

ارائه یک سیستم توصیه گر وب شخصی سازی شده بر اساس خوشه بندی ترتیبی و خودرمزگذار عمیق

مستوره معینی، علی برومندیا و منا مرادی

چکیده: حجم اطلاعات انتشار یافته در وب، استفاده از سیستم های توصیه گر را اجتناب ناپذیر کرده است. سیستم های توصیه گر وب به کاربران پیشنهادهای دقیق و مرتبط بر اساس علاقه ها و سلیقه هایشان ارائه می دهند. این توصیه ها می توانند به کاربران کمک کنند تا به سرعت به اطلاعات مورد نیاز خود دسترسی پیدا کنند و زمان جستجو را کاهش دهند. در این مقاله یک سیستم توصیه گر ترکیبی بر اساس ترکیب خوشه بندی ترتیبی فازی و شبکه خودرمزگذار عمیق بر اساس اطلاعات پروفایل کاربری و رتبه بندی وبسایت ها توسط کاربران ارائه شده است. در این سیستم توصیه گر، ابتدا کاربران بر اساس شباهت نظرات خود به صورت ترتیبی خوشه بندی می شوند. سپس رتبه بندی جدید برای کاربران با توجه به تابع عضویت فازی پیش بینی می شود. در نهایت اطلاعات موجود در پروفایل کاربری و رتبه بندی جدید کاربران به هر وبسایتی به منظور پیش بینی رتبه بندی وبسایت ها توسط کاربران، به عنوان ورودی شبکه خودرمزگذار عمیق ارائه شده مورد استفاده قرار می گیرد. در نهایت پس از یافتن کاربران مشابه، اقدام به ارائه توصیه بازدید و شخصی سازی صفحه وب کاربران جدید بر اساس وبسایت های مورد علاقه کاربران مشابه می نماید. روش پیشنهادی با توجه به لایه های آموزش عمیق و تکمیل فرایند آموزش در لایه میانی در مقایسه با سایر روش های طبقه بندی از نظر دقت آماری حدود ۴۲٪، نسبت توصیه های موفق به توصیه های مفید حدود ۴٪ و دقت تشخیص کاربران مشابه حدود ۲۰٪ بهبود داشته است.

از پرکاربردترین روش های مورد استفاده در سیستم های توصیه گر در شبکه های اجتماعی می توان به روش های مبتنی بر فیلتر مشارکتی اشاره کرد [۳]. در فیلتر مشارکتی، وظیفه این است که آیتم ها را با توجه به ترجیحات و سایر کاربران به یک کاربر فعال توصیه کنیم. در بیشتر موارد، ترجیحات به صورت نمرات ارزیابی عددی بیان می شوند [۴]. رویکردهای مبتنی بر فیلتر مشارکتی به دو نوع تقسیم می شوند: رویکردهای مبتنی بر حافظه و رویکردهای مبتنی بر مدل. رویکردهای مبتنی بر حافظه با در نظر گرفتن ترجیحات همسایگان، موارد جدیدی را توصیه می کنند. آنها از ماتریس ترجیحات مستقیماً برای پیش بینی استفاده می کنند. در این رویکرد، مرحله اول ساخت یک مدل است. مدل معادل یک تابع است که ماتریس ترجیحات را به عنوان ورودی می گیرد. سپس توصیه ها بر اساس تابعی که مدل و پروفایل کاربر را به عنوان ورودی می گیرد، صورت می گیرد. در اینجا تنها می توان به کاربرانی که پروفایل کاربری آنها متعلق به ماتریس ترجیحات است، توصیه کرد. بنابراین برای توصیه به یک کاربر جدید، پروفایل کاربر باید به ماتریس ترجیحات اضافه شود و ماتریس شباهت باید محاسبه مجدد شود که موجب افزایش پیچیدگی محاسباتی خواهد شد [۵].

در حال حاضر، سیستم های توصیه گر در تعداد زیادی از برنامه ها استفاده می شوند؛ بنابراین ساخت سیستم های توصیه گر با کیفیت بالا و منحصربه فرد برای ارائه توصیه های شخصی سازی شده به کاربران در برنامه های مختلف بسیار مهم است. با وجود پیشرفت های مختلف در سیستم های توصیه گر، نسل فعلی سیستم های توصیه گر نیازمند بهبودهای بیشتری است تا توصیه های کارآمدتری را که قابل استفاده در یک دامنه گسترده تر از برنامه ها هستند، ارائه دهند. سیستم های توصیه گر کاربردهای متنوعی دارند که نیاز است تحقیقات بیشتر برای بررسی مزایای آنها صورت گیرد [۲].

از پرکاربردترین روش های مورد استفاده در سیستم های توصیه گر در شبکه های اجتماعی می توان به روش های مبتنی بر فیلتر مشارکتی اشاره کرد [۳]. در فیلتر مشارکتی، وظیفه این است که آیتم ها را با توجه به ترجیحات و سایر کاربران به یک کاربر فعال توصیه کنیم. در بیشتر موارد، ترجیحات به صورت نمرات ارزیابی عددی بیان می شوند [۴].

رویکردهای مبتنی بر فیلتر مشارکتی به دو نوع تقسیم می شوند: رویکردهای مبتنی بر حافظه و رویکردهای مبتنی بر مدل. رویکردهای مبتنی بر حافظه با در نظر گرفتن ترجیحات همسایگان، موارد جدیدی را توصیه می کنند. آنها از ماتریس ترجیحات مستقیماً برای پیش بینی استفاده می کنند. در این رویکرد، مرحله اول ساخت یک مدل است. مدل معادل یک تابع است که ماتریس ترجیحات را به عنوان ورودی می گیرد. سپس توصیه ها بر اساس تابعی که مدل و پروفایل کاربر را به عنوان ورودی می گیرد، صورت می گیرد. در اینجا تنها می توان به کاربرانی که پروفایل کاربری آنها متعلق به ماتریس ترجیحات است، توصیه کرد. بنابراین برای توصیه به یک کاربر جدید، پروفایل کاربر باید به ماتریس ترجیحات اضافه شود و ماتریس شباهت باید محاسبه مجدد شود که موجب افزایش پیچیدگی محاسباتی خواهد شد [۵].

سیستم های مبتنی بر مدل از الگوریتم های داده کاوی و یادگیری ماشین مختلف برای توسعه یک مدل برای پیش بینی امتیاز کاربر برای یک مورد بدون امتیاز استفاده می کنند. زمانی که توصیه ها محاسبه می شوند، آنها به داده های کامل وابستگی ندارند، بلکه ویژگی ها را از داده ها استخراج کرده

کلیدواژه: سیستم توصیه گر، پروفایل کاربری، شبکه های خودرمزگذار، فیلتر مشارکتی.

۱- مقدمه

توسعه های اخیر در فناوری همراه با شیوع خدمات آنلاین، توانایی بیشتری برای دسترسی به حجم عظیمی از اطلاعات آنلاین را با سرعت بیشتری فراهم کرده است. کاربران می توانند بررسی ها، نظرات و امتیازها را برای انواع مختلفی از خدمات و محصولات آنلاین ارسال کنند. با این حال، توسعه های اخیر در محاسبات همه جا موجب بروز مشکل اضافه بار داده آنلاین شده است. این اضافه بار داده، فرایند یافتن محتوای مرتبط و

این مقاله در تاریخ ۲۵ خرداد ماه ۱۴۰۳ دریافت و در تاریخ ۹ مهر ماه ۱۴۰۳ بازنگری شد.

مستوره معینی، دانشکده کامپیوتر، واحد تهران جنوب، دانشگاه آزاد اسلامی، تهران، ایران، (email: st_m_moini@azad.ac.ir).

علی برومندیا، دانشکده کامپیوتر، واحد تهران جنوب، دانشگاه آزاد اسلامی، تهران، ایران، (email: Ali.Broumandnia@iau.ac.ir).

منا مرادی، دانشکده کامپیوتر، واحد تهران مرکز، دانشگاه آزاد اسلامی، تهران، ایران، (email: mona.moradi@iau.ac).

به همدیگر پیشنهاد شوند. روش اجرایی این مقاله به صورت زیر خلاصه شده است:

- (۱) همجوشی داده‌ها در قالب یک سیستم تصمیم یار
 - (۲) استفاده از خوشه‌بندی ترتیبی فازی به منظور یافتن کاربران مشابه
 - (۳) استفاده از شبکه عصبی رمزگذار خودکار عمیق به منظور پیش‌بینی رتبه‌بندی اعطایی کاربران جدید به وب‌سایت‌ها
 - (۴) محاسبه شباهت بین کاربران بر اساس شبکه‌های عصبی خودرمزگذار عمیق^۳ مبتنی بر فیلتر مشارکتی
 - (۵) یافتن K آیتم برتر بر اساس کاربران مشابه در شبکه اجتماعی
- ساختار ادامه مقاله به شرح زیر است. در بخش دوم کارهای مرتبط بررسی گردیده و در بخش سوم جزئیات روش پیشنهادی آمده است. در بخش چهارم نتایج پیاده‌سازی و در بخش پنجم نتیجه‌گیری مقاله بیان خواهد شد.

۲- کارهای مرتبط

با توجه به اهمیت سیستم‌های توصیه‌گر، اخیراً کارهای متنوعی در این زمینه ارائه شده که در ادامه به برخی از آنها اشاره خواهیم کرد. در [۹] یک چارچوب توصیه خبری شخصی به نام HYPNER^۴ پیشنهاد شده است. HYPNER هر دو روش فیلتر مشترک مبتنی بر فیلتر و فیلتر مبتنی بر محتوا را ترکیب می‌کند. هدف چارچوب پیشنهادی بهبود دقت توصیه‌های خبری از طریق حل مسائل مقیاس‌پذیری به دلیل مجموعه خبری بزرگ، غنی‌سازی نمایه کاربر، نمایش ویژگی‌ها و ویژگی‌های دقیق اخبار و توصیه مجموعه‌ای از اخبار متنوع است. در [۱۰] یک الگوریتم مبتنی بر یادگیری عمیق برای سیستم‌های توصیه‌گر چندمعیاره پیشنهاد شده که در آن رمزگذارهای خودکار عمیق برای بهره‌برداری از روابط غیر پیش پا افتاده، غیرخطی و پنهان بین کاربران با توجه به اولویت‌های چندمعیاره و ایجاد توصیه‌های دقیق‌تر استفاده می‌شوند. در [۱۱] تلاش شده تا یک تکنیک خوشه‌بندی جدید به نام الگوریتم مبتنی بر خوشه‌بندی مرتب^۵ (OCA) با هدف کاهش تأثیر شروع سرد و مشکلات پراکندگی داده در سیستم‌های توصیه تجارت الکترونیکی بررسی شود. در [۱۲] یک سیستم توصیه‌کننده مبتنی بر شبکه عصبی کانولوشن با استفاده از ماتریس محصول بیرونی ویژگی‌ها و فیلترهای متقابل کانولوشن پیشنهاد شده است. روش پیشنهادی می‌تواند با انواع مختلف ویژگی‌ها سروکار داشته باشد و تعاملات معنی‌دار مرتبه بالاتر بین کاربران و آیتم‌ها را به تصویر بکشد و به ویژگی‌های مهم وزن بیشتری بدهد. علاوه بر این، می‌تواند مشکل اضافه برازش^۶ را کاهش دهد؛ زیرا روش پیشنهادی شامل میانگین جهانی یا حداکثر ادغام به‌جای لایه‌های کاملاً متصل در ساختار است. در [۱۳] از یک رمزگذار خودکار که مدلی قدرتمند در کاهش ابعاد داده، استخراج ویژگی و بازسازی داده‌ها است، برای یادگیری و پیش‌بینی اولویت‌های دانش‌آموز در یک سیستم توصیه آموزش الکترونیکی مبتنی بر فیلترکردن مشارکتی استفاده شده است. در [۱۴] یک سیستم توصیه برای فروشگاه‌های خرده‌فروشی مد، بر اساس رویکرد چند خوشه‌بندی اقلام و نمایه‌های کاربران در فروشگاه‌های آنلاین و فیزیکی پیشنهاد شده است. راه‌حل پیشنهادی بر تکنیک‌های استخراج تکیه می‌کند که امکان

و مدل محاسبه می‌شود. بنابراین این روش به نام روش مبتنی بر مدل شناخته می‌شود. این تکنیک‌ها نیز برای پیش‌بینی به دو مرحله نیاز دارند: مرحله اول ساخت مدل است و مرحله دوم پیش‌بینی امتیازها با استفاده از یک تابع f است که مدل تعریف شده در مرحله اول و پروفایل کاربر را به عنوان ورودی می‌گیرد. تکنیک‌های مبتنی بر مدل نیازی به اضافه‌کردن پروفایل کاربر جدید به ماتریس خدمات عمومی قبل از انجام پیش‌بینی ندارند. ما می‌توانیم حتی به کاربرانی که در مدل حضور ندارند توصیه دهیم. سیستم‌های مبتنی بر مدل برای توصیه‌های گروهی بهتر عمل می‌کنند و می‌توانند به سرعت یک گروه از موارد را با استفاده از مدل پیش‌آموزش داده شده توصیه کنند. دقت این تکنیک به طور قابل توجهی بستگی به کارایی الگوریتم یادگیری دارد که برای ایجاد مدل استفاده می‌شود. تکنیک‌های مبتنی بر مدل قادر به حل برخی از مشکلات سنتی سیستم‌های توصیه‌گر مانند پراکندگی و قابلیت مقیاس‌پذیری با استفاده از تکنیک‌های کاهش بعد و یادگیری مدل هستند [۶].

یک تکنیک ترکیبی مجموعه‌ای از دو یا چند تکنیک است که با هم برای رسیدگی به محدودیت‌های تکنیک‌های توصیه‌کننده فردی به کار می‌روند. ترکیب تکنیک‌های مختلف می‌تواند به روش‌های مختلفی انجام شود. یک الگوریتم ترکیبی ممکن است نتایج به‌دست‌آمده از تکنیک‌های جداگانه را ترکیب کند یا می‌تواند از فیلتر مبتنی بر محتوا در یک روش مشارکتی استفاده کند یا از یک تکنیک فیلتر مشارکتی در یک روش مبتنی بر محتوا استفاده کند. این تکنیک ترکیبی از تکنیک‌های مختلف به طور کلی منجر به افزایش عملکرد و افزایش دقت در بسیاری از برنامه‌های توصیه‌کننده می‌شود [۷] و [۸].

