

Expert Detection In Question Answer Communities

Hamed Salimian¹, Mohammad Amin Fazli², Jafar Habibi³

¹ PhD Student of Computer Engineering, Sharif University of Tehran, Iran
hamed@ce.sharif.edu

² Assistant Professor, Sharif University of Tehran, Iran
fazli@sharif.edu

³ Associate Professor, Sharif University of Tehran, Iran
jhabibi@sharif.edu

Abstract

Community Question Answering has a high place in the lives of all members of a society nowadays. It is important for the owners of an community to be able to make this community better and more reliable. One way to achieve this is to find users who have more knowledge, expertise, experience and skills in the field and can well share their knowledge with others (which we call experts and aim to encourage them to be more active in the system).

One method to be used is to identify expert users, and whenever a new question is asked, we suggest this question to them to check and answer if its in their scope of expertise. One way to encourage users to post replies is to use gameplay techniques such as assigning points and badges to users. But as we will discuss, this method does not always reflect expert users well, because some users will try to have small and insignificant but numerous activities that will make them gain a lot of points, however they are not experts. In this study, we examine the methods by which experts in a question-and-answer system can be found, and try to evaluate and compare these methods, use their ideas and positive points, and add our own new ideas to a new way of finding them.

Keywords: Community Question Answering, Experts, Comment, Selected Answers, Positive Votes, Negative Votes, Point, Credits, Recommending System, Neural Networks, Machine Learning

یافتن کاربران خبره در انجمن‌های پرسش و پاسخ

حامد سلیمیان^۱، محمدامین فضلی^۲، جعفر حبیبی^۳

^۱ دانشجوی دکتری، دانشکده مهندسی کامپیوتر، دانشگاه صنعتی شریف، تهران
hamed@ce.sharif.edu

^۲ استادیار، گروه مهندسی نرم‌افزار، دانشکده مهندسی کامپیوتر، دانشگاه صنعتی شریف، تهران
fazli@sharif.edu

^۳ دانشیار، گروه مهندسی نرم‌افزار، دانشکده مهندسی کامپیوتر، دانشگاه صنعتی شریف، تهران
jhabibi@sharif.edu

چکیده

انجمن‌های پرسش و پاسخ امروزه جایگاه بالایی در زندگی تمام افراد یک جامعه دارند. اما نکته مهم برای صاحبان یک انجمن این است که بتوانند این انجمن را بهتر و قابل‌اعتمادتر کنند. یکی از راه‌های رسیدن به این مهم این است که کاربرانی را که دانش، تخصص، تجربه و مهارت بیشتری در زمینه‌های موردنظر دارند و به‌خوبی می‌توانند دانش خود را با دیگران به اشتراک بگذارند (که ما به آن‌ها خبره می‌گوییم) بیابند و سعی کنند با روش‌هایی آن‌ها را ترغیب کنند که بیشتر در سامانه فعالیت داشته باشند. یکی از راه‌هایی که می‌تواند مورد استفاده باشد این است که کاربران خبره^۱ شناسایی شوند و هرگاه سوال جدیدی پرسیده می‌شود این پرسش‌ها را به آن‌ها پیشنهاد دهیم که زودتر سوال‌های حوزه تخصصی خود را ببینند و پاسخ دهند. یکی از راه‌های ترغیب کاربران به ارسال پاسخ‌ها، استفاده از روش‌های بازی‌گونه‌سازی^۲ مانند تخصیص امتیاز و نشان^۳ به کاربران است. اما همان‌گونه که در ادامه خواهیم دید، این روش همیشه به‌خوبی کاربران خبره را نمایان نمی‌کند، زیرا برخی کاربران سعی بر این خواهند داشت تا فعالیت‌های کوچک و کم‌اهمیت اما پرشماری داشته باشند که باعث خواهد شد امتیاز و نشان‌های زیادی از آن خود کنند اما در اصل جزو دایره خبرگان نیستند. در این پژوهش روش‌هایی را که می‌توان با آن‌ها خبره‌های یک سامانه پرسش و پاسخ را یافت بررسی می‌کنیم و سعی می‌کنیم تا با ارزیابی و مقایسه این روش‌ها، استفاده از ایده‌ها و نقاط مثبت آنها و اضافه کرده ایده‌های جدید خودمان، به یک روش جدید برای یافتن خبره‌های واقعی در سامانه StackOverflow برسیم.

کلمات کلیدی

انجمن‌های پرسش و پاسخ، خبره، دیدگاه، پاسخ برگزیده، رای مثبت، رای منفی، امتیاز، اعتبار، سامانه توصیه‌گر، شبکه‌ها عصبی، یادگیری ماشین

1- مقدمه

در سال‌های اخیر، سایت‌ها و انجمن‌های پرسش و پاسخ (CQA^4) بی‌شماری در سطح اینترنت در دسترس عموم قرار گرفته‌اند. برخی از این سامانه‌ها، کاربردهای عمومی دارند و برخی دیگر به موضوع یا موضوعات خاصی اختصاص دارند. سامانه‌های پرسش و پاسخ امروزه نقش پررنگی در زندگی ما ایفا می‌کنند. کاربران به این سامانه‌ها رجوع می‌کنند تا پاسخ سوالات خود را بگیرند، سوالی بپرسند، دانش خود را در اختیار دیگران قرار دهند یا حداقل با ارسال نظر و رای مثبت و منفی به پیش‌برد آن انجمن کمک کنند. در این سامانه‌ها معمولاً تعداد بسیار زیادی کاربر عضو هستند، اما کاربران فعال در آن‌ها پرشمار نیستند. از بین همین کاربران فعال نیز نمی‌توان تمام کاربران را کاربر متخصص در نظر گرفت. ضمناً برخی کاربران هستند که فعالیت پرشماری ندارند اما تخصص و خبرگی بالایی دارند.

کاربر متخصص یا خبره را می‌توان کاربری در نظر گرفت که تخصص و مهارت زیادی (در یک یا چند زمینه) دارد و می‌تواند دانش خود را با دیگران به اشتراک بگذارد. تمام این سامانه‌ها بر مبنای اشتراک دانش بنا شده‌اند که نیازمند تخصص‌های افراد خبره هستند، پس نیاز است تا این افراد خبره کشف شده و به‌گونه‌ای آن‌ها را ترغیب کنیم تا مشارکت بیشتری در سامانه داشته باشند، مثلاً [25]، [26] و [27] روی این ایده به صورت مستقیم کار کرده‌اند که خبره‌هایی که در یک سامانه مشخص هستند را تشخیص داده و وظایف را به صورت دقیق به آن‌ها تخصیص داد. اما مشارکت در برخی از انجمن‌ها مانند *stackoverflow* نیازمند دانشی خاص و مهارتی خاص در یک زمینه مشخص است که کار را پیچیده‌تر می‌کند. به این منظور برخی از انجمن‌ها به کاربران امتیاز یا اعتبار نسبت می‌دهند که با هر فعالیتی که در سایت می‌کنند این امتیاز افزایش می‌یابد. برخی دیگر از سایت‌ها به کاربران نشان‌هایی اهدا می‌کنند که تعدد این نشان‌ها در یک نمایه^۵ نشان‌دهنده خبرگی بیشتر صاحب آن نمایه است.

اما آیا این امتیاز یا نشان همیشه نشان‌دهنده خبرگی کاربر است؟ پاسخ خیر است [1][24]. در ادامه و با بررسی برخی از کارهای پیشین سعی بر آن است که پاسخ این پرسش را مورد بررسی قرار دهیم، ببینیم چرا این معیار مناسب نیست و در نهایت در پی یافتن روش‌هایی دیگر و موثرتر در راه یافتن کاربران خبره خواهیم بود.

