

Automatic Intent Classification in a Task-Oriented Chatbot in Education

Elham Salehi

PhD Candidate; Faculty of Linguistics; Institute for Humanities and Cultural Studies; Tehran, Iran Email: e.salehi@ihcs.ac.ir

Masood Ghayoomi*

PhD in Computational Linguistics; Associate Professor; Faculty of Linguistics; Institute for Humanities and Cultural Studies; Tehran, Iran Email: m.ghayoomi@ihcs.ac.ir

Atoosa Rostambeik Tafreshi

PhD in Linguistics; Associate Professor; Faculty of Linguistics; Institute for Humanities and Cultural Studies; Tehran, Iran; Email: a.rostambeik@ihcs.ac.ir

**Iranian Journal of
Information
Processing and
Management**

Received: 04, Dec. 2025

Accepted: 06, Jun. 2026

Abstract: Natural language constitutes a fundamental medium through which humans perform social actions, construct interpersonal relations, and negotiate cultural identities. Far beyond the mere transmission of propositional content, linguistic interaction embodies a structured system of practices —such as turn-taking, adjacency pairing, sequencing, and repair— which collectively enable coordinated social conduct. Since the emergence of conversation analysis in the 1960s, these interactional structures have been systematically examined, while Speech Act Theory, articulated by Austin and subsequently advanced by Searle, has underscored the intrinsically action-oriented nature of utterances. Within this theoretical landscape, the identification of a speaker's communicative intent is a foundational requirement for meaningful interpretation. The advent of artificial intelligence and natural language processing has motivated efforts to computationally replicate this interpretive capacity within task-oriented dialogue systems, where intent detection serves as a core operational component.

Responding to the scarcity of annotated resources in Persian, the present study develops and evaluates a dedicated intent detection system for a Persian language-learning chatbot. To this end, a novel corpus comprising 3,607 sentences was constructed across four semantic relation types —synonymy, antonym, collocation, and idioms— derived from a systematic content analysis of authoritative Persian pedagogical materials and subsequently validated by domain experts. The dataset was thoroughly preprocessed, normalized, and annotated using the IOB scheme to capture the internal semantic structure of utterances. To model

**Iranian Research Institute
for Information Science and Technology
(IranDoc)**

ISSN 2251-8223

eISSN 2251-8231

Indexed by SCOPUS, ISC, & LISTA

Vol. 41 | No. 4 | pp. ???-???

Summer 2026

<https://doi.org/10.22034/ijpm.2026.?????.?????>



* Corresponding Author

lexical meaning numerically, two static embedding methods grounded in distributional semantics, Word2Vec and FastText, were employed, generating 300-dimensional vector representations for all tokens.

Six machine learning algorithms —Support Vector Machine, K-Nearest Neighbors, Random Forest, Decision Tree, Single-Layer Perceptron, and Multi-Layer Perceptron— were trained and assessed using five-fold cross-validation. The experimental results demonstrate that model performance is substantially shaped by the choice of embedding method. FastText, owing to its subword-level encoding, facilitated markedly improved semantic discrimination and resulted in superior performance across most algorithms. The Multi-Layer Perceptron combined with FastText achieved the highest accuracy (99.45%), indicating its strong capacity to model non-linear semantic patterns in Persian. Conversely, Word2Vec exhibited weaker separability for closely related semantic categories, though it yielded strong outcomes in conjunction with distance-based models such as KNN.

Keywords: Intent Detection, Machine Learning Algorithm, Task-Oriented Dialogue System, Vectors, Dialogue Analysis

دسته‌بندی خود کار قصدها در یک چت‌بات وظیفه‌محور در حوزه آموزش

الهام صالحی

دانشجوی دکتری زبان‌شناسی؛ پژوهشکده زبان‌شناسی؛
پژوهشگاه علوم انسانی و مطالعات فرهنگی؛
تهران، ایران e.salehi@ihcs.ac.ir

مسعود قیومی

دکتری زبان‌شناسی رایانشی؛ دانشیار؛
پژوهشکده زبان‌شناسی؛ پژوهشگاه علوم انسانی و
مطالعات فرهنگی؛ تهران، ایران؛
پدیده‌آور رابط m.ghayoomi@ihcs.ac.ir

آرتوسا رستم‌بیگ تفرشی

دکتری زبان‌شناسی؛ دانشیار؛ پژوهشکده زبان‌شناسی؛
پژوهشگاه علوم انسانی و مطالعات فرهنگی؛
تهران، ایران a.rostambeik@ihcs.ac.ir



مقاله برای اصلاح به مدت ۳۲ روز نزد پدیدآوران بوده است.

پذیرش: ۱۴۰۵/۰۳/۱۶

دریافت: ۱۴۰۴/۰۹/۱۳

تشریفات علمی | رتبه بین‌المللی
پژوهشگاه علوم و فناوری اطلاعات ایران
(ایرانداک)

شاپا (چاپی) ۲۳۵۱-۸۲۲۳

شاپا (الکترونیکی) ۸۲۳۱-۲۳۵۱

نمایه در SCOPUS، ISI، و LISTA

jipm.irandoc.ac.ir

دوره ۴۱ | شماره ۴ | صص ۱۱۴۶-۱۱۱۹
تابستان ۱۴۰۵

<https://doi.org/jipm.41.4>



چکیده: گفتار به‌عنوان یکی از بنیادی‌ترین سازوکارهای کنش اجتماعی زبان از دهه ۱۹۶۰، در حوزه تحلیل مکالمه و در کانون توجه زبان‌شناسی قرار گرفت. در این دیدگاه تحلیل مکالمه، گفتار تنها محمل انتقال پیام نیست، بلکه ابزاری برای انجام عمل است. بنابراین، شناسایی «قصد»ی که در پیش‌زمینه مکالمه وجود دارد، بخشی ضروری از درک معنا در هر تعامل زبانی است. با گسترش پردازش زبان طبیعی و افزایش توجه به تعامل ماشین-انسان در قالب سامانه‌های مکالمه وظیفه‌محور و افزایش توجه به آن در هنگام استفاده از مدهای زبانی بزرگ، مسئله «تشخیص قصد کاربر» همچنان به‌عنوان یک چالش مطرح است.

پژوهش حاضر با هدف طراحی و ارزیابی یک سامانه تشخیص قصد در حوزه آموزش زبان فارسی، یک پیکره تخصصی شامل ۳۶۰۷ جمله در چهار حوزه معنایی «هم‌معنایی»، «تضاد معنایی»، «باهم‌آیی» و «اصطلاحات» تدوین شده است. برای این منظور، منابع آموزشی معتبر فارسی مورد تحلیل محتوایی قرار گرفت و الگوهای زبانی پرتکرار با مشارکت متخصصان آموزش زبان استخراج شد. داده‌ها پس از پیش‌پردازش، پاک‌سازی و برچسب‌گذاری با الگوی IOB، برای بردارسازی بر پایه نظریه معناشناسی توزیعی آماده شدند. دو مدل «وردتووک» و «فست‌تکست» با بردارهای ۳۰۰ بُعدی برای بازنمایی معنایی واژگان به کار گرفته شدند.

در ادامه، شش الگوریتم یادگیری ماشین شامل ماشین بردار پشتیبان، K-نزدیک‌ترین همسایه، جنگل تصادفی، درخت تصمیم، پرسپترون تک‌لایه، و پرسپترون چندلایه در قالب اعتبارسنجی متقابل پنج‌تایی مورد ارزیابی قرار گرفتند و نتایج نشان داد که عملکرد مدل‌ها به‌طور چشمگیری تحت تأثیر نوع بردارسازی قرار دارد. «فست‌تکست» با بهره‌گیری از اطلاعات زیرواژه‌ای امکان استخراج مؤثرتر روابط معنایی را فراهم کرده و موجب افزایش چشمگیر عملکرد الگوریتم‌ها شده است. بهترین نتیجه مربوط به پرسپترون چندلایه همراه با «فست‌تکست» با میانگین دقت ۹۹/۴۵ درصد بود. در مقابل، «وردتووک» عملکرد ضعیف‌تری در تمایز مقوله‌های نزدیک معنایی نشان داد.

یافته‌های پژوهش بیانگر آن است که ترکیب بردارسازی‌های غنی و مدل‌های یادگیری مناسب می‌تواند نقش تعیین‌کننده‌ای در تشخیص دقیق قصد کاربران در سامانه‌های مکالمه‌محور فارسی ایفا کند. این پژوهش با ارائه یک پیکره تخصصی و تحلیل جامع الگوریتم‌ها، گامی مهم در توسعه چت‌بات‌های آموزشی زبان فارسی فراهم می‌آورد.

کلیدواژه‌ها: تشخیص قصد، الگوریتم یادگیری ماشین، سامانه مکالمه‌محور، بردارسازی، تحلیل گفتمان

۱. مقدمه

زبان طبیعی نه تنها ابزاری برای انتقال پیام، بلکه سازوکاری بنیادین برای انجام کنش‌های اجتماعی است و گفت‌وگو یکی از اساسی‌ترین شیوه‌های تعامل انسانی است که تداوم جوامع و تبادل فرهنگی را ممکن می‌سازد. «گودوین و هریتیج» بیان می‌کنند که تعاملات اجتماعی بستر اصلی تحقق کنش‌های انسانی همچون تجارت، همکاری و انتقال دانش هستند و از رهگذر همین تعاملات است که هویت و فرهنگ اجتماعی بازتولید می‌شود (Goodwin and Heritage 1990, 283). اهمیت تعاملات زبانی موجب شکل‌گیری شاخه‌ای مستقل در مطالعات زبان‌شناسی به‌نام تحلیل مکالمه^۱ در دهه ۱۹۶۰ شد که هدف آن واکاوی سازوکار درونی گفت‌وگو و سازماندهی تعاملات زبانی است (Liddicoat 2022, 7). این حوزه با رویکردی جامعه‌شناسانه به بررسی این مسئله می‌پردازد که افراد چگونه از طریق زبان، دنیای اجتماعی و سلسله‌مراتب اجتماعی را شکل می‌دهند و تعاملات هدفمند اجتماعی خود را پیش می‌برند (Dema 2021). در چارچوب تحلیل مکالمه، عناصر بنیادینی

1. conversation analysis

چون نوبت‌گیری^۱، جفت‌های همجوار^۲، توالی گفتار، همپوشی^۳، ترمیم^۴ گفتار، و تلویح^۵ به‌عنوان سازه‌های اصلی ارتباط انسانی شناخته می‌شوند (Jurafsky and Martin 2025, 567-576). در امتداد این رویکرد، نظریه کنش گفتاری^۶ که توسط Austin (1962) و سپس Searle (1969) بسط یافت، بر این نکته تأکید دارد که گفتار نه تنها وسیله‌ای برای انتقال معنا، بلکه ابزاری برای انجام عمل است. بر پایه این نظریه، هر جمله دربرگیرنده سه سطح کنش است: کنش بیانی^۷ یا صورت زبانی جمله، کنش منظورشناسی^۸ یا هدف گوینده، و کنش پس‌بیانی^۹ یا اثری که گفتار بر شنونده می‌گذارد. بنابراین، درک معنا در گفت‌وگو مستلزم تشخیص قصد ارتباطی گوینده است؛ زیرا آنچه در ظاهر یک پرسش ساده به نظر می‌رسد، ممکن است در واقع، نوعی درخواست یا دستور غیرمستقیم باشد. برای نمونه، جمله «آیا می‌توانی نمک را به من بدهی؟» از نظر نحوی پرسشی است، اما در سطح کنش گفتاری، به‌معنای صدور یک درخواست است (Yule 2000).