در این مقاله یک سیستم توصیه‌گر ترکیبی بر اساس پروفایل کاربری ارائه شده که در آن از خوشه‌بندی ترتیبی فازی^۱ FOCA بر اساس شباهت نظرات کاربران استفاده شده و بر این اساس به پیش‌بینی رتبه‌بندی کاربران جدید با استفاده از شبکه عصبی خودرمزگذار پرداخته است. بر اساس ترکیب این دو روش می‌توان دقت تشخیص کاربران مشابه و دقت پیش‌بینی رتبه‌بندی نظرات سایر کاربران را افزایش داد. روش ارائه شده یک رویکرد سلسله‌مراتبی است که در مرحله اول از اطلاعات پروفایل کاربری به استخراج کاربران مشابه در مجموعه داده آموزشی و وزن شباهت کاربران آموزشی استفاده می‌شود. از این اطلاعات به‌منظور خوشه‌بندی ترتیبی فازی کاربران استفاده خواهد شد. در نهایت وزن شباهت کاربران به همراه اطلاعات موجود در پروفایل کاربری به عنوان ویژگی برای کاربران به شبکه عصبی رمزگذاری خودکار ارائه می‌شود تا به‌جای استفاده از وزن اولیه تصادفی برای سوگیری^۲، از وزن شباهت کاربران استفاده شود. در نهایت ویژگی‌های کاربران در پروفایل کاربری و وزن شباهت آنها به عنوان ورودی شبکه عصبی رمزگذار خودکار عمیق ارائه شده مورد استفاده قرار می‌گیرد تا رتبه‌بندی اعطایی کاربران جدید به وب‌سایت‌ها پیش‌بینی شود. پس از تعیین رتبه‌بندی اعطایی کاربران جدید به وب‌سایت‌ها، تشابه بین کاربران جدید و کاربران آموزشی محاسبه می‌شود. در نهایت تصمیم برای ارائه توصیه آیتم به یک کاربر بر اساس شباهت بین کاربران و آیتم‌های مورد علاقه کاربران مشابه اتخاذ می‌شود. کاربرانی که تشابه ویژگی‌های آنها بالاست و وزن شباهت بیشتری دارند، می‌توانند به عنوان کاربران مشابه در نظر گرفته شوند و آیتم‌های مورد علاقه آنها که با بالاترین نمره به آنها نظر داده‌اند، به عنوان K آیتم برتر

3. Auto Encoder Deep Neural Network

4. Hybrid Personalized News Recommendation

5. Ordered Clustering-Based Algorithm

6. Overfitting

1. Fuzzy Ordered Clustering-Based Algorithm

2. Bias

جدول ۱: مقایسه روش های پیشین.

نویسنده و سال انتشار	روش مورد استفاده	ارائه توصیه شخصی	استفاده از فیلتر مشارکتی	تحلیل معنایی پروفایل کاربر	استفاده از یادگیری عمیق
Meghanathan (۲۰۱۶)	شخصی سازی شبکه های اجتماعی با استفاده از قوانین انجمنی توسعه یافته	✓	✗	✓	✗
Li (۲۰۱۸)	شخصی سازی نتایج موتورهای جستجو با شبکه های مفهومی فازای اتوماتیک سازگار	✓	✗	✓	✗
de Araújo (۲۰۱۹)	مدل سازی رفتار کاربران شبکه های اجتماعی با روش های مبتنی بر ساختار توصیه گر پویا	✓	✓	✓	✗
Xiao (۲۰۲۱)	آموزش آگاهانه کاربر برای توصیه	✓	✗	✓	✗
Mohanty (۲۰۲۱)	چارچوب صفحه وب مبتنی بر فازای با نمایه کاربر و طراحی هستی شناسی	✓	✓	✓	✗
Vullam (۲۰۲۳)	سیستم توصیه شخصی سازی شده چند نماینده	✓	✓	✗	✗
Bedi (۲۰۲۲)	بررسی معنایی مسیر جستجو و ارائه یک جستجوی سفارشی شده	✓	✗	✓	✗
Naumov (۲۰۱۹)	مدل توصیه مبتنی بر یادگیری عمیق (DLRM)	✓	✗	✓	✓
Doonan (۲۰۱۹)	خوشه بندی کاربران بر اساس معیار شباهت جدید	✓	✗	✗	✗
Cui (۲۰۲۰)	توصیه جدید مبتنی بر ضریب همبستگی زمانی	✓	✓	✓	✗
George (۲۰۲۰)	ساختار توصیه گر پویا مبتنی بر کاربرد و محتوا	✓	✓	✓	✗
Kong (۲۰۱۹)	با استفاده از داده های کاربرد و محتوا و ساختار	✓	✗	✓	✗
Darvishy (۲۰۲۰)	ارائه توصیه های خبری شخصی شده ترکیبی	✓	✓	✓	✗
Lathabai (۲۰۱۷)	استفاده از قواعد انجمنی و شاخص گذاری	✓	✓	✗	✗
Parish (۲۰۱۸)	ساختار توصیه گر پویا مبتنی بر محتوا	✓	✗	✓	✗
Tóth (۲۰۲۱)	بهبود نتایج جستجوی کاربر به صورت ضمنی	✓	✓	✓	✗
Logesh (۲۰۱۹)	رویکرد ترکیبی برای پیش بینی توصیه های POI متقاعد کننده	✓	✓	✓	✗

مختلف در آزمایش های این تحقیق گنجانده شده است. تحقیق دیگری با این هدف که به مشتریان خود توصیه های شخصی بر اساس آنچه افراد مشابه با نمونه کارهای مشابه دارند ارائه دهد، پیشنهاد شده است. این سیستم از ویژگی های مشتری علاوه بر داده های نمونه کارهای مشتری استفاده می کند. از آنجا که تعداد محصولات پیشنهادی احتمالی کمتر از محصولات موجود در دامنه سیستم های توصیه گر است و یا اطلاعات محصولات پرتکرار ممکن است در دسترس نباشد، تصمیم گرفته شده که از شبکه های بیزی برای مدل سازی سیستم استفاده شود. همچنین یک رویکرد مبتنی بر یادگیری عمیق برای ارائه توصیه هایی به بالقوه (مشتریان بالقوه) ارائه شده که در آن تنها داده های بازاریابی خارجی در زمان پیش بینی در دسترس است. در [۱۸]، یک شبکه عصبی کانولوشنال (CNN) با الگوریتم بهینه سازی ازدحام ذرات (PSO)^۲ یک پارچه شده تا توصیه های بی نظیری ایجاد کند. روش پیشنهادی مبتنی بر نقاط تغییر تمرکز است که شامل پارامترهای غیرمنتظره و مرتبط است. " در جدول ۱ مقایسه روش های پیشین ارائه شده است.

۳- روش پیشنهادی

همان طور که اشاره شد، در این مقاله یک رویکرد توصیه صفحات وب بر اساس پروفایل کاربر در شبکه های اجتماعی با استفاده از ترکیب خوشه بندی ترتیبی فازای و شبکه های رمزگذاری خودکار ارائه شده است. معماری روش پیشنهادی در شکل ۱ آمده است. روش ارائه شده با مجموعه داده بازبینی کاربران آغاز به کار می کند. این

پیش بینی رفتار خرید مشتریان جدید را فراهم می کند؛ بنابراین مشکلات شروع سرد را که نمونه ای از سیستم های پیشرفته است، حل می کند. در [۱۵] یک سیستم توصیه گردشگری معرفی شده که ترجیحات کاربران را به منظور ارائه توصیه های شخصی استخراج می کند. برای این منظور، نظرات کاربران در شبکه های اجتماعی گردشگری به عنوان منبعی غنی از اطلاعات برای استخراج ترجیحات آنها استفاده می شود. سپس نظرات از قبل پردازش شده، از نظر معنایی خوشه بندی می شوند و برای شناسایی ترجیحات یک گردشگر، تحلیل احساساتی می شوند. به طور مشابه، همه نظرات جمع آوری شده کاربران در مورد یک جاذبه برای استخراج ویژگی های این نقاط مورد علاقه استفاده می شود. در [۱۶] این سیستم توصیه از فیلتر مشارکتی مبتنی بر مدل استفاده می کند و محصولات را بر اساس رتبه بندی و سابقه خرید قبلی کاربران قدیمی توصیه می کند. همچنین کاربران جدید توصیه هایی از محصولات جدید، محصولات پرتعداد و محصولات موجود در فروش دریافت خواهند کرد. کاربران موجود بر اساس محصولات اخیراً مشاهده شده، محصولات مکمل و غیره توصیه هایی دریافت خواهند کرد. با توجه به این که این تحقیق بر اساس راه اندازی یک وبسایت تجارت الکترونیکی جدید انجام شده است، در ابتدا هیچ رتبه بندی کاربر برای محصولات مختلف وجود ندارد؛ بنابراین در این مورد، توصیه ها بر اساس تجزیه و تحلیل خوشه بندی متنی توضیحات محصول فیلتر مشارکتی مبتنی بر مدل همراه با خوشه بندی متنی، در بهبود دقت و هدف قرار دادن انواع کاربران، ارائه شده است. در [۱۷] عملکرد الگوریتم های مورد استفاده در سیستم های توصیه گر بر اساس مجموعه داده های مختلف و معیارهای ارزیابی، مقایسه شده است. یک رویکرد یادگیری عمیق به نام GRU4REC و روش های ساده تر در مقایسه گنجانده شده اند. مجموعه داده های دنیای واقعی از سه حوزه

1. Convolutional Neural Network

2. Partical Swarm Optimization

نامعتبری تولید می‌کند. تطبیق نحوی ممکن است نتواند با مشکل تفاوت معنایی مطلوب روبه‌رو شود. برای فرایند بعدی، مانند سیستم بازبایی اطلاعات، ممکن است نتایج نادرستی تولید کند و باعث کاهش کارایی آن شود. رابطه معیار شباهت کسینوسی به صورت زیر بیان می‌شود

$$Sim(q, d) = \frac{\vec{q} \cdot \vec{d}}{|\vec{q}| |\vec{d}|} = \frac{\sum_{k=1}^t w_{qk} \times w_{dk}}{\sqrt{\sum_{k=1}^t (w_{qk})^2} \sqrt{\sum_{k=1}^t (w_{dk})^2}} \quad (1)$$

که در آن q و d به متن‌ها اشاره دارند و t تعداد کلمات در متن و w وزن هر کلمه است. ماتریس شباهت به عنوان ورودی روش خوشه‌بندی ترتیبی ارسال می‌شود. در خوشه‌بندی ترتیبی ابتدا مقادیر مشابهت به ترتیب نزولی مرتب‌سازی می‌شوند. سپس اولین مقدار که نشان‌دهنده بیشترین شباهت بین کاربران است، به عنوان اولین خوشه انتخاب می‌شود. کاربرانی که بر اساس این مقدار با هم شباهت دارند درون اولین خوشه قرار می‌گیرند. پس از آن مقدار دوم در ماتریس مرتب‌شده شباهت به عنوان خوشه دوم در نظر گرفته می‌شود. در آخر برای هر مقدار شباهت در ماتریس همسایگی یک خوشه در نظر گرفته می‌شود. سپس رتبه‌بندی کاربران این خوشه با توجه به تابع عضویت فازای پیش‌بینی می‌شود. تابع عضویت فازای در روش پیشنهادی به صورت زیر تعریف می‌شود

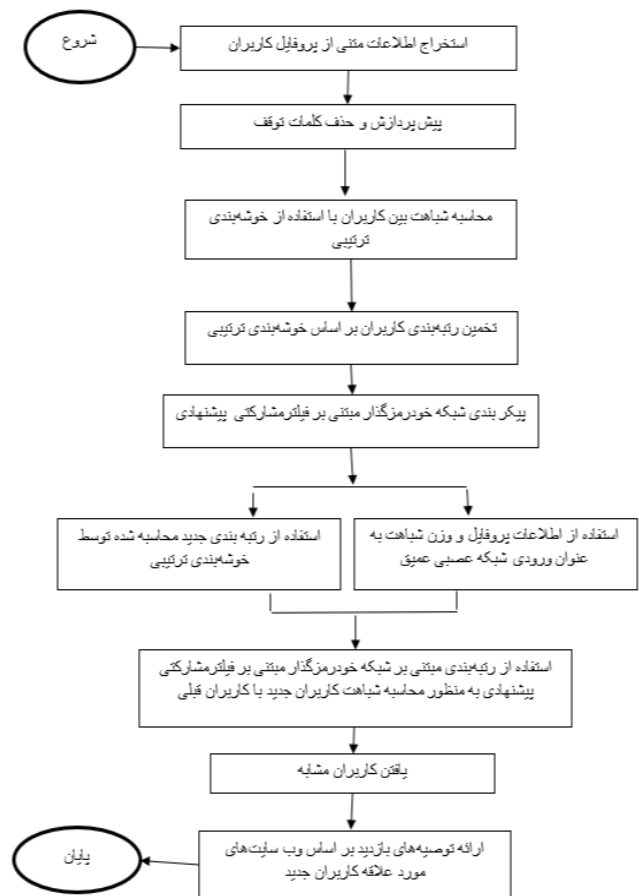
$$FMEM(i, k) = \frac{\sum_{i=1}^n \sum_{j=1}^k Rate_i - F_j \times Sim_{qd}}{n} \quad (2)$$

که در آن n تعداد کاربران در خوشه و k به عنوان کلاس یا رتبه است که می‌تواند مقداری بین بازه حقیقی یک تا پنج باشد. $Rate_i$ مقدار واقعی رتبه اعطاشده توسط کاربر و Sim_{qd} شباهت بین دو کاربر است. بر این اساس مقدار عضویت نظرات کاربران به هر یک از ۵ کلاس که شامل رتبه‌های ۱ تا ۵ است، مشخص می‌شود. پس از تعیین مقدار جدی رتبه‌بندی کاربران، این مقادیر به همراه نظرات کاربران به عنوان ورودی شبکه‌های عصبی خودرمزگذار پیشنهادی مورد استفاده قرار می‌گیرند. در ادامه به بررسی معماری شبکه‌های عصبی خودرمزگذار عمیق ارائه‌شده خواهیم پرداخت.