البته ناگفته نماند که مسئله یافتن خبره در حوزه‌های دیگر دنیای مهندسی نرم‌افزار هم مسئله جذابی است و اقدامات زیادی در این حوزه‌ها هم انجام می‌شود. برای مثال [8] و [10] روی این ایده کار کرده‌اند که از فعالیت افراد در *GitHub* بتوانند متخصصین را تشخیص دهند.

2- بررسی کارهای پیشین

در این بخش ابتدا نگاهی مختصر به به تاریخچه یافتن خبره می‌اندازیم، سپس الگوریتم‌ها و روش‌های یافتن خبره را مورد بررسی قرار می‌دهیم و در نهایت روش‌های نمایه‌محور را که یکی از ایده‌های اصلی پیاده‌شده در این مقاله هستند را بررسی خواهیم کرد. در انتها هم نگاهی اجمالی روی روش‌های شبکه‌ها عصبی، یادگیری ماشین و تحلیل گراف موجود خواهیم داشت.

1-2- تاریخچه یافتن خبره

در نخستین پژوهش [1] بیان می‌شود که اکثر روش‌های موجود برای یافتن خبره را می‌توان به دو گروه طبقه‌بندی کرد:

1. روش‌های مبتنی بر اعتبار^۶، که بر اساس تجزیه و تحلیل لینک فعالیت‌های موضوعی خبره در گذشته است.
2. روش‌های مبتنی بر موضوع که مبتنی بر تکنیک‌های مدل‌سازی موضوع و معنای نهفته است.

یکی از مدل‌هایی که در این مقاله مورد بررسی قرار می‌گیرد روش مبتنی بر MF^۷ هستند که به شباهت کاربران توجه می‌کند که با استفاده از این شباهت بتواند کاربرانی را که احتمالاً می‌توانند پاسخ یک سوال را داشته باشند، بیابد.

البته یافتن خبره در حوزه‌های دیگر هم کاربرد زیادی دارد مانند [8]، [10] و [12] اما ما در این پژوهش تنها روی یافتن خبره در فضای انجمن‌های پرسش و پاسخ تمرکز داریم.

2-2- الگوریتم‌های یافتن خبره

در پژوهش [5]، مسئله یافتن خبره را با ترکیب یادگیری ساختار شبکه انجمن پرسش و پاسخ و نمایش معنایی سوال^۸ با یک معیار رتبه‌بندی با نام $RMNL^9$ حل کرده‌اند. سپس با استفاده از شبکه RNN^{10} عمیق که بر اساس گشت تصادفی^{۱۱} توسعه یافته بود، رتبه‌بندی را برای کاربران و سوالات در یک شبکه انجمن پرسش و پاسخ ایجاد کردند.

در پژوهش [6] دیگر که در سال 2020 منتشر شده است، سعی نوبری و دیگران بر آن بود که هر حوزه مهارتی را به تعدادی واژه مربوط ترجمه^{۱۲} کنند تا بتوانند تطابق این واژه‌های ترجمه‌شده با سوال‌های پرسیده‌شده و پاسخ‌های داده‌شده را بهبود دهند. برای مثال واژه‌ای مانند *java-ee* به واژه‌های *session.http*، *request*، *controller* و *ejb* ترجمه می‌شود. مدل اول ترجمه آن‌ها *MI* نام داشت و مدل دوم آن‌ها *WE*. مدل اول روی اطلاعات مشترک تمرکز داشت و مدل دوم روی روش‌های جاسازی واژه‌ها^{۱۳}. هر دوی این مدل‌ها در نهایت نشان دادند که می‌توانند در یافتن خبره کمک‌کننده باشند.

در این پژوهش [7] که سومانث و دیگران در سال 2018 آن را منتشر کرده‌اند، با استفاده از داده‌های منتشرشده توسط *stackoverflow* و با استفاده از کتابخانه *SNAP* پایتون، یک گراف ایجاد کردند. راس‌های این گراف جهت‌دار کاربران هستند و هر یار جهت‌دار از کاربری است که پرسشی ارسال کرده به کاربری که پاسخ پذیرفته‌شده را ارسال کرده است. نتیجه حائز اهمیتی که این پژوهش برای ما دارد این است که برخی از کاربران در یک زمینه خاص مثلاً زبان برنامه‌نویسی *Python* میزان خبرگی بالایی دارند اما در کل سامانه *StackOverflow* از اعتبار بالایی برخوردار نیستند. در پژوهش‌های دیگری نیز همین ایده که خبره‌ها در هر حوزه به صورت جداگانه در نظر گرفته شوند پیاده شده است.

در پژوهش دیگری [9] به بررسی پاسخ کاربران به سوالات سخت‌تر می‌پردازد و بیان می‌شود که کاربرانی که سوالات سخت‌تر را پاسخ می‌دهند، خبرگی بالاتر دارند و برخی معیارها را برای یافتن سوالات سخت‌تر ارائه می‌دهند. آثار مشخصی در حوزه شناسایی سوالات دشوار وجود دارد. لیو

رویکردی را برای تعیین تاثیر کیفیت سؤال در تعیین کیفیت پاسخ در خدمات انجمن‌های پرسش و پاسخ پیشنهاد کرد [11] با این وجود، به غیر از کیفیت سؤال، در برچسب زدن به سؤالات دشوار، ویژگی‌های دیگری مانند تعداد نظرات را نیز در نظر می‌گیرد.

2-3- روش‌های یافتن خبره

پژوهش‌های زیادی مانند [3] و [4] وجود دارند که روی ایده ساخت نمایه کار کرده‌اند. پژوهش دیگری [13] نیز روی ایجاد نمایه برای کاربران و سؤالات و سپس تشخیص خبرگان با استفاده از یک سامانه توصیه‌گر از روی آن نمایه‌ها پرداخته است. ورودی این سامانه‌های توصیه‌گر ویژگی‌های کاربران و محتوایی است که توسط کاربران ایجاد شده است. یانگ و همکاران [18] یک مدل تخصصی یادگیری ماشین براساس موضوع را ارائه می‌دهند که موضوعات و تخصص مشترک را با یکپارچه‌سازی مدل محتوای متنی و تجزیه و تحلیل ساختار، پیوند می‌دهد. همچنین برخی از روش‌ها بر اساس تکمیل ماتریس برای توصیه متخصص^{۱۴} وجود دارد. در یک مطالعه [19]، تخصص کاربر توسط برچسب‌ها نشان داده شده است. تخصص کاربران به نمرات (کاربر، برچسب) ترجمه شده است، که به کیفیت برچسب‌ها بسیار وابسته است.

علاوه بر این، مدل‌های یادگیری عمیق در حال ظهور، با روش‌های فوق‌الذکر، برای بهبود بیشتر عملکرد روش‌های یافتن متخصص، ادغام می‌شوند مانند [2]. این روش‌ها قادر به یادگیری مؤثر در ابعاد بالا از اطلاعات تخصصی، اطلاعات موضوعی و تعامل موضوعات تخصصی هستند. پژوهش‌های دیگری نیز مانند [20] و [21] نیز روی ایده استفاده از یادگیری ماشین کار کرده‌اند و نتایج خوبی به دست آورده‌اند.

برخی از پژوهش‌ها هم مانند [15]، [16] و [17] هم از ایده استفاده از تحلیل گراف استفاده کرده‌اند. در این پژوهش‌ها، افراد به عنوان گره‌های گراف در نظر گرفته شده و روابط آن‌ها نیز یال‌ها هستند. در این پژوهش‌ها از روش‌های خوشه‌بندی، استفاده از ابرگراف یا ایده‌های ابتکاری خود استفاده کرده‌اند که همگی نتایج حائز اهمیتی داشته‌اند.

