر روش‌های بازنمایی متن کارآیی این روش به اثبات رسیده‌است. در وظیفه تحلیل احساسات، به‌کارگیری ساز و کار توجه بر روی هسته سخن بر اساس درخت *rst* ارزیابی، و نشان داده شد که چگونه ایجاد تمایز در سهم هسته و پیرو، بازنمایی متن و در نتیجه، کارایی وظیفه تحلیل احساسات را بهبود می‌دهد.

نوآوری این مقاله در استفاده از ساختار درختی *rst* دودویی در زیرشبکه‌های شبکه عصبی عمیق دوقلو جهت غنی‌سازی بازنمایی مستندات در وظیفه بازبینی اطلاعات و همچنین، استفاده از شبکه‌های عصبی عمیق مبتنی بر

- 1 coherence
- 2 Nucleus
- 3 Satellite

اطلاعات سلسله‌مراتبی موجود در درخت *rst* جهت ایجاد سازوکار توجه بر بخش‌های هسته‌ای متن و غنی‌سازی بازنمایی در وظیفه تحلیل احساسات است. ادامه مقاله به شکل زیر سازمان یافته است. نخست، نظریه ساختار بلاغی معرفی و توصیف می‌شود. سپس، کارهای مرتبط در زمینه بازنمایی متون و روش‌های به‌کارگیری ساختاری بلاغی عنوان شده‌است. در ادامه، روش‌های ارائه‌شده جهت به‌کارگیری نظریه ساختار بلاغی در ایجاد بازنمایی غنی‌تر در وظایف بازبینی اطلاعات و تحلیل احساسات ارائه می‌شود. در توصیف هر وظیفه به‌صورت جداگانه معماری الگو، دادگان و نتایج حاصل‌شده شرح داده می‌شود. در پایان، درباره نتیجه بحث شده و کارهای آینده عنوان می‌شود.

۲- نظریه ساختار بلاغی

نظریه ساختار بلاغی یک چارچوب سلسله‌مراتبی برای توصیف متن ارائه می‌کند. این نظریه رابطه بین بخش‌های مختلف متن را با عباراتی بامعنی بیان کرده و یک تحلیل جامع از متن ارائه می‌دهد. همچنین، قابل اعمال بر روی متون با طول‌های مختلف است. من و تامپسون، ۳۲ رابطه بلاغی مختلف بین واحدهای پایه‌ای متن ارائه داده‌اند [10].

برای تعریف واحدهای پایه‌ای سخن نیاز به شناسایی نشان‌گرهای سخن است. نشان‌گرهای سخن که انسان به‌طور مداوم استفاده می‌کند، می‌توانند رابطه‌ای منسجم بین عبارات و واحدهای بزرگتر متن باشند. این نشان‌گرها در هر سطحی از متن شامل عبارت، جمله، پاراگراف و متن استفاده می‌شوند.

پایه ریاضیاتی الگوریتم تجزیه سخن قاعده‌مندسازی درجه نخست یک ساختار است. فرضیات فرموله‌سازی در نظر گرفته‌شده در ادامه آمده‌است: (۱) واحدهای پایه‌ای عبارات غیرهم‌پوشان متن هستند. (۲) روابط بلاغی منسجمی بین واحدهای متنی با اندازه‌های مختلف وجود دارد. (۳) دو دسته رابطه بین عبارت‌هایی با اهمیت یکسان نگهداری وجود دارد. (۴) ساختار انتزاعی بیشتر متنها درخت دودویی است. (۵) اگر رابطه‌ای بین دو عبارت متنی در یک ساختار درختی برقرار باشد، این رابطه بین زیرعبارت‌های تشکیل دهنده مهم آن نیز برقرار است.

در یک متن معمولی یک نشان‌گر برای ایجاد یک ساختار بلاغی غنی کفایت می‌کند. فرض بر این است که

متن مورد پردازش دارای ساختار نحوی مناسبی است. اگرچه هنوز ابهام‌های فراوانی در زمینه تشخیص عبارات و نشان‌گر سخن وجود دارد. گاهی نشان‌گرها به‌عنوان بخشی از جمله به کار رفته و در تشخیص طول عبارت ابهام ایجاد می‌کنند.

ایجاد یک پیکره و تحلیل آن برای فراهم کردن اطلاعاتی درباره انواع روابط بلاغی، وضعیت بلاغی (هسته یا پیرو) و اندازه عبارت متنی که هر نشان‌گر می‌تواند مشخص کند، ضروری به نظر می‌رسد. در این پیکره مجموعه‌ای از نمونه‌های متنی برای هر نشان‌گر جمع‌آوری شده‌است. هر نمونه از متن دارای یک پنجره ۲۰۰ کلمه‌ای با رخداد نشان‌گر است. برای نشان‌گرهای مبهم مانند «و» نمونه‌های بیشتری جمع‌آوری شده‌است.

قطعات متنی با تعاریفی که در ادامه می‌آیند، هم‌راستا می‌شود: (۱) ساختاری که استفاده بالقوه از نشان‌گر را مشخص می‌کند؛ شامل رخداد نقطه، کاما، دونقطه، سمیکالون و غیره. (۲) نوع استفاده از نشان‌گر؛ استفاده در جمله یا ساختار بلاغی و یا هر دو. (۳) موقعیت نشان‌گر در واحدی که به آن تعلق دارد. (۴) موقعیت نسبی واحدهای متنی که واحد متنی حاوی نشان‌گر به آن متصل است؛ قبلی و بعدی. (۵) رابطه بلاغی که نشان‌گر مشخص می‌کند. (۶) نوع متن واحدهایی که از طریق نشان‌گر به هم متصل شده‌اند؛ عبارت تا چندین پاراگراف. (۷) وضعیت بلاغی هر واحد درگیر در رابطه؛ هسته یا پیرو. متن خام وارد تجزیه‌گر شده و با استفاده از نشان‌گرهای ذکرشده تجزیه و به‌صورت درخت بلاغی نشان داده می‌شود. همچنین، نشان‌گرها سلسله‌مراتب واحدهای پایه‌ای سخن و روابط بین آن‌ها را مشخص می‌کند. بسیاری از مطالعات، *rst* را به‌عنوان چارچوبی جهت بررسی مسائل زبانی به‌کار برده‌اند. برخی از این مطالعات در ادامه آمده‌اند:

نظریه ساختار بلاغی رویکردی جهت توصیف روابط بین عبارات در متن فراهم می‌کند. این روابط مستقل از رابطه نحوی یا لغوی هستند. در نتیجه، این نظریه ساختاری مفید برای مرتبط کردن معانی عبارات را فراهم می‌کند. *rst* به‌عنوان یک ابزار تحلیلی برای انواع مختلف متن به کار می‌رود. برای مثال نوئل در [11] از *rst* برای توصیف متون خبری بهره برده‌است. فاکس [12] نشان می‌دهد که چگونه انتخاب بین ضمیر و یا عبارت اسمی در متون تفسیری انگلیسی می‌تواند از ساختار منتج از *rst* حاصل شود.

در این بخش به معرفی یکی از روابط تعریف‌شده خواهیم پرداخت.^۱ بدون شک روابط دیگری نیز می‌تواند در این نظریه وجود داشته باشد که برخی از آن‌ها در پیوست آمده‌است. در اینجا موارد پرکاربرد معرفی شده‌اند. این تعاریف بر نشانه‌های ریخت‌شناسی^۲ یا نحوی استوار نبوده و تشخیص رابطه همواره بر قضاوت کاربردی و معنایی تکیه دارد. برای مثال، تشخیص رابطه شرطی^۳ به وجود کلمه "اگر" در جمله وابسته نیست. در واقع هیچ نشانه قابل‌اتکا و غیرمبهمی برای هیچ‌کدام از روابط وجود ندارد. در طرح‌واره‌های روابط، بخش‌های پایه‌ای متن شامل عبارات و جملات کامل با خطوط افقی نشان داده شده‌اند. خطوط مورب واحدهای سخن را با روابط بلاغی به یکدیگر متصل می‌کنند و کل پیغام به‌صورت سلسله‌مراتبی شکل می‌گیرد. اگرچه ساختار کلی مشابه تجزیه نحوی سخن است، اطلاعات منتقل‌شده، معنایی و در سطح عبارت است.

۲-۱- رابطه شرحی^۴

اطلاعی زیر از یک خبرنامه علمی رابطه شرحی را توصیف می‌کند:

۱. همایش بین‌المللی زبان‌شناسی رایانشی در شهر sanga-saby-Kursgard سوئد، در تاریخ یکم تا چهارم سپتامبر ۱۹۶۹ میلادی برگزار خواهد شد.
۲. انتظار می‌رود ۲۵۰ زبان‌شناس از آسیا، اروپای غربی و شرقی شامل روسیه و همچنین، ایالات متحده آمریکا در همایش شرکت کنند.
۳. محورهای همایش شامل کاربردهای ریاضیات و فنون کامپیوتری در مطالعه پردازش زبان طبیعی، توسعه برنامه‌های کامپیوتری جهت مطالعات زبانی، و کاربرد زبان‌ها برای توسعه سیستم تعامل انسان-ماشین است.