۳-۱ معماری شبکه‌های عصبی خودرمزگذار عمیق ارائه‌شده

در این بخش، یک رمزگذار خودکار برای یادگیری یک نمایش فشرده از ویژگی‌های ورودی برای یک مسئله مدل‌سازی پیش‌بینی طبقه‌بندی ایجاد می‌شود. ابتدا یک مسئله مدل‌سازی پیش‌بینی طبقه‌بندی تعریف می‌شود. برای تعریف یک کار طبقه‌بندی باینری (۲ کلاسه)، ویژگی‌های پروفایل کاربران در شبکه‌های اجتماعی با استفاده از ماتریس ورودی با ستون‌های ویژگی F و ردیف‌هایی از n کاربر در ماتریس ورودی اولیه مدل‌سازی می‌شوند. همچنین یک ماتریس $n \times n$ به عنوان وزن تشابه اولیه وارد رمزگذار خودکار می‌شود. در خودرمزگذارها وزن اولیه به صورت تصادفی برای ورودی‌ها تولید می‌شود و پس از چندین بار تکرار این وزن به‌روز شده و به سمت وزن ایده‌آل کاربران همگرا می‌شود. به شرطی که وزن اولیه بر اساس شباهت نظرات کاربران ارائه شود که می‌تواند روند همگرایی به سمت وزن‌های بهینه را تسریع بخشد.

وقتی شبکه رمزگذار خود را به عنوان یک استخراج‌کننده ویژگی استفاده می‌کنیم، لازم است برخی از داده‌های ورودی را در لایه میانی



شکل ۱: معماری روش پیشنهادی.

مجموعه داده شامل ویژگی‌های کاربران و نیز رأی کاربران به آیتم‌های موجود در مجموعه داده است. در مجموعه داده Yelp [۱۹] نظرات کاربران بر اساس علاقه آنها به مکان‌هایی است که مشاهده کرده‌اند و در مجموعه داده TripAdvisor [۲۰] نظرات کاربران نسبت به محل اقامت یا هتل‌هایی است که در آن بوده‌اند. در روش پیشنهادی در مجموعه داده نظرات کاربران در مورد اقامت خود نیز اشاره شده که از آن به عنوان اطلاعات پروفایل کاربری استفاده می‌شود.

در روش پیشنهادی، در ابتدا با توجه به معیار شباهت کسینوسی، ماتریس شباهت کاربران با توجه به نظرات آنها استخراج می‌شود. تشابه کسینوسی یک معیار شناخته‌شده در بازبایی اطلاعات و مطالعات مرتبط است. این معیار یک سند متنی را به عنوان یک بردار از اصطلاحات مدل‌سازی می‌کند. با استفاده از این مدل، تشابه بین دو سند می‌تواند با محاسبه مقدار کسینوس بین بردارهای اصطلاحی دو سند، مشخص شود. اجرای این معیار می‌تواند بر روی هر دو متن (جمله، پاراگراف یا کل سند) اعمال شود. در مورد موتور جستجو، مقدار تشابه بین پرس‌وجوی کاربر و سندها از بالاترین تا کمترین مرتب می‌شود. امتیاز تشابه بالاتر بین بردار اصطلاحی سند و بردار اصطلاحی پرس‌وجو به معنای وجود تطابق بیشتر بین سند و پرس‌وجو است.

تشابه کسینوسی برای اندازه‌گیری تشابه بین سند و پرس‌وجوی کاربر باید با معنای کلمات سازگار باشد. با این حال، تشابه کسینوسی هنوز نمی‌تواند به‌خوبی معنای معنایی متن را کنترل کند. اجرای اندازه‌گیری تشابه کسینوسی بین دو بردار اصطلاحی از نظر نحوی گاهی نتایج

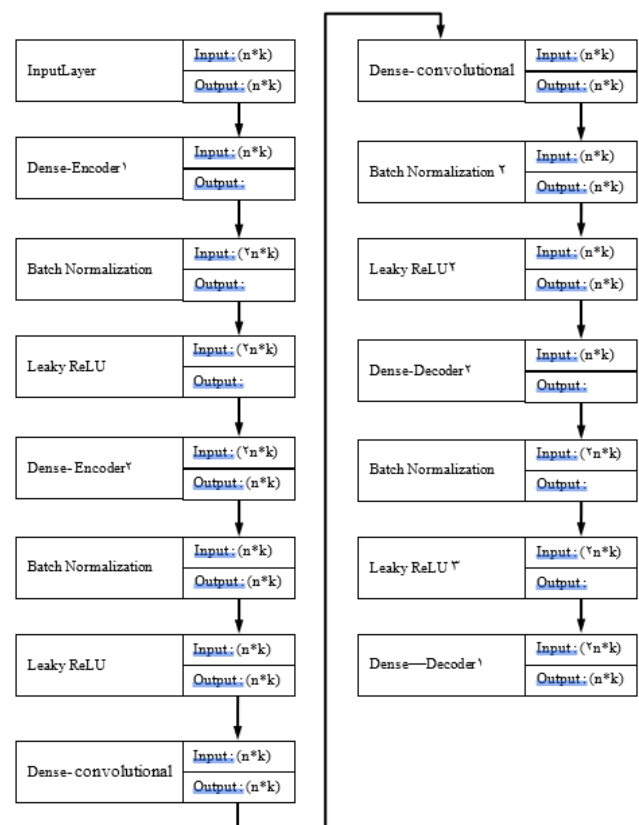
با نزول تصادفی گرادیان^۲، با کمینه کردن میانگین مربعات خطا آموزش داده می شود. شکل ۲ معماری رمزگذار ارائه شده را نشان می دهد. مطابق با شکل، معماری شبکه رمزگشایی بر اساس لایه های تعریف شده، ارائه گردیده است. خروجی این مرحله به ماژول توصیه داده می شود. تفاوت بین شبکه های عصبی ارائه شده و شبکه های قبلی در تعداد لایه های میانی است که اجازه آموزش عمیق بر روی ویژگی های پروفایل کاربران را می دهد و دقت پیش بینی رتبه بندی به کاربران جدید را افزایش می دهد. در ادامه به توضیح ماژول توصیه ارائه شده خواهیم پرداخت.

۳-۲ ماژول توصیه ارائه شده

ماژول توصیه در روش ارائه شده، آخرین بخشی است که در آن توصیه های بازدید صفحات به کاربران بر اساس نتایج شبکه رمزگذار ارائه شده اتفاق می افتد. خروجی شبکه رمزگذار شامل ویژگی های مهم استخراج شده از پروفایل کاربران و وزن شباهت کاربران است که در بخش قبلی توضیح داده شد. در این مرحله، بر اساس وزن شباهت کاربران در خوشه ها و رتبه بندی جدید بر اساس خوشه بندی ترتیبی فازی، برای هر کاربر، k کاربر برتر مشابه انتخاب می شود. مقدار k در چند سناریوی مختلف، مقادیر مختلفی را می گیرد و در نهایت عملکرد سیستم بر اساس تعداد k کاربر مشابه نیز سنجیده می شود. این k کاربر به عنوان کاربرانی در نظر گرفته می شوند که از نظر روش ارائه شده نظرات و سلاقی مشابهی می توانند داشته باشند. بنابراین برای هر کاربر، آیت مهایی که k کاربر مشابه بیشترین رتبه را به آنها اختصاص داده اند، می تواند به عنوان گزینه های مناسبی برای توصیه استفاده شوند. به عبارت دیگر، اگر کاربر i و j بر اساس مقادیر وزن شباهت، مشابه هم در نظر گرفته شده باشند، آیت مهایی که از نظر کاربر j بیشترین رتبه یا ۵ را دریافت کرده اند، می توانند به عنوان آیت مهایی مناسبی برای پیشنهاد به کاربر i در نظر گرفته شوند و همچنین آیت مهایی که از نظر کاربر i بیشترین رتبه یا ۵ را دریافت کرده اند، می توانند به عنوان آیت مهایی مناسبی برای پیشنهاد به کاربر j در نظر گرفته شوند.

۴- پیاده سازی روش ارائه شده

همان طور که اشاره شده است، در روش ارائه شده به منظور پیاده سازی سیستم توصیه گر شخصی سازی شده بر مبنای پروفایل کاربر از دو مجموعه داده زمان واقعی مقیاس بزرگ Yelp و TripAdvisor استفاده شده است. این مجموعه داده ها شامل اطلاعاتی مستخرج از پروفایل کاربران در مورد محصولات به صورت متنی و رتبه بندی کاربران از محصولات موجود می باشد. کاربران با توجه به شناسه منحصر به فرد در مجموعه داده مشخص شده اند. روش ارائه شده با توجه به این داده ها در مجموعه داده سعی در شخصی سازی صفحه وب بر اساس رفتار کاربران دارد. رفتار کاربران بر اساس اطلاعات موجود در پروفایل شخصی کاربر استخراج شده و شباهت بین کاربران بر اساس آن محاسبه می شود. در جدول ۲ مشخصات مربوط به مجموعه داده های مورد استفاده آمده است. در این مجموعه، داده ها به دو قسمت داده های آموزشی با نسبت ۷۰٪ کل داده ها و داده های تست با ۳۰٪ کل داده ها تقسیم می شود. در مجموعه داده آموزشی Yelp در کنار رتبه بندی کاربران نسبت به محصولات، نظرات و کاربران و اطلاعات پروفایل آنها به صورت متنی درج شده است. اطلاعات متنی درج شده در پروفایل کاربر در واقع



شکل ۲: معماری شبکه خود رمزگذار ارائه شده.

حذف کنیم تا ویژگی های اساسی را از ورودی ها استخراج کنیم. بنابراین مسئله به گونه ای تعریف می شود که شامل برخی از متغیرهای ورودی اضافی باشد و این امکان را به رمزگذار می دهد که در ادامه یک نمایش فشرده مفید را یاد بگیرد. شبکه های رمزگذار معمولاً به گونه ای محدود می شوند که فقط به تقریب و کپی کردن ورودی هایی که شبیه داده های آموزشی هستند، اجازه داده می شود. از آنجایی که مدل مجبور است اولویت بندی کند که کدام جنبه از ورودی باید کپی شود، اغلب ویژگی های مفید داده ها را یاد می گیرد. سپس داده های ورودی می توانند از دامنه به مدل ارائه شوند و خروجی مدل در گلوگاه می تواند به عنوان یک بردار ویژگی در یک مدل یادگیری نظارت شده، برای تجسم یا به طور کلی برای کاهش ابعاد استفاده شود. همچنین وزن هر لایه به عنوان وزن شباهت کاربران در نظر گرفته می شود و این وزن در هر لایه به روزرسانی می شود تا در نهایت وزن واقعی شباهت کاربران از شبکه رمزگذار استخراج شود.

رمزگذار به گونه ای تعریف شده که دو لایه پنهان داشته باشد: لایه اول با دو برابر تعداد ورودی ($2n$) و لایه دوم با همان تعداد ورودی (n) و پس از آن لایه گلوگاه با همان تعداد ورودی به عنوان مجموعه داده ورودی (n) است. برای اطمینان از یادگیری بهتر مدل، از نرمال سازی دسته ای و فعال سازی Leaky ReLU^۱ استفاده می شود.

رمزگشا در یک ساختار مشابه، به صورت معکوس، تعریف شده و دو لایه پنهان دارد: لایه اول با تعداد ورودی های مجموعه داده (n) و لایه دوم با دو برابر تعداد ورودی ها ($2n$). لایه خروجی همان تعداد گره های ستون های داده ورودی را دارد و از یک تابع فعال سازی خطی برای خروجی مقادیر عددی استفاده می کند. مدل با استفاده از بهینه ساز Adam



شکل ۳: ماتریس شباهت بین کاربران.

کاربری آنها در شبکه‌های اجتماعی به منظور ایجاد سیستم‌های توصیه‌گر صفحات وب هستند که در روش ارائه‌شده بر روی آن تمرکز شده است. بنابراین در روش ارائه‌شده کاربرانی که در مورد یک یا بیشتر آیتم با هم نظرات مشابهی دارند به عنوان کاربران مشابه در نظر گرفته می‌شوند. آیتم‌هایی که این کاربران در مورد آنها با هم تشابه دارند به شکل مقداری نرمال‌شده به عنوان وزن اولیه در روش ارائه‌شده در نظر گرفته می‌شود. این وزن به عنوان یک ویژگی با عنوان وزن شباهت در روش ارائه‌شده به همراه سایر ویژگی‌های موجود در پروفایل کاربری به شبکه عصبی ارائه‌شده وارد می‌شود. پیش‌پردازش مجموعه داده، کاربران با رتبه‌بندی کمتر در مکان‌ها و نتایج با ۳۹۱۰۴ مکان، ۲۰۱۶۶ کاربر منحصربه‌فرد و در کل ۵۸۶۲۷۴ رتبه‌بندی را حذف می‌کند. TripAdvisor یک وبسایت توصیه سفر است که از بررسی‌ها و بازخوردهای مکان‌ها تشکیل شده است. مشابه پیش‌پردازش مجموعه داده Yelp، مجموعه داده‌های TripAdvisor نیز از قبل پردازش شده و کاربران با رتبه‌بندی کمتر حذف می‌شوند. مجموعه داده فیلترشده تریپ‌ادایزر شامل ۹۱۴۹ مکان، ۱۳۴۱۰ کاربر منحصربه‌فرد و در کل ۱۵۲۷۲۱ رتبه‌بندی است.

پس از پیش‌پردازش، شباهت بین نظرات کاربران بر اساس معیار شباهت کسینوسی استخراج می‌شود. در شکل ۳ ماتریس شباهت کسینوسی برای ۱۰ کاربر از مجموعه داده آموزشی به عنوان نمونه نشان داده شده است. با توجه به شکل می‌توان دید ماتریس شباهت برای کاربران در سیستم توصیه‌گر ارائه‌شده بر اساس اطلاعات موجود در پروفایل کاربری محاسبه شده است. همان‌طور که دیده می‌شود مقدار شباهت برای هر زوج کاربر مقداری در محدوده [۰، ۱] است که میزان شباهت بین دو کاربر از نظر اطلاعات پروفایل کاربری را نشان می‌دهد. هرچه این مقدار بیشتر باشد نشان‌دهنده شباهت بیشتر دو کاربر از نظر اطلاعات پروفایل کاربری است. از این رو در روش ارائه‌شده برای یافتن کاربران مشابه یک آستانه تعریف می‌شود که به صورت رابطه زیر است

$$\alpha = \frac{\sum_{i=1}^n \sum_{j \neq i}^n \text{sim}(i, j)}{n} \geq 1 \quad (3)$$

در (۳) i و j اندیس کاربران و n تعداد کل کاربران در مجموعه داده

جدول ۲: مشخصات مربوط به مجموعه داده‌های مورد استفاده.

