

استفاده از ارتباطات بین برچسبها در ایجاد زنجیره ردهبندها برای بهبود ردهبندی چندبرچسبی

خلیل غفوری پور^۱، دانشجوی کارشناسی ارشد؛ زهرا میرزامومن^۲، استادیار

۱- دانشکده مهندسی کامپیوتر - دانشگاه تربیت دبیر شهیدرجایی - تهران - ایران - ghafoori_khalil@yahoo.com

۲- دانشکده مهندسی کامپیوتر - دانشگاه تربیت دبیر شهیدرجایی - تهران - ایران - mirzamomen@sru.ac.ir

چکیده: در این مقاله، با فرض وجود ارتباطات معنادار بین برچسبها در مسائل چندبرچسبی، دو روش جدید برای بهبود روش پایه زنجیره ردهبندها (CC) در ردهبندی چندبرچسبی پیشنهاد شده است. در این مقاله، برای اولین بار از قوانین انجمنی برای تعیین ترتیب ردهبندها در روش CC استفاده شده است. در روشهای پیشنهادی این مقاله، ابتدا از قوانین انجمنی برای مدل سازی ارتباطات بین برچسبها استفاده می شود و سپس با استفاده از قوانین انجمنی استخراج شده، یک گراف ارتباط ساخته می شود و در نهایت، این گراف مبنای تعیین ترتیب زنجیره ردهبندها قرار می گیرد. از آنجا که در مسائل واقعی چندبرچسبی مانند ردهبندی متون، تصاویر و کاربردهای پزشکی ارتباطات معناداری بین برچسبها وجود دارد، روشهای پیشنهادی به بهبود ردهبندی در این حوزهها منجر خواهد شد. آزمایشهای تجربی گسترده انجام شده روی مجموعه دادههای استاندارد و رایج در حوزه ردهبندی چندبرچسبی نشان می دهند، استفاده از ارتباطات بین برچسبها در ساخت زنجیره ردهبندها باعث بهبود معیارهای مهم ارزیابی در ردهبندی مبتنی بر زنجیره ردهبندها می شود.

واژههای کلیدی: ردهبندی چندبرچسبی، ارتباط بین برچسبها، قوانین انجمنی، زنجیره ردهبندها.

To Use the Relationships between Class Labels in Creating Classifier Chains to Improve Multi-label Classification

K. Ghafouri pour, MSc Student¹; Z. Mirzamomen, Assistant Professor²

1- School of Computer Engineering, Shahid Rajaee Teacher Training University, Tehran, Iran, Email: ghafoori_khalil@yahoo.com

2- Faculty of Computer Engineering, Shahid Rajaee Teacher Training University, Tehran, Iran, Email: mirzamomen@sru.ac.ir

Abstract: In this paper, we have supposed that there is meaningful relationships between the class labels in the multi-label classification problems and based on it, we have proposed two novel methods to improve the base classifier chains (CC) method for multilabel classification. In this paper, association rules are employed to determine the order of classifiers in the CC method for the first time. In the proposed methods, association rule mining is first employed to model the relationships between the class labels and then, an association graph is built based on the extracted rules and finally, the classifier chains is built based on the obtained graph. As there is meaningful relationships between the class labels in the real multi-label problems such as classifying the images and texts and medical applications, the proposed methods will improve the classification results in such contexts. Extensive experimental evaluations conducted on the benchmark datasets in the multi-label classification context show that to use the associations between the labels in constructing the classifier chains improves the results obtained by the main evaluation measures.

Keywords: Multi-label classification, relationship between labels, association rules, classifier chains.

تاریخ ارسال مقاله: ۱۳۹۵/۰۹/۱۴

تاریخ اصلاح مقاله: ۱۳۹۵/۱۲/۱۸ و ۱۳۹۵/۱۲/۲۳ و ۱۳۹۵/۱۱/۱۲

تاریخ پذیرش مقاله: ۱۳۹۶/۰۶/۲۹

نام نویسنده مسئول: زهرا میرزامومن

نشانی نویسنده مسئول: ایران - تهران - دانشگاه تربیت دبیر شهید رجایی - دانشکده مهندسی کامپیوتر.

۱- مقدمه

رده‌بندی یکی از وظایف مهم یادگیری ماشین است که کاربردهای فراوانی دارد [۱، ۲]. رده‌بندی چندبرچسبی، نوعی مسئله رده‌بندی است که در آن هر نمونه می‌تواند هم‌زمان متعلق به بیش از یک کلاس باشد. این درحالی‌است که در رده‌بندی چندکلاسه، هر نمونه تنها متعلق به یکی از کلاس‌های مسئله است. بسیاری از مسائل جهان امروز به‌گونه‌ای هستند که در قالب رده‌بندی تک‌برچسبی قرار نمی‌گیرند، زیرا نمونه‌ها بیش از یک برچسب دارند. رده‌بندی کردن متون، یادداشت‌گذاری روی تصاویر و ویدئوها و تشخیص بیماری‌ها همه از کاربردهای رده‌بندی چندبرچسبی است [۳]. به‌عنوان کاربردی دیگر از این مسئله در حوزه وب‌کاوی، می‌توان به رده‌بندی موضوعی صفحات وب اشاره کرد.