2-4- روش‌های اعتباردهی

یکی دیگر از روش‌هایی که می‌توان بر اساس آن خبره‌های یک سامانه را مشخص کرد، روشی بازیگوشانه‌سازی مانند اعطای امتیاز یا اعتبار و نشان به کاربران است. در ادامه نگاهی روی پژوهش‌هایی خواهیم داشت که با در نظر گرفتن اعتبار سعی در یافتن خبره‌ها دارند.

در اولین پژوهشی [13] که این دسته بررسی می‌کنیم بیان می‌شود که روش‌هایی که بر مبنای اعتبار هستند، به خوبی نمی‌توانند خبرگان را تشخیص دهند. زیرا یکی از روش‌هایی که کاربران اعتبار دریافت می‌کنند پاسخ به سؤالات است. حال اگر کاربری روی مفاهیم زیادی آشنایی داشته باشد (حتی اگر در هیچ‌یک از آن‌ها خبره نباشد) چون پاسخ‌های زیادی ارسال می‌کند احتمالاً اعتبار بالایی خواهد داشت در صورتی که ممکن است فردی تنها در چند زمینه محدود خبره باشد که چون نمی‌تواند پاسخ‌های زیادی ارسال کند، پس احتمالاً اعتبار پایین‌تری هم خواهد داشت.

2-5- رویکردهای توصیه‌ی^{۱۵} خبره

در این پژوهش [22] هم یک دسته‌بند دودویی^{۱۶} توسعه دادند که ویژگی^{۱۷}‌های نمایه‌محور^{۱۸} (مانند مدت زمانی که از ثبت‌نام کاربر گذشته است یا تعداد دنبال‌کنندگان و غیره) با ویژگی‌های پست‌های ارسالی ترکیب کردند که این ویژگی‌ها هم بر اساس پرسش‌گران بود و هم پاسخ‌دهندگان. در نهایت و با اعمال این مدل روی داده‌ها، به این نتیجه رسیدند که کاربرانی که مدت زمان کمتری از پیوستنشان به سامانه می‌گذرد و کاربرانی که فعال هستند، علاقه بیشتری به پاسخ دادن به سؤالی‌هایی دارند که به آن‌ها پیشنهاد شده است.

در مقاله‌ای دیگر [23]، دایک و دیگران به روی تحلیل نتیجه‌های الگوریتم‌های دسته‌بندی مطالعاتی داشته‌اند. در این پژوهش مشخص شد که از بین الگوریتم‌ها، جنگل تصادفی^{۱۹} از نظر معیار F1-score از Gaussian Naïve Bayes بهتر عمل کرده است. همچنین این الگوریتم از Linear Support Vector Classification بهتر عمل کرده است.

برای ارزیابی برای هر سوال که پاسخ‌دهندگان خبره آن‌ها مشخص بود، الگوریتم‌ها و روش‌های مورد نظر را اجرا می‌کردند و لیستی از N کاربر پیشنهادی که ممکن بود به این سوال پاسخ دهند را استخراج می‌کردند. سپس به ازای هر یک از کاربرانی که پاسخ داده بودند و در این لیست بودند یک واحد به متغیری اضافه می‌کردند و در نهایت میانگین این متغیر را گرفتند. نتیجه این پژوهش این شد که روش LDA از روش TF-IDF و مدل زبانی در یافتن بهترین خبرگان برای یک پرسش بهتر عمل کرد.

3- تعریف‌های اصلی

در ادامه پژوهش برخی از تعاریف به دفعات مورد استفاده قرار خواهند گرفت.

به این منظور در این بخش به توصیف دقیق این مفاهیم می‌پردازیم.

- سامانه‌های پرسش و پاسخ: سامانه‌هایی هستند که کاربران با عضویت در آنها می‌توانند در یک حوزه خاص و تخصصی یا در حوزه‌های عمومی، سؤالی از دیگر کاربران بپرسند یا پاسخ سوال کاربر دیگری را ارسال کنند. برخی از سامانه‌ها پرسش و پاسخ ویژگی‌های بیشتر مانند ارسال دیدگاه روی پاسخ‌ها و سؤالات دیگر کاربران، ویرایش پاسخ دیگر کاربران، اعلام تکراری بودن سوال و غیره را شما می‌شوند.
- خبره: در سامانه‌های پرسش و پاسخ، خبره به شخصی گفته می‌شود که در یک زمینه مشخص، از مهارت، توانایی، دانش و تجربه کافی برخوردار است و می‌تواند نظرات تخصصی خود را در یک زمینه به دیگران منتقل کند. شخص خبره لزوماً از اعتبار و امتیاز بالایی در سامانه‌های پرسش و پاسخ برخوردار نیست (همان‌گونه که در بهش کارهای پیشین دیدیم). البته ما خبره را شخصی در نظر می‌گیریم که با تمام توانایی‌هایی که دارد بتواند پاسخ سوال را بدهد. یعنی اگر کاربری خبره باشد اما به واسطه عدم فعالیت قادر به پاسخ‌گویی سؤالی نباشد، می‌توان وی را از لیست خبرگان خارج کرد. البته این خروج در مدل‌هایی در نظر گرفته خواهد شد که زمان آخرین فعالیت کاربر را در نظر گرفته‌ایم.

- نمایه کاربر: در این پژوهش ما برای هر کاربر یک نمایه ایجاد می‌کنیم. این نمایه سه بخش اصلی دارد:

1. زمان پیوستن کاربر به سامانه: در بخش کارهای پیشین دیدیم که کاربرانی که تازه وارد سامانه شده‌اند تعامل بیشتری برقرار می‌کنند و این مورد می‌تواند به ما در کشف کاربرانی که احتمالا سوالی را پاسخ دهند یاری کند.
2. مدت‌زمان سپری‌شده از آخرین تعامل کاربر در سامانه: اگر کاربری (هر چند خبره) ماه‌هاست در سامانه تعاملی نداشته است، پس احتمالا نمی‌توان روی وی برای پاسخ‌گویی به سوالات حساب کرد.
3. امتیاز کاربر روی هر برچسب: هر کاربر با توجه به تعاملاتی که روی سوال‌ها دارد، با توجه به میزان سختی سوال و نوع فعالیتی که انجام داده است، به ازای تمام برچسب‌های سوال امتیازی دریافت می‌کند که نشان‌دهنده میزان خبرگی وی روی هر کدام از آن برچسب‌هاست.

- پرسش سخت: در یکی از پژوهش‌های پیشین مشاهده شد که کاربرانی هستند که تنها به سوالات ساده پاسخ می‌دهند و در این راه امتیاز زیادی کسب می‌کنند، در صورتی که مشاهدات نشان می‌داد کاربران خبره کاربرانی هستند که بیشتر به سمت سوالات سخت گرایش دارند که این امر باعث می‌شود از روی اعتبار به تنهایی نتوان کاربران خبره را تشخیص داد. به این منظور باید سوالات سخت را از دیگر سوالات تمییز داد. به این منظور ما ویژگی‌های زیر را دلالی بر سختی سوال در نظر گرفته‌ایم و هر چه این موارد پررنگ‌تر باشند، میزان سختی سوال بیشتر است:

- پاسخ خوب: پاسخ خوب یا به عبارتی پاسخ خبره لزوماً پاسخ منتخب نیست. البته که پاسخ منتخب می‌تواند به عنوان پاسخ خبره در نظر گرفته شود، اما نکاتی است که باید به آنها توجه شود. برای مثال ممکن است پاسخی بهتر از پاسخ منتخب بعدا ارسال شده باشد، یا کاربری که سوال پرسیده است فراموش کرده باشد پاسخ منتخب را مشخص کند. همچنین ممکن است ذیل یک پرسش، چندین پاسخ خبره وجود داشته باشد که باید همگی را در نظر بگیریم. به این منظور فاکتورهای زیر را برای تعیین خبرگی یک پاسخ در نظر گرفته و از ترکیب آنها می‌توانیم عددی به میزان خبرگی آن نسبت دهیم:

1. انتخاب به عنوان پاسخ منتخب
2. تعداد دیدگاه ثبت‌شده روی پاسخ
3. نسبت امتیاز به تعداد مشاهده سوال
4. میزان خبرگی پاسخ‌دهنده در برچسب‌های سوال
5. نسبت امتیاز پاسخ به امتیاز سوال
6. تعداد دیدگاه‌های روی پاسخ (دست کم سه مورد)

- پاسخ منتخب: کاربری که سوالی را پرسیده است می‌تواند از بین پاسخ‌های ارسال‌شده، یک مورد را به عنوان پاسخ صحیح که نیاز وی را مرتفع می‌کند انتخاب کند. البته پاسخ‌های منتخب همواره

بهترین پاسخ نیستند و گاه مشاهده می‌شود در بین پاسخ‌های ارسالی، پاسخی است که از پاسخ برگزیده به مراتب بهتر است (با توجه به معیارهایی که در بند پیش بیان شد).

- دیدگاه: هر کاربر می‌تواند روی هر پرسش یا پاسخ دیدگاه خود را ثبت کند. ما تا حد کمی در این پژوهش، دیدگاه روی یک پاسخ را نشانه‌ای مبنی بر خبرگی ارسال‌کننده پاسخ دانسته‌ایم. البته اگر تعداد دیدگاه‌ها بسیار کم باشد ممکن است اتفاقاً به معنی نامناسب بودن پاسخ باشد، مثلاً کاربری بیان کرده باشد که به این دلیل پاسخ اشتباه است یا نیاز به ویرایش دارد.
- زمان آخرین فعالیت: ممکن است کاربری در یک زمینه بسیار خبره باشد اما ماه‌ها باشد که در سامانه فعالیتی نداشته باشد. ما این شخص را دیگر خبره در نظر نخواهیم گرفت، زیرا ما به دنبال خبرگانی هستیم که بتوانند سوالاتی را پاسخ دهند و نه تنها خبرگانی که دانش بالایی دارند.
- برچسب: کاربری که سوالی می‌پرسد باید به این منظور که کاربران دیگر این سوال را در دسته‌بندی‌های مشخص راحت‌تر پیدا کنند و همچنین دسته سوال مشخص باشد، روی سوال بر حسب نیازهایی که مطرح می‌شود، برچسب‌هایی اضافه کند. مبنای کار ما در این پژوهش استفاده از این برچسب‌ها و یافتن خبره‌های هر برچسب است.

4- طرح مسئله

همان‌گونه که در بخش مقدمه بررسی شد، یافتن افراد خبره در سامانه‌های پرسش و پاسخ اهمیت بالایی دارد. به این منظور و با توجه به اینکه بیشتر پژوهش‌های پیشین مطالعه‌شده تمام جوانب را به صورت کامل در نظر نگرفته‌اند، بر آن شدیم تا با ارائه چندین روش و مقایسه آنها، یک روش بهینه و با کارایی بالاتر با در نظر گرفتن بسیاری از جوانب حاصل از تجمیع کارهای پیشین مهم‌تر، ایجاد کنیم.

در بخش‌های پیشین دیدیم هر مورد از کارهای پیشین دارای نقاط قوت و نقاط ضعفی هستند که ما سعی بر آن خواهیم داشت تا از تجمیع این نکات مثبت استفاده کنیم. همچنین سعی می‌کنیم تا با ایجاد روش امتیازدهی شخصی خودمان و استفاده از برخی ایده‌های جدید مانند بررسی زوج برچسب‌های پرتکرار، منطق کار را تقویت ببخشیم. در بخش آتی به جزئیات اقدامات و ایده‌های جدیدی که استفاده کرده‌ایم می‌پردازیم.

هدف نهایی در این پژوهش این خواهد بود که برای سوال‌های جدیدی که وارد سامانه می‌شوند یا سوالاتی که مدتی است ارسال شده‌اند اما پاسخ مناسبی دریافت نکرده‌اند، کاربران خبره‌ای که ممکن است بتوانند با احتمال زیادی به آن سوال پاسخ دهند را بیابیم.

5- راه‌حل پیشنهادی

برای انجام فعالیت‌ها از مجموعه داده‌های سامانه StackOverflow که در kaggle در اختیار است استفاده کرده‌ایم. همان‌گونه که در بخش کارهای پیشین مشاهده شد، یکی از رویکردها موفق جهت بررسی خبره‌ها،

ایجاد نمایه برای کاربران و پرسش‌ها است. در این پژوهش ما نیز از همین رویکرد موفق استفاده خواهیم کرد.

از طرفی در برخی دیگر از مقاله‌های معتبر دیدیم که کاربران خبره لزوماً در کل سامانه خبره در نظر گرفته نمی‌شوند و در برخی از حوزه‌ها به عنوان کاربر خبره هستند و نه در کل سامانه. به این منظور ما نیز در پیاده‌سازی‌های خود برچسب‌هایی که روی سوال‌ها گذاشته می‌شوند را معیار خبرگی در حوزه برای کاربران در نظر گرفته‌ایم. البته با توجه به تعداد زیادی برچسب که در سامانه StackOverflow وجود دارد، نیاز به انجام پردازش‌هایی است که در ادامه به آن‌ها اشاره شده است.

زمان از آخرین فعالیت هر کاربر نیز یکی از ملاک‌هایی است که در بیشتر کارهای پیشین روی آن تمرکز نشده است و تنها برخی از مقاله‌ها که اتفاقاً خروجی‌های نسبتاً قابل قبولی داشتند روی زمان آخرین فعالیت کاربر تمرکز کرده‌اند. ما نیز برای بررسی اینکه هر کاربر برای هر سوال مشخص چه میزان خبرگی دارد، آخرین زمانی که کاربر در سامانه فعالیتی داشته است را در نظر گرفته‌ایم. دلیل این کار این است که ممکن است کاربری بر اساس تمام معیارها خبره باشد، اما چون ماه‌هاست در سامانه فعالیتی نداشته است، پس نمی‌توان روی پاسخ احتمالی وی برای سوال حساب کرد و به این منظور باید وی را از لیست خبرگان احتمالی خارج کرد.

زمان گذشته از ثبت‌نام کاربر در سامانه نیز یکی دیگر از مواردی است که باید در نظر گرفته شود. همان‌گونه که در کارهای پیشین مشاهده شد، کاربرانی که به تازگی به سامانه پیوسته‌اند، تمایل بیشتری برای پاسخ‌گویی به سوالات و ایجاد تعامل دارند. پس با در نظر گرفتن تعدادی محدودیت (مثلاً حداقل تعداد تعامل یا حداقل تعداد پاسخ درست ارسال‌شده) می‌توان کاربران تازه‌واردی که فعالیت زیادی در سامانه ندارند را نیز در لیست احتمالی خبرگان در نظر گرفت.

عدم تکیه کامل بر اعتبار خود سامانه نیز یکی دیگر از مواردی است که ما در نظر داشته‌ایم. پیش‌تر دیدیم که کاربرانی هستند که سوالات ساده و زیادی را پاسخ می‌دهند یا سوالات ساده‌ای می‌پرسند که تنها به افزایش تعامل آنها منجر می‌شود که این‌ها باعث می‌شوند شخص امتیاز بالایی در سامانه بگیرد. این به این معنی است که کاربرانی که امتیاز بالاتری در سامانه دارند، احتمالاً کاربران فعال‌تری هستند اما نمی‌توان لزوماً آن‌ها را خبره در نظر گرفت.