عنوان مجموعه داده	تعداد نمونه‌ها	تعداد فیلد	عناوین فیلدها
TripAdvisor	۲۰۴۹۲	۲	Review Rating business_id date review_id stars
Yelp	۱۰۰۰۰	۱۰	Text Type user_id Cool useful funny

نشان‌دهنده احساسات و عواطف کاربران نسبت به محصول هستند و در کنار مقادیر رتبه اعطاشده به محصول توسط کاربر، می‌توانند برای تحلیل رفتار کاربران مورد استفاده قرار گیرند. سیستم‌های توصیه‌گر سنتی، تحلیل اطلاعات متنی موجود در پروفایل کاربران را نادیده می‌گیرند و برای محاسبه شباهت بین کاربران از رتبه‌بندی آنها استفاده می‌کنند. حال آن که اطلاعات متنی درج‌شده در پروفایل کاربران حاوی دانش نهفته ارزشمندی است که بر اساس آن می‌توان شباهت کاربران را با دقت بالاتری سنجید. از این رو در روش ارائه‌شده تحلیل اطلاعات موجود در پروفایل کاربری ارائه خواهد شد.

اطلاعات موجود در پروفایل کاربری به صورت متن پیوسته است که قبل از اعمال تغییرات بر روی آن نیاز به پیش‌پردازش و حذف کلمات توقف^۱ است.

ماتریس کاربر-کاربر یک ماتریس مقارن است که بر اساس شباهت نظرات اعطاشده توسط کاربران تنظیم شده است. در این ماتریس هر درایه می‌تواند مقدار بین ۰ تا m دریافت کند که نشان‌دهنده اندیس آیتمی است که دو کاربر در مورد آن شباهت دارند. همچنین می‌توان دید در برخی درایه‌ها بیش از یک عدد وجود دارد که نشان می‌دهد دو کاربر به بیش از یک آیتم نظر مشابهی اعطا نموده‌اند. چنین کاربرانی که در بیش از یک مورد با هم شباهت دارند، گزینه‌های مناسبی برای بررسی پروفایل

جدول ۳: خوشه بندی ترتیبی کاربران بر اساس شباهت نظرات.

رتبه پیش بینی شده	کاربران مشابه	مقدار شباهت	شماره خوشه
۵	۶۸۰ و ۲۱۶	۰,۴۶۴۰	۱
۲	۹۹۸ و ۵۵۴	۰,۴۴۸۰	۲
۴	۴۵۷ و ۱۸	۰,۴۳۷۰	۳
۵	۹۳۵ و ۲۴۵	۰,۴۲۴۰	۴
۵	۵۴۲ و ۱۳۳	۰,۳۸۵۰	۵

در این مرحله نتایج به دست آمده از شبکه عصبی عمیق برای کاربران در مجموعه داده تست بر اساس رتبه های اعطاشده به محصولات توسط کاربران تست سنجیده می شود. بدین منظور کاربران تستی که به وبسایت حداکثر رتبه را اعطا کرده باشند و بر اساس پیش بین شبکه عصبی عمیق نیز رتبه اعطایی این کاربران به وبسایت حداکثر باشد، به عنوان نمونه های مثبت صحیحی^۱ (TP) در نظر گرفته می شوند. کاربران تستی که به وبسایت حداکثر رتبه را اعطا نکرده باشند و بر اساس پیش بین شبکه عصبی عمیق نیز رتبه اعطایی این کاربران به وبسایت حداکثر نباشد، به عنوان نمونه های منفی صحیحی^۲ (TN) در نظر گرفته می شوند. کاربران تستی که به وبسایت حداکثر رتبه را اعطا کرده باشند ولی بر اساس پیش بین شبکه عصبی عمیق، رتبه اعطایی این کاربران به وبسایت حداکثر نباشد، به عنوان نمونه های مثبت کاذب^۳ (FP) در نظر گرفته می شوند. در نهایت کاربران تستی که به وبسایت حداکثر رتبه را اعطا نکرده باشند ولی بر اساس پیش بین شبکه عصبی عمیق، رتبه اعطایی این کاربران به وبسایت حداکثر باشد، به عنوان نمونه های منفی کاذب^۴ (FN) در نظر گرفته می شوند. این چهار پارامتر با هم به عنوان پارامترهای ماتریس آشفتگی^۵ تعیین می شوند. در شکل ۶ ماتریس آشفتگی مربوط به شبکه های خودرمزگذار عمیق ارائه شده برای کاربران تست در مجموعه داده TripAdvisor آمده است.

همان طور که در شکل ۶ نشان داده شده است، ماتریس آشفتگی مربوط به روش شبکه های خودرمزگذار عمیق پیشنهادی ارائه شده است. در ادامه ارزیابی روش پیشنهادی ارائه خواهد شد.

۴-۱- ارزیابی روش ارائه شده

کیفیت عملکرد سیستم توصیه گر می تواند توسط معیارهای زیادی مورد سنجش قرار گیرد. نوع معیار مورد استفاده به روش فیلتر و ارائه به کاررفته بستگی دارد. دقت عملکرد سیستم توصیه گر بر اساس دقت توصیه های ارائه شده به کاربران تعیین می شود. در این تحقیق با توجه به ترکیب در روش فیلتر مشارکتی و یادگیری عمیق، از دو معیار برای سنجش دقت عملکرد روش ارائه شده ترکیبی استفاده شده است. معیار اول معیار دقت آماری برای سنجش میزان دقت رتبه بندی پیش بینی شده برای کاربران در روش فیلتر مشارکتی و دقت در ارائه توصیه به کاربران مشابه در میان کاربران جدید است. معیار دوم ارزیابی عملکرد شبکه های عصبی عمیق به منظور پیش بینی رتبه ارائه شده به محصول بر اساس اطلاعات موجود در پروفایل کاربر و مقایسه آن با سایر روش های طبقه بندی است. در ادامه ابتدا ارزیابی روش پیشنهادی بر اساس دقت آماری ارائه خواهد شد و بعد از آن دقت روش های طبقه بندی با هم مقایسه می گردد.

۴-۱-۱- ارزیابی عملکرد روش ارائه شده بر اساس دقت آماری

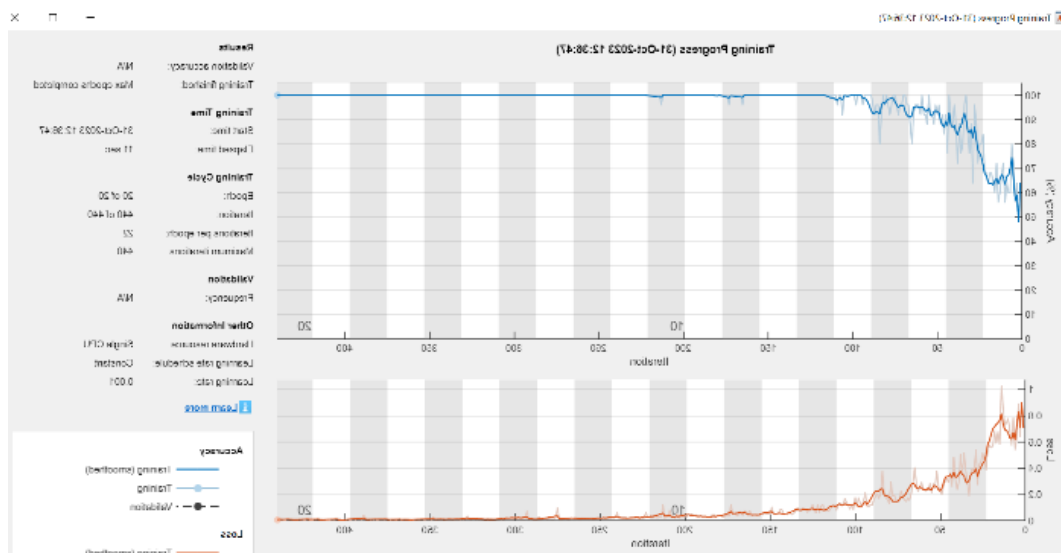
همان طور که اشاره شد، معیارهای دقت آماری، دقت از روش فیلتر مشارکتی را با مقایسه به طور مستقیم رتبه های پیش بینی شده برای کاربران را با امتیاز واقعی شده آنها مورد بررسی قرار می دهد. میانگین خطای مطلق^۶ (MAE) و خطای میانگین مربعات^۷ ($RMSE$) معمولاً

آموزشی است. پس از محاسبه شباهت بین نظرات کاربران، در روش ارائه شده خوشه بندی مبتنی ترتیب فازی انجام می شود که در آن کاربران به هر خوشه با توجه به مقدار شباهت اختصاص می یابند. پس از اختصاص کاربران به خوشه ها، مقدار رتبه بندی جدید کاربران با توجه به تابع عضویت فازی محاسبه می شود. در جدول ۳ نمونه ای از خوشه بندی ترتیبی کاربران بر اساس شباهت نظرات آنها نشان داده شده است.

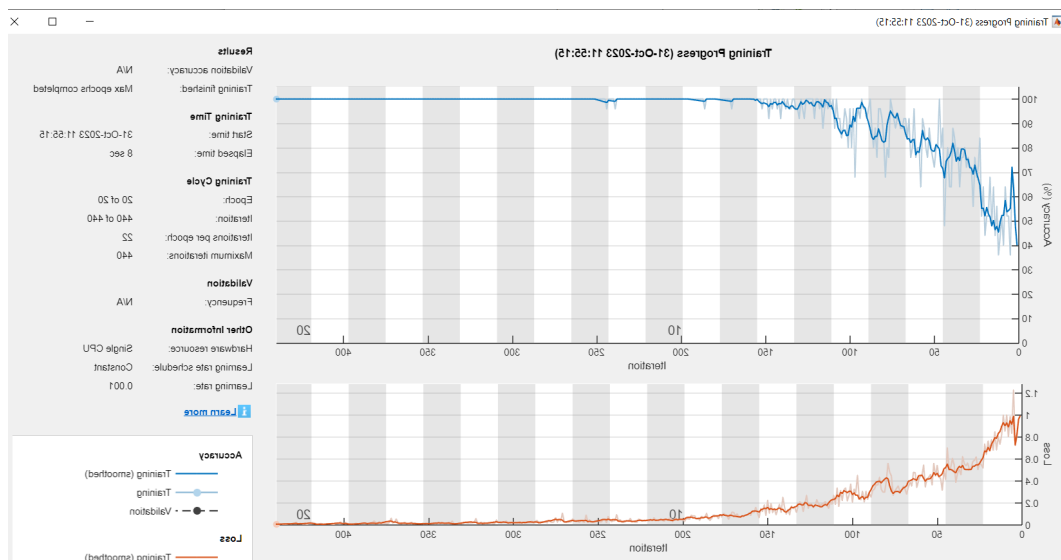
با توجه به جدول می توان دید کاربران بر اساس شباهت نظرات در خوشه ها قرار می گیرند. با توجه به این که هر کاربر فقط در یک خوشه می تواند جا گیرد، ممکن است در برخی از خوشه ها تنها یک کاربر وجود داشته باشد. بر اساس روش ارائه شده کاربران اختصاص یافته به یک خوشه، بالاترین احتمال رفتارهای مشابه را دارند که بر اساس آن توصیه های بازدید وب پیشنهاد شده است. در واقع کاربرانی با شباهت بیشتر از نظر اطلاعات پروفایل کاربری، می توانند گزینه مناسبی برای ورودی روش فیلتر مشارکتی بر اساس یادگیری عمیق ارائه شده باشد. در روش ارائه شده به منظور شناسایی دقیق کاربران مشابه از یادگیری عمیق استفاده شده است. شبکه های خودرمزگذار عمیق ارائه شده با توجه به کاربردهای متنوع، گزینه مناسبی برای یافتن کاربران مشابه از نظر اطلاعات پروفایل کاربری هستند. در این نوع از شبکه ها اطلاعات پروفایل کاربری و ماتریس شباهت بین کاربران به عنوان ویژگی های ورودی شبکه در نظر گرفته می شود. همچنین رتبه بندی کاربران به محصولات به عنوان متغیر هدف در نظر گرفته می شود. شبکه های عصبی ارائه شده با آموزش عمیق سعی می کنند ویژگی های هر کاربر را یاد گرفته و کاربران مشابه از نظر رتبه های اعطاشده به وبسایت ها توسط کاربران را پیدا کنند. در نهایت به منظور اعتبارسنجی آموزش خود از شبکه تشابه بین کاربران به عنوان معیار محک و ارزیابی استفاده می کنند. در صورتی که دقت شبکه های عصبی ارائه شده در حد مطلوبی باشد، نتایج آموزش به صورت وزن ویژگی ها و کلمات موجود در پروفایل کاربری استخراج شده و به عنوان الگویی برای یافتن کاربران مشابه جدید مورد استفاده قرار می گیرند. از این رو هرچه آموزش شبکه های عصبی عمیق دقیق تر باشد، کاربران مشابه با دقت بیشتری یافته می شود. یافتن کاربران مشابه می تواند موجب ارائه توصیه های دقیق تری به کاربران و علی الخصوص کاربران جدید باشد. در شکل ۴ و ۵ روند آموزش شبکه عصبی ارائه شده بر روی مجموعه داده های آموزشی TripAdvisor و Yelp نشان داده شده است. همچنین در جداول ۴ و ۵ میزان دقت شبکه های عصبی ارائه شده در یافتن کاربران مشابه در طی مراحل تکرار در مجموعه داده آموزشی TripAdvisor و Yelp نشان داده شده است.

با توجه به شکل ۴ و ۵ و جداول ۴ و ۵ می توان دید شبکه های عصبی ارائه شده در ابتدا در تنظیم وزن ها و آموزش ویژگی های مربوط به کاربران مقدار خطای بالایی داشته که با افزایش تعداد تکرارها این خطا به مقدار نزدیک به صفر رسیده است.

1. True Positive
2. True Negative
3. False Positive
4. False Negative
5. Confusion Matrix
6. Mean Absolute Error
7. Root Mean Square Error



شکل ۴: روند آموزش شبکه عصبی ارائه شده بر روی مجموعه داده آموزشی TripAdvisor.