در مسائل واقعی چندبرچسبی، ارتباط و همبستگی بین برچسب‌ها معمولاً به‌نحو چشمگیری وجود دارد. برای مثال در حوزه پزشکی، بیماری‌های بسیاری دارای نشانه‌ها و عوارض مشترک هستند و بنابراین برای تشخیص بیماری یک فرد بر اساس عوارض و مشاهدات انجام‌شده، چندین برچسب قابل‌تصور است. همچنین برای رده‌بندی موضوعی متون و رسانه‌های تصویری، می‌توان از جنبه‌ها و دیدگاه‌های مختلف، برچسب‌های مرتبط ولی متفاوت از هم را به یک متن، تصویر یا ویدیو نسبت داد.

بنابراین چنان‌چه ارتباطات بین برچسب‌ها به طرز مناسبی مدل شود، می‌تواند نقش بسزایی در بهبود رده‌بندی چندبرچسبی در مراحل یادگیری و رده‌بندی ایفا کند.

روش‌های ارائه‌شده در حوزه رده‌بندی چندبرچسبی را می‌توان به دو دسته کلی تقسیم کرد: روش‌هایی که وجود هرگونه ارتباط بین برچسب‌ها را نادیده می‌گیرند و روش‌هایی که به نحوی سعی می‌کنند از این ارتباطات استفاده کنند. ازسوی دیگر، دو رویکرد عمده انتقال مسئله و الگوریتم‌های تطبیقی در روش‌های ارائه‌شده در این حوزه قابل تشخیص است [۴، ۳]. الگوریتم‌های مبتنی بر رویکرد انتقال مسئله را می‌توان به سه زیردسته انتقال به مسئله رده‌بندی دودویی، انتقال به مسئله رتبه‌بندی برچسب‌ها و انتقال به مسئله رده‌بندی چندکلاسه تقسیم کرد [۴].

در رویکرد انتقال مسئله، مسئله L برچسبی به L مسئله تک‌برچسبی تبدیل می‌شود و هر یک از این مسائل توسط رده‌بندی مجزا حل می‌شود. درحالی‌که در رویکرد الگوریتم‌های تطبیقی، الگوریتم‌های رده‌بندی تک‌برچسبی به حوزه چندبرچسبی گسترش یافته و توسعه می‌یابند [۴].

از جمله اولین روش‌های ارائه‌شده مبتنی بر رویکرد اول، روش BR^1 است که در آن به ازای هر برچسب در فضای برچسب‌ها، یک رده‌بند در نظر گرفته می‌شود، به‌نحوی که هر رده‌بند مسئول تشخیص نمونه‌های یک برچسب از بین سایر نمونه‌ها است. از آن‌جاکه در این روش ارتباط بین برچسب‌ها در نظر گرفته نمی‌شود، روش دیگری به نام زنجیره

رده‌بندها یا CC^2 [۳] پیشنهاد شده است. در روش CC ، ارتباط بین برچسب‌ها بدین‌طریق مورد استفاده قرار می‌گیرد که یک زنجیره (ترتیب) از رده‌بندها در نظر گرفته می‌شود و بر اساس آن (هم در مرحله آموزش و هم در مرحله رده‌بندی) خروجی رده‌بند اول به فضای ویژگی رده‌بند دوم اضافه می‌گردد. این روند برای سایر رده‌بندها نیز ادامه می‌یابد. این روش از جمله روش‌های موفق در رده‌بندی چندبرچسبی است، ولی از چالش‌های مهم آن، یافتن ترتیب مناسب برای زنجیره رده‌بندها است که نقش عمده‌ای در عملکرد این روش رده‌بندی دارد.

در کارهای انجام‌شده گذشته روش‌هایی برای پاسخ به چالش تعیین ترتیب مناسب برای زنجیره رده‌بندها پیشنهاد شده است که در بخش ۲ به‌اختصار معرفی شده‌اند. همچنین، در بخش بزرگی از کارهای انجام‌شده در این حوزه که تحت عنوان رده‌بندی چندبرچسبی سلسله مراتبی^۳ شناخته شده است، فرض براین است که سلسله مراتب ارتباطات بین برچسب‌ها در مسئله به‌صورت درخت یا گراف از ابتدا موجود است و می‌تواند مبنای تعیین ترتیب مناسب برای زنجیره رده‌بندها قرار گیرد [۵].