با ظهور فناوری‌های نوین هوش مصنوعی و پردازش زبان طبیعی، تلاش برای شبیه‌سازی همین فرایندهای انسانی در تعامل با ماشین آغاز شد. نخستین نمونه از چنین تلاشی را می‌توان در سامانه «الیسا»^{۱۰} مشاهده کرد که در دهه ۱۹۶۰ توسط «وایزباوم»^{۱۱} طراحی شد. «الیسا» به‌عنوان یک چت‌بات ساده، مکالمه‌ای شبه‌درمانی بر پایه الگوی «درمان راجرزی»^{۱۲} شبیه‌سازی می‌کرد و در ظاهر می‌توانست به پرسش‌ها و اظهارات کاربر پاسخ دهد (Jurafsky and Martin 2025, 145). نمونه‌ای از این نوع مکالمه در شکل ۱، آمده است. با این حال، پرسش بنیادین همچنان باقی ماند: چگونه ممکن است یک سامانه هوشمند، بدون درک انسانی، بتواند قصد و منظور کاربر را شناسایی کند؟

-
1. turn taking
 2. adjacency pairs
 3. overlapping
 4. repair
 5. implicature
 6. speech act theory
 7. locutionary act
 8. illocutionary act
 9. perlocutionary act
 10. ELIZA
 11. Weizenbaum
 12. Rogerian

User: I need some help, that much seems certain.
 ELIZA: WHAT WOULD IT MEAN TO YOU IF YOU GOT SOME HELP
 User: Perhaps I could learn to get along with my mother.
 ELIZA: TELL ME MORE ABOUT YOUR FAMILY
 User: My mother takes care of me.
 ELIZA: WHO ELSE IN YOUR FAMILY TAKES CARE OF YOU
 User: My father.
 ELIZA: YOUR FATHER
 User: You are like my father in some ways.

Weizenbaum (1966)

شکل ۱. نمونه‌ای از گفت‌وگوی «الیزا» با یک کاربر (Jurafsky and Martin 2025: 4)

در پاسخ به این پرسش، پژوهش‌های جدید به بهره‌گیری از روش‌های آماری و یادگیری ماشین روی آورده‌اند تا فرایند «درک معنا» را از طریق الگوهای زبانی و بافتی مدل‌سازی کنند. در این میان، نظریه معناشناسی توزیعی^۱ با تکیه بر فرضیه Harris (1954) مبنی بر اینکه «واژگانی که در بافت‌های مشابه به کار می‌روند، معانی مشابهی دارند»، بستری نظری برای بازنمایی معنای واژگان ب شکل عددی فراهم کرده است (قیومی ۱۴۰۱). این نظریه اساس مدل‌های بردارسازی مانند «وردتوک»^۲ (Mikolov et al. (2013 و «فست‌تکست»^۳ (Bojanowski et al. (2017 را تشکیل می‌دهد که در آن معنا از طریق هم‌وقوعی آماری کلمات در بافت زبانی استخراج می‌شود.

ترکیب دیدگاه‌های زبان‌شناختی (مانند تحلیل مکالمه و نظریه کنش گفتاری) با ابزارهای محاسباتی معناشناسی توزیعی، راهی نوین برای درک قصد کاربر در سامانه‌های مکالمه وظیفه‌محور^۴، از جمله چت‌بات‌ها فراهم می‌کند. هدف از انجام این پژوهش طراحی و پیاده‌سازی یک چت‌بات وظیفه‌محور در حوزه آموزش زبان فارسی است که در فرایند اجرا تلاش می‌شود به صورت تجربی به مقایسه نظام‌مند الگوریتم‌های یادگیری ماشین پرداخته شود. تمرکز اصلی بر بررسی تعامل میان نوع بازنمایی معنایی ایستاده و بررسی عملکرد انواع دسته‌بندها در این سامانه است. سرانجام، این پژوهش به دنبال پاسخ به این مسئله است که کدام‌یک از روش‌های به کار رفته، بهترین عملکرد را در تشخیص قصد کاربران دارد.

1. distributional semantics

2. word2vec

3. fasttext

4. task-oriented

۵. static vectorization: به روش بردارسازی گفته می‌شود که برای هر کلمه یک بردار ثابت تولید می‌کند (Jurafsky and Martin 2025, 105)

۲. پیشینه مطالعاتی

مطالعات تحلیل مکالمه ریشه در پژوهش‌های کلاسیک «ساکس، شکلف و جفرسون» دارد که برای نخستین بار نظام نوبت‌گیری و ساختار قاعده‌مند گفت‌وگوهای انسانی را تبیین کردند و نشان دادند که مکالمات طبیعی بر اساس الگوهای منظم مانند نوبت‌ها، جفت‌های همجوار و فرایندهای ترمیم سامان می‌یابند (Sacks, Schegloff & Jefferson 1974). این دیدگاه، بنیان نظری بسیاری از پژوهش‌های بعدی را در تحلیل مکالمه و طراحی سامانه‌های مکالمه و چت‌بات‌ها در هوش مصنوعی فراهم کرده است. در ادامه، پژوهش «گونزالس-لورت» با تلفیق تحلیل مکالمه و نظریه کنش گفتاری نشان می‌دهد که تشخیص قصد و معنا در تعامل زبانی تنها با توجه همزمان به ساختار گفت‌وگو و نوع کنش گفتاری ممکن است؛ رویکردی که مبنای مهمی برای طراحی سامانه‌های خودکار تشخیص قصد کاربر به شمار می‌رود (González-Lloret 2010). در سال‌های اخیر، پژوهش‌هایی با رویکرد فلسفی و معناشناختی به بررسی کنش گفتاری در سامانه‌های هوشمند پرداخته‌اند. «ویلیامز و بین» با نگاهی انتقادی بررسی می‌کنند که آیا چت‌بات‌ها به‌واقع، توان انجام کنش‌های گفتاری اظهاری را دارند یا فقط به تقلید صوری از آنها می‌پردازند (Williams and Bayne 2024).

در زمینه تشخیص قصد نیز مطالعاتی صورت گرفته که در ادامه آورده شده است. «کیم» در پژوهشی بر پایه شبکه‌های عصبی پیچشی^۱ نشان داد که نوع بردارسازی (ایستا در برابر غیرایستا)^۲ و تنوع معماری شبکه عصبی پیچشی بر دقت مدل در طبقه‌بندی جمله‌ها تأثیر دارد. او در پیکره نقد فیلم^۳ با استفاده از بردارسازی غیرایستا به دقت ۸۱/۵ درصد و در پیکره مقالات خبری با بردارسازی ایستا به دقت ۸۹/۶ درصد دست یافت (Kim 2014). در ادامه، «لیو و لین» با استفاده از شبکه‌های عصبی بازگشتی^۴ مبتنی بر توجه^۵ توانستند

1. convolutional Neural Network

۲. non-static: به دیگر بردارسازی گفته می‌شود که در این مدل‌ها بردارهای هر واژه وابسته به بافت است (Jurafsky and Martin, 2025: 105).

3. film review

4. recurrent neural network(RNN)

5. attention-based

عملکرد مدل را در داده‌های سنجه^۱ «ای تی آی اس»^۲ Hemphill, Godfrey and Doddington (1990) بهبود داده و خطای تشخیص قصد و پر کردن جای خالی را کاهش دهند (Lio and Lane 2016). «چن، ژو و وانگ» نیز با معرفی مدلی مشترک بر پایه «برت»^۳ برای دو وظیفه یادشده بر روی داده‌های سنجه «اسنپس»^۴ (Coucke et al. 2018) و «ای تی آی اس» توانستند دقتی نزدیک به ۹۸ درصد در تشخیص قصد و ۹۷ درصد در پر کردن جای خالی به دست آورند (Chen, Zhuo & Wang 2019).

«دی جنارو» و همکاران با بهره‌گیری از معماری شبکه عصبی حافظه طولانی کوتاه‌مدت^۵ در حوزه تشخیص قصد به دقت ۹۹ درصد در داده‌های آموزش و ۸۷ درصد در داده‌های آزمون رسیدند (Di Gennaro et al. 2020). همچنین «لووان و مگینی» در مطالعه‌ای مروری، سه نوع معماری را بررسی کردند: مدل‌های مستقل برای هر وظیفه^۶، مدل‌های مشترک^۷، و مدل‌های یادگیری انتقالی^۸ که عملکرد بهتری در تطبیق با حوزه‌های جدید داشتند (Louvan and Magnini 2020). «احمد» و همکاران با استفاده از داده‌های سنجه زبان انگلیسی به اسم پیکره «پالسی آی ای»^۹ نیز دو مدل «توالی مشترک» و «توالی به توالی»^{۱۰} را مقایسه کرده و گزارش کردند که مدل دوم دقت بالاتری در هر دو وظیفه دارد (Ahmad et al. 2021).