شکل ۵: روند آموزش شبکه عصبی ارائه شده بر روی مجموعه داده آموزشی Yelp.

جدول ۵: میزان دقت شبکه عصبی ارائه شده در یافتن کاربران مشابه در طی مراحل تکرار در مجموعه داده آموزشی YELP.

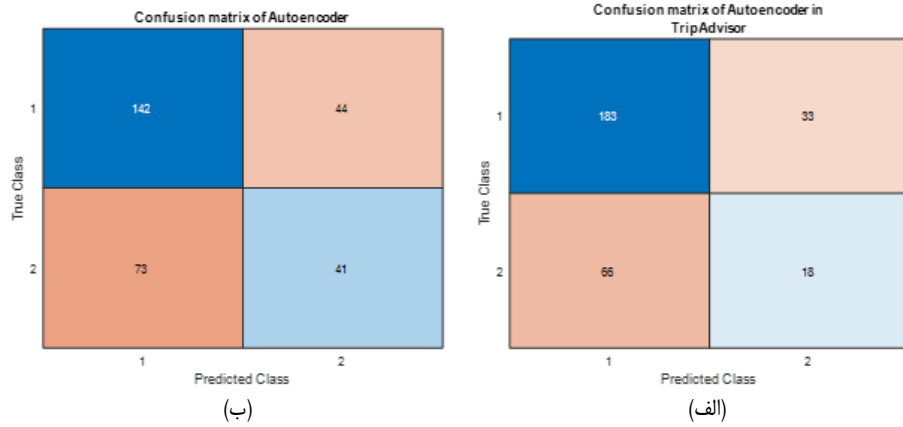
Epoch	Iteration	Time Elapsed (hh:mm:ss)	Mini-batch Accuracy	Mini-batch Loss	Base Learning Rate
۱	۱	۰۰:۰۰:۰۳	%۴۰/۰۰	۰/۰۹۹۱۸	۰/۰۰۱۰
۳	۵۰	۰۰:۰۰:۰۳	%۷۲/۰۰	۰/۰۵۳۹۴	۰/۰۰۱۰
۵	۱۰۰	۰۰:۰۰:۰۴	%۹۶/۰۰	۰/۰۲۰۶۰	۰/۰۰۱۰
۷	۱۵۰	۰۰:۰۰:۰۴	%۱۰۰/۰۰	۰/۰۲۲۱۴	۰/۰۰۱۰
۱۰	۲۰۰	۰۰:۰۰:۰۵	%۱۰۰/۰۰	۰/۰۵۲۴	۰/۰۰۱۰
۱۲	۲۵۰	۰۰:۰۰:۰۶	%۱۰۰/۰۰	۰/۰۳۷۶	۰/۰۰۱۰
۱۴	۳۰۰	۰۰:۰۰:۰۶	%۱۰۰/۰۰	۰/۰۲۲۰	۰/۰۰۱۰
۱۶	۳۵۰	۰۰:۰۰:۰۷	%۱۰۰/۰۰	۰/۰۱۱۵	۰/۰۰۱۰
۱۹	۴۰۰	۰۰:۰۰:۰۷	%۱۰۰/۰۰	۰/۰۰۸۱	۰/۰۰۱۰
۲۰	۴۴۰	۰۰:۰۰:۰۸	%۱۰۰/۰۰	۰/۰۰۶۲	۰/۰۰۱۰

جدول ۴: میزان دقت شبکه عصبی ارائه شده در یافتن کاربران مشابه در طی مراحل تکرار در مجموعه داده آموزشی TRIPADVISOR.

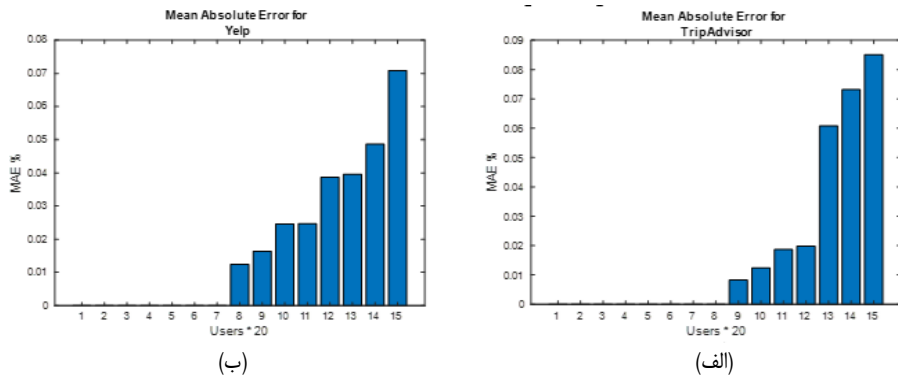
Epoch	Iteration	Time Elapsed (hh:mm:ss)	Mini-batch Accuracy	Mini-batch Loss	Base Learning Rate
۱	۱	۰۰:۰۰:۰۴	%۶۴/۰۰	۰/۰۷۰۴۷	۰/۰۰۱۰
۳	۵۰	۰۰:۰۰:۰۵	%۹۲/۰۰	۰/۰۲۲۵۷	۰/۰۰۱۰
۵	۱۰۰	۰۰:۰۰:۰۶	%۱۰۰/۰۰	۰/۰۱۷۶۳	۰/۰۰۱۰
۷	۱۵۰	۰۰:۰۰:۰۶	%۱۰۰/۰۰	۰/۰۴۹۵	۰/۰۰۱۰
۱۰	۲۰۰	۰۰:۰۰:۰۷	%۱۰۰/۰۰	۰/۰۳۳۸	۰/۰۰۱۰
۱۲	۲۵۰	۰۰:۰۰:۰۸	%۱۰۰/۰۰	۰/۰۱۴۷	۰/۰۰۱۰
۱۴	۳۰۰	۰۰:۰۰:۰۹	%۱۰۰/۰۰	۰/۰۲۷۳	۰/۰۰۱۰
۱۶	۳۵۰	۰۰:۰۰:۱۰	%۱۰۰/۰۰	۰/۰۰۸۷	۰/۰۰۱۰
۱۹	۴۰۰	۰۰:۰۰:۱۰	%۱۰۰/۰۰	۰/۰۰۷۳	۰/۰۰۱۰
۲۰	۴۴۰	۰۰:۰۰:۱۱	%۱۰۰/۰۰	۰/۰۰۵۹	۰/۰۰۱۰

هرچه مقدار این معیار کمتر باشد، رتبه بندی پیش بینی شده برای کاربران با امتیاز واقعی آنها نزدیک تر است. این معیارها به صورت (۴) و (۵) محاسبه می شوند

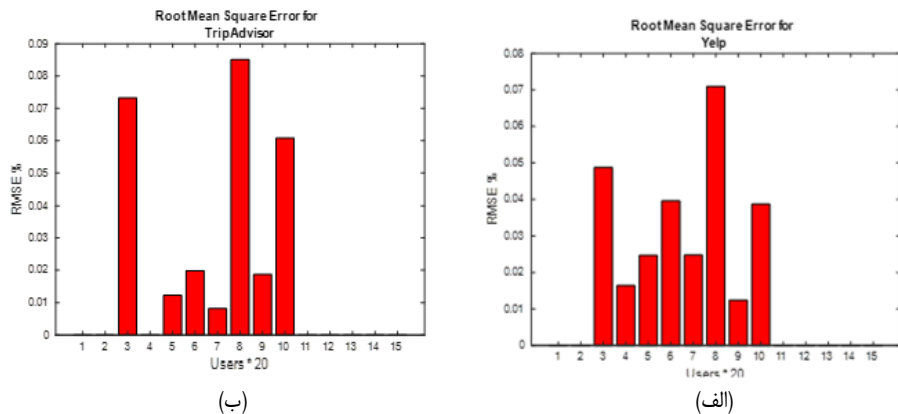
به عنوان معیارهای دقت آماری استفاده می شود. MAE از محبوب ترین و معمول ترین روش هایی است که در بین بیشتر روش ها مورد استفاده قرار گرفته است. طبیعتاً هرچه مقدار این معیار در ازای کاربران، کمتر باشد، دقت عملکرد روش فیلتر مشارکتی بیشتر خواهد بود. به عبارت دیگر



شکل ۶: ماتریس آشفتگی مربوط به شبکه‌های خودرمزگذار عمیق در مجموعه داده‌های (الف) TripAdvisor و (ب) Yelp.



شکل ۷: مقدار خطای مطلق در ازای کاربران تست.



شکل ۸: مقدار خطای میانگین مربعات در ازای کاربران تست.

تجمعی خطاهای مطلق و میانگین مربعات به ازای همه کاربران تست در مجموعه داده‌های TripAdvisor و Yelp محاسبه خواهد شد. شکل ۹ نمودار تجمعی مقدار خطای مطلق در ازای همه کاربران تست در مجموع داده TripAdvisor را نشان می‌دهد. همچنین در شکل ۱۰ نمودار تجمعی خطای میانگین مربعات در ازای همه کاربران تست در مجموع داده TripAdvisor نشان داده شده است.

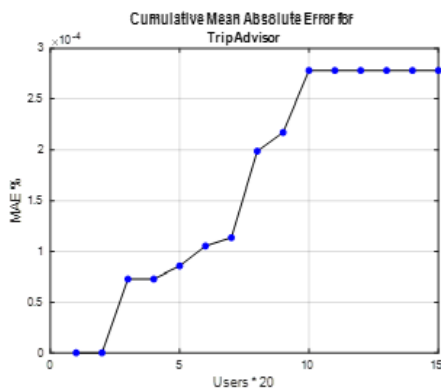
همان طور که در شکل ۹ و ۱۰ دیده می‌شود، نمودار تجمعی مقادیر خطای مطلق و خطای میانگین مربعات برای کاربران تست در هر دو مجموعه داده با شیب ملایمی افزایش می‌یابد که نشان‌دهنده دقت بالای پیش‌بینی رتبه اعطاشده به کاربر جدید بر اساس شبکه‌های خودرمزگذار عمیق ارائه شده است. در نهایت مقدار خطای مطلق تجمعی در مجموعه داده‌های TripAdvisor و Yelp به ترتیب به حداکثر مقدار ۰/۰۰۴۲۲ و ۰/۰۰۲۲۸ و مقدار خطای میانگین مربعات تجمعی به ترتیب به حداکثر مقدار ۰/۰۰۲۷۲ و ۰/۰۰۲۶۷ می‌رسد.

$$MAE = \frac{1}{N} \sum_{u,i} |p_{u,i} - r_{u,i}| \quad (۴)$$

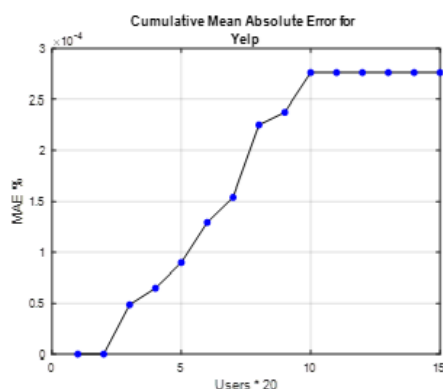
$$RMSE = \sqrt{\frac{\sum_{u,i} (p_{u,i} - r_{u,i})^2}{N}} \quad (۵)$$

که $P_{u,i}$ به عنوان رتبه پیش‌بینی‌شده برای کاربر u روی محصول i ، $r_{u,i}$ رتبه واقعی کاربر u روی محصول i و N کل تعداد رتبه‌بندی در مجموعه گزینه‌ها تعریف شده است. شکل ۷ مقدار خطای مطلق در ازای کاربران تست در مجموعه داده‌های TripAdvisor و Yelp را نشان می‌دهد. همچنین در شکل ۷ مقدار خطای میانگین مربعات در ازای کاربران تست در مجموعه داده TripAdvisor و Yelp آمده است.

همان طور که در شکل ۷ و ۸ نشان داده شده است، مقادیر خطای مطلق و خطای میانگین مربعات برای هر یک از کاربران در مجموعه داده‌های تست TripAdvisor و Yelp محاسبه شده است. حال مقادیر

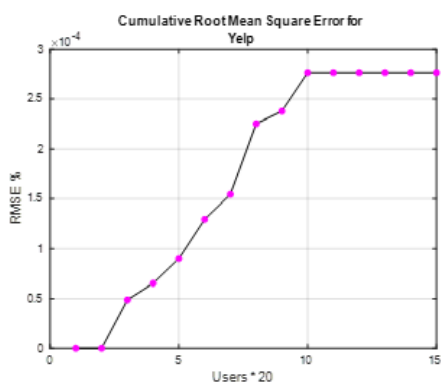


(ب)

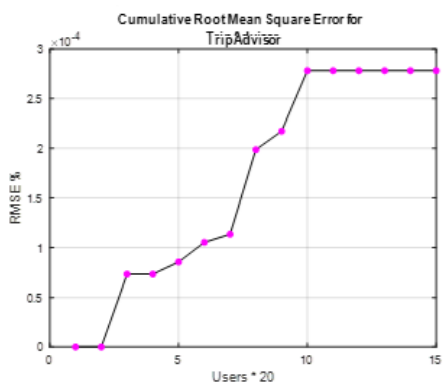


(الف)

شکل ۹: نمودار تجمعی خطای میانگین مربعات.

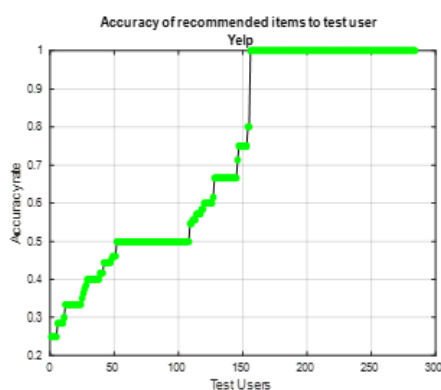


(ب)

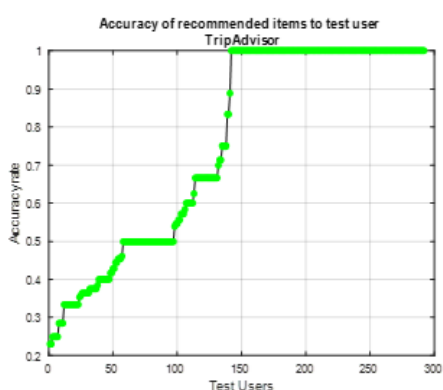


(الف)

شکل ۱۰: نمودار تجمعی مقدار خطای میانگین مربعات.