در این مقاله، دو روش جدید برای تعیین ترتیب مناسب برای زنجیره رده‌بندها برای مسائلی که ارتباطات بین برچسب‌ها از ابتدا مشخص نیست، پیشنهاد شده است. بااینکه در کارهای گذشته، از قوانین انجمنی در رده‌بندی چندبرچسبی به شکل‌های مختلف استفاده شده است (رجوع کنید به پاراگراف آخر بخش ۲) ولی تابحال ایده‌ای برای استفاده از قوانین انجمنی برای تعیین ترتیب مناسب برای زنجیره رده‌بندها مطرح نشده‌است. در روش‌های پیشنهادی، ابتدا به‌منظور مدل‌سازی ارتباطات بین برچسب‌ها از کاوش قوانین انجمنی [۶، ۷] استفاده شده است و سپس بر مبنای قوانین انجمنی استخراج‌شده، گرافی که نشان‌دهنده ارتباط بین برچسب‌ها است، ساخته شده است. درنهایت، این گراف مبنای تعیین ترتیب برچسب‌ها در زنجیره رده‌بندها قرار گرفته است.

به‌طور خلاصه، نوآوری ما در این مقاله، پیشنهاد روش‌هایی برای تعیین ترتیب زنجیره رده‌بندها مبتنی بر قوانین انجمنی استخراج‌شده از برچسب‌ها است. به‌عبارت دیگر، در این مقاله ترتیب زنجیره رده‌بندها در روش CC از حالت تصادفی خارج گردیده و ارتباط بین برچسب‌ها با استفاده از قوانین انجمنی مثبت، منفی و ترکیبی موجود بین برچسب‌ها، به‌صورت یک گراف مدل شده است. درواقع بین قوانین استخراج‌شده و گراف ارتباطات برچسب‌ها نداشت به‌وجود آمده‌است.

در این مقاله، برای نگاشت قوانین انجمنی بین برچسب‌ها به گراف ارتباط، دو روش پیشنهاد داده‌شده و هر دو مورد ارزیابی قرار گرفته است و درنهایت، با استفاده از گراف ارتباط، ترتیب زنجیره رده‌بندها تعیین شده است. از آن‌جا که در این مقاله از گراف (به‌جای درخت) برای مدل‌سازی ارتباط بین برچسب‌ها استفاده شده است، زنجیره رده‌بندها به‌شکل یک توالی پی‌درپی از تک‌برچسب‌ها نیست و بنابراین

استفاده می‌شود که در [۱۴] PS^V نامیده شده است و در آن، تنها مجموعه برچسب‌های با طول حداقل c در نظر گرفته می‌شوند.

در [۱۵] با گسترش ایده اصلی روش نزدیک‌ترین همسایگی KNN، الگوریتم تطبیقی $MLkNN^A$ برای حل مسائل چندبرچسبی ارائه شده است که در آن، برچسب‌گذاری هر نمونه جدید بر اساس احتمالات تعلق به کلاس‌های مختلف، مستخرج از k نزدیک‌ترین همسایه آن، انجام می‌شود. درخت تصمیم از جمله روش‌های مهم در رده‌بندی تک‌برچسبی است که در [۱۶] برای چندبرچسبی توسعه داده شده است و بدین‌منظور، روابط محاسبه بهره اطلاعاتی در گره‌های داخلی درخت و رده‌بندی نمونه آزمون بر اساس یک حد آستانه در برگ‌ها تغییر یافته و گسترش یافته‌اند.

در [۱۷] روشی برای رده‌بندی چندبرچسبی بر اساس یادگیری عمیق و ماشین بولتزمن ارائه شده است که در آن رویکرد کاهش ابعاد مسئله و پیش‌بینی‌کننده برچسب‌ها در لایه آخر در نظر گرفته شده است. روش پیشنهادی در [۱۸] به‌ازای هر برچسب یک پرسپترون را آموزش می‌دهد و برای این‌که ارتباط بین برچسب‌ها در نظر گرفته شود، پرسپترون‌ها با یکدیگر آموزش داده می‌شوند. [۱۹] به‌صورت تصادفی k برچسب را با جایگزینی یا بدون جایگزینی انتخاب کرده و با ترکیب‌های مختلف از آن‌ها، به روش LP به حل مسئله می‌پردازد.