در حوزه زبان فارسی، مطالعات مرتبط با تشخیص قصد و پر کردن جای خالی هنوز در مراحل آغازین خود قرار دارند، ولی روندی روبه‌رشد را طی می‌کنند. نخستین تلاش‌ها را می‌توان در پژوهش «قائم و کاهانی» مشاهده کرد که با به کارگیری رویکردی ترکیبی از دو دسته‌بند یادگیری ماشین (ماشین بردار پشتیبان و نمایش پراکنده) در کنار

1. benchmark

۲. ATIS (Air Travel Information System): یک پیکره معروف انگلیسی در حوزه سامانه‌های مکالمه وظیفه‌محور است که در اوایل دهه ۱۹۹۰ برای پژوهش در زمینه درک زبان طبیعی طراحی شد

3. Bidirectional Encoder Representations from Transformers (BERT)

4. SNIPS

5. longShort Term Memory (LSTM)

6. task

7. joint

8. transfer learning

9. PolicyIE

10. seq2Seq

یک دسته‌بند قاعده‌بنیان^۱، به بهبود چشمگیری در دقت طبقه‌بندی پرسش‌ها دست یافتند (۱۳۹۵). به‌دنبال آن، «زادکمالی، ممتازی و زینالی» با استفاده از مدل‌های چندزبانۀ ازپیش‌آموزش‌دیده مانند «اکس‌ال‌ام-روبرت»^۲ (Conneau et al. (2020) بر روی داده‌های سنجه «ای‌تی‌آی‌اس» انگلیسی و ترجمه آن به فارسی نشان دادند که یادگیری بین‌زبانی^۳ می‌تواند در مقایسه با مدل‌های تک‌زبانۀ منجر به عملکرد بهتری در تشخیص قصد و استخراج اطلاعات معنایی شود (Zadkamali, Momtazi and Zeinali 2024).

«قهرمانی و شمسی‌فرد» مدلی بین‌زبانی برای فارسی ارائه کردند که از انتقال دانش از زبان‌های پُرمنبع به فارسی بهره می‌گیرد و توانستند دقت سیستم را در هر دو وظیفه تشخیص قصد و پر کردن جای خالی به‌میزان قابل توجهی افزایش دهند (Ghahramani & Shamsfard 2023). همچنین، «کریمی و میرزایی» با ایجاد نخستین پیکره معیار فارسی برای این دو وظیفه و ارزیابی مدل‌های مختلف مبتنی بر «برت و ام‌برت»^۴ (Pires, Schlinger and Garrette 2019) نشان دادند که یادگیری مشترک دو وظیفه به‌صورت هم‌زمان می‌تواند دقت مدل را تا بیش از ۹۳ درصد ارتقا دهد (Karimi & Mirzae 2024). سرانجام، «حاجی رمضانعلی» و همکاران با استفاده از «پارس‌برت»^۵ (Farahani et al. 2020) و افزودن لایه‌های توجه^۶، موفق شدند میانگین نمره اف ۱^۷ مدل‌های فارسی را حدود ۴ درصد بهبود دهند (HajiRamezanAli, Deldar and Homayounpour 2025).

در مجموع، مرور پژوهش‌های انجام‌شده نشان می‌دهد که پیشرفت‌های قابل توجه در حوزه تشخیص قصد و تحلیل مکالمه بیشتر در زبان‌هایی حاصل شده است که از نظر حجم پیکره‌های زبانی، تنوع داده، و در دسترس بودن معماری‌های عمیق و مدل‌های ازپیش‌آموزش‌دیده غنی هستند. در مقابل، پژوهش‌های فارسی در این حوزه به‌دلیل محدودیت منابع داده‌ای و نبود پیکره‌های سنجه استاندارد، اغلب یا بر توسعه ابزارهای اولیه مکالمه متمرکز بوده‌اند یا به‌صورت انتقالی از مدل‌های چندزبانۀ بهره گرفته‌اند.

-
1. rule-based
 2. XLM-Roberta
 3. cross-lingual
 4. mBERT
 5. ParsBERT
 6. attention layers
 7. f1-score

هرچند در سال‌های اخیر بهره‌گیری از مدل‌های چندزبانه و رویکرد انتقال یادگیری گسترش یافته، اما اغلب این مطالعات یا فاقد طراحی پیکره اختصاصی متناسب با یک حوزه کاربردی مشخص بوده‌اند، یا ارزیابی جامعی از مقایسه نظام‌مند الگوریتم‌های مختلف یادگیری ماشین در تعامل با انواع بازنمایی‌های معنایی ارائه نکرده‌اند. بر این اساس، شکاف پژوهشی در سه محور قابل طرح است: ۱) نبود پیکره تخصصی در حوزه آموزش لغات زبان فارسی، ۲) فقدان مقایسه نظام‌مند الگوریتم‌های مرسوم یادگیری ماشین و شبکه عصبی در تعامل با بازنمایی‌های معنایی، و ۳) کمبود تحلیل خطای دقیق در تمایز مقولات قصد.

نوآوری پژوهش حاضر در پاسخ به این شکاف‌ها قابل تبیین است. نخست، یک پیکره تخصصی و ساخت‌یافته در حوزه آموزش واژگان زبان فارسی طراحی و برچسب‌گذاری شده است که همزمان امکان تشخیص قصد و پر کردن جای خالی را فراهم می‌آورد. دوم، عملکرد تعدادی از الگوریتم‌های یادگیری ماشین مرسوم و شبکه‌های عصبی در تعامل با دو روش بردارسازی ایستا و به‌صورت نظام‌مند و تحت اعتبارسنجی متقابل ارزیابی و مقایسه شده است. سوم، تحلیل ماتریس‌های درهم‌ریختگی و الگوی خطاها، تصویری دقیق از چالش‌های معنایی در تمایز مقولات واژگانی فارسی ارائه می‌دهد. بدین ترتیب، این پژوهش افزون‌بر ارائه یک چارچوب تجربی برای ارزیابی مدل‌ها، به توسعه منابع داده‌ای و تحلیل روش‌مند عملکرد الگوریتم‌ها در زبان فارسی نیز یاری می‌رساند.

۳. چارچوب نظری

در راستای هدف پژوهش، این بخش شامل مباحث تحلیل مکالمه، نظریه کنش گفتار، و بازنمایی معنایی واژه‌ها در چارچوب «معناشناسی توزیعی» است.

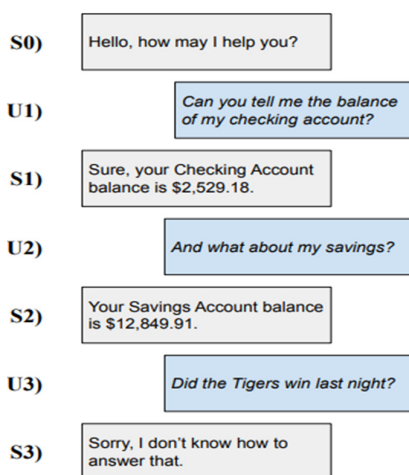
۳-۱. تحلیل مکالمه و نظریه کنش گفتار^۱

تحلیل مکالمه نشان می‌دهد که معنا و قصد در گفت‌وگوهای طبیعی تنها از طریق جمله‌های منفرد قابل درک نیست، بلکه در چارچوب نوبت‌گیری، توالی پاسخ‌ها، جفت‌های همجوار و شیوه‌های ترمیم^۲ سوء تفاهم‌ها شکل می‌گیرد. از سوی دیگر،

1. speech act theory

2. repair

نظریه کنش گفتاری «آستین و سرل» بر این اصل تأکید دارد که گفتار همواره نوعی عمل است و می‌تواند کارکردهایی مانند پرسش، درخواست یا اظهار نظر داشته باشد (Jurafsky and Martin 2025, 569-576). ترکیب این دو دیدگاه امکان شناسایی دقیق‌تر قصد کاربر و تعیین نوع پاسخ مناسب در سامانه‌های مکالمه و چت‌بات‌ها را فراهم می‌کند (Jurafsky and Martin 2025, 569-576). در تصویر ۲، نمونه‌ای از مکالمه بین یک کاربر و یک سامانه مکالمه وظیفه-محور^۱ یا خاص‌منظوره در حوزه بانکداری است که بر اساس تحلیل مکالمه، جمله سوم کاربر نشان از تخطی کاربر از ویژگی جفت همجوار به‌وضوح قابل رؤیت است.



شکل ۲. مکالمه بین یک کاربر و یک سامانه مکالمه وظیفه-محور در حوزه بانکداری (Jurafsky and Martin 2023, 15)

معماری کلی چت‌بات‌ها در تصویر ۳، به نمایش گذاشته شده است. به‌طور کلی، چت‌بات‌ها از چهار بخش اصلی تشکیل شده‌اند: (۱) بخش پردازش زبان طبیعی^۲، (۲) پایگاه داده^۳، (۳) پایگاه دانش^۴، و (۴) تولید طبیعی زبان^۵ (Ghayoomi 2022). همان‌طور

1. task-oriented dialogue system

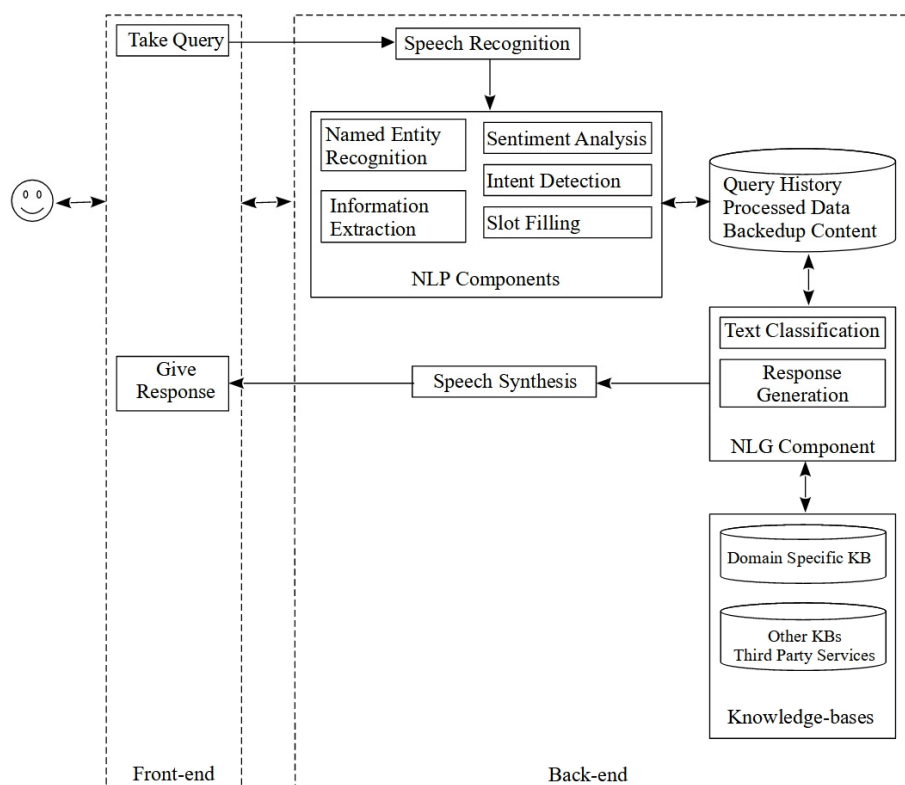
2. NLP

3. database

4. knowledge base

5. NLG

که در تصویر ۳، مشاهده می‌شود، هر یک از این بخش‌ها نقشی مکمل در فرایند دریافت، تحلیل و پاسخ‌گویی به ورودی‌های کاربر دارد. در این میان، بخش پردازش زبان طبیعی هسته اصلی سامانه محسوب می‌شود، زیرا مسئول تبدیل ورودی‌های زبانی انسان به داده‌های قابل درک برای ماشین است. این بخش شامل چندین زیرفرایند کلیدی از جمله تشخیص موجودیت‌ها، تحلیل احساسات، استخراج اطلاعات و به‌ویژه تشخیص قصد کاربر است. از میان این وظایف، تشخیص قصد مهم‌ترین مرحله محسوب می‌شود، زیرا نقطه آغاز درک معنای تعامل انسان و ماشین است.