(ب)



(الف)

شکل ۱۱: دقت توصیه‌های ارائه‌شده.

توصیه‌های مفید، نسبت توصیه‌های مفید در میان کل توصیه‌های ارائه‌شده مشخص گردیده است. در واقع دقت توصیه‌های ارائه‌شده برابر با نسبت توصیه‌های مفید ارائه‌شده نسبت به کل توصیه‌های ارائه‌شده به کاربران تست می‌باشد. در شکل ۱۱ دقت توصیه‌های ارائه‌شده در مجموعه داده‌های TripAdvisor و Yelp نشان داده شده است.

با توجه به شکل ۱۱ می‌توان دید دقت توصیه‌های ارائه‌شده در مجموعه داده TripAdvisor و Yelp به صورت صعودی افزایش می‌یابد که این مورد نشان‌دهنده پیش‌بینی دقیق رتبه اعطاشده به کاربران جدید و یافتن کاربران مشابه بر اساس اطلاعات پروفایل کاربری است که به توصیه‌های بازدید وب دقیق منجر شده است. میانگین دقت ارائه توصیه بازدید وب در مجموعه داده‌های TripAdvisor و Yelp به ترتیب برابر با ۷۴/۵۳٪ و ۷۲/۲۲٪ است.

۴-۱-۲ ارزیابی عملکرد روش ارائه‌شده بر اساس دقت توصیه‌های ارائه‌شده

پس از محاسبه دقت آماری بر اساس خطای مطلق و میانگین مربعات برای کاربران تست در مجموعه داده‌های TripAdvisor و Yelp، حال نوبت به محاسبه دقت توصیه‌های ارائه‌شده به کاربران تست در این مجموعه داده‌ها می‌رسد. به منظور محاسبه دقت توصیه‌های ارائه‌شده، توصیه‌های مفید از میان تمام توصیه‌های ارائه‌شده در مجموعه داده شناسایی می‌شود. توصیه‌های مفید شامل آن دسته از وبسایت‌هایی است که به کاربر I توصیه شده و کاربر i در مجموعه داده اصلی به آن وبسایت‌ها حداکثر مقدار رتبه‌بندی را اعطا نموده است. با توجه به این که مقدار رتبه‌بندی کاربران تست در مجموعه داده اصلی مشخص است، به راحتی می‌توان توصیه‌های مفید را مشخص نمود. پس از تشخیص

۴-۱-۳ ارزیابی دقت روش های طبقه بندی

ارزیابی روش ارائه شده برای تجزیه و تحلیل و بررسی عملکرد روش ارائه شده بر روی نمونه های جدید که روند آموزش بر روی آنها اتفاق نیفتاده است، به منظور پیش بینی رتبه بندی این نمونه ها صورت می گیرد. به منظور بررسی عملکرد و ارزیابی روش های طبقه بندی که داده ها در کلاس ها مختلفی توزیع شده اند، معیارهای سنجش مختلفی وجود دارد که بیشتر این معیارها از تقابل رتبه بندی پیش بینی شده برای نمونه ها در مقابل رتبه بندی واقعی آنها که در ماتریس آشفتگی مشخص می شود، به دست می آیند. برای محاسبه معیارهای ارزیابی در ماتریس آشفتگی چهار پارامتر اصلی مثبت صحیح (TP)، مثبت کاذب (FP)، منفی صحیح (TN) و منفی کاذب (FN) باید از روش طبقه بندی به دست آید که این پارامترهای قبلاً معرفی شده اند. بر اساس این پارامترها می توان معیارهای ارزیابی زیر را تعریف کرد

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (6)$$

$$Recall = \frac{TP}{TP + FN} \quad (7)$$

$$Precision = \frac{TP}{TP + FP} \quad (8)$$

$$F - measure = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (9)$$

معیارهای ارزیابی فوق از معروف ترین معیارهایی هستند که با اتکا بر آنها عملکرد یک روش طبقه بندی مورد سنجش قرار گرفته است. این معیارها استاندارد هستند و بر اساس آنها می توان روش ارائه شده را با سایر روش ها در این زمینه مورد مقایسه قرار داد. در روش ارائه شده علاوه بر شبکه های خودرمزگذار عمیق ارائه شده از سه طبقه بند معروف K نزدیک ترین همسایه، ماشین پشتیبان بردار و طبقه بند بیزین به منظور پیش بینی رتبه بندی کاربران تست استفاده گردیده و با شبکه های خودرمزگذار عمیق ارائه شده مقایسه می شود. در روش پیشنهادی مقدار K در الگوریتم K نزدیک ترین همسایه به صورت پیش فرض ۵ در نظر گرفته شده است. همچنین برای الگوریتم بیزین از رابطه قضیه بیز که در (۱۰) نشان داده شده است، استفاده نموده ایم. قضیه بیز را می توان به صورت زیر نشان داد

$$P(Q|R) = \frac{P(R|Q)P(Q)}{P(R)} \quad (10)$$

در (۱۰)، احتمال Q با توجه به R محاسبه می شود. $P(Q)$ احتمال Q به طور مستقل، $P(R)$ احتمال R به طور مستقل و $P(R|Q)$ احتمال R به شرط تحقق Q است. در شکل ۱۲ ماتریس آشفتگی مربوط به سایر روش های طبقه بندی نشان داده شده است.

در شکل ۱۳ نمودار معیار دقت در روش های طبقه بندی مختلف در مجموعه داده TripAdvisor و Yelp نمایش داده شده است. همان طور که در شکل می توان دید، در روش ارائه شده بر اساس اطلاعات موجود در پروفایل کاربری و رتبه بندی کاربران نسبت به صفحات وب و ترکیب آن با شبکه عصبی عمیق و سایر طبقه بندها، روش شبکه عصبی عمیق نسبت به سایر روش های طبقه بندی مقدار بهتری برای معیار دقت را به دست آورده است. با توجه به شکل ۱۳ می توان دریافت تعداد کاربران تست ۳۰۰ کاربر می باشد که مقادیر دقت پیش بینی از هر ۱۰ تست یک بار در شکل ظاهر شده است. این کار برای جلوگیری از ازدحام نقاط بر روی نمودار

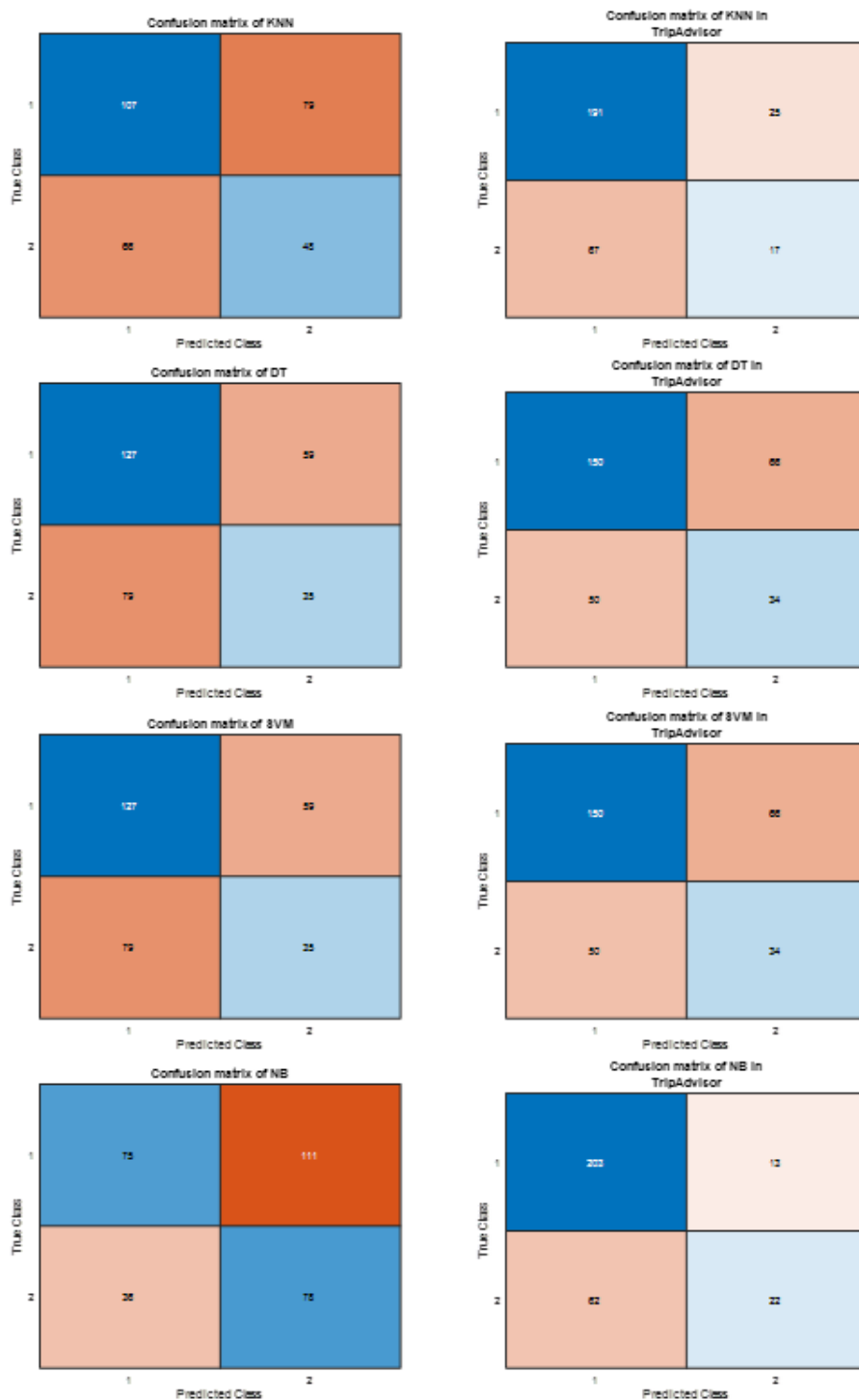
دقت و سایر نمودارها صورت گرفته است تا وضوح تصاویر مربوط به دقت و سایر معیارهای ارزیابی افزایش یابد.

معیار دیگری که در روش پیشنهادی مورد استفاده قرار گرفته است، معیار حساسیت است. معیار حساسیت به عنوان نسبت رتبه های حداکثری پیش بینی شده، تعریف می شود که به طور صحیح در میان تمام رتبه های حداکثری در مجموعه داده تست پیش بینی شده است. در شکل ۱۴ نمودار معیار حساسیت در روش های طبقه بندی مختلف در مجموعه داده TripAdvisor و Yelp نمایش داده شده است. همان طور که در شکل می توان دید، در روش ارائه شده بر اساس اطلاعات موجود در پروفایل کاربری و رتبه بندی کاربران نسبت به صفحات وب و ترکیب آن با شبکه عصبی عمیق و سایر طبقه بندها، روش شبکه عصبی عمیق نسبت به سایر روش های طبقه بندی مقدار بهتری را برای معیار حساسیت به دست آورده است.

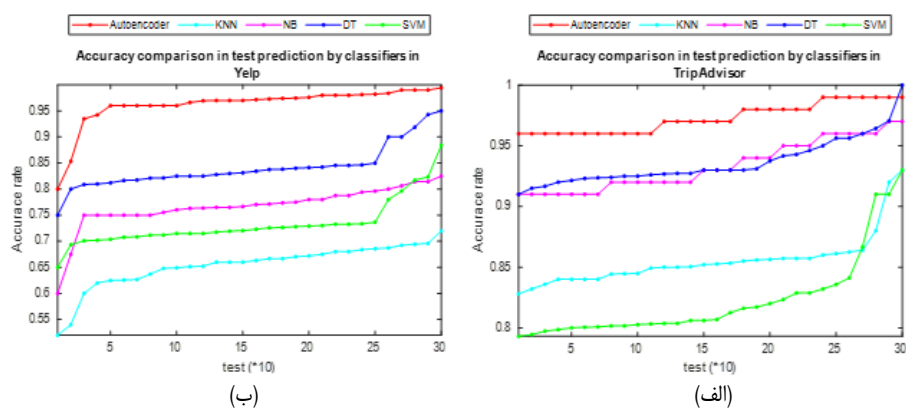
سومین معیار در روش ارائه شده، معیار صحت است. معیار صحت نشان دهنده درستی رتبه های حداکثری پیش بینی شده در میان رتبه های حداکثری کاربران تست کشف شده توسط روش های یادگیری عمیق است. در شکل ۱۵ نمودار معیار صحت در روش های طبقه بندی مختلف در مجموعه داده TripAdvisor و Yelp نمایش داده شده است. همان طور که در شکل می توان دید، در روش ارائه شده بر اساس اطلاعات موجود در پروفایل کاربری و رتبه بندی کاربران نسبت به صفحات وب و ترکیب آن با شبکه عصبی عمیق و سایر طبقه بندها، روش شبکه عصبی عمیق نسبت به سایر روش های طبقه بندی مقدار بهتری برای معیار صحت را به دست آورده است.

آخرین معیار استفاده شده، معیار F است که حاصل ترکیب صحت و حساسیت است. این معیار نشان دهنده عملکرد کلی مدل ارائه شده به ازای تشخیص رتبه بندی کاربران در مجموعه داده های استاندارد است. در شکل ۱۶ نمودار معیار F در روش های طبقه بندی مختلف در مجموعه داده TripAdvisor و Yelp نمایش داده شده است. همان طور که در شکل می توان دید، در روش ارائه شده بر اساس اطلاعات موجود در پروفایل کاربری و رتبه بندی کاربران نسبت به صفحات وب و ترکیب آن با شبکه عصبی عمیق و سایر طبقه بندها، روش شبکه عصبی عمیق نسبت به سایر روش های طبقه بندی مقدار بهتری برای معیار F را به دست آورده است. با توجه به نتایج به دست آمده می توان گفت روش ارائه شده با استفاده از اطلاعات موجود در پروفایل کاربری و رتبه بندی کاربران نسبت به وبسایت ها ورودی بسیار مناسبی برای روش طبقه بندی شبکه عصبی عمیق مهیا کند. روش شبکه عصبی عمیق با توجه به مناسب بودن داده های ورودی، قدرت آموزش و پیش بینی بالایی بر روی این مجموعه ویژگی ها به دست آورده اند. در میان روش های طبقه بندی، روش شبکه عصبی عمیق مناسب ترین نتایج را برای طبقه بندی نمونه های آموزشی و پیش بینی نمونه های رتبه کاربران در مجموعه داده آزمایشی به دست آورده است. در شکل ۱۷ نمودار میله ای مقایسه تفکیکی معیارهای ارزیابی برای روش های طبقه بندی مختلف در مجموعه داده TripAdvisor و Yelp نشان داده شده است. در شکل ۱۸ نمودار میله ای مقایسه تجمیعی معیارهای ارزیابی برای روش های طبقه بندی مختلف در مجموعه داده TripAdvisor و Yelp نشان داده شده است.