در تحقیقات گذشته روش‌هایی برای غیر تصادفی کردن ترتیب رده‌بندها در زنجیره CC پیشنهاد شده است. در روش $BCC^{[9,10]}$ زنجیره رده‌بندها مبتنی بر شبکه‌های بیزین شکل می‌گیرند. BCC ارتباط بین برچسب‌ها را به‌صورت یک گراف در نظر می‌گیرد که ساده‌ترین آن $TNBCC^{[10]}$ یا درخت ارتباط است. در [۱۱] روش $CT^{[1]}$ (روش داربستی) به‌عنوان روشی ابتکاری برای ساختن ترتیب رده‌بندها پیشنهاد شده است. در این روش از یک برچسب تصادفی شروع کرده و با توجه به اطلاعات متقابل آن با سایر برچسب‌ها، سمت چپ و پایین ساختار داربستی را می‌سازد. این روند تا پایان یافتن برچسب‌ها ادامه می‌یابد. در این ساختار هر برچسب با برچسب پایین و چپ خود در ارتباط است. در نسخه دیگری از این روش هر برچسب با سه برچسب در ارتباط است. برچسب سوم در سر دیگر قطر داربست قرار دارد. در [۱۲] با استفاده از روش BR دقت تک‌برچسبی‌ها سنجیده شده و به ترتیب نزولی دقت، زنجیره ساخته می‌شود. [۲۰] برای ساخت زنجیره رده‌بندها از شبکه‌های بیزین به‌صورت مؤثرتری استفاده کرده تا در برخورد با انتساب برچسب ناقص^{۱۲} مناسب باشد.

در [۲۱] روشی برای بهبود رده‌بندهای چندبرچسبی با استفاده از قوانین انجمنی از طریق کاهش تعداد برچسب‌ها ارائه شده است. در [۲۲] روشی برای اصلاح برچسب‌های پیش‌بینی‌نشده توسط رده‌بند بر اساس قوانین انجمنی استخراج شده از برچسب‌ها ارائه شده است و با توجه به این‌که تعداد قوانین استخراج شده در این روش زیاد بوده است، از روشی ابتکاری بر اساس یک حد آستانه و مقایسه قوانین برای جداسازی قوانین انجمنی باکیفیت استفاده شده است. در [۲۳] از

در هر مرحله، فضای ویژگی یک رده‌بند ممکن است به‌وسیله خروجی چند رده‌بند تکمیل گردد.

در ادامه مقاله به‌صورت زیر سازماندهی شده است: در بخش ۲ کارهای مرتبط بررسی شده است. در بخش ۳ روش پیشنهادی به‌صورت مفصل معرفی شده است. در بخش ۴، پس از ذکر شرایط آزمایش، مجموعه داده‌ها و معیارهای ارزیابی مورد استفاده، نتایج ارزیابی‌های انجام شده گزارش و تحلیل شده است. در پایان، بخش ۵ به نتیجه‌گیری و پیشنهادهایی برای تحقیقات آینده اختصاص داده شده است.

۲- کارهای مرتبط

در این بخش پژوهش‌های پیشین و روش‌های مطرح ارائه شده برای حل مسائل چندبرچسبی را مرور کرده‌ایم.

همان‌طور که در مقدمه آمده است، در روش مبتنی بر انتقال مسئله BR، مسئله رده‌بندی L برچسبی به L مسئله رده‌بندی تک‌برچسبی تبدیل می‌شود و هر یک از مسائل بوجودآمده مستقل از دیگری حل می‌شود [۳] بنابراین از ارتباطات موجود بین برچسب‌ها چشم‌پوشی می‌شود. در این روش ممکن است برای یک نمونه جدید، مجموعه تهی پیش‌بینی شود. در روش CC، با فرض وجود یک ارتباط ثابت بین برچسب‌ها، یک زنجیره از رده‌بندها به‌صورت تصادفی ایجاد می‌شود [۳] و با توجه به نحوه استفاده از این زنجیره، هر برچسب بر روی برچسب‌های بعدی خود در زنجیره تأثیر می‌گذارد. این امر باعث می‌شود که در فاز رده‌بندی، پیش‌بینی‌های درست رده‌بندهای ابتدایی زنجیره، به انجام پیش‌بینی‌های درست در رده‌بندهای بعدی کمک کند و از سوی دیگر، اشتباه رده‌بندهای اول در کل زنجیره گسترش یابد. بنابراین در این روش، ترتیب رده‌بندها از اهمیت ویژه‌ای برخوردار است. در $ECC^{[8]}$ چندین رده‌بند CC در یک رده‌بند جمعی قرار داده می‌شوند، به‌نحوی که ترتیب زنجیره رده‌بندها در هر یک از آن‌ها با دیگری متفاوت است.

در روش جفت‌جفت (PW^5) به ازای هر دو برچسب، یک رده‌بند ایجاد می‌شود تا یکی از این دو برچسب را برای هر نمونه پیش‌بینی کند. بنابراین در این روش از پیش‌بینی شدن مجموعه برچسب تهی برای نمونه‌های آزمون اجتناب می‌شود. با این حال، زیاده‌شدن تعداد رده‌بندها از چالش‌های این روش محسوب می‌شود [۴].