شکل ۳. معماری کلی چت‌بات (Ghayoomi 2022)

تشخیص قصد نخستین و حیاتی‌ترین گام در چت‌بات‌های وظیفه‌محور است؛ زیرا خطا در این مرحله، کل مکالمه را منحرف می‌سازد.

۳-۲. معناشناسی توزیعی و بردارسازی کلمات

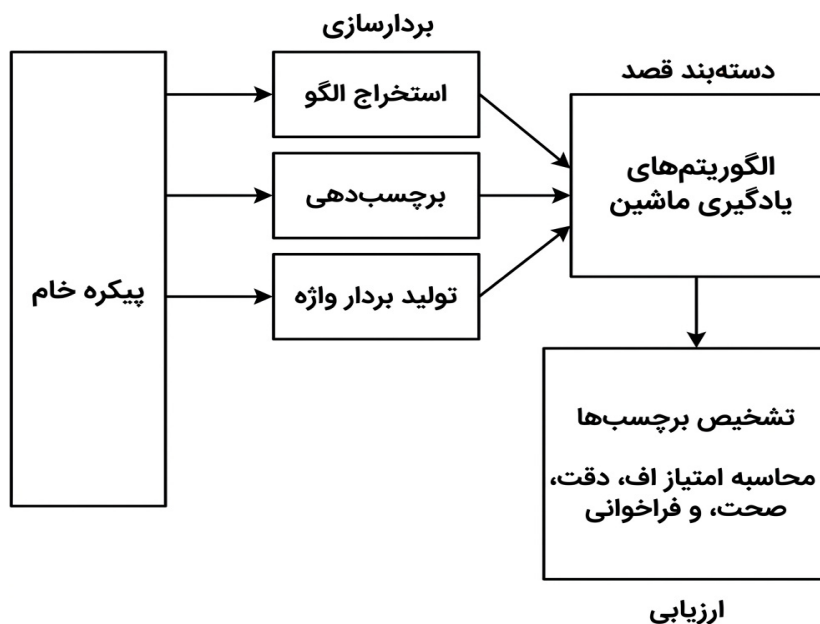
در معناشناسی توزیعی معنای یک کلمه از طریق بافت تعیین می‌شود. «هریس» معتقد است واژه‌هایی که در بافت‌های مشابه به کار می‌روند، بیشتر وقت‌ها معنای مشابهی دارند (Harris 1954). این ارتباط بین مشابهت در بافت توزیعی کلمات و معنای کلمات اساس «فرضیه توزیعی» را تشکیل داد. این فرضیه را «هریس، فرث، و جوز»^۱ برای اولین بار ارائه دادند (Jurafsky and Martin 2025, 95). در این چارچوب، «میکولو» و همکاران مدلی ارائه کردند که امکان نمایش معنای واژه‌ها را به صورت بردار فراهم می‌آورد (Mikolov et al. 2013). در این رویکرد، به جای استفاده از شکل ظاهری کلمه، از بردار واژه بهره گرفته می‌شود که بر اساس الگوی هم‌وقوع و موقعیت آن واژه در یک پیکره زبانی استخراج شده است. افزون‌بر این، معناشناسی توزیعی به عنوان ابزار محاسباتی مکمل این دو نظریه عمل می‌کند. با مدل کردن بردارهای معنایی واژگان و عبارات در متن و بافت مکالمه، می‌توان نشانه‌هایی از قصد را استخراج کرد و توالی‌ها را کمی‌تر تحلیل نمود؛ برای نمونه، واژگانی که بیشتر وقت‌ها، در کنش‌های گفتاری خاص استفاده می‌شوند یا ساختارهایی که در جفت همجوار خاص دیده می‌شوند.

۴. مدل پیشنهادی و داده‌های پژوهش

۴-۱. مدل پیشنهادی

مدل پیشنهادی دسته‌بندی قصد در پژوهش حاضر، در شکل ۴، نمایش داده شده است. روند اجرا به چهار مرحله اصلی تقسیم شده است که شامل تهیه پیکره، بردارسازی، آموزش، و ارزیابی است. بعد از تهیه پیکره، الگوها استخراج و برچسب‌گذاری شده و بردار واژه‌ها تولید می‌شود. به منظور آموزش دسته‌بند، برچسب‌ها، الگوها و داده‌های بردارسازی شده به ماشین داده می‌شود و پس از آموزش، به مرحله ارزیابی می‌رسیم و در این مرحله عملکرد ماشین در زمینه برچسب‌های قابل پیش‌بینی مورد ارزیابی قرار می‌گیرد.

1. Harris, firth and Joes



شکل ۴. مدل پیشنهادی دسته‌بند قصد

۴-۲. داده‌های پژوهش

در نخستین گام به‌منظور گردآوری داده‌های پژوهش، از آنجا که زبان فارسی نسبت به زبان انگلیسی با محدودیت مواجه است و به‌دلیل در دسترس نبودن پیکره‌ای یا داده سنجه‌ای در حوزه تشخیص قصد و پر کردن جای خالی در حوزه آموزش زبان، این فرایند طراحی، و تدوین یک پیکره تخصصی آغاز شد. برای این منظور، ابتدا مجموعه‌ای از کتاب‌های آموزش زبان فارسی برای غیرفارسی‌زبانان از جمله کتاب گام اول «صحرايي» و همکاران (۱۳۹۸)، آموزش کاربردی واژه (همان) و واژه آموزشی زبان فارسی «عسگری» (۱۳۸۴) مورد مطالعه و تحلیل محتوایی قرار گرفتند. پس از مشورت با متخصصان آموزش زبان فارسی و مدرسان حوزه واژه آموزشی، الگوها و موقعیت‌های زبانی متداولی که در تعاملات یادگیری زبان رخ می‌دهند، استخراج شدند تا پرسگان و ساختارهای مورد نیاز چت‌بات بر اساس آنها شکل گیرد.

در ادامه، با الگوبرداری از این منابع و با در نظر گرفتن نیازهای واقعی زبان‌آموز در محیط یادگیری، مجموعه‌ای از جملات و پرسگان طراحی شد که پاسخگوی اهداف آموزشی در حوزه یادگیری واژگان باشد. این پرسگان در چهار دسته معنایی اصلی شامل

«هم‌معنایی»، «تضاد معنایی»، «باهم‌آیی» و «اصطلاحات» سازماندهی شدند. داده‌های نهایی پس از بازبینی و پالایش دستی توسط یک ارزیاب متخصص، در قالب یک فایل متنی با فرمت CSV ذخیره و برای مراحل بعدی پردازش زبان طبیعی آماده شدند. پیکره مورد استفاده شامل حدود ۳۶۰۷ جمله و در مجموع، ۱۵۳۰۵۴ واژه است. لازم به ذکر است که الگوهای زبانی مربوط به چهار قصد ارتباطی، از قبیل «هم‌معنایی»، «باهم‌آیی»، «تضاد معنایی» و «اصطلاح»، به ترتیب، شامل ۸۸، ۶۳، ۶۹ و ۵۸ الگوی یکتا هستند که در مجموع، ۲۷۸ الگوی یکتا در پیکره وجود دارد. نمونه‌هایی از این پیکره در جدول ۱، آورده شده است.

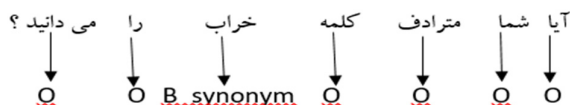
جدول ۱. نمونه‌ای از داده‌های پیکره برای قصدهای هم‌معنایی، تضاد، باهم‌آیی و اصطلاح

الگو	برچسب	قصد	حوزه معنایی	جای خالی
هم‌معنای پوشیدن را بگو.	o b_synonym o o	هم‌معنایی	آموزش زبان	پوشیدن
برنده با چه کلماتی هم‌معناست؟	b_synonym o o o o o	هم‌معنایی	آموزش زبان	برنده
آیا شما مخالف کلمه درشت را می‌دانید؟	o o o o b_antonym o o o	متضاد	آموزش زبان	درشت
چه کلمه متضادی واسه خندان پیشنهاد میدی؟	o o o o b_antonym o o o	متضاد	آموزش زبان	خندان
کلمه سلاح با چه کلماتی بیشتر میاد؟	o o b_collocation o o o o o	باهم‌آیی	آموزش زبان	سلاح
واژه‌های باهم‌آیند با واژه فیلم چیست؟	o o o o b_collocation o o	باهم‌آیی	آموزش زبان	فیلم
با کلمه استخوان اصطلاح به من بده، لطفا.	o o b_idiom o o o o o o o	اصطلاح	آموزش زبان	استخوان