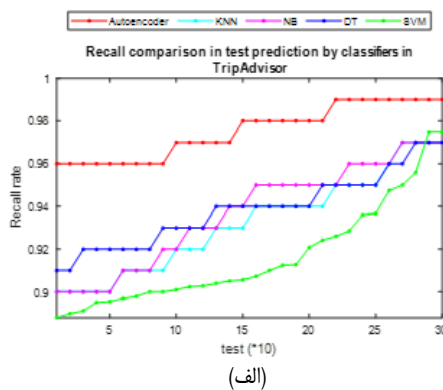
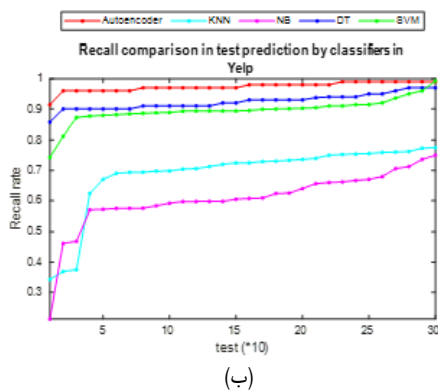
با توجه به شکل ۱۸ می توان دید که شبکه های خودرمزگذار عمیق ارائه شده با توجه به یادگیری عمیق بر روی اطلاعات موجود در پروفایل کاربری در مجموعه داده TripAdvisor و Yelp، توانسته اند با بالاترین دقت رتبه بندی محصولات برای کاربران جدید را پیش بینی نماید. روش



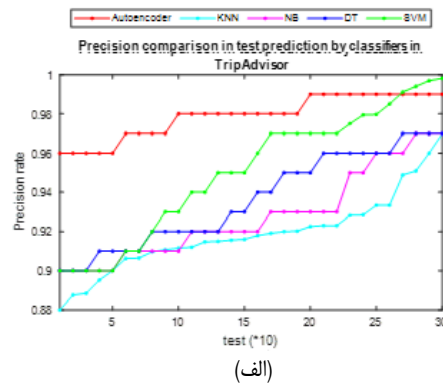
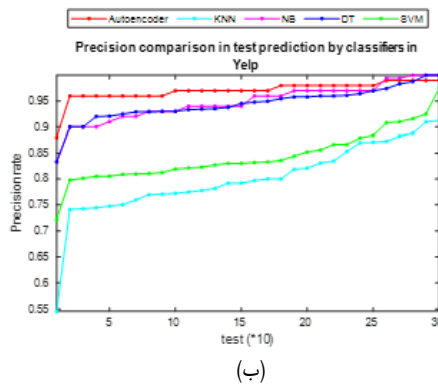
شکل ۱۲: ماتریس آشفته‌گی سایر روش‌های طبقه‌بندی.



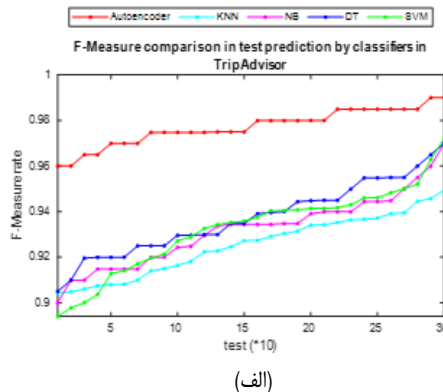
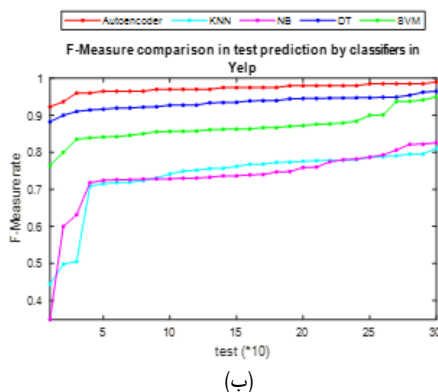
شکل ۱۳: نمودار معیار دقت در روش‌های طبقه‌بندی.



شکل ۱۴: نمودار معیار حساسیت در روش های طبقه بندی.



شکل ۱۵: نمودار معیار صحت در روش های طبقه بندی مختلف در مجموعه داده Yelp.



شکل ۱۶: نمودار معیار F در روش های طبقه بندی.

مشابه استفاده کرد. در شکل ۱۹ مقایسه روش پیشنهادی با روش های پیشین از نظر معیار میانگین خطای مطلق آمده است.

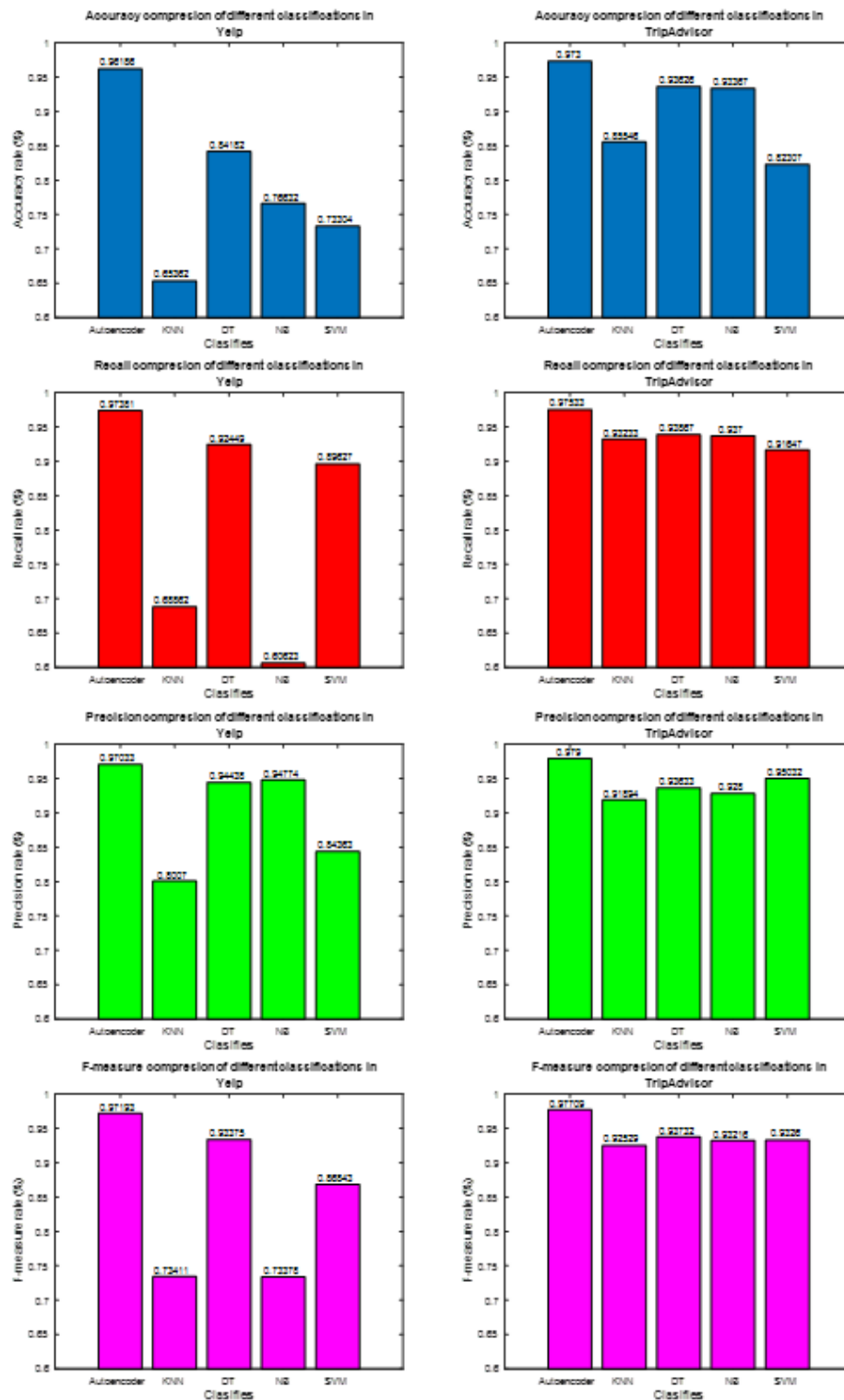
همان طور که در شکل ۱۹ قابل مشاهده است، روش پیشنهادی میانگین خطای مطلق کمتری نسبت به روش های پیشین به دست آورده است. دلیل این امر آموزش عمیق ویژگی های پروفایل کاربر و یافتن کاربران مشابه با دقت بالایی توسط شبکه عصبی خودرمزگذار پیشنهادی است. از این رو انتظار می رود با توجه به مقدار کم خطای مطلق، دقت توصیه های ارائه شده که به عنوان نرخ برخورد نیز در مقالات یاد می شود، در روش پیشنهادی نسبت به سایر روش های موجود، مقدار بهتری داشته باشد. در شکل ۲۰ مقایسه روش پیشنهادی با روش های پیشین از نظر نرخ برخورد یا دقت توصیه های ارائه شده به کاربران آمده است.

با توجه به شکل ۲۰ می توان دید روش پیشنهادی با تکیه بر آموزش عمیق در شبکه های خودرمزگذار پیشنهادی توانسته است، ویژگی های کاربران را به خوبی آموزش دهد که همین امر موجب کاهش خطای مطلق در ارائه توصیه و افزایش نرخ برخورد یا دقت توصیه های ارائه شده

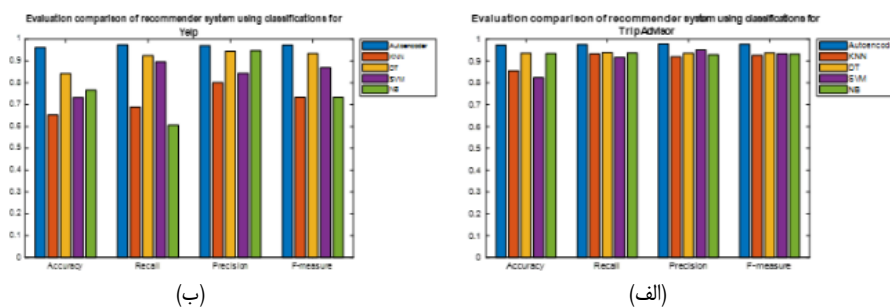
ارائه شده با توجه به لایه های آموزش عمیق و تکمیل فرایند آموزش در لایه میانی توانسته است از سایر روش های طبقه بندی از نظر معیارهای ارزیابی بهتر عمل کند.

۴-۲ مقایسه روش پیشنهادی با روش های پیشین

پس از پیاده سازی و ارزیابی به منظور اعتبارسنجی، روش پیشنهادی با روش های پیشین شامل [۲۱] مورد مقایسه قرار گرفته است. آنچه مسلم است این است که مقایسه بر روی نتایج روش های پیشین بر روی مجموعه داده های یکسان شامل مجموعه داده Yelp و TripAdvisor در شرایط یکسانی صورت می گیرد. از این رو نتایج به دست آمده در روش پیشنهادی در شرایط یکسان و مجموعه داده یکسان با روش های دیگری مقایسه می شود تا میزان بهبود حاصل از ایده پیشنهادی در این مقاله مشخص شود. با توجه به این که در روش پیشنهادی از ترکیب یادگیری عمیق و فیلتر مشارکی استفاده شده است، می توان برای مقایسه از معیارهایی مانند میانگین خطای مطلق، نرخ برخورد و دقت یافتن کاربران



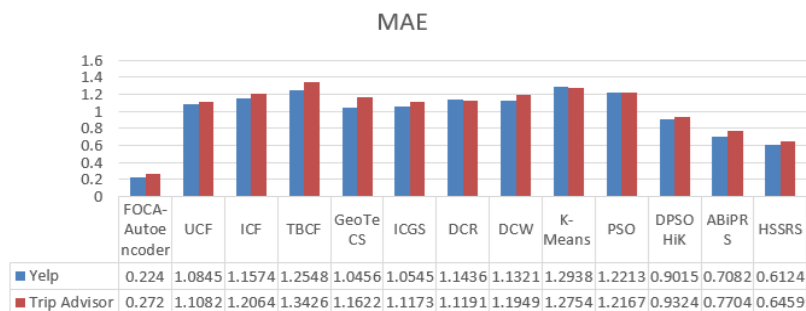
شکل ۱۷: نمودار میله‌ای مقایسه تفکیکی معیارهای ارزیابی.



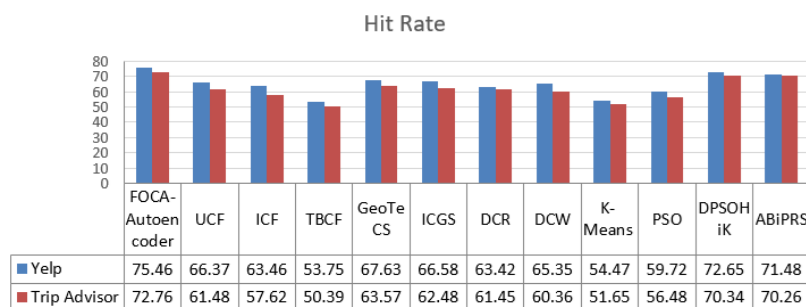
(ب)

(الف)

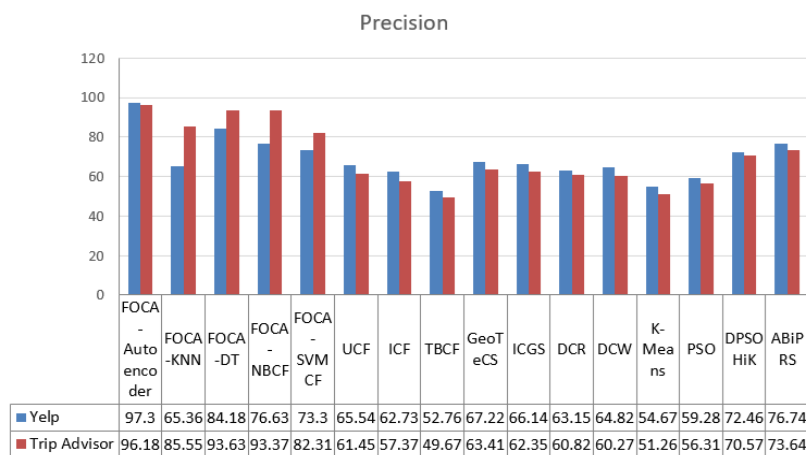
شکل ۱۸: نمودار میله‌ای مقایسه تجمیعی معیارهای ارزیابی.