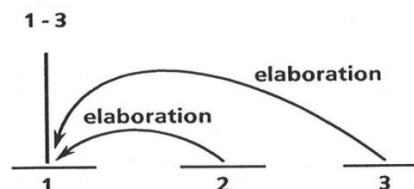
هدف از متن اطلاع‌رسانی به خوانندگان در رابطه با یک همایش است. واحد یک به‌عنوان هسته در نظر گرفته می‌شود. واحدهای پیرو دو و سه اطلاعات جزئی‌تری در رابطه با همایش نام‌برده‌شده در واحد یک فراهم می‌کنند. ساختار *rst* متن در شکل (۱) **Error! Reference source not found.** نشان داده شده‌است.

^۱ چند مثال دیگر از روابط بلاغی پرکاربرد در پیوست آمده‌است.

^۲ Morphological

^۳ Condition

^۴ Elaboration



(شکل - ۱): طرح‌واره رابطه شرحی در متن مرتبط با همایش

Figure- 1 : RST diagram for "conference" text

توصیف رابطه شرحی در جدول (۱) آمده‌است.

(جدول - ۱): رابطه شرحی

Table ۱ : Elaboration relation

نام رابطه	شرحی
قیود بر هسته	ندارد
قیود بر پیرو	ندارد
قیود بر ترکیب هسته و پیرو	پیرو جزئیات بیشتری در رابطه با موقعیت ارائه شده، یا برخی از عناصر موجود در هسته، یا مواردی که می‌توانند به طرق زیر استنباط شوند، ارائه می‌کند. فهرست زیر مواردی از هسته و پیرو را نشان می‌دهد: ■ مجموعه: عضو ■ انتزاع: مثال ■ کل: جزء ■ فرآیند: مرحله ■ شی: ویژگی ■ عام: خاص
اثر	خواننده تشخیص می‌دهد که موقعیت ارائه شده در پیرو جزئیات بیشتری به هسته می‌افزاید. همچنین، عناصر این جزئیات را تشخیص می‌دهد.
محل اثر	هسته و پیرو

با توجه به توسعه نظریه ساختار بلاغی، وجود مجموعه داده برچسب‌گذاری شده، و تجزیه سلسله‌مراتبی آن در سطح معنا، اکنون زمان به‌کارگیری نظریه ساختار بلاغی جهت ایجاد بازنمایی متون طولانی در پردازش زبان طبیعی فرا رسیده‌است. همچنین، این روابط می‌توانند به‌طور مشخص برای زبان فارسی تعریف شوند.

۳- کارهای مرتبط

برخی از روش‌های ایجاد بازنمایی متن به کمک شبکه‌های عمیق در ادامه آمده‌است. سالاخو و هینتون [13] یک الگوی گرافیکی عمیق بر اساس شمارش کلمات آموزش داده‌اند. این الگوی گرافیکی مانند یک هش‌کننده معنایی به کار می‌رود زیرا آخرین لایه تعداد کمی اعداد دودویی تولید می‌کند. این بازنمایی معنایی اسناد مشابه را به نشانی‌هایی نزدیک در حافظه نگاشت می‌دهد. ونگ و دیگران [14] از تگ و الگوسازی موضوعی^۱ برای تولید یک

¹ Topic Modeling

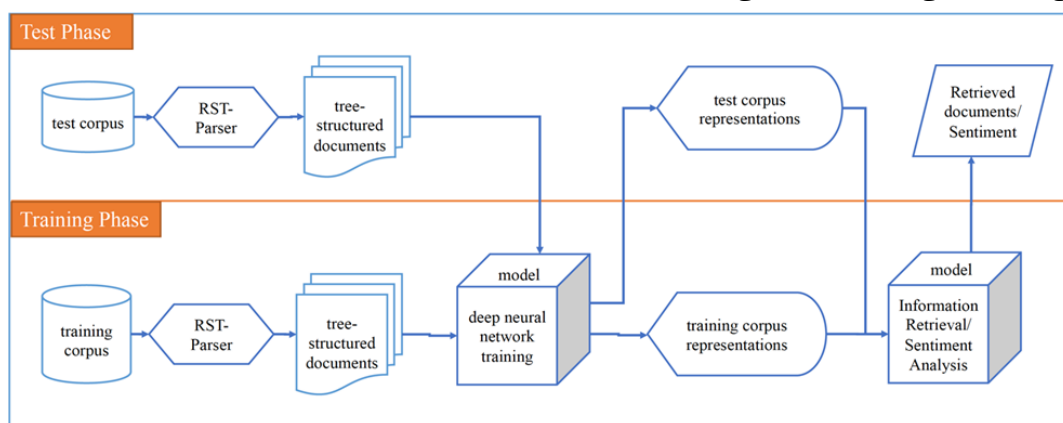
بازنمایی معنایی بهره می‌برند. لیورمور و دیگران [15] ترکیبی از الگوسازی موضوعی را به همراه شبکه استناد^۲ برای بازنمایی اطلاعات معنایی در هنگام بازیابی اطلاعات بهره برده‌اند. این روش‌ها ترتیب کلمات و عبارات را حفظ نمی‌کنند.

هانگ و دیگران [16] بازنمایی درخواست و اسناد را در فضای برداری معنا به‌طور توأمان آموزش می‌دهند. در این رویکرد درخواست و سند مرتبط دارای امتیاز بالا و درخواست و سند غیرمرتبط دارای امتیاز پایین هستند. مولر [18] یک شبکه دوقلو مبتنی بر شبکه حافظه کوتاه‌مدت ماندگار (LSTM) را برای ارزیابی تشابه معنایی بین جملات به کار برده‌اند. لیوما و دیگران [17] روابط بلاغی را در وظیفه بازیابی متن به کار برده و به بهبود قابل‌ملاحظه‌ای (>۱۰٪ در میانگین دقت) در مقابل روشهای مبنا دست یافته‌اند. جی و اسمیت در [19] از شبکه عصبی بازگشتی و سازوکار توجه برای دستیابی به بازنمایی متن متمرکز بر هسته متن استفاده کرده‌اند. یین و دیگران [20] یک شبکه پیچشی به همراه سازوکار توجه برای الگوسازی جفت جملات ارائه کرده‌اند. ونگ و می در [21] از یک شبکه عصبی برای بازنمایی متن به صورت بردار توزیع شده استفاده کرده و با خوشه‌بندی داده‌های دارای برچسب به ارزیابی بازنمایی‌ها پرداخته‌اند. بومن و دیگران [22] یک شبکه خودرمزگذار مبتنی بر شبکه عصبی بازگشتی را برای بازنمایی توزیع پنهان کل جمله معرفی کرده‌اند.

لیوما و دیگران [17] روابط بلاغی را در وظیفه بازیابی متن به کار برده و به بهبود قابل‌ملاحظه‌ای (>۱۰٪ در میانگین دقت) در مقابل روش‌های مبنا دست یافته‌اند. اما از مزایای شبکه‌های عصبی بهره نبرده‌اند. بهیوتا و دیگران [23] از شبکه عصبی بازگشتی مبتنی بر rst برای ترکیب edu بهره برده‌اند. آن‌ها از وزن‌های ازپیش‌تعریف شده برای تغییر اثربخش‌های مختلف درخت وابستگی بهره برده‌اند. در مقالات [8] [24] [7]، بخش-های مختلف متن با اعداد ثابت تغییر وزن یافته‌اند. لیو و لاپتا [25] رویکردی مبتنی بر ساختار متنی بدون استفاده از تجزیه‌گر سخن ارائه کرده‌اند. کراس و فریگل [26] با افزودن سازوکار توجه مبتنی بر ضرب با استفاده از sentiWordNet و محاسبه امتیاز قطبیت برای هر برگ در درخت ساختار به استخراج بخش‌های مهم متن پرداخته‌اند. آن‌ها از شبکه حافظه کوتاه‌مدت مبتنی بر rst برای ترکیب edu بهره برده‌اند. بسیاری از شبکه‌های

² Citation

عمیق به بازنمایی جملات و پاراگراف بسنده کرده و توانایی ارائه بازنمایی برای متون طولانی را ندارند.



(شکل - ۲): شمای کلی الگو پیشنهادی

Figure 2 : proposed model schema

بازیابی متن بر اساس یافتن شباهت بین متون است [13]. علاوه بر در نظر گرفتن ساختار نحوی، جهت تسهیل در جستجوی اطلاعات، درخواست و مجموعه مستندات باید دارای بازنمایی‌های قابل قیاس باشند. یکی از روش‌های قابل قیاس‌سازی بازنمایی مستندات، یادگیری توأمان بازنمایی سند درخواست و مجموعه اسناد با استفاده از شبکه‌های دوقلو است.

شبکه‌های دوقلو، نخست، برای وظیفه تصدیق امضا به کار گرفته شد. تصدیق امضا تأیید می‌کرد که آیا دو امضا به یک شخص تعلق دارد یا خیر. شبکه دوقلو شامل دو شبکه یکسان است که در لایه خروجی به یکدیگر متصل می‌شوند و در هر زیرشبکه از ورودی‌های قابل قیاس ویژگی استخراج می‌شود. این ویژگی‌های استخراج شده از ورودی‌ها با یکدیگر الحاق شده و در لایه خروجی امتیازی به آن‌ها انتصاب داده می‌شود. این شبکه زمانی به کار می‌رود که شباهت و یا ارتباطی بین دو ورودی وجود داشته باشد.