۴-۳. پیش‌پردازش و برچسب‌گذاری

مرحله بعدی که نقش بسیار تعیین‌کننده‌ای در بهبود عملکرد مدل دارد، پیش‌پردازش و برچسب‌گذاری داده‌هاست. این مرحله بیشتر وقت‌ها یکی از مهم‌ترین گام‌های پروژه‌های داده‌کاوی و پردازش زبان طبیعی محسوب می‌شود. هدف از پیش‌پردازش، آماده‌سازی داده‌ها برای مدل‌های یادگیری ماشین از طریق هنجارسازی، پاک‌سازی و غنی‌سازی متون است. در این پژوهش، جملات ابتدا از نظر نویسه‌گذاری، فاصله و نیم‌فاصله، علائم نگارشی، املا و واژه‌ها و کاراکترهای غیرمجاز اصلاح و یکدست شدند تا از بروز خطا در مرحله بردارسازی جلوگیری شود. از آنجا که پیکره طراحی شده شامل جملاتی با جای خالی بود، برای نمایش ساختار درونی جملات و تعیین مرز عناصر معنایی از الگوی

شیوه برچسب‌گذاری IOB استفاده شد. این روش یکی از شناخته‌شده‌ترین الگوهای برچسب‌گذاری در وظایف تشخیص موجودیت‌های نامدار و پر کردن جای خالی است (Tjong, Sang and De Meulder 2003). این الگو شامل سه برچسب است: B شروع جای خالی، I ادامه جای خالی، و O هر واژه‌ای که جای خالی نیست. به عنوان نمونه، در جمله زیر که از پیکره استخراج شده، از این نوع برچسب‌گذاری استفاده شده است:



شکل ۵. نمایش شیوه برچسب‌گذاری IOB در داده‌های پژوهش

این ساختار برچسب‌گذاری باعث می‌شود مدل بتواند مرز دقیق اجزای زبانی مربوط به قصد کاربر را بیاموزد و در مراحل بعدی، هنگام پردازش ورودی‌های جدید، به درستی تشخیص دهد که کدام بخش از جمله نقش معنایی اصلی را ایفا می‌کند.

۴-۴. بردارسازی داده‌ها

در مرحله بعد، تمامی داده‌های متنی و جملات در چارچوب نظریه معناشناسی توزیعی در قالب بردار بازنمایی شدند تا برای پردازش توسط الگوریتم‌های یادگیری ماشین قابل استفاده باشند. به طور کلی، در پردازش زبان طبیعی سه رویکرد عمده برای بازنمایی متون وجود دارد: رویکرد کیسه واژگان^۱، رویکرد بازنمایی توزیع شده^۲، و رویکرد انتقال یادگیری^۳. هر یک از این رویکردها به گونه‌ای متفاوت معنا را در قالب فضای برداری مدل‌سازی می‌کنند. در پژوهش حاضر، به منظور استخراج ویژگی‌های معنایی دقیق‌تر از داده‌های فارسی، از دو مدل «ورد تو وک» و «فست تکست» برای بردارسازی استفاده شد. این دو روش از خانواده مدل‌های بردارسازی ایستا هستند و بر پایه فرضیه معناشناسی توزیعی «هریس» عمل می‌کنند که بیان می‌دارد: واژگانی که در بافت‌های مشابه به کار می‌روند، از نظر معنایی نیز مشابه‌اند (Harris 1954). در این چارچوب هر واژه یا جمله به یک بردار n-بعدی تبدیل می‌شود که فاصله معنایی میان واژگان را در فضای برداری

1. bag-of-words

2. distributed representation

3. transfer learning

بازتاب می‌دهد (Mikolov et al. 2013). بدین ترتیب، مشابهت معنایی میان واژگان و ساختار جمله‌ها به صورت عددی قابل تحلیل و یادگیری توسط مدل‌های طبقه‌بندی قصد و پر کردن جای خالی خواهد بود. در هر دو مدل، تمامی واژگان موجود در پیکره پژوهش به بردارهای ۳۰۰ بعدی تبدیل شدند تا در مراحل بعدی برای آموزش الگوریتم‌های یادگیری ماشین مورد استفاده قرار گیرند. شایان ذکر است که به منظور رفع مشکل واژه‌های ناشناخته^۱ در مراحل بعدی پژوهش، یک پیکره پشتیبان نیز با روش «وردتووک» آموزش داده شد. با این حال، این پیکره تاکنون در مراحل اولیه تحقیق مورد استفاده قرار نگرفته و فقط به عنوان منبع کمکی برای فازهای بعدی پژوهش در نظر گرفته شده است.

۴-۵. الگوریتم‌های یادگیری ماشین و تنظیمات شبکه عصبی

در این پژوهش به منظور شناسایی و دسته‌بندی روابط واژگانی در زبان از روش‌های مختلف بردارسازی واژگان و الگوریتم‌های یادگیری ماشین مرسوم از جمله ماشین بردار پشتیبان^۲ (Vapnik 1995)، جنگل تصادفی^۳ (Breiman 2001)، درخت تصمیم^۴ (Breiman et al. 1984) و K_ نزدیک‌ترین همسایه^۵ (Cover and Hart 1967) و دو مدل مبتنی بر شبکه عصبی شامل پرسپترون تک‌لایه^۶ (Rosenblatt 1958) و پرسپترون چندلایه^۷ (Rumelhart, Hinton and Williams 1986) استفاده شده است.

لازم به ذکر است که در این پژوهش از یک معماری پرسپترون چندلایه مبتنی بر شبکه عصبی پیش‌خور استفاده شده است که شامل یک لایه ورودی متناظر با ابعاد بردارهای واژگانی و دو لایه پنهان است و به ترتیب از ۱۲۸ و ۶۴ نورون تشکیل شده‌اند. برای هر دو لایه پنهان از تابع فعال‌ساز ReLU استفاده شد تا شبکه بتواند الگوهای غیرخطی و وابستگی‌های معنایی پیچیده را در فضای برداری استخراج کند.

به منظور کاهش خطر بیش‌برازش^۸، پس از هر لایه پنهان یک لایه حذف و

-
1. out-of-vocabulary
 2. support vector machine
 3. random forest
 4. decision tree
 5. K-nearest neighbor
 6. single layer perceptron
 7. multilayer perceptron
 8. overfitting

منظم‌سازی^۱ با نرخ ۰/۳ قرار داده شد. این فرایند در هر تکرار آموزش، به‌صورت تصادفی بخشی از نورون‌های لایه را غیرفعال کرده و بدین ترتیب، موجب بهبود قابلیت تعمیم مدل می‌شود. لایه خروجی شبکه از یک واحد چگال با تعداد نورون برابر با تعداد مقولات قصد (چهار دسته معنایی) تشکیل شده و با تابع فعال‌ساز بیشینه‌نرم^۲ احتمال تعلق جمله ورودی به هر دسته را برآورد می‌کند. این ساختار به مدل امکان می‌دهد که افزون‌بر تفکیک دقیق الگوهای زبانی، توزیع احتمالاتی مناسبی برای پیش‌بینی قصد کاربر ارائه کند.

به‌طور کلی، استفاده از دو لایه پنهان با ظرفیت مناسب همراه با حذف و منظم‌سازی برای کنترل پیچیدگی مدل سبب شده است که این معماری بتواند روابط معنایی موجود در داده‌های آموزشی را به‌خوبی یاد بگیرد و در مقایسه با بسیاری از الگوریتم‌های مرسوم یادگیری ماشین، عملکردی به‌مراتب دقیق‌تر و پایدارتر در وظیفه تشخیص قصد ارائه دهد.

۴-۶. ارزیابی

پس از انتخاب روش‌های بردارسازی مناسب، مرحله بعد به آموزش مدل‌ها و ارزیابی عملکرد آنها اختصاص یافت. برای اطمینان از پایداری و قابلیت تعمیم مدل‌ها، از روش اعتبارسنجی متقابل پنج‌تایی^۳ استفاده شد. در این روش، داده‌ها به‌صورت تصادفی به پنج بخش تقسیم می‌شوند. در هر تکرار، چهار بخش (معادل ۸۰ درصد داده‌ها) برای آموزش مدل و بخش باقی‌مانده (۲۰ درصد) برای آزمون به کار می‌رود و تمامی بخش‌ها یک‌بار نقش داده آزمون را ایفا می‌کنند.

ویژگی مهم این روش در آن است که داده‌های آموزشی و آزمون هیچ‌گونه همپوشانی ندارند و بدین ترتیب، عملکرد مدل به‌شکل منصفانه و دقیق ارزیابی می‌شود. تکرار مکرر فرایند آموزش و آزمون با ترکیب‌های متفاوت داده، دو مزیت اساسی به‌همراه دارد: نخست اینکه برآوردی پایدارتر از توانایی تعمیم مدل بر داده‌های نادیده فراهم می‌آورد؛ و دوم اینکه خطر بیش‌برازش در داده‌های آموزشی را کاهش می‌دهد.

در مرحله آموزش، هر یک از الگوریتم‌های مورد نظر به‌صورت جداگانه اجرا و تنظیم شدند. تمامی مراحل پیاده‌سازی با استفاده از زبان برنامه‌نویسی «پایتون» و کتابخانه

1. dropout

2. softmax

3. 5-fold cross validation

استاندارد scikit-learn انجام گرفت. این محیط به دلیل پشتیبانی از طیف گسترده‌ای از الگوریتم‌های یادگیری ماشین، از جمله مدل‌های مرسوم و شبکه‌های عصبی عمیق، امکان بهینه‌سازی و مقایسه دقیق عملکرد مدل‌ها را فراهم ساخت.

برای ارزیابی عملکرد مدل‌های طبقه‌بندی در وظایف تشخیص قصد و پر کردن جای خالی از چندین شاخص متداول در یادگیری ماشین استفاده شد. این شاخص‌ها شامل دقت^۱، صحت^۲، فراخوانی^۳ و میانگین متوازن این دو معیار، یعنی امتیاز اف^۴ هستند. هر یک از این معیارها ابعاد متفاوتی از عملکرد مدل را اندازه‌گیری می‌کنند.

الف) دقت

فرمول ۱، نسبت تعداد پیش‌بینی‌های درست مدل به کل نمونه‌ها را نشان می‌دهد و از رابطه زیر محاسبه می‌شود:

$$accuracy = \frac{tp + tn}{tp + fp + tn + fn} \quad (۱)$$

که در آن،

- ◇ TP، تعداد نمونه‌های مثبت قصد که درست پیش‌بینی شده‌اند.
- ◇ TN، تعداد نمونه‌های منفی قصد که درست پیش‌بینی شده‌اند.
- ◇ FP، تعداد نمونه‌های منفی قصد که به اشتباه مثبت پیش‌بینی شده‌اند.
- ◇ FN، تعداد نمونه‌های مثبت قصد که به اشتباه منفی پیش‌بینی شده‌اند.