ایجاد رویکرد اعتباردهی شخصی خودمان در کنار استفاده از سیستم اعتباردهی سامانه StackOverflow، راه‌حل نهایی ما برای اعتباردهی به کاربران است. درست است که گفتیم رویه اعتباردهی خود سامانه شاید به صورت مستقیم باعث کشف کاربران خبره نشود، اما به هر حال نمی‌توان آن را نادیده گرفت. اما برای رفع نقاط ضعف این رویه، ما رویه اعتباردهی خودمان را نیز اضافه کرده‌ایم. در این بخش پس از پیاده‌سازی برخی از ایده‌ها، در نهایت بهترین ایده‌ای که به آن رسیدیم این بود که چهار موضوع را با هم ترکیب کنیم: 1. ایجاد نمایه برای کاربران 2. برچسب‌محور کردن ایده کلی 3. ایجاد رویه امتیازدهی جدید 4. در نظر گرفتن سختی سوالات. به این منظور، هر کاربری که روی سوالی تعاملی ایجاد می‌کند (اعم از درج دیدگاه، درج پاسخ و انتخاب پاسخ به عنوان پاسخ منتخب) به برچسب‌های موردنظر روی سوال روی نمایه وی با وزن عددی را اضافه می‌کنیم. این وزن حاصل ضرب اهمیت فعالیت (دیدگاه کمتر از پاسخ و پاسخ کمتر از پاسخ منتخب) در سختی سوال

است. سختی سوال هم همان‌گونه که در بخش دوم گفتیم، با معیارهایی همانند تعداد پاسخ‌های ارسال‌شده، تعداد دیدگاه‌های ارسال‌شده، نسبت امتیاز سوال به تعداد مشاهده‌های سوال، مدت زمان بین زمان ارسال سوال و ارسال پاسخ منتخب و غیره در نظر گرفته می‌شود.

در برخی از کارهای پیشین مشاهده شد که برخی از پژوهش‌گران به دلایل گوناگون سعی در نمونه‌برداری از بخشی از داده‌ها بودند. در این پژوهش ما نیز روی سوالاتی که انتخاب کرده‌ایم، یک نمونه‌برداری پیاده‌کردیم. در مجموعه داده‌ها تعداد حدود 20 میلیون سوال بود که به دو دلیل تنها سوالات سال 2018 به بعد را در نظر گرفتیم:

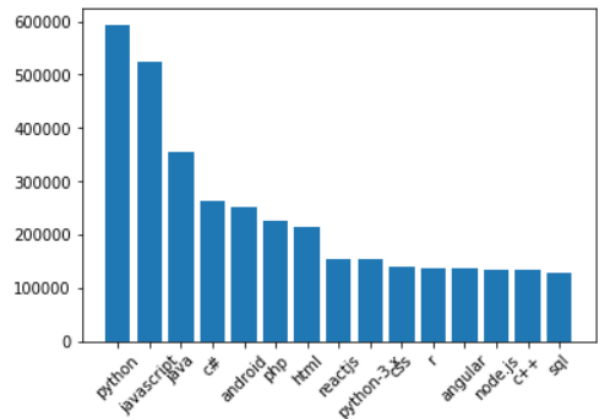
1. این داده‌ها حدود پنج میلیون سطر داده را شامل می‌شوند که کسر مناسبی از کل داده‌ها است.

2. برخی از زبان‌های برنامه‌نویسی یا در کل موضوعاتی که راجع به آنها سوال پرسیده می‌شود در سال‌های اخیر بیشتر مورد بحث قرار گرفته‌اند (مانند پایتون در برابر جاوا) و همچنین برخی از برچسب‌ها در سال‌های دورتر بیشتر مورد استفاده بوده‌اند و در حال حاضر افت زیادی داشته‌اند (مانند سوالات اسمبلی). با توجه به اینکه ما می‌خواهیم تشخیص خبرگان را در داده‌های جدید بررسی کنیم، لازم دانستیم که داده‌های قدیمی‌تر را حذف کنیم.

یکی دیگر از نمونه‌برداری‌هایی که انجام دادیم روی برچسب‌ها بود. در مجموعه داده‌ای که روی آن کار کردیم بیش از 58 هزار برچسب متمایز وجود داشت که در مجموع حدود 58 میلیون بار استفاده شده بودند. در داده‌های سال 2018 به بعد این عدد حدوداً 14.4 میلیون بود که ما برای کوچک کردن اندازه مسئله و کاهش پردازش‌ها، 100 برچسب پر استفاده را از دیگر برچسب‌ها جدا کردیم و تنها روی آنها تمرکز کردیم. این 100 برچسب روی هم رفته حدود هفت میلیون بار استفاده شده بودند که حدود نیمی از کل استفاده‌ها بود که تقریب بدی نیست. همچنین این امر نشان می‌دهد دیگر برچسب‌ها تعداد استفاده خیلی کمی داشته‌اند (یعنی دقیقاً 57653 برچسب تکرار استفاده‌ای حدود هفت میلیون داشته‌اند در برابر تنها صد برچسب که همین تکرار استفاده را داشته‌اند). در برچسب‌های گلچین‌شده، بیشترین تکرار به ترتیب برای برچسب‌های python، javascript و java بودند با حدود 600 هزار، 520 هزار و 355 هزار استفاده و کم‌استفاده‌ترین برچسب‌ها در این لیست گلچین‌شده هم unit-testing با حدود 17 هزار استفاده است. نمودار توزیع تعداد برچسب‌ها در شکل (1) قابل مشاهده است.

اصلاح نمایه‌ها نیز آخرین اقدامی است که روی نمایه‌ها صورت دادیم. به این منظور بررسی کردیم که کدام زوج برچسب‌ها تکرار زیادی دارند. این کار با الگوریتم A-Priori انجام شد. این ایده به این منظور اتخاذ شد که برای مثال کسی که خبرگی در زمینه nodeJS دارد بدون شک باید در زمینه javascript هم خبره باشد. حال ممکن است سوالی پرسیده شود که برچسب nodeJS داشته باشد اما برچسب javascript نداشته باشد. در این مرحله ما یک گام پیش رفته و در کنار استفاده مستقیم از برچسب‌های مورد استفاده، به احتمال تسلط روی برچسب‌های جامانده هم پرداخته‌ایم. این مورد هم یکی دیگر از ایده‌های جدیدی بودی که در این پژوهش از آن استفاده نمودیم.

شکل (1): تعداد استفاده برچسب‌های پرتکرار



با توجه به ترکیب روش‌ها نمی‌توان از یکی از مقاله‌ها به عنوان معیار استفاده نمود. به این منظور ما در ادامه چندین روش ارائه می‌دهیم و این روش‌ها را با هم مقایسه می‌نماییم. یکی از این روش‌ها ساده‌ترین روش ممکن یعنی کاملاً تصادفی است که منطقاً تمام روش‌های ارائه‌شده باید کارایی بهتری نسبت به این روش داشته باشند. در روش بعدی تصادفی کمی هوشمندانه‌تر.

روش انجام کار نیز به این‌گونه است که سعی می‌کنیم اطلاعات سوال و هر کاربر را در مدل خود لحاظ کنیم و از مدل انتظار عددی بین صفر و یک داشته باشیم. هر چه این عدد به یک نزدیک‌تر باشد به معنی احتمال بیشتر پاسخ از سمت آن کاربر برای آن سوال است و برعکس، هر چه این عدد به صفر میل کند، احتمالاً کاربر نمی‌تواند به آن سوال پاسخ دهد.