۴-۱-۱- شبکه پیشنهادی

هدف اصلی ارائه این روش برای بازنمایی متن در وظیفه بازیابی اطلاعات، یادگیری بازنمایی معنایی حاوی اطلاعات نحوی است. به طوری که هم‌زمان اطلاعات وابستگی جهت جستجوی شباهت الگو شود. به منظور حصول به این هدف از شبکه‌های دوقلو مبتنی بر دو شبکه بازگشتی بهره برده شد. آموزش این شبکه با یک وظیفه دسته‌بندی دوگانه، که سندهای مرتبط را از سندهای غیر مرتبط تمیز می‌دهد، انجام می‌شود. پس از آموزش شبکه،

علاوه بر موارد مذکور، از نظریه ساختار بلاغی در کاربردهای دیگری مانند تحلیل گفتمان [37] و خلاصه‌سازی [38] نیز استفاده شده است.

مشکلات روش‌های ارائه‌شده، عدم نگهداری ترتیب عبارات و کلمات، عدم بهره‌گیری از بردارهای بازنمایی معنایی کلمه و شبکه‌های عمیق، عدم در نظر گرفتن بخش‌های کلیدی متن و یا ایجاد سازوکار توجه از پیش تعریف‌شده، عدم توانایی ارائه بازنمایی برای متون طولانی و عدم در نظر گرفتن تشابهات متنی در زمان ایجاد بازنمایی است. در این مقاله تلاش بر آن است تا با روش ارائه‌شده این موارد مرتفع شود.

۴- الگو پیشنهادی

شمای کلی الگو پیشنهادی در شکل (۲) آمده است. همان‌طور که مشاهده می‌شود، نخست، مستندات آموزش و آزمون با تجزیه‌گر به شکل ساختاریافته درخت دودویی تبدیل می‌شوند؛ سپس، آموزش شبکه عصبی عمیق بازگشتی بر حسب وظیفه مدنظر بر روی دادگان آموزش انجام می‌شود. الگوی آموزش داده‌شده جهت ایجاد بازنمایی اسناد استفاده می‌شود. بازنمایی‌های آموخته شده نیز برای آموزش الگوی وظیفه مدنظر به کار گرفته شده و خروجی مناسب تولید می‌شود.

۴-۱- وظیفه مورد ارزیابی نخست: بازیابی اطلاعات

هدف از وظیفه بازیابی اطلاعات، یافتن متون مرتبط با درخواست اطلاعات است [28]. اساس الگوریتم‌های

در این شکل بردار بازنمایی واحدهای پایه a و b و c با V_a و V_b و V_c نمایش داده شده‌اند. W_1 تابع ترکیب برای ساختن بازنمایی اسناد است. در این مقاله سند مورد جستجو و مجموعه اسناد ماهیت یکسانی دارند، بنابراین، تابع ترکیب برای هر دو یکسان است. در هر گره از درخت rst تابع ترکیب یکسانی بازنمایی گره والد را شکل می‌دهد. بازنمایی حاصل در ریشه درخت، بازنمایی کل سند است.

پ. شبکه دوقلو برای الگو کردن ارتباط مستندات

شبکه دوقلو ارائه شده شامل دو زیر شبکه RNN است که وزن‌های یکسانی دارند. همان‌طور که اشاره شد، بازنمایی‌های edu با هم ترکیب شده و بازنمایی والد را ایجاد می‌کنند. در ریشه درخت rst ، بازنمایی‌های دو سند با یکدیگر ترکیب شده و با استفاده از یک شبکه چندلایه به خروجی متصل می‌شوند. لایه خروجی میزان ارتباط دو سند را محاسبه می‌کند. **Error! Reference source not found.** معماری شبکه ارائه شده را نمایش می‌دهد.

شبکه برای کمینه‌سازی میانگین مربّع خطا بر روی خروجی واقعی و خروجی مورد انتظار بر روی داده ارزیابی آموزش می‌بیند. مشتق تابع خطا برای تمام وزن‌های موجود محاسبه شده و خطا در طول ساختار درخت پس‌انتشار می‌یابد و وزن‌ها با گرادین کاهشی به‌روزرسانی می‌شوند. برای یکسان نگه‌داشتن وزن‌ها از میانگین گرادین وزن‌های زیر شبکه بازنمایی سند درخواست و زیر شبکه بازنمایی مجموعه اسناد برای به‌روزرسانی استفاده می‌شود. ΔW_1 با معادله زیر محاسبه می‌شود:

$$\Delta W_{Q/doc} = \frac{1}{2} \sum \left(\sum_{i=0}^{i=d_1} \delta_{W_{1Q}}(i), \sum_{j=0}^{j=d_2} \delta_{W_{1doc}}(j) \right) \quad (2)$$

در این معادله $\delta_{W_{1Q}}(i)$ و $\delta_{W_{1doc}}(j)$ به ترتیب، خطای لایه Δm از درخت سند درخواست و خطای لایه Δm از درخت سند موجود در مجموعه اسناد است. d_1 و d_2 عمق درخت rst از دو زیر درخت را نشان می‌دهد.

۴-۱-۲- پیکربندی آزمایش‌ها

برای تجزیه متن به درخت rst از تجزیه‌گر متن‌باز DPLP [31] بهره برده شده‌است. انتخاب این تجزیه‌گر به دلیل کارایی بهتر آن نسبت به سایر تجزیه‌گرهای موجود است. DPLP در تشخیص روابط بین edu ‌های دارای دقت ۶۱ درصد است.

عمیق‌ترین لایه شبکه عصبی بازگشتی به‌عنوان بازنمایی متن استفاده قرار می‌شود.

در این قسمت به توضیح نحوه بازنمایی بخش‌های کوچک متن، چگونگی ترکیب بخش‌های کوچک متن جهت دستیابی به بازنمایی متون طولانی و همچنین، نحوه محاسبه امتیاز ارتباط میان اسناد در لایه آخر شبکه دوقلو می‌پردازیم.

آ. بازنمایی واحد پایه‌ای سخن

هر برگ از درخت rst یک edu را نمایش می‌دهد. برای محاسبه بازنمایی edu ، از میانگین بردار کلمات [29] موجود در هر edu استفاده شده‌است:

$$v_{edu} = \frac{\sum_{j=1}^n w_j}{n} \quad (1)$$

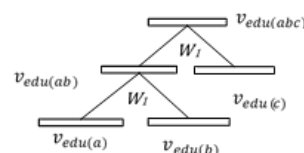
در اینجا n تعداد کلمات موجود در edu و w_j بردار کلمه j امین کلمه است. بردار کلمات پیش‌آموزش‌شده¹ [30] Glove¹ با ابعاد ۱۰۰ در این وظیفه استفاده شد. ایست واژگان حذف شدند و جهت نرمال‌سازی، بردار از تابع تانژانت هایپربولیک^۲ می‌گذرد. این تابع ترکیب، بازنمایی مناسبی برای edu ارائه می‌کند. بهبود این بازنمایی به کارهای آینده واگذار می‌شود.

ب. بازنمایی سند

در اینجا ما به بررسی چگونگی بازنمایی سند و محاسبه امتیاز تشابه آن‌ها می‌پردازیم.

شبکه عصبی بازگشتی

ایده استفاده از شبکه عصبی بازگشتی به‌عنوان تابع ترکیب برای ترکیب برگ‌های درخت نحوی نخست، توسط [32] ارائه شده‌است. در اینجا شبکه عصبی بازگشتی^۳ (RNN)، مشابه [19] الگو شده‌است. شبکه عصبی بازگشتی از برگ‌ها که edu ها هستند آغاز و تا ریشه درخت به‌طور بازگشتی بازنمایی‌ها را می‌سازد. ساختار این درخت در شکل (۳) نمایش داده شده‌است.



شکل ۱: شبکه عصبی بازگشتی

Figure2 : Recursive composition function

1 <https://nlp.stanford.edu/projects/glove/>

2 tanh

3 Recursive Neural Network

به دلیل ماهیت درختی دادگان در شبکه عصبی دوقلو، این شبکه در زبان برنامه نویسی پایتون پیاده سازی و اجرا شد. در فاز پیش پردازش تمام نشانه ها و اعداد حذف شده و حروف بزرگ به کوچک تبدیل شدند. وزن های شبکه دوقلو با مقادیر تصادفی در بازه $[-0.01, 0.01]$ مقداردهی اولیه شده اند. تابع واحد یک سوسده خطی نیست^۱ با معادله زیر به عنوان تابع فعال سازی استفاده شد:

$$f(x) = \max(\varepsilon x, x) \quad (3)$$

در این بررسی مقدار ε برابر 0.1 در نظر گرفته و در لایه ادغام تابع فعال سازی $\tanh$ استفاده شد. نرخ یادگیری با 0.0005 مقداردهی اولیه شد و با معادله (۴) به تدریج و در هر تکرار کاهش یافت.

$$lr_i = \frac{lr_0}{1 + \frac{i}{s}} \quad (4)$$

lr_i نرخ یادگیری در تکرار i ام و lr_0 نرخ یادگیری اولیه است. s شاخصی است که کاهش نرخ یادگیری را نشان می دهد. برای جلوگیری از بیش برآزش روش حذف تصادفی وزن ها^۲ با نرخ ۰.۵ استفاده شد.

۴-۱-۳- دادگان

دو مجموعه داده خبری برای ارزیابی بازنمایی ارائه شده استفاده شده است: دادگان خبری رویترز^۳ و دادگان خبری بی بی سی^۴. از مجموعه دادگان رویترز هشت موضوع شامل earn, ship, acquisition, grain, money, trade, crude و interest که پرکاربردترین موضوعات هستند، انتخاب شد.