ب) صحت

این شاخص در فرمول ۲، نشان می‌دهد که از میان تمام پیش‌بینی‌های صحیح قصد، چه نسبتی به‌واقع درست بوده است.

$$precision = \frac{tp}{tp + fp} \quad (۲)$$

صحت بالا بیانگر این است که مدل در شناسایی نمونه‌های هدف (برای نمونه، قصد خاص یا نوع اسلات^۵ مشخص) اشتباهات اندکی دارد.

1. accuracy
2. precision
3. recall
4. f1-Score
5. solot

ج) فراخوانی

فراخوانی در فرمول ۳، نشان می‌دهد که از میان تمام نمونه قصدهای واقعی صحیح، چند مورد توسط مدل به درستی شناسایی شده‌اند:

$$recall = \frac{tp}{tp+fn} \quad (۳)$$

مقدار بالای فراخوانی به این معناست که مدل توانایی شناسایی بخش بزرگی از قصدها را به صورت صحیح دارد؛ حتی اگر برخی پیش‌بینی‌ها نادرست باشند.

د) معیار امتیاز اف

برای دستیابی به توازن میان فراخوانی و صحت، از میانگین امتیاز اف (فرمول ۴) استفاده می‌شود که به صورت زیر محاسبه می‌شود:

$$f1 = \frac{(1+\beta^2).precision.recall}{(\beta^2.precision)+recall} \quad (۴)$$

امتیاز اف شاخصی جامع‌تر برای ارزیابی مدل در شرایطی است که داده‌ها نامتوازن باشند و تنها دقت کلی نتواند کارایی مدل را به درستی بازتاب دهد (Jurafsky and Martin, 2025, 82). مقدار بتا در این پژوهش ۱ ($\beta=1$) گرفته شده است تا بین فراخوانی و صحت توازن برقرار شود.

۵. یافته‌ها

۵-۱. نتایج ارزیابی الگوریتم‌ها

نتایج حاصل از ارزیابی الگوریتم‌ها نشان می‌دهد که عملکرد روش‌های یادگیری ماشین مرسوم و مدل‌های مبتنی بر شبکه عصبی به‌طور چشمگیری تحت تأثیر نوع بردارسازی معنایی قرار دارد. در بردارسازی «فست‌تکست»، که از اطلاعات زیرواژگانی بهره می‌گیرد، روش‌های یادگیری ماشینی K نزدیک‌ترین همسایه و جنگل تصادفی توانسته‌اند دقتی نزدیک به ۹۹ درصد ارائه دهند و مدل ماشین بردار پشتیبان نیز عملکردی مطلوب، اما پایین‌تر از این دو مدل ثبت کرده است. این الگوریتم‌ها آن است که غنای اطلاعاتی «فست‌تکست» مرزهای معنایی میان مقوله‌ها را برای مدل‌های مرسوم به‌خوبی روشن می‌سازد. با این حال، در بردارسازی «وردتووک»، وابستگی مدل‌ها به ویژگی‌های معنایی ساده‌تر آشکار می‌شود؛ به‌طوری که جنگل تصادفی و K نزدیک‌ترین همسایه همچنان

عملکردی بسیار قوی دارند، اما ماشین بردار پشتیبان و درخت تصمیم با افت قابل توجه دقت مواجه می‌شوند. در مقابل، مدل‌های شبکه عصبی رفتاری متفاوت نشان می‌دهند: پرسپترون چندلایه در ترکیب با «فست‌تکست» بالاترین عملکرد کل پژوهش را ثبت کرده (۹۹/۴۵ درصد)، اما در بردارسازی «وردتووک» هر دو مدل شبکه عصبی تک‌لایه و چندلایه با افت چشمگیر دقت مواجه شده‌اند.

جدول ۲. متوسط عملکرد الگوریتم‌ها (عددها به درصد)

الگوریتم یادگیری	بردارسازی فست‌تکست				بردارسازی وردتووک			
	متوسط دقت	متوسط فراخوانی	متوسط امتیاز f	متوسط صحت	متوسط دقت	متوسط فراخوانی	متوسط امتیاز f	متوسط صحت
ماشین بردار پشتیبان	۹۳/۷۶	۹۳/۷۶	۹۳/۷۱	۹۳/۷۳	۷۹/۹۰	۸۰/۲۴	۸۳/۶۰	۸۰/۰۳
K_ نزدیک‌ترین همسایه	۹۹/۰۰	۹۹/۰۰	۹۹/۰۰	۹۹/۰۰	۹۷/۹۹	۹۷/۹۷	۹۷/۹۶	۹۷/۹۵
جنگل تصادفی	۹۸/۶۴	۹۸/۶۱	۹۸/۶۱	۹۸/۶۱	۹۹/۳۷	۹۹/۳۶	۹۹/۳۶	۹۹/۳۶
درخت تصمیم	۸۶/۰۲	۸۵/۹۵	۸۵/۹۵	۸۵/۸۸	۹۷/۳۴	۹۷/۳۳	۹۷/۳۲	۹۷/۳۰
پرسپترون تک‌لایه	۷۹/۹۲	۷۹/۷۹	۷۹/۷۷	۷۹/۷۰	۷۳/۶۰	۷۰/۶۷	۶۸/۰۵	۷۰/۴۷
پرسپترون چندلایه	۹۹/۴۵	۹۹/۴۶	۹۹/۴۵	۹۹/۴۴	۹۵/۷۴	۷۴/۷۶	۷۴/۴۲	۷۴/۶۸

۲-۵. نتایج تحلیل خطا

به‌منظور درک بهتر الگوی خطاهای مدل، ماتریس‌های درهم‌ریختگی^۱ نیز مورد بررسی قرار گرفتند. براساس جدول ۳، مدل جنگل تصادفی در تشخیص مقوله‌های باهم‌آیی (۷۴۱ مورد صحیح) و اصطلاح (۷۱۷ مورد صحیح) عملکرد قابل قبولی نشان داده است و در مجموع، مدل توانسته است حدود ۷۴/۲ درصد کل نمونه‌ها را به‌درستی پیش‌بینی کند. این امر بیانگر آن است که این مدل می‌تواند الگوهای بسامدی و به‌نسبت ثابت موجود در فضای برداری «فست‌تکست» را برای این دو دسته تشخیص دهد. باین حال، ضعف اساسی آن در تفکیک هم‌معنایی و تضاد معنایی آشکار است.

1. confusion matrices

جدول ۳. ماتریس درهم‌ریختگی جنگل تصادفی و «فست تکست»

برچسب صحیح				
برچسب پیش‌بینی شده	تضاد معنایی	باهم‌آیی	اصطلاح	هم‌معنایی
تضاد معنایی	۵۹۳	۰	۶	۳۱۸
باهم‌آیی	۰	۷۴۱	۶۳	۱۱۳
اصطلاح	۱۰	۳۸	۷۱۷	۹۲
هم‌معنایی	۱۹۹	۵۳	۳۹	۶۲۵

پیش‌بینی اشتباه ۳۱۸ نمونه هم‌معنایی به‌عنوان تضاد معنایی در جدول ۴ و ۱۹۹ نمونه تضاد معنایی به‌عنوان هم‌معنایی نشان می‌دهد که مرزبندی بین این دو مفهوم در فضای برداری زیرواژه‌محور «فست تکست» به‌اندازه کافی تفکیک‌پذیر نبوده است. از سوی دیگر، ماهیت درختی و نیمه‌خطی جنگل تصادفی قدرت لازم برای استخراج الگوهای غیرخطی و ظریف مورد نیاز برای تفکیک مترادف و متضاد را فراهم نمی‌کند. شباهت سطحی ساختار پرسش‌ها نیز این مشکل را تشدید کرده است.

جدول ۴. نمونه‌ای از اشتباهات ماشین در پیش‌بینی برچسب‌ها

برچسب پیش‌بینی شده	برچسب صحیح
تضاد معنایی	هم‌معنایی
هم‌معنایی خراب؟	هم‌معنایی خراب؟
مترادف سالم چیست؟	مترادف سالم چیست؟

مدل K- نزدیک‌ترین همسایه در جدول ۵، نیز الگویی مشابه، اما با شدت بیشتر نشان داده است. این مدل سراسر به ساختار هندسی فضای برداری وابسته است. بنابراین، هر جا که دسته‌ها در بردارسازی به هم نزدیک یا همپوشان باشند، عملکرد آن به شدت افت می‌کند. خطاهای قابل توجه به‌ویژه در تمایز میان هم‌معنایی و تضاد معنایی، از جمله پیش‌بینی ۳۳۷ نمونه هم‌معنایی به‌عنوان تضاد معنایی نشان می‌دهد که این مدل قادر به جداسازی این روابط معنایی نزدیک نیست. افزون‌بر این، حساسیت بالای این مدل به نوفه و به‌ویژه به تنوع ساختارهای زبانی، ضعف آن را برجسته‌تر کرده است. با وجود این، عملکرد به‌نسبت قابل قبول آن در برخی مقولات مانند باهم‌آیی نشان می‌دهد که

«فست‌تکست» برای برخی روابط معنایی خاص ساختارهای برداری مناسبی تولید کرده است. لازم به ذکر است که مدل توانسته است حدود ۷۰ درصد کل نمونه‌ها را به‌درستی پیش‌بینی کند.

جدول ۵. ماتریس درهم‌ریختگی K- نزدیک‌ترین همسایه و «فست‌تکست»

بر چسب صحیح				
بر چسب پیش‌بینی شده	تضاد معنایی	باهم‌آیی	اصطلاح	هم‌معنایی
تضاد معنایی	۵۱۹	۲۴	۳۷	۳۳۷
باهم‌آیی	۱۵	۷۶۱	۶۷	۷۴
اصطلاح	۱۳	۱۰۱	۶۵۹	۶۶
هم‌معنایی	۲۵۹	۴۹	۳۹	۵۶۹

در مقابل، پرسپترون چندلایه با «فست‌تکست» در جدول ۶، عملکردی تقریباً بی‌نقص از خود نشان داده است. دقت بالای این مدل در تمامی کلاس‌ها شامل پیش‌بینی صحیح ۹۱۲ نمونه تضاد معنایی، ۹۱۷ نمونه باهم‌آیی، ۸۵۷ نمونه اصطلاح و ۸۹۴ نمونه هم‌معنایی بیانگر توانایی شبکه عصبی در استخراج الگوهای غیرخطی و بهره‌گیری کامل از بازنمایی زیرواژه‌ای «فست‌تکست» است. تعداد خطا در تمام کلاس‌ها به‌طور چشمگیری محدود بوده و بیشتر در طبقه هم‌معنایی که کمتر از ۲۵ نمونه بود، مشاهده شده است و نشان‌دهنده پایداری و تعمیم‌پذیری بسیار بالای مدل است. در مجموع، این مدل حدود ۹۹/۳ درصد نمونه‌ها را درست پیش‌بینی کرده است.