برای یافتن کاربران خبره نیز تمام نمایه کاربر را در نظر نمی‌گیریم. از آنجایی که در پژوهش‌های پیشین هم دیدیم که یک کاربر نمی‌تواند در تعداد زیادی زمینه خبره باشد، برای هر کاربر در نهایت پنج برچسبی که بیشترین امتیاز را شامل می‌شوند در نظر گرفته و از مابقی برچسب‌ها صرف نظر می‌کنیم.

روش انتخاب پاسخ بهینه را هم این‌گونه در نظر گرفتیم که می‌دانیم لزوماً پاسخی که به عنوان پاسخ منتخب انتخاب می‌شود پاسخ مناسب نیست و ممکن است پاسخی بهتر وجود داشته باشد. هم‌چنین ممکن است پاسخی منتخب نشده باشد و پاسخ مناسبی باشد. به این منظور اگر پاسخی امتیاز بیشتری از سوال گرفته باشد، روی آن بیش از پنج دیدگاه وجود داشته باشد، به عنوان پاسخ منتخب انتخاب شده باشد، نسبت امتیاز به تعداد مشاهده آن بالا باشد، به عنوان پاسخ از سمت یک خبره در نظر گرفته می‌شود. یعنی ممکن است در یک پرسش چندین پاسخ به عنوان پاسخ خبره در نظر گرفته شود.

به این منظور هر دیدگاهی که کاربر ارسال می‌کند یک واحد، پاسخی که می‌دهد سه واحد و هر پاسخ صحیحی که ارسال کرده است پنج واحد مثبت برای وی در نظر گرفته می‌شود. سپس این امتیاز را در ضریبی ضرب می‌کنیم که این ضریب اینگونه محاسبه می‌شود: نسبت امتیاز به مشاهده با وزن 0.5، تعداد دیدگاه‌ها (تا سقف 10 مورد) با وزن 0.2، مدت زمان گذشته از پرسش (تا سقف 72 ساعت) با ضریب 0.1 و امتیاز پاسخ تا سقف 2500 با ضریب 0.2.

از طرفی برای هر سوال میزان سختی را عددی بین صفر و یک در نظر گرفتیم. فاصله زمانی بین ارسال سوال و ارسال پاسخی که بعداً به عنوان

برگزیده انتخاب شده است، تعداد پاسخ‌ها (دست کم سه مورد)، تعداد دیدگاه‌های روی سوال (دست کم سه مورد)، امتیاز سوال به تنهایی (تا سقف 2500 که بیشترین امتیاز بین تمام سوالات از سال 2018 به بعد است) و نسبت تعداد مشاهده سوال به امتیاز سوال همگی معیارهای ما برای امتیازدهی به سختی سوال هستند. البته در این روش ما ابتدا فرض می‌کنیم تمام سوال‌ها کاملاً ساده هستند (یعنی با سختی صفر) و سپس به ازای هر یک از معیارهای نام‌برده، به میزان سختی آن در صورت لزوم اضافه می‌کنیم. سپس این ضریب سختی را در امتیازی که در ابتدای این پاراگراف گفتیم ضرب می‌کنیم و به تمام برچسب‌های سوال این میزان را در نمایه کاربر اضافه می‌کنیم. به این صورت نمایه شخصی‌سازی‌شده ما برای کاربر ایجاد می‌شود. برای مثال اگر سوالی دارای سختی 0.75 باشد و سودمندی پاسخی نیز 0.5 باشد، در نمایه ارسال‌کننده آن پاسخ به مقدار تمام برچسب‌ها میزان 0.75×0.5 اضافه می‌کنیم. لازم به ذکر است در صورتی که میزان سختی سوال یا سودمندی پاسخ ممکن است برخی مواقع منفی شود (مثلاً پاسخ یا پرسشی که امتیاز منفی زیادی دارد) در این صورت ما آن را صفر در نظر می‌گیریم که همواره این دو عدد بین صفر و یک باشند.

در ادامه برای مثال بخشی از نمایه یکی از کاربران را مشاهده می‌کنیم:

```
{'python': 44445.2994120573,
'javascript': 40785.373228161494,
'java': 24850.98611114744,
'c#': 20163.131018360233,
'android': 17469.807392388288,
'php': 12274.59503363403,
'html': 13432.882015758008,
'reactjs': 12154.339674976654,
'python-3.x': 10858.816594901491,
'css': 9782.591122652624,
'r': 10856.262039384912,
'angular': 12430.032216431942,
'node.js': 7853.19516844874,
'c++': 18396.492872628103,
'sql': 9905.570381176232,
'jquery': 7010.051779561444}
```

تا کنون و با این روش ایجاد نمایه، هم تخصص کاربر در هر حوزه را جداگانه در نظر گرفته‌ایم، هم سختی سوالات را دخیل کرده‌ایم. این دو مورد از اصلی‌ترین اصولی است که در این پژوهش مورد استفاده قرار داده‌ایم که مورد دوم ایده جدیدی بود که با این جزئیات به سختی سوالات توجه می‌شود.

1-5- استفاده از روش فاصله کسینوسی

روش نخستی که انجام داده‌ایم، تشابه نمایه کاربری است که سوالی پرسیده است با کاربر مورد بررسی. می‌توانیم امیدوار باشیم که اگر دو کاربر تخصص‌های مشابهی داشته باشند، احتمالاً بتوانند سوالات یکدیگر را پاسخ دهند. برای این منظور نمایه هر کاربر را به صورت یک بردار 100 تایی در نظر می‌گیریم از امتیازاتی که به وی بر اساس عملکردش تا کنون داده‌ایم و برای محاسبه تشابه نمایه‌ها، فاصله کسینوسی بردار 100 تایی هر دو نمایه را لحاظ می‌کنیم. این روش موفق شد 48 درصد از پاسخ‌های صحیح را به درستی تشخیص دهد.

5-2-2- روش رگرسیون لاجیستیکی

دومین روشی که در روش‌های یادگیری ماشین اتخاذ کردیم، روش رگرسیون لاجیستیکی بود. این روش همان گونه که در ادامه خواهیم دید نسبت به روش پیش بسیار بهتر عمل کرده است. در اینجا هم داده‌های آموزش را 75 درصد داده‌ها و داده‌های آزمون را 25 درصد داده‌ها در نظر گرفتیم. در مجموع در این روش روی مجموعه آزمون حدود 650 هزار TN و 667 هزار TP ایجاد شد و هیچ نمونه FN و FP مشاهده نشد.

با توجه به عدم حضور هیچ FN یا FP در این روش، مساحت زیر نمودار ROC برابر یک خواهد شد و نمایش نمودار آن موضوعیتی ندارد و به این منظور از رسم آن صرف نظر می‌کنیم.

5-2-3- روش شبکه‌های عصبی

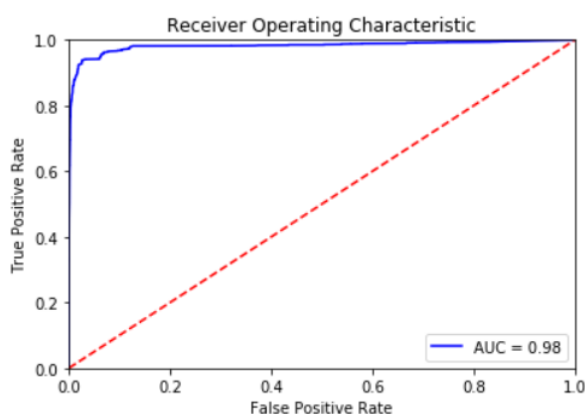
در این روش چندین معماری گوناگون شبکه‌های عصبی مورد بررسی قرار گرفتند که در نهایت یک شبکه عصبی عمیق با معماری زیر بهترین پاسخ را ارائه داد:

- لایه اول با هشت نورون و تابع فعال‌ساز ReLU
- لایه دوم با 16 نورون و تابع فعال‌ساز ReLU
- لایه سوم با 16 نورون و تابع فعال‌ساز ReLU
- لایه چهارم با هشت نورون و تابع فعال‌ساز ReLU
- لایه پنجم با تک نورون جهت پیش‌بینی نهایی با تابع فعال‌ساز sigmoid

ضمناً پس از هر لایه Dropout به میزان 0.2 در نظر گرفته شد تا از بیش‌برازش^{۲۰} پیش‌گیری شود.