شکل ۱۹: مقایسه روش پیشنهادی با روش‌های پیشین از نظر میانگین خطای مطلق.



شکل ۲۰: مقایسه روش پیشنهادی با روش‌های پیشین از نظر نرخ برخورد.



شکل ۲۱: مقایسه روش پیشنهادی با روش‌های پیشین از نظر دقت تشخیص کاربران مشابه.

۵- نتیجه‌گیری

سیستم توصیه‌گر بر اساس الگوریتم‌های پیچیده و مدل‌های یادگیری ماشین، با تحلیل رفتار کاربران و استفاده از اطلاعات موجود درباره آنها، توصیه‌هایی ارائه می‌دهد. این توصیه‌ها ممکن است بر اساس سابقه مشاهده و علاقه‌های کاربران، محتوای مرتبط و جذاب برای آنها پیشنهاد شود. در صورتی که کاربر در حال جستجوی یک محصول خاص است، محصولات مشابه و مرتبط با آن به او پیشنهاد می‌شود. علاوه بر این بر اساس رفتار کاربرانی که سلیقه و علاقه‌های مشابهی دارند، به کاربران پیشنهاد می‌شود تا با دیدن و استفاده از تجربه‌های آنها، بهترین تجربه را به دست آورند. همچنین ممکن است بر اساس زمان و فعالیت کاربران، توصیه‌هایی در زمان‌های مناسب و بهینه به آنها ارائه شود. با استفاده از این سیستم‌ها، کاربران قادرند به راحتی به اطلاعات مورد نیاز خود دسترسی پیدا کنند و از تجربه کاربری بهتری برخوردار شوند. همچنین وبسایت‌ها نیز با ارائه توصیه‌های دقیق می‌توانند بازدیدکنندگان بیشتری جذب کنند و سطح رضایت آنها را افزایش دهند. از این رو در این تحقیق یک سیستم توصیه‌گر ترکیبی بر اساس ترکیب خوشه‌بندی ترتیبی فازی و

گردیده است. آنچه در شکل ۲۰ قابل مشاهده است این است که روش پیشنهادی نرخ برخورد بالاتری نسبت به روش‌های پیشین بر روی مجموعه داده‌های یکسان به دست آورده است.

با توجه به استفاده از یادگیری عمیق و سایر روش‌های طبقه‌بندی در زمینه تشخیص کاربران مشابه بر اساس اطلاعات پروفایل کاربران، روش پیشنهادی را می‌توان از لحاظ دقت تشخیص کاربران مشابه با روش‌های پیشین مقایسه کرد. در روش پیشنهادی علاوه بر شبکه عصبی خودرمزگذار، از سایر روش‌های طبقه‌بندی مانند k نزدیک‌ترین همسایه (KNN)، بیزین ساده (NB) و ماشین پشتیبان بردار (SVM) نیز به منظور یافتن کاربران مشابه بر اساس اطلاعات پروفایل کاربری استفاده شده است. مقایسه روش پیشنهادی با روش‌های طبقه‌بندی دیگر و روش‌های پیشین در نشریات از نظر دقت یافتن کاربران مشابه در شکل ۲۱ نشان داده شده است. همان طور که در شکل آمده است، روش پیشنهادی با توجه به استفاده از یادگیری عمیق بر روی ویژگی‌های کاربران که از پروفایل آنها استخراج می‌شود، توانسته است دقت بالاتری در درستی تشخیص کاربران مشابه نسبت به این روش‌های طبقه‌بندی و روش‌های پیشین به دست آورد.

- [11] Y. Gulzar, A. A. Alwan, R. M. Abdullah, A. Z. Abualkashik, and M. Oumrani, "OCA: Ordered clustering-based algorithm for e-commerce recommendation system," *Sustainability*, vol. 15, no. 4, Article ID: 2947, Feb.-2 2023.
- [12] S. Lee and D. Kim, "Deep learning based recommender system using cross convolutional filters," *Information Sciences*, vol. 592, pp. 112-122, May 2022.
- [13] L. El Youbi El Idrissi, I. Akharraz, and A. Ahaitouf, "Personalized e-learning recommender System Based on Autoencoders" *Applied System Innovation*, vol. 6, no. 6, Article ID: 102, Dec. 2023.
- [14] P. Bellini, L. Alessandro Ipsaro Palesi, P. Nesi, and G. Pantaleo, "Multi clustering recommendation system for fashion retail," *Multimedia Tools and Applications*, vol. 82, pp. 9989-10016, 2023.
- [15] Z. Abbasi-Moud, H. Vahdat-Nejad, and J. Sadri, "Tourism recommendation system based on semantic clustering and sentiment analysis," *Expert Systems with Applications*, vol. 167, Article ID: 114324, Apr. 2021.
- [16] H. Khaterr, S. Arif, U. Singh, S. Mathur, and S. Jain, "Product recommendation system for e-commerce using collaborative filtering and textual clustering," in *Proc. 3rd. Int. Conf. on Inventive Research in Computing Applications*, pp. 612-618. Coimbatore, India, 2-4 Sept. 2021.
- [17] T. K. Dang, Q. P. Nguyen, and V. S. Nguyen, "A study of deep learning-based approaches for session-based recommendation systems," *SN Computer Science*, vol. 1, Article ID: 216, 2020.
- [18] R. J. Ziarani and R. Ravanmehr, "Deep neural network approach for a serendipity-oriented recommendation system," *Expert Systems with Applications*, vol. 185, Article ID: 115660, Dec. 2021.
- [19] Yelp Inc., *Yelp Open Dataset*, Accessed 2024. <https://business.yelp.com/data/resources/open-dataset/>
- [20] M. H. Alam, W. -J. Ryu, and S. Lee, *TripAdvisor Dataset*. 2016, Accessed 2024. <https://brightdata.com/products/datasets/tripadvisor>
R. Logesh, V. Subramaniaswamy, V. Vijayakumar, and X. Li, "Efficient user profiling based intelligent travel recommender system for individual and group of users," *Mobile Networks and Applications*, vol. 24, pp. 1018-1033, 2019.

مستوره معینی تحصیلات کارشناسی خود را در سال ۱۳۸۹ از دانشگاه آزاد اسلامی واحد ماهشهر در حوزه IT، مدرک کارشناسی ارشد خود را در سال ۱۳۹۵ از دانشگاه آزاد اسلامی واحد علوم تحقیقات، در رشته مهندسی کامپیوتر گرایش نرم افزار، و مدرک دکترای خود را در رشته مهندسی کامپیوتر گرایش نرم افزار در سال ۱۴۰۳ از دانشگاه آزاد اسلامی واحد تهران جنوب دریافت کرد. در این سال‌ها او چندین کتاب و مقاله در حوزه کامپیوتر منتشر کرده است. زمینه‌های تحقیقاتی مورد علاقه ایشان عبارتند از: هوش مصنوعی، داده کاوی و سیستم‌های توصیه‌گر.

علی برومندینیا متولد خطه اصفهان، ایران است. او مدرک کارشناسی خود را در سال ۱۳۷۱ از دانشگاه صنعتی اصفهان، مدرک کارشناسی ارشد خود را در سال ۱۳۷۴ از دانشگاه علم و صنعت ایران، هر دو در رشته مهندسی معماری کامپیوتر، و مدرک دکترای خود را در رشته هوش مصنوعی و مهندسی کامپیوتر از دانشگاه آزاد اسلامی واحد علوم و تحقیقات تهران در سال ۱۳۸۵ دریافت کرد. از سال ۱۳۷۱ تا ۱۳۷۴، او بر روی کنترل هوشمند حمل و نقل با پردازش تصویر کار کرد و سیستم تشخیص خودکار پلاک خودرو را برای شرکت کنترل ترافیک تهران طراحی نمود. ایشان بیش از صد کتاب، مجله و مقاله کنفرانسی در حوزه کامپیوتر منتشر کرده است. نام‌برده به پنهان‌سازی اطلاعات، امنیت چندرسانه‌ای، تشخیص و قطعه‌بندی کاراکترهای فارسی و عربی، قطعه‌بندی اسناد فارسی و عربی، تصویربرداری پزشکی، پردازش سیگنال و تصویر و تحلیل مویک علاقه‌مند است. ایشان داور برخی از مجلات و کنفرانس‌های بین‌المللی است و در حال حاضر دانشیار دانشکده هوش مصنوعی دانشگاه آزاد اسلامی واحد تهران جنوب است.

مونا مرادی تحصیلات خود را در مقاطع کارشناسی و کارشناسی ارشد کامپیوتر به‌ترتیب در سال‌های ۱۳۷۹ و ۱۳۸۳ در دانشگاه آزاد اسلامی به پایان رسانده است. مدرک دکترای ایشان نیز در همین حوزه می‌باشد. وی اکنون استادیار دانشکده مهندسی کامپیوتر دانشگاه آزاد اسلامی می‌باشد. زمینه‌های تحقیقاتی مورد علاقه ایشان عبارتند از: هوش مصنوعی، معماری کامپیوتر، طراحی مدارهای محاسباتی و نانو الکترونیک.

شبکه خودمرزگذار عمیق بر اساس اطلاعات پروفایل کاربری و رتبه‌بندی وبسایت‌ها توسط کاربران ارائه شده است. در این سیستم توصیه‌گر، ابتدا کاربران بر اساس شباهت نظرات خود به صورت ترتیبی خوشه‌بندی می‌شوند. سپس رتبه‌بندی جدید برای کاربران با توجه به تابع عضویت فازی پیش‌بینی می‌شود. در نهایت اطلاعات موجود در پروفایل کاربری و رتبه‌بندی جدید کاربران به هر وبسایتی به‌منظور پیش‌بینی رتبه‌بندی وبسایت‌ها توسط کاربران، به عنوان ورودی شبکه خودمرزگذار عمیق ارائه‌شده مورد استفاده قرار می‌گیرد. شبکه خودمرزگذار عمیق به پیش‌بینی رتبه‌بندی کاربران جدید نسبت به وبسایت‌ها اقدام نموده و از این رتبه‌بندی برای یافتن کاربران مشابه بر اساس روش فیلتر مشارکتی اقدام می‌کند. در نهایت پس از یافتن کاربران مشابه، اقدام به ارائه توصیه بازدید و شخصی‌سازی صفحه وب کاربران جدید بر اساس وبسایت‌های مورد علاقه کاربران مشابه می‌نماید. روش ارائه‌شده با استفاده از شبکه‌های عصبی عمیق و با توجه به یادگیری عمیق بر روی اطلاعات موجود در پروفایل کاربری در مجموعه داده TripAdvisor و Yelp، توانسته است با بالاترین دقت رتبه‌بندی محصولات برای کاربران جدید را پیش‌بینی نماید. روش ارائه‌شده با توجه به لایه‌های آموزش عمیق و تکمیل فرایند آموزش در لایه میانی توانسته است از سایر روش‌های طبقه‌بندی از نظر معیارهای ارزیابی بهتر عمل کند. نتایج شبیه‌سازی نشان می‌دهد روش ارائه‌شده از نقطه‌نظر دقت آماری و نسبت توصیه‌های موفق به توصیه‌های مفید، عملکرد مناسبی دارد.

مراجع

- [1] H. Ko, S. Lee, Y. Park, and A. Choi, "A survey of recommendation systems: recommendation models, techniques, and application fields," *Electronics*, vol. 11, no. 1, Article ID: 141, Jan.-1 2022.
- [2] D. Roy and M. Dutta, "A systematic review and research perspective on recommender systems," *Journal of Big Data*, vol. 9, Article ID: 59, 2022.
- [3] S. Ephina. Thendral and C. Valliyammai, "Understanding personalization of recommender system: A domain perspective," *International Journal of Applied Engineering Research*, vol. 13, no. 15, pp. 2422-12428, 2018.
- [4] A. Papagrigoriou, C. Panagiotakis, E. Kosmas, and P. Fragopoulou, "Collaborative filtering recommender systems taxonomy," *Knowledge and Information Systems*, vol. 64, pp. 35-74, 2022.
- [5] S. Eliyas and P. Ranjana, "Recommendation systems: Content-based filtering vs collaborative filtering" in *Proc. 22nd. Int. Conf. on Advance Computing and Innovative Technologies in Engineering*, pp. 1360-1365, Greater Noida, India, 28-29 Apr. 2022.
- [6] M. Yu, T. Quan, Q. Peng, X. Yu, and L. Liu, "A model-based collaborate filtering algorithm based on stacked AutoEncoder," *Neural Computing and Applications*, vol. 34, no. 4, pp. 2503-2511, Feb. 2022.
- [7] F. O. Isinkaye, Y. O. Folajimi, and B. A. Ojokoh, "Recommendation systems: Principles, methods and evaluation," *Egyptian Informatics Journal*, vol. 16, no. 3, pp. 261-273, Nov. 2015.
- [8] G. Behera and N. Nain, "Trade-off between memory and model-based collaborative filtering recommender system," in *Proc. Int. Conf. on Paradigms of Communication, Computing and Data Sciences*, pp. 137-146, Kurukshetra, India, 7-9 May 2022.
- [9] A. Darvishy, H. Ibrahim, F. Sidi, and A. Mustapha, "HYPNER: A hybrid approach for personalized news recommendation," *IEEE Access*, vol. 8, pp. 46877-46894, 2020.
- [10] Q. Shambour, "A deep learning based algorithm for multi-criteria recommender systems," *Knowledge-Based Systems*, vol. 211, Article ID: 106545, Jan. 2021.