جدول ۶. ماتریس درهم‌ریختگی پرسپترون چندلایه و «فست‌تکست»

بر چسب صحیح				
بر چسب پیش‌بینی شده	تضاد معنایی	باهم‌آیی	اصطلاح	هم‌معنایی
تضاد معنایی	۹۱۲	۰	۰	۵
باهم‌آیی	۰	۹۱۷	۰	۰
اصطلاح	۰	۰	۸۵۷	۰
هم‌معنایی	۱۰	۹	۳	۸۹۴

۶. بحث

در مقایسه پژوهش‌های حوزه تشخیص قصد، باید توجه داشت که تفاوت‌های بنیادین در متغیرهای پژوهش از جمله نوع پیکره آموزشی، تعداد قصدها، حجم داده، ماهیت تک‌زبانه یا چندزبانه بودن مدل و استفاده از بردارسازی ایستا یا پویا موجب می‌شود نتایج این مطالعات به‌صورت مستقیم و منصفانه قابل مقایسه نباشند. برای نمونه، مدل‌هایی که بر پیکره‌های کوچک و تک‌زبانه آموزش دیده‌اند، با چالش‌هایی متفاوت از مدل‌های مبتنی بر پیکره‌های بزرگ چندزبانه مواجه‌اند، و مدل‌هایی که از بردارسازی پویا مانند «برت» یا «روبرتا» بهره می‌برند، در اصل ظرفیت‌های بازنمایی عمیق‌تری دارند و نمی‌توان عملکرد آنها را در کنار مدل‌های کلاسیک یا بردارسازی‌های ایستا با یک معیار یکنواخت ارزیابی کرد. همچنین، اختلاف در تعداد قصدها و پیچیدگی معنایی هر پیکره سبب می‌شود بار پردازشی و دشواری وظیفه برای هر مدل متفاوت باشد و همین امر تحلیل تطبیقی نتایج را محدود می‌سازد. در جدول ۷، همین ناهمگونی‌ها دیده می‌شود. پژوهش‌های بین‌المللی بیشتر بر پیکره‌های سنجه استاندارد مانند «ای تی آی اس» و «اسنپس» با تعداد قصدهای ثابت و بردارسازی‌های پویا متمرکزند؛ در حالی که پژوهش‌های انجام‌شده در داخل کشور از پیکره‌هایی متفاوت از نظر حجم، تنوع زبانی و روش برجسب‌گذاری استفاده کرده‌اند و در برخی موارد اطلاعات دقیقی درباره حجم داده یا تعداد قصدها ارائه نشده است. افزون‌بر این، تفاوت در نوع مدل‌های زبانی و انتخاب معماری‌های یادگیری مشترک یا انتقالی، به‌طور مستقیم بر کیفیت پیش‌بینی اثر می‌گذارد. از این‌رو، هرچند این مطالعات تصویری کلی از روند پیشرفت در تشخیص قصد ارائه می‌دهند، اما به‌دلیل ناهمگونی در متغیرهای کلیدی، مقایسه کمی نتایج آنها تنها در چارچوب‌های هم‌تراز و کنترل‌شده امکان‌پذیر است.

جدول ۷. پژوهش‌های انجام‌شده در تشخیص قصد

پژوهش	پیکره داده	تعداد قصدها	حجم داده	نوع مدل زبانی	نوع بردارسازی	مدل یادگیری
Kim (2014)	پیکره فیلم	-	-	تک‌زبانه	ایستا و پویا	عصبی پیچشی
Liu & Lane (2016)	ای تی آی اس	حدود ۲۱	حدود پنج‌هزار جمله	تک‌زبانه	پویا	عصبی بازگشتی مبتنی بر توجه

پژوهش	پیکره داده	تعداد قصدها	حجم داده	نوع مدل زبانی	نوع بردارسازی	مدل یادگیری
Chen, Zhuo & Wang (2019)	ای تی آی اس و اسنپس	۲۱ از ای تی آی اس و ۷ از اسنپس	حدود پنج هزار جمله و سیزده هزار از پیکره دوم	تک‌زبان	پویا (برت)	مدل مشترک بر پایه برت
Di Gennaro et al. (2020)	داده مکالمه‌ای	-	-	تک‌زبان	پویا	شبکه عصبی حافظه طولانی کوتاه مدت
Ahmad et al. (2021)	پالسی آی ای	-	-	تک‌زبان	پویا	توالی مشترک
پناه‌اک و یمناق (1395)	پرسش‌های فارسی	-	-	تک‌زبان	بردارسازی کلاسیک	بردار ماشین پشتیبان و رویکرد قاعده‌بنیان
Zadkamali, Momtazi & Zeinali (2024)	ای تی آی اس ترجمه شده فارسی	۲۱	پنج هزار جمله	چندزبان	پویا (اکس‌ال‌ام روبرتا)	مدل مشترک برت و سی آر اف
Ghahramani & Shamsfard (2023)	داده بین زبانی	-	-	چندزبان	پویا	رویکرد انتقال یادگیری
Karimi & Mirzaei (2024)	پیکره معیار فارسی	-	-	تک‌زبان و چندزبان	پویا (برت و ام‌برت)	یادگیری مشترک

نتیجه‌گیری و پیشنهادات

هدف این پژوهش طراحی و پیاده‌سازی یک چت‌بات وظیفه‌محور در حوزه آموزش زبان فارسی و مقایسه تجربی الگوریتم‌های مختلف یادگیری ماشین در تشخیص قصد کاربران بود. تمرکز اصلی بر بررسی تعامل میان بازنمایی معنایی ایستا و عملکرد دسته‌بندهای گوناگون قرار داشت. نتایج به‌دست آمده از ارزیابی شش الگوریتم مختلف نشان داد که نوع روش بردارسازی و نوع معماری مدل یادگیری ماشین تأثیر تعیین‌کننده‌ای بر کیفیت نهایی تشخیص قصد دارد.

به‌طور مشخص، «فست‌تکست» به‌دلیل ماهیت زیرواژه‌محور خود توانسته است ساختارهای صرفی و واژگانی زبان فارسی را به‌صورت مؤثرتری نسبت به «وردتووک» بازنمایی کند. این ویژگی در کنار مدل‌های غیرخطی مانند پرسپترون چندلایه موجب شده است که این ترکیب بهترین عملکرد را با میانگین دقت نزدیک به ۹۹/۴۵ درصد

ارائه دهد. این نتایج نشان می‌دهد که مدل‌های مبتنی بر شبکه‌های عصبی، زمانی که از بردارهای غنی و بافت‌مدار استفاده می‌کنند، قادرند مرزهای معنایی ظریف میان مقاصد زبانی را با دقتی بسیار بالا شناسایی کنند.

در مقابل، عملکرد الگوریتم‌های درخت‌محور، به‌ویژه در ترکیب با «وردتووک» به‌نسبت ضعیف‌تر بوده است. نتایج ماتریس‌های درهم‌ریختگی نشان داد که بیشترین خطاها در شناسایی دو مقوله «هم‌معنایی» و «تضاد معنایی» رخ می‌دهد؛ خطایی که می‌تواند ناشی از شباهت ساختاری الگوهای پرسشی، همپوشانی معنایی و محدودیت‌های «وردتووک» در بازنمایی واحدهای زیرواژه‌ای در فارسی باشد. در همان حال، مدل K_{-} نزدیک‌ترین همسایه در ترکیب با «وردتووک»، برخلاف انتظار، به عملکردی بسیار بالا دست یافته است. در مجموع، یافته‌های پژوهش به چند نکته کلیدی اشاره دارند:

۱. بردارسازی نقش محوری در موفقیت مدل دارد و استفاده از روش‌های غنی مانند «فست‌تکست» برای زبان فارسی بسیار مؤثر است.

۲. مدل‌های غیرخطی و لایه‌عمق‌دار مانند پرسپترون چندلایه توانایی بالاتری در استخراج الگوهای معنایی پیچیده نشان می‌دهند.

۳. الگوریتم‌های مبتنی بر فاصله، مانند K_{-} نزدیک‌ترین همسایه، زمانی که ساختار برداری داده‌ها مناسب باشد، می‌توانند عملکردی رقابتی و حتی فوق‌العاده ارائه دهند.

ایجاد پیکره معیار برای زبان فارسی گامی ضروری است و پژوهش حاضر با طراحی یک پیکره تخصصی با هدف کاربرد در سامانه‌های مکالمه‌محور، زیرساختی مهم برای توسعه چت‌بات‌های فارسی‌زبان فراهم کرده است.

برای ادامه مسیر، چند جهت‌گیری پژوهشی پیشنهاد می‌شود: (۱) ارزیابی مدل‌های زبانی از پیش آموزش‌دیده مبتنی بر ترنسفورمر^۱ مانند «برت»، «ام‌برت» و «پارس‌برت» به‌منظور مقایسه آنها با روش‌های بردارمحور، (۲) گسترش حجم و تنوع پیکره جهت افزایش توان تعمیم مدل‌ها، (۳) بررسی همزمان تشخیص قصد و استخراج مؤلفه‌های معنایی در قالب مدل‌های یادگیری مشترک، و (۴) ارزیابی عملکرد مدل‌ها در محیط‌های واقعی مانند چت‌بات‌ها و سامانه‌های مکالمه‌محور فارسی. این مسیرها می‌توانند پایه‌ای برای توسعه سامانه‌های درک زبان طبیعی کارآمدتر در حوزه زبان فارسی فراهم آورند.