شبکه‌های عصبی با حدود 614 هزار TN، 38 هزار FN، 39 هزار FP و 629 هزار TP نتیجه نسبتاً قابل قبولی ارائه داد. البته مساحت زیر نمودار ROC برای این روش عدد 0.98 را نشان می‌دهد که عدد بسیار مناسبی است و در شکل (3) قابل مشاهده است.

شکل (۳): نمودار ROC روش شبکه‌های عصبی



6- نتیجه‌گیری

در این پژوهش ابتدا به بررسی دلیل اهمیت یافتن کاربران خبره پرداختیم. سپس برخی از فعالیت‌های پیشین که روی یافتن خبره یا دست کم بررسی کاربران خبره بودند را بررسی کرده و نقاط قوت و ضعف آنها را بیان کردیم. سپس ایده‌های پیشین را با ایده‌های جدیدی از خودمان مانند ایجاد نمایه

در ادامه تنها به بررسی دقت این روش با پاسخ‌های صحیح بسنده نکردیم و روی تمام پاسخ‌ها این روش را مورد بررسی قرار دادیم. در روش پیش تنها می‌توانستیم روی True positive ها و False Negative ها بحث کنیم، اما برای بررسی کارایی نیاز است False Positive ها و True Negative ها را هم بررسی کنیم. به این منظور دو ستون به مجموعه داده اضافه کردیم که همان سختی سوال و سودمندی پاسخ بود. سپس هر پاسخی که سودمندی بیشتر از 0.5 داشت را پاسخ خبره و با برچسب یک و در غیر این‌صورت آن پاسخ را پاسخ غیر خبره و با برچسب صفر در نظر گرفتیم. سپس مجدداً روش فاصله کسینوسی بردار نمایه کاربر صاحب سوال و کاربر صاحب پاسخ را اعمال کردیم. در این حالت حدود دو میلیون TN، 1.8 میلیون FN، 600 هزار FP و 840 هزار TP ایجاد شد.

5-2- استفاده از روش‌های یادگیری ماشین

در روش بعدی سراغ روش‌های یادگیری ماشین رفتیم. سه روش در این بخش مورد بررسی قرار خواهند گرفت که در ادامه به بررسی آنها می‌پردازیم.

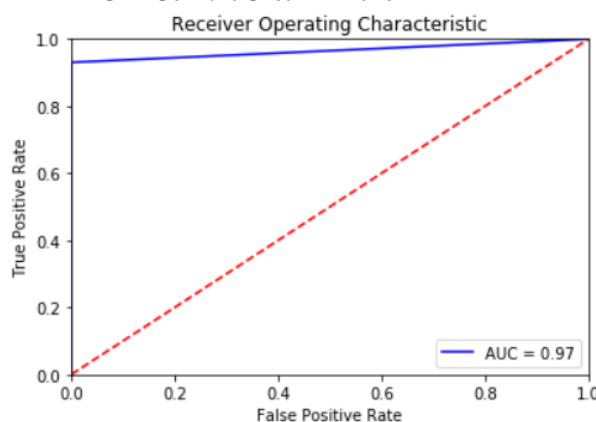
5-2-1- روش رگرسیون خطی

نخستین روشی که در این بخش پیاده کردیم، روش رگرسیون خطی بود. در این روش که از بیان جزئیات آن می‌پرهیزیم، مجموعه داده خود را به عنوان ورودی در نظر گرفته و خروجی میزان سودمندی سوال را رند کرده و به عنوان خروجی در نظر گرفتیم.

البته شایان ذکر است که در این بین برخی از ویژگی‌های مجموعه داده را به عنوان ورودی در نظر نگرفتیم، برای مثال میزان امتیاز پاسخ یا تعداد مشاهده پاسخ. زیرا مدل نهایی ما قرار است فارغ از پاسخ و تنها از روی ویژگی‌های شخص سوال‌کننده و شخص پاسخ‌دهنده پیش‌بینی کند.

در این روش مجموعه آزمون را 25 درصد مجموعه آموزش در نظر گرفتیم. در مجموع در این روش روی مجموعه آزمون حدود 650 هزار TN، 46 هزار FN و 620 هزار TP ایجاد شد و هیچ نمونه FP مشاهده نشد. نمودار ROC این روش در شکل (2) قابل مشاهده است.

شکل (۲): نمودار ROC روش رگرسیون خطی



نتایج خوبی دست یافت. همچنین اگر کاربری پست دیگری را ویرایش می کند نیز می تواند حاوی نکات مهمی برای ما باشد.

9- مراجع

- [1] S Yuan, Y Zhang, J Tang, W Hall, JB Cabotà, "Expert finding in community question answering: a review.," *Artificial Intelligence Review*, vol. 53, no. 2, pp. 843 - 874, 2020.
- [2] J Wei, J He, K Chen, Y Zhou, Z Tang, "Collaborative filtering and deep learning based recommendation system for cold start items.," *Expert Systems with Applications*, , vol. 69, pp. 29 - 39, 2017.
- [3] Tian, Y., Ng, W., Cao, J. and McIntosh, S. Geek talents: Who are the top experts on github and stack overflow?. *CMC-COMPUTERS MATERIALS & CONTINUA*, 61(2), pp.465-479, 2019.
- [4] Ghasemi, N., Fatourechi, R. and Momtazi, S., 2021. User embedding for expert finding in community question answering. *ACM Transactions on Knowledge Discovery from Data (TKDD)*, 15(4), pp.1-16.
- [5] Z Zhao, Q Yang, D Cai, X He, Y Zhuang, "Expert Finding for Community-Based Question Answering via," in *Ijcai* , 2016.
- [6] AD Nobari, M Neshati, SS Gharebagh, "Quality-aware skill translation models for expert finding on StackOverflow.," *Information Systems*, vol. 87, p. 101413, 2020.
- [7] P Sumanth, K Rajeshwari, " Discovering Top Experts for Trending Domains on Stack Overflow.," *Procedia computer science*, vol. 143, pp. 333- 340, 2018.
- [8] A Di Sorbo, G Canfora, S Panichella, "Won't We Fix this Issue? Qualitative characterization and automated identification of wontfix issues on GitHub," *Information and Software Technology*, 2021.
- [9] M Bhanu, J Chandra, "Exploiting response patterns for identifying topical experts in stackoverflow.," in *Eleventh International Conference on Digital Information Management (ICDIM) (pp. 139-144). IEEE.*, 2016.
- [10] JA Oliveira, M Viggiano, E Figueiredo, " Mining Experts from Source Code Analysis: An Empirical Evaluation," *Journal of Software Engineering Research and Development*, 2021.
- [11] J Liu, H Shen, L Yu, "Question quality analysis and prediction in community question answering services with coupled mutual reinforcement.," *IEEE Transactions on Services Computing.*, vol. 10, no. 2, pp. 286 - 301, 2015.
- [12] Bozzon, A., Brambilla, M., Ceri, S., Silvestri, M. and Vesca, G., 2013, March. Choosing the right

شخصی، ایجاد اعتبار جدید برای هر کاربر، بررسی حوزه محور کاربران، در نظر گرفتن ویژگی های زمانی، استفاده از روش هایی مانند A-priori برای اصلاح نمایه و غیره ترکیب کردیم و مجموعه داده نهایی را ایجاد نمودیم. سپس با چهار روش فاصله کسینوسی، رگرسیون خطی، رگرسیون لاجیستیکی و شبکه های عصبی سعی بر ایجاد مدل های گوناگون نمودیم. از آن جایی که روش ما با روش های پایه ای که از آنها ایده گرفته بودیم تفاوت های زیادی داشتند و حتی رویکرد نهایی و هدف اصلی ما با بسیاری از آنها عینا همسو نبود، به ناچار برای اعتبار سنجی روش های خود را با یکدیگر مقایسه نمودیم که در نهایت روش رگرسیون لاجیستیکی بهترین نتیجه را داد و پس از آن روش شبکه های عصبی بهترین نتیجه را داشت.