دادگان بی بی سی شامل مقالاتی در بازه زمانی ۲۰۰۴-۲۰۰۵ هستند و هر مقاله با یکی از پنج موضوع زیر برچسب گذاری شده است:

tech و business, entertainment, politics, sport,

داده به سه بخش آموزش، آزمون و ارزیابی تقسیم شده است. جدول (۲) آماره دادگان را نمایش می دهد:

(جدول-۲): آماره پیکره

Table 2: Corpus statistics

تعداد نمونه ها	رویتز	بی بی سی
آموزش	2584	1803
آزمون	782	221
اعتبارسنجی	534	201

¹ Leaky Relu

² dropout

³ <http://archive.ics.uci.edu/ml/datasets/reuters-21578+text+category+collection>

⁴ <http://mlg.ucd.ie/datasets/bbc.html>

برای آموزش شبکه نیاز به جفت سندهای مرتبط و غیرمرتبط است. در این مقاله، مستنداتی که دارای موضوع یکسانی هستند، مرتبط قلمداد می شوند. جفت اسناد موجود در هر موضوع نمونه های مرتبط و سند موجود در هر موضوع با سندهای موجود در سایر موضوعات که به صورت تصادفی انتخاب می شوند، نمونه های غیرمرتبط را ایجاد می کنند. آماره این نمونه ها در **Error!** **Reference source not found.** موجود است. امتیاز تشابه برای نمونه های مرتبط (۱) و برای نمونه های غیرمرتبط (۰) است.

(جدول-۳): آماره نمونه ها

Table 3: Pairs statistics

تعداد نمونه ها	رویتز	بی بی سی
آموزش	100000	100000
آزمون	800	800
اعتبارسنجی	200	200

شبکه با بیش از ۲۰۰۰۰۰ نمونه مرتبط و غیرمرتبط آموزش داده می شود. نقطه اتمام آموزش شبکه کمینه خطا بر روی داده ارزیابی است.

بعد از پایان آموزش شبکه، ایجاد بازنمایی متن از طریق تابع ترکیب و در ساختار درختی صورت می پذیرد. وزن های لایه نخست، به عنوان ترکیب کننده ساختار درختی، یک سند را به یک بردار صد بعدی نگاشت می دهد. بازنمایی ایجاد شده با محاسبه و مقایسه کارایی در وظیفه بازیابی اطلاعات ارزیابی می شود.

۴-۱-۴- معیار ارزیابی

برای ارزیابی سامانه پیشنهادی از معیار دقت در k سند بازیابی شده استفاده شده است. این معیار میزان ارتباط در تعدادی از اسناد بازیابی شده را با سند درخواست ارزیابی می کند. همان طور که اشاره شد، ارتباط در اینجا به معنای داشتن موضوع مشترک است.

از معیار شباهت کسینوسی برای هر سند در داده آزمون با همه اسناد آموزش بهره برده شده و k سند با بیشترین میزان شباهت بازیابی شدند. با استفاده از معادله (۱) از $P@k$ محاسبه شده برای تمام اسناد درخواست در داده آزمون میانگین گرفته می شود:

$$P@k = \frac{\sum_{q=1}^n \left(\frac{\text{the number of retrieved relevant document}}{k} \right)_q}{n} \quad (5)$$

K تعداد اسناد بازیابی شده و n تعداد اسناد درخواست است. در این مقاله دقت بازیابی در $k=50$ گزارش شده است.

۵-۱-۴- نتایج

جهت مقایسه روش ارائه شده با روش پایه‌های پایه، برخی از روش‌های سنتی پیاده‌سازی و ارزیابی شدند. روش TF-IDF یکی از روش‌های پرکاربرد در زمینه بازنمایی متون طولانی است. این بازنمایی هر کلمه را با فراوانی آن و میزان اهمیت آن در کل مجموعه داده وزن دهی می‌کند [33]. این بازنمایی ترتیب کلمات را در نظر نگرفته و شباهت بین کلمات را بازتاب نمی‌دهد. روش LDA^1 نخست، توسط [34] معرفی شده است. این روش هر سند را به صورت مجموعه‌ای از موضوعات در نظر می‌گیرد. LSA [35] مجموعه‌ای از مفاهیم مرتبط با سند و کلمات تولید می‌کند. روش دیگر محاسبه میانگین بردار کلمات موجود در سند بعد از حذف ایست واژگان است.

روش‌های مبنا در بازنمایی متن، ساختار بلاغی متن را در نظر نگرفته و سند را به عنوان سبد کلمه در نظر می‌گیرند. در اینجا می‌توان تأثیر شبکه عصبی عمیق مبتنی بر نظریه ساختار بلاغی در بازنمایی متن را مشاهده کرد. **Error! Reference source not found.** (۴) نتایج این الگوها را بر روی دادگان روترز و بی‌بی‌سی با معیار $P@k$ نمایش می‌دهد.

(جدول ۴): معیار $P@k$ ($k=50$)

روش	ابعاد	رویترز	بی‌بی‌سی
TF-IDF	100	0.6086	0.6176
TF-IDF (بهترین نتایج)	Reuters: 500 BBC: 2600	0.6206	0.7737
میانگین بردار کلمات	100	0.6677	0.8653
LSA	100	0.7122	0.8664
LDA	100	0.5958	0.6578
LDA (بهترین نتایج)	Reuters: 60 BBC: 40	0.6591	0.7380
SDS-RNN (الگوی پیشنهادی)	100	0.7403	0.9289

نتایج نشان می‌دهد بازنمایی ارائه شده نسبت به LDA بیش از ده درصد بهبود داشته است و عملکرد آن در داده روترز سه درصد و در داده بی‌بی‌سی شش درصد بهتر از بازنمایی LSA بوده است. این مقادیر نسبت به بازنمایی TF-IDF، ۱۵٪ و ۱۱٪ و نسبت به بازنمایی میانگین بردار کلمات شش درصد و هفت درصد هستند.

¹ Latent Dirichlet Allocation

قابل توجه است که ابعاد بازنمایی ارائه شده در این مقاله ۱۰۰ بعد بوده که با بهترین ابعاد سایر روش‌ها نیز مقایسه شده است.

در شکل (۷)، یک بازنمایی دوبعدی از بازنمایی داده آزمون نشان داده شده است. این بازنمایی توسط t-sne ایجاد شده است. t-sne داده با ابعاد بالا را در فضای دو یا سه بعدی نمایش می‌دهد. در اینجا هر رنگ به یک موضوع انتساب داده شده است. هر موضوع نشان داده شده در شکل به یک سند آزمون برمی‌گردد.

همان‌طور که در شکل دیده می‌شود، بازنمایی ارائه شده با روش SDS-RNN سندهایی با موضوعات یکسان را به فضای مشابهی نگاشت می‌کند و مانند سایر بازنمایی‌ها در کل فضای بازنمایی پخش نشده است.

۴-۲- وظیفه مورد ارزیابی دوم: تحلیل

احساسات

وظیفه تحلیل احساسات [39] به دسته‌بندی قطبیت متون می‌پردازد. از کاربردهای این وظیفه می‌توان به تعیین میزان مثبت یا منفی بودن نظرات بینندگان فیلم اشاره کرد.

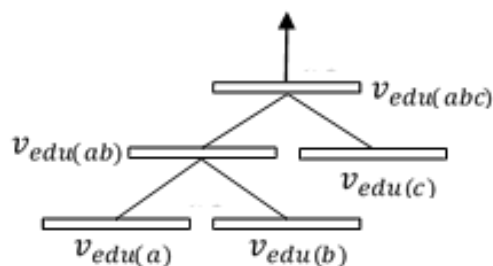
۴-۲-۱- الگوی پیشنهادی

در اینجا یک سازوکار توجه از طریق شبکه بازگشتی عمیق جهت الگوسازی اهمیت بخش‌های مختلف متن ارائه شده است. این کار به رویکرد مقاله [19] شبیه بوده، اما در استفاده از یک لایه عمیق جهت اعمال توجه به بخش‌های مختلف متن متفاوت است.

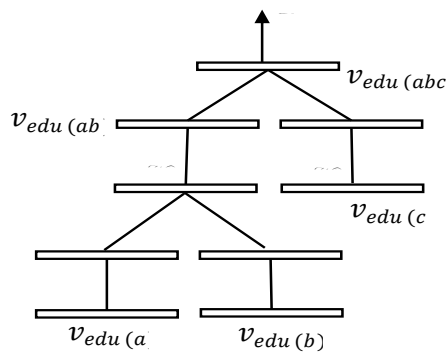
در اینجا برخلاف [26] ما از بردار کلمات، و نه امتیاز قطبیت نظرات برای بازنمایی eduها بهره بردیم. همچنین، علاوه بر استفاده از وزن‌های ازپیش‌تعیین شده جهت مقایسه، از وزن‌های آموزش دیده شبکه عصبی عمیق بهره بردیم.