در [۱۳] روش مجموعه کامل برچسب (LP^6) پیشنهاد شده است که در آن، مجموعه برچسب‌های هر نمونه به‌عنوان یک برچسب مجزا در یک مسئله جدید در نظر گرفته می‌شود و مسئله به‌صورت تک‌برچسبی درمی‌آید. بنابراین، رده‌بند حاصل از این روش، قادر به پیش‌بینی ترکیب‌هایی از برچسب‌ها که در مرحله آموزش مشاهده نشده‌اند، نیست. از سوی دیگر، با توجه به این‌که ممکن است تعداد مجموعه‌های مجزای برچسب‌ها زیاد شود، برای اجتناب از پیچیدگی‌های محاسباتی مدل LP، از مجموعه نمونه‌های هرس شده

متعددی ارائه شده که نشان از اهمیت و توجه به این روش است [۱۲-۱۸].

در این پژوهش، دو روش برای ساخت ترتیب زنجیره بر اساس قوانین انجمنی پیشنهاد شده است که ما آن‌ها را زنجیره رده‌بندهای قوانین انجمنی یا $ARCC^1$ شماره ۱ و ۲ نامیده‌ایم. قابل ذکر است که در روش CC، ترتیب زنجیره به صورت تصادفی ایجاد می‌شود و روش‌های پیشنهادی این ترتیب را از حالت تصادفی خارج می‌کنند. روش CC زنجیره رده‌بندها را به شکل درخت و روش‌های پیشنهادی ما زنجیره رده‌بندها را به شکل گراف می‌سازند.

در بخش‌های ۳-۱ و ۳-۲ روش‌های پیشنهادی ارائه شده‌اند. در روش اول ($ARCC^1$)، ترتیب صعودی تعداد تکرار هر برچسب در سمت تالی قوانین انجمنی به عنوان ملاک تعیین برچسب‌های آغازکننده گراف ارتباط پیشنهاد شده است و در روش دوم ($ARCC^2$)، ترتیب نزولی صحت تک برچسب‌ها بدین منظور پیشنهاد شده است. سپس در هر دو روش، از قوانین انجمنی استخراج شده از ارتباطات بین برچسب‌ها برای تکمیل گراف ارتباط استفاده می‌شود.

۳-۱- روش اول: $ARCC^1$

در این روش، ابتدا برچسب‌ها بر اساس فراوانی‌شان در سمت تالی همه قوانین انجمنی استخراج شده، مرتب می‌شوند. هر چه فراوانی یک برچسب در سمت تالی قوانین کمتر باشد، به این معنی است که این برچسب، اثرپذیری کمتری از دیگر برچسب‌ها دارد. اگر این فراوانی برای یک برچسب صفر باشد، به این معنی است که با توجه به شرایط استخراج قوانین انجمنی، از هیچ برچسبی نمی‌توان وجود/عدم وجود آن را برای یک نمونه نتیجه گرفت. بنابراین انتظار می‌رود این گونه برچسب‌ها، برای شروع ترتیب زنجیره مناسب باشند.

به این ترتیب، در روش پیشنهادی $ARCC^1$ اولین لایه زنجیره با برچسب‌های دارای کمترین فراوانی در سمت تالی قوانین ساخته می‌شود. سپس زنجیره بر اساس اینکه برچسب‌های لایه اول روی کدام برچسب‌های دیگر اثر می‌گذارند (از طریق ظاهر شدن در سمت مقدم یک قانون) تکمیل می‌شود. برای ساخت ادامه زنجیره، همین روال تکرار می‌شود. الگوریتم ۱ نحوه کار روش پیشنهادی را نشان می‌دهد.

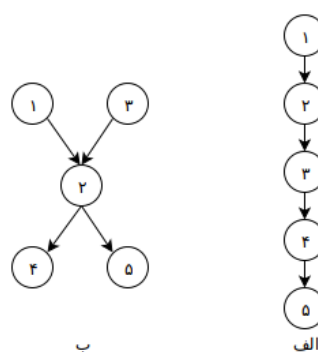
۳-۱-۱- تشریح الگوریتم

ورودی الگوریتم ۱، یک مجموعه داده چندبرچسبی است و خروجی آن، یک آرایه به اندازه تعداد برچسب‌ها است که مقدار عنصر i ام آن، مجموعه‌ای است که اعضای آن، گره‌های پدر برچسب i در گراف ارتباط هستند و حداکثر اندازه آن، یکی کمتر از تعداد برچسب‌های مسئله است (چون یک گره نمی‌تواند پدر خودش در گراف ارتباط باشد). درواقع در الگوریتم پیشنهادی، با استفاده از یک آرایه از مجموعه‌ها، گراف زنجیره رده‌بند نمایش داده شده است. برای مثال، آرایه متناظر با گراف شکل ۱-ب به شکل زیر خواهد بود:

قوانین انجمنی برای استخراج قانون تصمیم (Decision rule) استفاده شده است که این قوانین مستقیماً برای تعیین برچسب نمونه‌های جدید استفاده شده‌اند. در [۲۴] روشی برای تبدیل یک مجموعه داده چندبرچسبی به یک مجموعه داده تک‌برچسبی با استفاده از قوانین انجمنی پیشنهاد شده است و سپس دو الگوریتم کلاس‌بندی فازی مبتنی بر نزدیک‌ترین همسایگی برای اعمال روی مجموعه داده حاصل ارائه شده است. در [۲۵] روشی پس‌پردازشی مبتنی بر ترکیب استفاده از قوانین انجمنی و آنتولوژی کلاس‌ها برای اصلاح پیش‌بینی‌های رده‌بندها پیشنهاد شده است و ادعا شده است که این روش، برای اصلاح برچسب‌های مربوط به کلاس‌های کوچک مجموعه داده نیز مؤثر خواهد بود. همچنین، ما قبلاً در [۲۶] روشی پس‌پردازشی برای اصلاح برچسب‌های پیش‌بینی نشده و اشتباه پیش‌بینی شده با استفاده از قوانین انجمنی ترکیبی پیشنهاد کرده‌ایم.

۳- روش پیشنهادی

همان‌طور که در بخش‌های قبل بیان شد، چالش عمده در روش CC، تعیین ترتیب مناسب برچسب‌ها در زنجیره رده‌بندها است. به طور کلی روش‌های ارائه شده برای ساخت ترتیب زنجیره، در دو دسته کلی قرار می‌گیرند [۸]. در دسته اول که ساده‌تر است، ترتیب به شکل یک درخت نگاشت می‌شود. یعنی در هر مرحله، هر برچسب در هر جای این زنجیره که قرار داشته باشد، نتیجه را فقط از یک برچسب پدر دریافت و به فضای ویژگی خود اضافه می‌کند. در دسته دوم، این نگاشت به صورت گراف است و بنابراین، از حالت درخت عمومی‌تر است. یعنی در هر مرحله بسته به شکل گراف، ممکن است نتیجه یک یا چند رده‌بند، به فضای ویژگی برچسب بعدی اضافه شود. شکل ۱-الف یک مثال از درخت زنجیره و ۱-ب یک مثال از گراف زنجیره نشان می‌دهد.



شکل ۱: درخت و گراف ترتیب زنجیره

در حالت الف، فضای ویژگی برای رده‌بند برچسب ۲، از فضای ویژگی اولیه مسئله به اضافه خروجی رده‌بند برچسب ۱ تشکیل شده است، ولی در حالت ب، فضای ویژگی برای رده‌بند برچسب ۲ از فضای ویژگی اولیه مسئله به اضافه خروجی رده‌بندهای برچسب‌های ۱ و ۳ تشکیل شده است. برای ساخت ترتیب زنجیره رده‌بندها تاکنون روش‌های

For each rule $r \in R$ (such that $r: A \rightarrow B$)	:۱۰
if $label \in A$ then	:۱۱
For each label $k \in B$ do	:۱۲
$ARChain_out[k] = ARChain_out[k] \cup label$	:۱۳
End For	:۱۴
End if	:۱۵
End For	:۱۶

۳-۱-۲- تبیین روش با یک مثال

به‌منظور روشن‌شدن نحوه عملکرد روش پیشنهادی، در این بخش با ذکر یک مثال مراحل کار تشریح می‌شود. فرض کنید در یک مسئله تعداد برچسب‌ها شش باشد (یعنی $q=6$). هم‌چنین در نظر بگیرید پس از استخراج و ثبت قوانین و شمارش برچسب‌های سمت تالی قوانین، آرایه شمارش به شکل زیر باشد:

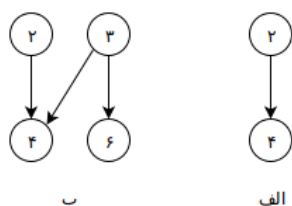
$$labelFrequency_right[q] = [12, 5, 10, 19, 28, 22]$$

بدین معنی که برچسب اول ۱۲ مرتبه در سمت تالی قوانین استخراج‌شده آمده‌است و برچسب‌های دوم تا ششم به ترتیب ۵، ۱۰، ۱۹، ۲۸ و ۲۲ مرتبه در سمت تالی قوانین استخراج‌شده آمده‌اند.

برای این مثال مطابق الگوریتم ۱، در تکرار اول حلقه آغازشده در خط ۸، برچسب دوم (در خط ۹) انتخاب می‌شود، زیرا در آرایه $labelFrequency_right$ کمترین تکرار (برابر با ۵) را در سمت تالی تمامی قوانین داشته‌است. در حلقه درونی در هنگام مرور همه قوانین، قانونی که برچسب دوم در سمت مقدم آن است ($I_p \rightarrow I_m$) انتخاب و بر اساس آن، برچسب دوم به‌عنوان والد برچسب چهارم در نظر گرفته‌شده و آرایه $ARChain_out$ به شکل رابطه (۱) اصلاح می‌گردد که معادل گراف نشان‌داده‌شده در نمودار شکل ۲-الف است.