7- تهدید اعتبار

چند مورد در این پژوهش می توانند به عنوان تهدیدی بر اعتبار در نظر گرفته شوند. مثلاً ما داده ها را از سال 2018 به بعد جدا کردیم که البته برای این کار استدلال کافی داشتیم، اما برخی از زبان های برنامه نویسی یا برخی حوزه های مورد بحث در StackOverflow هستند که سالیان زیادی است جزو مباحث داغ هستند (مانند Java) هر چند که ممکن است در سال های گذشته کمی از محبوبیت آنها کاسته شده باشد، اما به هر حال همچنان جای بررسی دارند.

همچنین ما برچسب هایی که تکرار کمی داشتند را از دور خارج کردیم. اما چه بسا کاربرانی که در آن حوزه ها خبره باشند از کاربرانی که در حوزه های داغ تر خبره اند، تخصص بیشتری داشته باشند.

نکته بعدی این است که اگر پرسش یا پاسخی امتیاز منفی در رویه های ما پیدا کند، ما آن را صفر در نظر می گیریم، اما سوالی که منفی است با سوالی که صفر است واقعا تفاوت دارند که این تفاوت در پژوهش ما دیده نشده است.

8- کارهای آتی

به عنوان یکی از پیشنهادهای کارهای آتی می توان روی داده های پیش از سال 2018 پژوهش کرد. البته ما گفتیم که به علت تغییرات زیاد در فناوری و گرایش توسعه دهنده ها و دیگر فعالان حوزه های کامپیوتری، ما داده های سال های پیش از 2018 را حذف نموده ایم، اما همین بررسی که چقدر سلیقه ها و رویکردها و گرایش ها از سال 2018 به بعد تغییر یافته اند، خود می تواند یک موضوع جالب برای اقدامات پژوهشی باشد. همچنین در صورتی که برخی موضوعات یافت شدند که از دیرباز تا کنون مورد اقبال همگانی بوده اند، همین روش های ما و ایده های ما را می توان روی آن مجموعه داده هم در نظر گرفت.

اقدام بعدی که می تواند در نظر گرفته شود، توجه به رای های صادره از خود کاربر روی سوال ها و پاسخ های دیگران است. اینکه کاربر به سوالات مناسب و سخت رای خوب داده باشد یا پاسخ های سودمند را با رای مثبت خود همراه کرده باشد، نشان از خبرگی کاربر است و برعکس.

یکی دیگر از مواردی که می توان در نظر گرفت ویرایش های پرسش یا پاسخ است. مطمئناً با دقیق شدن در بحث ویرایش پست ها، ارتباط آنها با دیدگاه های ارسالی، ارتباط ویرایش ها با پاسخ های ارسالی و غیره می توان به

International Conference on Software Maintenance (ICSM), Cleveland, Ohio, USA, 2021.

- [25] MS Faisal, A Daud, AU Akram, RA Abbasi, "Expert ranking techniques for online rated forums," *Computers in Human Behavior*, vol. 176, p. 168, 2019.
- [26] J Lock, P Redmond, "Embedded experts in online collaborative learning: A case study," *The Internet and Higher Education*, vol. 48, 2021.
- [27] Pradeep Kumar Roy, Jyoti Prakash Singh, Amitava Nag, "Finding Active Expert Users for Question Routing in Community Question Answering Sites," *Machine Learning and Data Mining in Pattern Recognition*, 2018.

پانویس‌ها

- ¹ Expert
- ² Gamification
- ³ Badge
- ⁴ Community-based Question Answer
- ⁵ Profile
- ⁶ authority-based methods
- ⁷ matrix factorization (MF)
- ⁸ Semantic Representation
- ⁹ Ranking Metric Network Learning
- ¹⁰ Recurrent Neural Networks
- ¹¹ Random walk
- ¹² Translate
- ¹³ Word Embedding
- ¹⁴ matrix completion for expert recommendation
- ¹⁵ Recommend
- ¹⁶ Binary Classifier
- ¹⁷ feature
- ¹⁸ Profile-based
- ¹⁹ Random Forest
- ²⁰ Overfit

- crowd: expert finding in social networks. In *Proceedings of the 16th international conference on extending database technology* (pp. 637-648).
- [13] I Srba, M Grznar, M Bielikova, " Utilizing non-qa data to improve questions routing for users with low qa activity in cqa," in *he 2015 IEEE/ACM international conference on advances in social networks analysis and mining 2015* (pp. 129-136)., 2015.
- [14] J Wang, J Sun, H Lin, H Dong, S Zhang, "Convolutional neural networks for expert recommendation in community question answering.," *Science China Information Sciences*, vol. 60, no. 11, pp. 102 - 110, 2017.
- [15] Cifariello, P., Ferragina, P. and Ponza, M. Wiser: A semantic approach for expert finding in academia based on entity linking. *Information Systems*, 82, pp.1-16, 2019.
- [16] Amato, F., Cozzolino, G. and Sperli, G. A hypergraph data model for expert-finding in multimedia social networks. *Information*, 10(6), p.183, 2019.
- [17] Zhao, N., Cheng, J., Chen, N., Xiong, F. and Cheng, P. A Novel Expert Finding System for Community Question Answering. *Complexity*, 2020.
- [18] Baoguo Yang, Suresh Manandhar, "Tag-based expert recommendation in community question answering.," in *EEE/ACM International Conference on Advances in Social Networks Analysis and Mining (ASONAM 2014)* (pp. 960-963). *IEEE.*, 2014.
- [19] Z Zhao, L Zhang, X He, W Ng, " Expert finding for question answering via graph regularized matrix completion.," *IEEE Transactions on Knowledge and Data Engineering*, vol. 27, no. 4, pp. 993 - 1004, 2014.
- [20] Yuan, S., Zhang, Y., Tang, J., Hall, W. and Cabotà, J.B. Expert finding in community question answering: a review. *Artificial Intelligence Review*, 53(2), pp.843-874, 2020.
- [21] Miri, M. and Beigy, H., Expert Finding in Community Question Answering. *8th Iran Web conf*, 2020.
- [22] Z Liu, BJ Jansen, " Predicting potential responders in social Q&A based on non-QA features.," *CHI'14 Extended Abstracts on Human Factors in Computing Systems*, pp. 2131 - 2136, 2014.
- [23] D van Dijk, M Tsagkias, M de Rijke, "Early detection of topical expertise in community," in *the 38th international ACM SIGIR conference on research and development in information retrieval* (pp. 995-998)., 2015.
- [24] S Wang, DM German, TH Chen, Y Tian, "Is reputation on Stack Overflow always a good indicator for users' expertise? No!," in