الف. بازنمایی متن

معماری شبکه پیشنهادی در شکل (۴) ارائه شده است:



شکل ۴: طرح‌واره شبکه عصبی بازگشتی برای تحلیل احساسات



(شکل-۲): شبکه عصبی بازگشتی با سازوکار توجه عمیق

Figure 4: Recursive composition function with deep layer attention mechanism

قبل از ترکیب دو edu بر اساس نوع سلسله مراتب آنها سازوکار توجه اعمال می شود. بردارهای حاصل الحاق شده و به گره والد ارسال می شوند و این روند تا ریشه درخت ادامه می یابد. در نهایت به لایه $softmax$ وارد شده و مشتق تابع خطا نسبت به وزن های $\{W_1, W_2, W_N, W_S\}$ گرفته شده و خطا در ساختار پسانتشار [36] می یابد و شاخص ها با گرادین کاهشی به روزرسانی می شوند.

وزن های سازوکار توجه هسته و پیرو با استفاده از معادلاتی که در ادامه آمده اند، به روزرسانی می شوند:

(۷)

$$\Delta_{W_{N/S}} = \sum_{i=0}^{i=d} \delta_{W_{N/S}}(i)$$

$$\delta_{W_{N/S}}(i) = W_1[:, i] * \delta_{d-1} \odot dReLU(W_{N/S} * v_{edu})$$

در اینجا $\Delta_{W_{N/S}}$ خطای کل برای به روزرسانی وزن های سازوکار توجه است و $\delta_{W_{N/S}}$ خطا در هر سطح درخت است. $dReLU$ مشتق تابع $ReLU$ و δ_{d-1} خطای لایه قبل است.

۴-۲-۲- پیکربندی آزمایش ها

تنظیمات این آزمایش مشابه تنظیمات ارائه شده در وظیفه بازیابی اطلاعات بوده و علاوه بر آن این آزمایش بر روی ابعاد مختلف بردار کلمه $\{50, 100, 200, 300\}$ انجام شده است.

۴-۲-۳- دادگان

برای ارزیابی الگوی ارائه شده در وظیفه تحلیل احساسات از دادگان اینترنتی نظرات بینندگان فیلم^۳ بهره بردیم. دادگان imdb دارای ۲۵۰۰۰ داده آموزش و ۲۵۰۰۰ داده آزمون است. در هر بخش از دادگان آموزش و آزمون

³ Internet movie database (imdb)

Figure 3 : Recursive composition function for sentiment analysis

در اینجا نیز واحدهای پایه متن a و b و c با v_a و v_b و v_c نشان داده شده اند. W_1 تابع ترکیب و W_2 لایه بیشینه هموار^۱ است. در هر گره از تابع ترکیب یکسانی برای ترکیب edu ها بهره برده می شود. در ریشه درخت، بردار برای دسته بندی به لایه $softmax$ فرستاده می شود. شبکه جهت کمینه سازی خطای آنتروپی^۲ احتمال تعلق بردار به رده مربوط به آن آموزش می بیند.

ا. بازنمایی متن با استفاده از سازوکار توجه حاصل

از درخت rst

همان طور که در مقدمه اشاره شد edu هسته حاوی اطلاعات بیشتری نسبت به پیرو است. اما شبکه عصبی بازگشتی تفاوتی بین بخش های مختلف متن ایجاد نمی کند. یکی از روش های پیشنهاد شده جهت بهبود بازنمایی متن، تغییر سهم edu های هسته و پیرو در بازنمایی نهایی بر اساس این ویژگی است.

در اینجا دو روش برای افزایش تأکید بر هسته و کاهش تأیید بر پیرو به کار رفته است. در روش نخست، بازنمایی edu در وزن های عددی ثابت ضرب شده اند. در روش دوم سازوکار توجه در حین آموزش دسته بند آموخته می شود.

ب. شبکه عصبی بازگشتی با وزن های عددی ثابت

در این روش سازوکار توجه از طریق ضرب بازنمایی edu در یک مقدار عددی ثابت (λ) ، اعمال می شود. دو بازه عددی متفاوت برای افزایش تأکید بر بردار هسته و کاهش تأکید بر پیرو در نظر گرفته شده است. در هر آزمایش، λ ۰.۱ افزایش یافته است.

(۶)

$$v_{edu} = \lambda \circ v_{edu} \quad \begin{cases} \lambda \in [1, 2] \text{ if } edu \text{ is nucleus} \\ \lambda \in [0, 1] \text{ if } edu \text{ is satellite} \end{cases}$$

ت. شبکه عصبی بازگشتی عمیق با سازوکار توجه

در این رویکرد، با افزودن یک لایه عمیق در هر برگ درخت سهم هر بخش از متن با توجه به سلسله مراتب آن، هسته یا پیرو بودن، تغییر می کند. در اینجا دو وزن متفاوت برای توجه بر هسته W_N و پیرو W_S در نظر گرفتیم. این وزن ها سازوکار توجه را الگوسازی می کنند. W_N وقتی edu هسته است و W_S وقتی edu پیرو است، اعمال می شود. شکل (۵) معماری سازوکار توجه ارائه شده را نشان می دهد.

¹ softmax

² cross-entropy

۱۲۵۰۰ داده با قطبیت منفی و ۱۲۵۰۰ داده با قطبیت مثبت وجود دارد.

۴-۲-۴- نتایج

نتایج تمامی آزمایش‌ها با معیار F1 در جدول (۵) آمده است. بهترین نتایج به صورت پررنگ مشخص شده‌اند. الگوی مبنای یک، شبکه عصبی بازگشتی بدون سازوکار توجه است. الگوی مبنای دو، نتایج اعمال سازوکار توجه عددی ثابت بر روی هسته و پیرو را نشان می‌دهد. DRNN نتایج حاصل از شبکه عصبی عمیق به همراه سازوکار توجه را نشان می‌دهد. به طور کلی اعمال سازوکار توجه نتایج بهتری نسبت به روش بدون این سازوکار ارائه می‌دهد. در سازوکار توجه عددی کاهش میزان سهم پیرو با وزن ۰.۹ باعث افزایش میزان F1 شده است. می‌توان نتیجه گرفت که کاهش توجه بر پیرو منجر به بهبود بازنمایی شده است. سازوکار توجه از طریق شبکه عمیق باعث افزایش سه درصد معیار F1 نسبت به شبکه عادی و دو درصد نسبت به سازوکار توجه عددی شده است. همان طور که نتایج نشان می‌دهد، لایه عمیق در شبکه آموزش بازگشتی به بهبود بازنمایی کمک می‌کند. با آزمودن فرضیه، تفاوت نتایج در روش ارائه شده و روش‌های مبنا $p < 0.05$ است، که نشان‌دهنده یک تفاوت معنی‌دار است. در ارزیابی مؤلفه‌ها در اغلب موارد افزایش تعداد بعد منتج به افزایش دقت الگو می‌شود.

۵- نتیجه‌گیری

در این مقاله رویکردی نوین جهت در نظر گرفتن ترتیب و اهمیت عبارات موجود در متن و همچنین، تشابهات متنی

در بازنمایی متون طولانی ارائه شد. این بازنمایی در وظیفه بازبایی متن و همچنین، تحلیل احساسات در زبان انگلیسی بررسی شد.

در اینجا یک شبکه عصبی دوقلو شامل دو زیرشبکه بازگشتی بر روی ساختار درخت rst، جهت تولید بازنمایی اسناد و ارزیابی ارتباط آن‌ها طراحی شده است. نتایج نشان می‌دهد به کارگیری این شبکه ضمن تسهیل یافتن تشابه متون، کارایی بازبایی اسناد را بر روی داده خبری رویترز به میزان شش درصد و بر روی داده خبری بی‌بی‌سی به میزان سه درصد بهبود داده است.

همچنین، در وظیفه تحلیل احساسات بررسی تأثیر به کارگیری اطلاعات درخت rst شامل سلسله‌مراتب گره در بهبود ایجاد بازنمایی متنی به اثبات رسید. به طوری که نتایج بر روی دادگان imdb به میزان سه درصد نسبت به روش‌های پایه بهتر شده است.

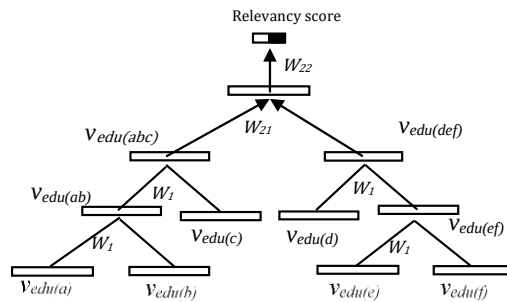
بهینه‌سازی شاخص‌هایی چون نرخ آموزش، ابعاد بردار کلمات و درصد حذف تصادفی وزن‌ها در بهبود الگو کارآمد هستند. یافتن روش‌های بهتر بازنمایی edu و هرس درخت rst باعث کاهش پیچیدگی الگو و بهبود نتایج می‌شود. علاوه بر آن به کارگیری اطلاعات موجود در نوع رابطه بین eduها و عمق گره‌ها در درخت rst می‌تواند منجر به بهبود بازنمایی شود. کارایی پایین تجزیه‌گر دلیل دیگری برای پایین بودن کارایی روش ارائه شده است که در [37] اثبات شده است. با اثبات کارایی روش ارائه شده و همچنین، عدم وجود تجزیه‌گر برای زبان فارسی، نخستین گام طراحی یک تجزیه‌گر دقیق خواهد بود. که در کارهای آینده به آن خواهیم پرداخت.