$$ARChain_out[4] = \{2\} \quad (1)$$

در تکرار دوم حلقه (خط ۸)، برچسب سوم با دومین کمترین تعداد تکرار (۱۰ مرتبه) در سمت تالی قوانین انتخاب می‌شود. در حلقه داخلی، چون قانون ($I_p \rightarrow I_m$) وجود دارد که برچسب سوم در مقدم آن است و در سمت تالی آن برچسب‌های چهارم و ششم آمده‌اند، نتیجه شکل ۲-ب به‌دست می‌آید. این روند تا تعیین تکلیف تمامی برچسب‌ها ادامه می‌یابد. باید در نظر داشت ممکن است در حلقه درونی، تعداد قوانینی که برچسب موردنظر در سمت مقدم آن‌ها آمده، زیاد باشد که تمامی این قوانین در ساخت گراف نقش خواهند داشت.



شکل ۲: گراف ارتباط از روی قوانین انجمنی

$$ARChain_out = [\{\}, \{1, 3\}, \{\}, \{2\}, \{2\}]$$

که در آن، عناصر اول و سوم آرایه، آرایه‌هایی تهی هستند، به این معنی که گره‌های متناظر با برچسب‌های اول و سوم در گراف، بدون والد هستند و هیچ یال ورودی ندارند. عنصر دوم آرایه، برابر مجموعه $\{1, 3\}$ است که بدین معنی است که گره متناظر با برچسب دوم در گراف خروجی، فرزند گره متناظر با برچسب اول و هم‌ین‌طور، فرزند گره متناظر با برچسب سوم است و بنابراین، در زنجیره رده‌بندها، خروجی رده‌بند متناظر با برچسب‌های ۱ و ۳ به رده‌بند برچسب ۲ داده خواهد شد. عناصر چهارم و پنجم این آرایه، حاوی برچسب ۲ هستند و بنابراین، خروجی رده‌بند برچسب ۲ به‌صورت مشترک به هر دو رده‌بند متناظر با برچسب‌های ۴ و ۵ داده می‌شود. به این ترتیب، گراف ارتباط به‌صورت آرایه‌ای از مجموعه‌ها نمایش داده شده است.

در الگوریتم پیشنهادی، مجموعه قوانین انجمنی استخراج‌شده از روابط بین برچسب‌های مسئله در R ذخیره می‌گردد و سپس، فراوانی هر برچسب در سمت تالی همه قوانین شمرده شده و در آرایه $labelFrequency_right$ ثبت می‌گردد. سپس در دو حلقه تودرتو، آرایه خروجی یا $ARChain_out$ که نشان‌دهنده گراف ارتباط است، ساخته می‌شود.

از بین تمام برچسب‌ها، هر باریک برچسب (L_i) انتخاب می‌شود که کمترین تکرار را در سمت تالی کلیه قوانین داشته‌است و پس از اتمام روند، این برچسب از مجموعه برچسب‌ها کنار گذاشته می‌شود. برچسب انتخابی (L_i) در سمت مقدم هر قانونی که باشد، به‌عنوان والد برچسب یا برچسب‌های سمت تالی آن قانون، شناخته می‌شود و در آرایه $ARChain_out$ ثبت می‌گردد. منطق کار این است که طبق آن قانون و به این دلیل که برچسب انتخابی در سمت مقدم آن قرارداد، بر روی برچسب یا برچسب‌های سمت تالی اثرگذار بوده است و بنابراین باید به‌عنوان برچسب والد در گراف ارتباط در نظر گرفته‌شود. برای جلوگیری از تکرار و حلقه در انتخاب برچسب‌های والد، این برچسب در دور بعد از مجموعه برچسب‌ها کنار گذاشته می‌شود.

الگوریتم ۱ الگوریتم نگاشت قوانین انجمنی به گراف ارتباط برای رده‌بندی چندبرچسبی با شمارش برچسب‌ها در سمت تالی قوانین

۱: ورودی: مجموعه داده (D).

۲: خروجی: گراف ارتباط در قالب آرایه $ARChain_out[q]$: اندیس i ام این آرایه، حاوی یک مجموعه است که والدین گره متناظر با برچسب i را در گراف ارتباط نشان می‌دهد. در ابتدا تمام عناصر آرایه حاوی مجموعه تهی هستند.

معرفی متغیرها:

۳: R مجموعه قوانین انجمنی استخراج شده از برچسب‌های مسئله

۴: q تعداد کل برچسب‌ها

۵: $labelFrequency_right[q]$: اندیس i ام آرایه، تعداد تکرار برچسب i در سمت تالی قوانین R را نشان می‌دهد.