1. transformer

فهرست منابع

- صحرای، رضامراد، فائزه مرصوص، و داوود ملک‌لو. ۱۳۹۸. گام اول. تهران: انتشارات آرفا بنیاد سعدی.
- صحرای، رضامراد، شهناز احمدی قادر، فائزه مرصوص، و لیلا بنفشه. ۱۳۹۸. آموزش کاربردی واژه. ویراست دوم. تهران: بنیادسعدی، انتشارات آرفا.
- قائم، هادی، و محسن کاهانی. ۱۳۹۵. دسته‌بندی پرسش‌ها با استفاده از ترکیب دسته‌بندها. فصل‌نامه علمی-پژوهشی پردازش علایم و داده‌ها. DOI: 10.18869/acadpub.jsdp.13.3.99
- قیومی، مسعود. ۱۴۰۱. ارزیابی ساختار هرم وارونه در پیکره بزرگ خبری فارسی: تحلیل گفتمان خبری براساس همبستگی میان عنوان و محتوای خبر. زبان و زبان‌شناسی ۱۸ (۳۵): ۲۱-۴۵. doi: 10.30465/lsi.2023.849
- عسگری، مهناز. ۱۳۸۴. واژه‌آموزی زبان فارسی. تهران: کانون زبان ایران.

References

- Ahmad, W. U., J. Chi, T. Le, T. Norton, Y. Tian, & K-W. Chang. 2021. *Intent classification and slot filling for privacy policies*. arXiv preprint arXiv:2101.00123. <https://arxiv.org/abs/2101.00123> (accessed Feb 22, 2021)
- Asgari, M. 2005. *Vocabulary Learning in Persian*. Tehran: Iran Language Institute (ILI). [In Persian]
- Austin, J. L. 1962. *How to do things with words*.: Oxford University Press.
- Bojanowski, P., E. Grave, A. Joulin, & T. Mikolov. 2017. Enriching word vectors with subword information. *Transactions of the Association for Computational Linguistics*, 5, 135–146. Doi:10.1162/tacl_a_00051.
- Breiman, L. 2001. *Random forests*. *Machine Learning*, 45 (1): 5–32. <https://doi.org/10.1023/A:1010933404324>
- Breiman, L., J.H. Friedman, R. A. Olshen, & C. J. Stone. 1984. *Classification and Regression Trees*. Monterey, CA: Wadsworth.
- Chen, Q., Z. Zhuo, & W. Wang. 2019. *BERT for joint intent classification and slot filling*. arXiv preprint arXiv:1902.10909.
- Conneau, A., K. Khandelwal, N. Goyal, V. Chaudhary, G. Wenzek, F. Guzmán, & V. Stoyanov. 2020. *Unsupervised cross-lingual representation learning at scale*. *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 8440–8451. Doi: 10.18653/v1/2020.acl-main.747.
- Coucke, A., A. Saade, A. Ball, T. Bluche, A. Caulier, D. Leroy, & J. Dureau. 2018. *Snips voice platform: an embedded spoken language understanding system for private-by-design voice interfaces*. arXiv preprint arXiv:1805.10190.
- Cover, T., & P. Hart. 1967. *Nearest neighbor pattern classification*. *IEEE Transactions on Information Theory* 13 (1): 21–27.
- Deldar, H., & M. Khajepour. 2025. Improvement in intent detection and slot filling using Persian BERT representations. *Journal of Intelligent Computing and Systems Engineering* 5 (2): 45–56.
- Dema, O. 2021. *Conversation and social interaction: A sociolinguistic perspective*. London: Routledge.
- Di Gennaro, G., A. Buonanno, A. Di Girolamo, A. Ospedale, & F. A. N. Palmieri. 2020. *Intent classification in question-answering using LSTM architectures*. arXiv preprint arXiv:2001.09330. <https://arxiv.org/abs/2001.09330> (accessed Jan. 25, 2020)

- Devlin, J., M.-W. Chang, K. Lee, & K. Toutanova. 2019. *BERT: Pre-training of deep bidirectional transformers for language understanding. Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT)*, 1, 4171–4186. <https://doi.org/10.48550/arXiv.1810.04805>
- Farahani, M., M. Gharachorloo, M. Farahani, & M. Manthouri. 2020. *ParsBERT: Transformer-based Model for Persian Language Understanding*. arXiv preprint arXiv:2005.12515.
- Ghaemi, Hadi, and Mohsen Kahani. 2016. *Question Classification Using Classifier Combination. Scientific-Research Quarterly of Signal and Data Processing* 13 (3): 99. DOI: 10.18869/acadpub.jsdp.13.3.99. [In Persian]
- Ghahramani, A., & M. Shamsfard. 2023. *Intent detection and slot filling for Persian: Cross-lingual training for low-resource languages*. arXiv preprint arXiv:2303.00408. <https://arxiv.org/abs/2303.00408> (accessed May 1, 2023)
- Ghayoomi, M. 2022. *Applications of Chatbots in Education. Trends, Applications, Challenges of Chatbot Technology*. Chapter 4, 80-118. DOI: 10.4018/978-1-6684-6234-8..
- Ghayoomi, M. 2023. Evaluating the Structure of the Inverted Pyramid in the Big Persian News Corpus: News Discourse Analysis based on the Correlation Coefficient between Title and News Content. *Language and Linguistics* 18 (35): 21-45. doi: 10.30465/lsi.2023.8497. [In Persian]
- González-Lloret, M. 2010. *Conversation analysis and speech act performance: Bridging two approaches to study interaction. Language Learning & Technology* 22 (3): 1–22.
- Goodwin, C., & J. Heritage. 1990. Conversation analysis. *Annual Review of Anthropology* 19: 283–307.
- Grice, H. P. 1975. *Logic and conversation*. In P. Cole & J. L. Morgan (Eds.), *Syntax and semantics: Vol. 3. Speech acts* (pp. 41–58). New York: Academic Press.
- Harris, Z. S. 1954. *Distributional structure*. *Word*, 10: 146–162. <https://doi.org/10.1080/00437956.1954.11659520>.
- Hemphill, C. T., J. J. Godfrey, & G. R. Doddington. 1990. *The ATIS spoken language systems pilot corpus*. In *Proceedings of the Workshop on Speech and Natural Language* (pp. 96–101). Hidden Valley Pennsylvania.
- Jurafsky, D., & J. H. Martin. 2025. *Speech and language processing*; 3rd ed.. Draft version.?: Prentice Hall.
- Karimi, R., & M. Mirzaei. 2024. *A Persian benchmark for joint intent detection and slot filling. Proceedings of the 6th Workshop on NLP for Low-Resource Languages (LoResMT)*. arXiv preprint arXiv:2401.01234.
- Kim, Y. 2014. *Convolutional neural networks for sentence classification*. arXiv preprint arXiv:1408.5882. <https://arxiv.org/abs/1408.5882> (accessed Sept. 3, 2014)
- Liddicoat, A. J. 2022. *An introduction to conversation analysis (2nd ed.)*. New York: Bloomsbury Academic.
- Liu, B., & I. Lane. 2016. Attention-based recurrent neural network models for joint intent detection and slot filling. *Interspeech?*: 685–689.
- Louvan, S., & B. Magnini. 2020. *Recent neural methods on slot filling and intent classification for task-oriented dialogue systems: A survey*. arXiv preprint arXiv:2011.00564. <https://arxiv.org/abs/2011.00564>
- Mikolov, T., I. Sutskever, K. Chen, G. Corrado, & J. Dean. 2013. *Distributed representations of words and phrases and their compositionality*. arXiv. <https://doi.org/10.48550/arXiv.1310.4546>
- Pires, T., E. Schlinger, & D. Garrette. 2019. *How multilingual is Multilingual BERT? In Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics (ACL)*. DOI: 10.18653/v1/P19-1493

- Rosenblatt, F. 1958. The perceptron: A probabilistic model for information storage and organization in the brain. *Psychological Review* 65 (6): 386–408.
- Rumelhart, D. E., G. E. Hinton, & R. J. Williams. 1986. Learning representations by back-propagating errors. *Nature* 323: 533–536.
- Sacks, H., E. A. Schegloff & G. Jefferson. 1974. A simplest systematics for the organization of turn-taking for conversation. *Language* 50 (4): 696–735.
- Sahraei, R. M., F. Marsous., & D. Malekloo. 2019. *First Step*. Tehran: Saadi Foundation AZFA Publications. [In Persian]
- Sahraei, R. M., S. Ahmadi Ghader., F. Marsous, & L. Banafsheh. 2019. *Applied Vocabulary Teaching*. 2nd ed. Tehran: AZFA Publications, Saadi Foundation. [In Persian]
- Tjong, Kim, E. F. Sang, & F. De Meulder. 2003. *Introduction to the CoNLL-2003 shared task: Language-independent named entity recognition*. In *Proceedings of the 7th Conference on Natural Language Learning at HLT-NAACL 2003* (pp. 142–147). Edmonton, Canada: Association for Computational Linguistics. <https://doi.org/10.3115/1119176.1119195>
- Vapnik, V. N. 1995. *The Nature of Statistical Learning Theory*. New York: Springer
- Williams, J., & S. Bayne. 2024. *Chatting with bots: AI, speech acts, and the edge of assertion*. *AI & Society*
- Yule, G. 2000. *Pragmatics*. Oxford: Oxford University Press.
- Zadkamali, R., S. Momtazi, & H. Zeinali. 2024. *Intent detection and slot filling for Persian: Cross-lingual training for low-resource languages*. Cambridge: Cambridge University Press.
doi:10.1017/nlp.2024.17

الهام صالحی

متولد سال ۱۳۶۱، دانشجوی دکتری زبان‌شناسی از پژوهشگاه علوم انسانی و مطالعات فرهنگی است.
زبان‌شناسی پیکره‌ای، زبان‌شناسی رایانشی و پردازش زبان طبیعی، سامانه‌های مکالمه و تحلیل گفتمان از جمله علایق پژوهشی وی است.



مسعود قیومی

متولد سال ۱۳۵۸، دارای مدرک تحصیلی دکتری در رشته زبان‌شناسی رایانشی از دانشگاه آزاد برلین، آلمان است. ایشان هم‌کنون دانشیار پژوهشگاه زبان‌شناسی در پژوهشگاه علوم انسانی و مطالعات فرهنگی است.
زبان‌شناسی رایانشی و پردازش زبان طبیعی، مدل‌سازی زبانی، یادگیری ماشینی، نحو و معناشناسی واژگانی از جمله علایق پژوهشی وی است.



آنوسا رستم‌بیک تفرشی

متولد سال ۱۳۵۷، دارای مدرک دکتری زبان‌شناسی از پژوهشگاه علوم انسانی و مطالعات فرهنگی است و در حال حاضر به‌عنوان دانشیار زبان‌شناسی در همان پژوهشگاه فعالیت می‌کند.

زبان‌شناسی کاربردی، تحلیل گفتمان انتقادی و رویکردهای پیکره‌محور و رایانشی در تحلیل زبان از جمله علایق پژوهشی وی است.