(جدول - ۵): کارایی تحلیل احساس بر روی دادگان IMDB

Table 5: Performance of sentiment analysis on IMDB test dataset

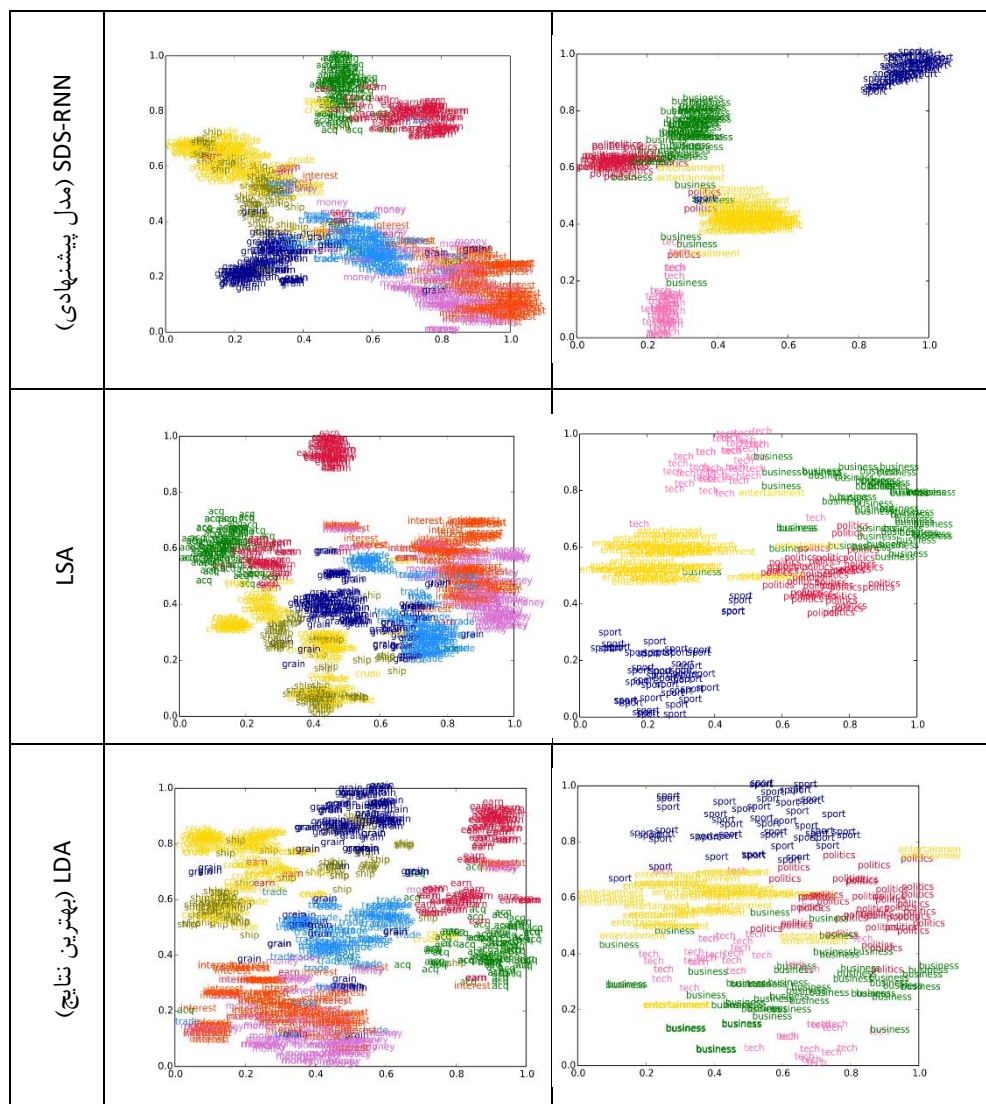
معیار F1	بازخوانی	دقت	صحت	ابعاد	عدد ثابت	سلسله‌مراتب توجه	سازوکار توجه	
73.62	77.74	69.93	72.14	50	1		-	روش مبنا (۱)
75.99	80.51	71.97	74.57	100	1		-	
77.03	82.31	72.4	75.46	200	1		-	
77.23	80.54	74.54	76.29	300	1		-	
73.9	79.75	68.85	71.83	200	1.2	هسته	عددی	روش مبنا (۲)
75.84	79.88	72.2	74.56	300	1.2	هسته	عددی	
77.00	82.32	72.32	75.41	200	1.1	هسته	عددی	
77.48	79.77	75.33	76.82	300	1.1	هسته	عددی	
73.86	84.15	65.81	70.22	50	0.9	پیرو	عددی	
75.95	79.85	72.41	74.71	100	0.9	پیرو	عددی	
77.22	81.85	73.08	75.85	200	0.9	پیرو	عددی	
78.06	83.24	73.48	76.6	300	0.9	پیرو	عددی	

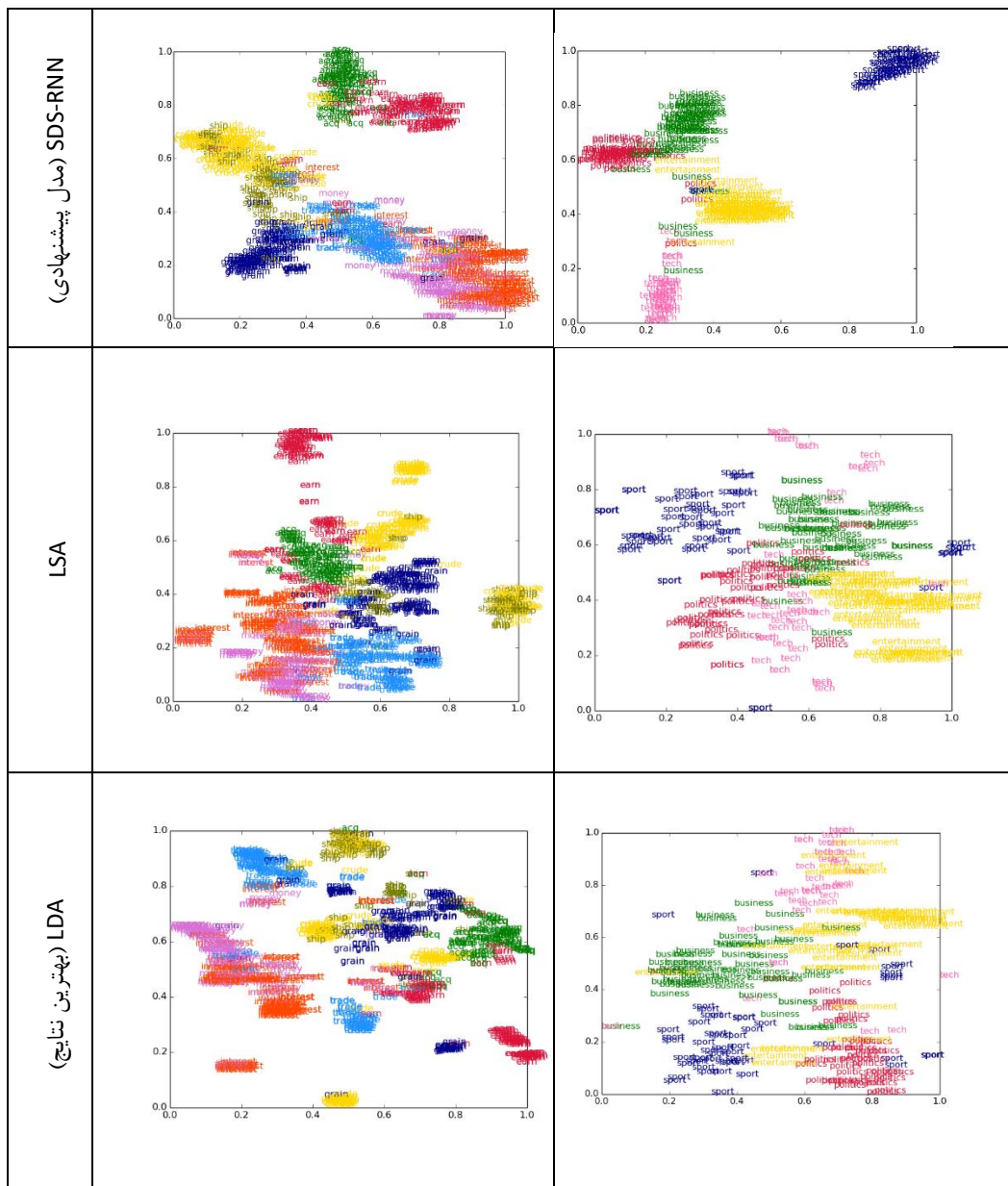
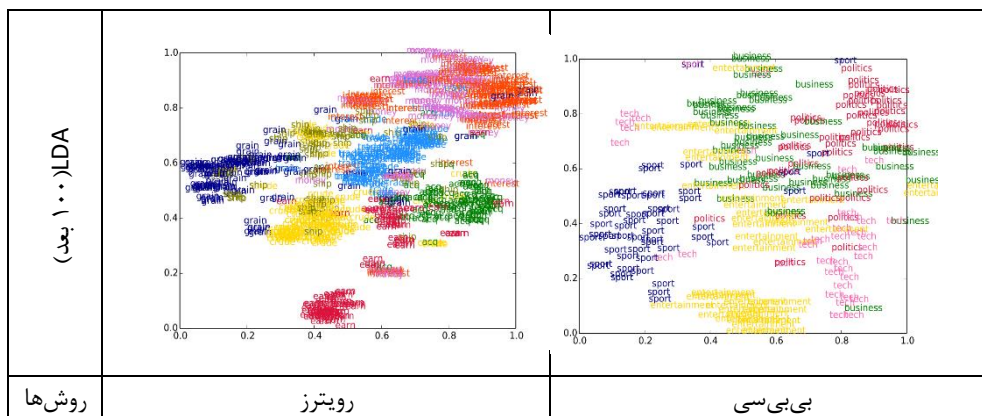
78.67	78.72	79.03	78.72	50	-	هسته و پیرو	DRNN	DRNN
77.88	77.88	77.91	77.88	100	-	هسته و پیرو	DRNN	
79.13	79.14	79.23	79.14	200	-	هسته و پیرو	DRNN	
80.07	80.08	80.18	80.08	300		هسته و پیرو	DRNN	

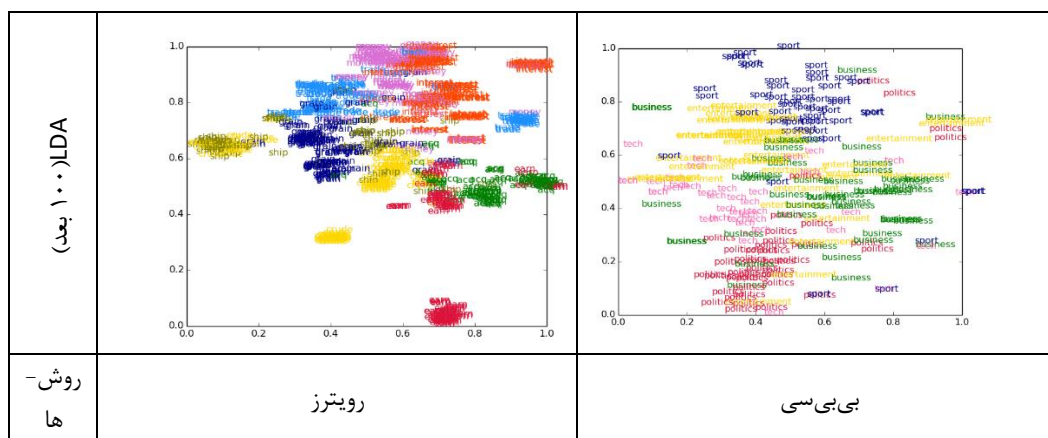


(شکل - ۶): طرح‌واره شبکه دوقلو بازگشتی

Figure 1: Schematic representation for SDS-RNN







شکل ۷: بازنمایی دوبعدی دادگان تست با استفاده از روش های مختلف به کمک t-sne.
برای نمایش بهتر به صورت رنگی مشاهده شود.

Figure 7: A 2-dimensional embedding for representations of test documents by different methods using t-sne. See in color for better visualization.

- sentiment analysis,” *Commun. ACM*, vol. 58, no. 7, pp. 69–77, 2015.
- [9] D. Marcu, “Discourse Trees are Good Indicators of Importance in Text,” in *Advances in Automatic Text Summarization*, 1999, pp. 123–136.
- [10] W. C. Mann and S. A. Thompson, “Rhetorical Structure Theory: Toward a functional theory of text organization,” *Text*, vol. 8, no. 3, pp. 243–281, 1988.
- [11] D. Noel, *Towards a functional characterization of the news of the BBC World Service*. 1986.
- [12] B. A. Fox, *Discourse Structure and Anaphora: Written and Conversational English*. Cambridge University Press, 1993.
- [13] R. Salakhutdinov and G. Hinton, “Semantic hashing,” *Int. J. Approx. Reason.*, vol. 50, no. 7, pp. 969–978, 2009.
- [14] Q. Wang, D. Zhang, and L. Si, “Semantic hashing using tags and topic modeling,” in *Proceedings of the 36th international ACM SIGIR conference on Research and development in information retrieval - SIGIR '13*, 2013, p. 213.
- [15] M. A. Livermore, F. Dadgostari, M. Guim, P. Beling, and D. Rockmore, “Law Search as Prediction,” *Virginia Public Law Leg. Theory Res. Pap.*, no. 2018–61, 2018.
- [16] P. Huang et al., “Learning Deep Structured Semantic Models for Web Search using Clickthrough Data,” *22nd ACM Int. Conf. Conf. Inf. Knowl. Manag.*, pp. 2333–2338, 2013.
- [17] J. Mueller, “Siamese Recurrent Architectures for Learning Sentence Similarity,” *Proc. 30th Conf. Artif. Intell. (AAAI 2016)*, no. 2012, pp. 2786–2792, 2016.
- [18] C. Lioma, B. Larsen, and W. Lu, “Rhetorical Relations for Information Retrieval,” in *Proceedings of the 35th International ACM*

6-Refrence

۶- مراجع

- [1] R. Socher, C. D. C. Manning, and A. Y. A. Ng, “Learning continuous phrase representations and syntactic parsing with recursive neural networks,” *Proc. NIPS-2010 Deep Learn. Unsupervised Featur. Learn. Work.*, pp. 1–9, 2010.
- [2] R. Socher, C. Manning, B. Huval, and A. Ng, “Semantic compositionality through recursive matrix-vector spaces,” in *EMNLP-CoNLL '12: Proceedings of the 2012 Joint Conference on Empirical Methods in Natural Language Processing and Computational Natural Language Learning*, 2012, pp. 1201–1211.
- [3] Y. Bengio, R. Ducharme, P. Vincent, and C. Janvin, “A Neural Probabilistic Language Model,” *J. Mach. Learn. Res.*, vol. 3, pp. 1137–1155, 2003.
- [4] K. S. Tai, R. Socher, and C. D. Manning, “Improved Semantic Representations From Tree-Structured Long Short-Term Memory Networks,” *Proc. 53rd Annu. Meet. Assoc. Comput. Linguist. 7th Int. Jt. Conf. Nat. Lang. Process.*, pp. 1556–1566, 2015.
- [5] N. Kalchbrenner, E. Grefenstette, and P. Blunsom, “A Convolutional Neural Network for Modelling Sentences,” *Acl*, pp. 655–665, 2014.
- [6] Q. V. Le and T. Mikolov, “Distributed Representations of Sentences and Documents,” vol. 32, pp. 1188–1196, 2014.
- [7] J. Märkle-Huß, S. Feuerriegel, and H. Prendinger, “Improving Sentiment Analysis with Document-Level Semantic Relationships from Rhetoric Discourse Structures,” in *Proceedings of the 50th Hawaii International Conference on System Sciences*, 2017, pp. 1142–1151.
- [8] A. Hogenboom, F. Frasinca, F. de Jong, and U. Kaymak, “Using rhetorical structure in

- Learn. Res., vol. 3, no. Jan, pp. 993–1022, 2003.
- [34] S. T. Dumais, G. W. Furnas, T. K. Landauer, S. Deerwester, and R. Harshman, "Using latent semantic analysis to improve access to textual information," in *Proceedings of the SIGCHI conference on Human factors in computing systems - CHI '88*, 1988, pp. 281–285.
- [35] C. Goller and A. Kuchler, "Learning task-dependent distributed representations by backpropagation through structure," *Proceedings of International Conference on Neural Networks (ICNN'96)*, vol. 1, pp. 347–352, 1996.
- [36] M. Morey, P. Muller, and N. Asher, "How much progress have we made on RST discourse parsing? A replication study of recent results on the RST-DT," *Emnlp*, pp. 1330–1335, 2017.



عرفانه غروی دانشجوی مقطع

دکترای رشته فناوری اطلاعات در

دانشکده علوم و فنون نوین دانشگاه

تهران می‌باشد. زمینه‌های پژوهشی

مورد علاقه ایشان پردازش زبان طبیعی، یادگیری عمیق، شبکه‌های عصبی مصنوعی و یادگیری ماشین می‌باشد. نشانی رایانامه ایشان عبارت است از:

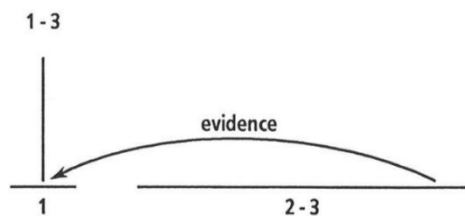
e.gharavi@ut.ac.ir

هادی ویسی دانش‌آموخته مقطع دکترا در رشته مهندسی کامپیوتر از دانشگاه شریف است. در حاضر استادیار گروه علوم و فناوری شبکه در دانشکده علوم و فنون نوین دانشگاه تهران می‌باشد. زمینه پژوهشی مورد علاقه ایشان شامل پردازش سیگنال، پردازش صوت، تشخیص الگو، شبکه‌های عصبی مصنوعی، یادگیری عمیق، منطق فازی و سیستم‌های فازی می‌باشد. نشانی رایانامه ایشان عبارت است از:

h.veisi@ut.ac.ir

- SIGIR Conference on Research and Development in Information Retrieval*, 2012, pp. 931–940.
- [19] Y. Ji and N. Smith, "Imported from Neural Discourse Structure for Text Categorization. (arXiv:1702.01829v1 [cs.CL]) <http://arxiv.org/abs/1702.01829>," *Preprint*, 2017.
- [20] W. Yin, H. Schütze, B. Xiang, and B. Zhou, "Abcnn: Attention-based convolutional neural network for modeling sentence pairs," *arXiv Prepr. arXiv1512.05193*, 2015.
- [21] and A. I. Zhiguo Wang, Haitao Mi, "Semi-supervised clustering for short text via deep representation learning," in *The 20th SIGNLL Conference on Computational Natural Language Learning (CoNLL)*, 2016.
- [22] S. R. Bowman, L. Vilnis, O. Vinyals, A. M. Dai, R. Jozefowicz, and S. Bengio, "Generating sentences from a continuous space," *arXiv Prepr. arXiv1511.06349*, 2015.
- [23] P. Bhatia, Y. Ji, and J. Eisenstein, "Better Document-level Sentiment Analysis from RST Discourse Parsing," *Emnlp*, no. September, pp. 2212–2218, 2015.
- [24] M. Taboada, K. Voll, and J. Brooke, "Extracting sentiment as a function of discourse structure and topicality," *Tech. Rep.*, vol. 20, pp. 1–22, 2008.
- [25] Y. Liu and M. Lapata, "Learning Structured Text Representations," *arXiv Prepr. arXiv1705.09207*, 2017.
- [26] M. Kraus and S. Feuerriegel, "Sentiment analysis based on rhetorical structure theory: Learning deep neural networks from discourse trees," *arXiv Prepr. arXiv1704.05228*, 2017.
- [27] C. D. Manning, P. Ragahvan, and H. Schütze, *An Introduction to Information Retrieval*, no. c. 2009.
- [28] T. Mikolov, K. Chen, G. Corrado, and J. Dean, "Distributed Representations of Words and Phrases and their Compositionality arXiv : 1310 . 4546v1 [cs . CL] 16 Oct 2013," *arXiv Prepr. arXiv1310.4546*, pp. 1–9, 2013.
- [29] J. Pennington, R. Socher, and C. D. Manning, "GloVe: Global Vectors for Word Representation," *Proc. 2014 Conf. Empir. Methods Nat. Lang. Process.*, pp. 1532–1543, 2014.
- [30] R. Socher, "Recursive Deep Learning for Natural Language Processing and Computer Vision," *PhD thesis*, no. August, 2014.
- [31] Y. Ji and J. Eisenstein, "Representation Learning for Text-level Discourse Parsing," *Proc. 52nd Annu. Meet. Assoc. Comput. Linguist.*, pp. 13–24, 2014.
- [32] G. Salton and C. Buckley, "Term-weighting approaches in automatic text retrieval," *Inf. Process. Manag.*, vol. 24, no. 5, pp. 513–523, 1988.
- [33] D. M. Blei, A. Y. Ng, and M. I. Jordan, "Latent Dirichlet Allocation David," *J. Mach.*





شکل ۳: طرح‌واره رابطه گواهی در متن برنامه مالیاتی

Figure 8 : RST diagram for "Tax Program" text

رابطه رجحان^۵

یکی از واضح‌ترین راه‌ها برای نشان دادن رابطه رجحان گزاره اگرچه^۶ است. (جدول-۳) (۷) این رابطه را توصیف می‌کند. در ادامه یک مثال از خلاصه Scientific America آمده‌است:

عنوان: دیوکسین^۷

۱. نگرانی در مورد مضر بودن این مواد برای سلامتی و یا محیط زیست ممکن است نادرست باشد.

۲. اگرچه این ماده برای برخی جانوران سمی است،

۳. شواهد کافی بر تأثیر درازمدت آن بر انسان وجود ندارد.

در این متن نویسنده نشان می‌دهد که واحدهای ۲ و ۳ هم سازگار هستند و هم دارای ناسازگاری بالقوه هستند. بدین صورت که سمی بودن دیوکسین برای برخی از جانوران با فقدان شواهد کافی بر مضر بودن آن برای انسان سازگار است. اما از آنجا که سمی بودن برای جانوران، بر سمی بودن برای انسان دلالت دارد، بطور بالقوه با آن ناسازگار است. (شکل-۹) نمودار rst این متن را نمایش می‌دهد.

(جدول-۳): رابطه رجحان

Table 7: Concession relation

نام رابطه	رجحان
قیود بر هسته	نویسنده رویکرد مثبتی نسبت به موقعیت ارائه شده در هسته دارد.
قیود بر پیرو	نویسنده مدعی عدم درستی پیرو نیست.
قیود بر ترکیب هسته و پیرو	نویسنده به یک ناسازگاری بالقوه یا آشکار بین موقعیت ارائه شده در هسته و پیرو اذعان دارد. همچنین به موقعیت ارائه شده در

5 Concession

6 Although

7 Dioxin

(جدول ۲): رابطه گواهی

Table 3: Evidence relation

نام رابطه	گواهی
قیود بر هسته	خواننده ممکن است هسته را به میزان دلخواه نویسنده باور نکند. ^۲
قیود بر پیرو	خواننده پیرو را باور می‌کند و آن را معتبر ^۳ می‌یابد.
قیود بر ترکیب هسته و پیرو	درک خواننده از پیرو، باور وی از هسته را افزایش می‌دهد.
اثر	باور خواننده از هسته افزایش یافته است.
محل ^۴ اثر	هسته

مثال زیر از رابطه گواهی از یک نامه به سردبیر مجله BYTE استخراج شده است. نویسنده به تمجید از برنامه بازگشت مالیاتی فدرالی که در نسخه قبلی مجله منتشر شده‌است می‌پردازد:

۱- برنامه‌ای که برای سال مالی ۲۰۱۳ انتشار یافته است بخوبی کار می‌کند.

۲ و ۳. تمام فرم‌های بازگشت مالیاتی در عرض چند دقیقه وارد سیستم شده و نتایج حاصل با نتایج محاسبات دستی تا دقیق‌ترین ارقام قابل قیاس است.

نمودار rst در شکل ۳ (۸) واحدهای ۲ و ۳ را در رابطه گواهی با واحد ۱ نمایش می‌دهد. این دو واحد برای افزایش باور خواننده در ادعای مطرح شده در واحد یک ارائه شده‌اند.

¹ Evidence

^۲ در rst باور یک مفهوم درجه‌ای در نظر گرفته می‌شود. که اگرچه جزو ویژگی‌های مرکزی به شمار نمی‌رود اما برخی ویژگی‌های متن را توصیف می‌کند. برای مثال چند خط بودن گواه. تمام قضاوت‌ها از حالت خواننده و عکس‌العمل‌ها لزوماً از نگاه تحلیلگر از نگاه نویسنده ناشی می‌شود. زیرا بر اساس متن است.

³ credible

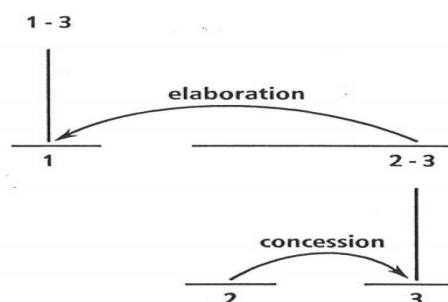
⁴ locus

مثالی از رابطه شرطی در مثال زیر از UCLA PERSONNEL News آمده است. پیرو شرطی به صورت مورب نوشته شده:

۱. از کارمندان درخواست شده بود فرم‌های جدید مزایای بیمه و یا بازنگشتگی وابستگان را پر کنند.
۲. هر زمان که تغییری در وضعیت تأهل و یا وضعیت خانوادگی وجود داشت.
۳. اخیراً مواردی مشاهده شده‌است که یک همسر مطلقه مزایا را دریافت می‌کرده‌است.
۴. زیرا کارمندان در زمینه تکمیل فرم‌های جدید غفلت می‌ورزند.

این متن رابطه شرطی‌ای را که توصیف می‌کند که توسط کلمه هر زمان مشخص شده‌است: پیرو یک موقعیت تحقق نیافته را نام می‌برد. تحقق موقعیتی که در هسته ذکر شده‌است، به تحقق موقعیت ارائه شده در پیرو وابسته است.

هسته و پیرو به صورت سازگار می‌نگرد. تشخیص سازگاری بین موقعیت‌های ارائه شده در هسته و پیرو رویکرد مثبت خواننده نسبت به موقعیت ارائه شده در هسته را افزایش می‌دهد.	
رویکرد مثبت خواننده به موقعیت ارائه شده در هسته افزایش یافته است.	اثر
هسته و پیرو	محل اثر



(شکل ۹): طرح‌واره رابطه رجحان در متن "دیوکسین"

Figure 8 : RST diagram for "Dioxin" text

رابطه شرطی

رابطه شرطی در زبان انگلیسی با گزاره شرطی^۱، قاعده-بندی شده‌است. هرچند، مانند تمام روابط ارائه شده در این بررسی، که روابط معنایی و نه نحوی هستند، رابطه شرطی نیازی به تصریح با عبارت "اگر" ندارد. (جدول این رابطه را شرح می‌دهد.

(جدول ۸): رابطه شرطی

Table 8: Condition relation

نام رابطه	شرطی
قیود بر هسته	ندارد
قیود بر پیرو	پیرو موقعیتی فرضی یا تحقق نیافته را توصیف می‌کند.
قیود بر ترکیب هسته و پیرو	تحقق موقعیت ارائه شده در هسته، به تحقق موقعیت ارائه شده در پیرو بستگی دارد.
اثر	خواننده متوجه می‌شود که چگونه تحقق موقعیت ارائه شده در هسته بستگی به تحقق موقعیت ارائه شده در پیرو دارد.
محل اثر	هسته و پیرو

¹ Conditional clause