شروع:

۶: استخراج قوانین انجمنی از مجموعه اقلام حاوی برچسب‌های نمونه‌ها $\leftarrow R$

۷: شمارش فراوانی برچسب‌ها در سمت تالی قوانین عضو $\leftarrow labelFrequency_right$

۸: For $i=1$ to q do

۹: انتخاب i امین کم‌تکرارترین برچسب بر اساس $label \leftarrow labelFrequency_right$

```

۶: استخراج قوانین انجمنی از مجموعه اقلام حاوی برچسب‌های نمونه‌ها  $R \leftarrow$ 
۷: محاسبه صحت رده‌بندی تک‌برچسبی همه برچسب‌ها  $\leftarrow \text{Accuracy\_SingleLabel}$ 
۸: For  $i = 1$  to  $q$  do
۹: انتخاب  $i$  امین باصحت‌ترین برچسب بر اساس  $\text{Accuracy\_SingleLabel}$ 
۱۰: For each rule  $r \in R$  ( such that  $r: A \rightarrow B$  )
۱۱: if  $label \in A$  then
۱۲: For each label  $k \in B$  do
۱۳:  $\text{ARChain\_out}[k] = \text{ARChain\_out}[k] \cup label$ 
۱۴: End For
۱۵: End if
۱۶: End For
۱۷: End For

```

برای این که یک برچسب به صورت تکراری به عنوان والد قرار نگیرد، قبل از اصلاح آرایه ARChain_out کنترل‌های لازم انجام می‌شود. مورد دیگری که لازم است کنترل گردد، جلوگیری از تشکیل حلقه در گراف است. برای مثال اگر برچسب ۱ والد برچسب ۲ در آرایه ARChain_out قرار گرفت و برچسب ۲ والد برچسب ۳ شد، امکان والد شدن برچسب ۳ برای برچسب ۱ نباشد. همان‌طور که در این مثال مشخص است، روش پیشنهادی ARCC1 گراف ارتباط را با توجه به قوانین استخراج‌شده از روابط بین برچسب‌ها می‌سازد. این گراف مبنای تعیین ترتیب زنجیره رده‌بندی قرار خواهد گرفت.

۳-۲- روش دوم: ARCC2

در این روش، معیار دیگری برای انتخاب برچسب شروع کننده زنجیره پیشنهاد شده است. به این ترتیب که در ابتدا، صحت تک‌تک برچسب‌ها توسط یک رده‌بند پایه سنجیده می‌شود. صحت هر کدام از برچسب‌ها در یک آرایه ذخیره و به صورت نزولی مرتب می‌شود. در روش پیشنهادی، زنجیره از برچسبی شروع خواهد شد که بیشترین صحت در رده‌بندی تک‌برچسبی را داشته باشد. منطق این روش این است که اگر زنجیره از برچسب یا برچسب‌های دارای بالاترین صحت شروع شود، از آنجا که این برچسب‌ها بهتر از سایرین یاد گرفته می‌شوند، نتایج رده‌بند متناظر با آن‌ها، نتایج با صحت بالاتری خواهد بود که در اختیار دیگر رده‌بندها در زنجیره قرار خواهد گرفت و در نتیجه، انتشار خطای رده‌بندها در طول زنجیره کمتر خواهد بود و از طرف دیگر، با اضافه شدن نتایج این رده‌بند با صحت بالا به فضای ویژگی سایر رده‌بندها، به رده‌بندی بهتر توسط آن‌ها کمک خواهد کرد.

در ادامه، مشابه روش ARCC1 عمل می‌شود، یعنی قوانینی که برچسب انتخابی (با بیشترین صحت) در سمت مقدم آن‌هاست بررسی می‌شود تا برچسب انتخابی به عنوان والد برچسب‌های سمت تالی آن‌ها در گراف قرار گیرد. الگوریتم ۲ روش پیشنهادی ARCC2 را نشان می‌دهد که تشریح الگوریتم آن مشابه با الگوریتم ۱ است، با این تفاوت که به جای برچسب دارای کمترین تکرار در سمت تالی قوانین، در خط ۹ الگوریتم ۲، برچسبی انتخاب می‌شود که بیشترین صحت رده‌بندی را داشته باشد.

الگوریتم ۲ نگاشت قوانین انجمنی به گراف ارتباط برای رده‌بندی چندبرچسبی با صحت تک‌برچسبی‌ها

۱: ورودی: مجموعه داده (D).
 ۲: خروجی: گراف ارتباط در قالب آرایه $\text{ARChain_out}[q]$: اندیس i ام این آرایه، حاوی یک مجموعه است که والدین گره متناظر با برچسب i را در گراف ارتباط نشان می‌دهد. در ابتدا تمام عناصر آرایه حاوی مجموعه تهی هستند.
 معرفت متغیرها:
 ۳: R مجموعه قوانین انجمنی استخراج‌شده از روابط بین برچسب‌ها
 ۴: q تعداد کل برچسب‌ها
 ۵: $\text{Accuracy_SingleLabel}[q]$: اندیس i ام آرایه حاوی صحت برچسب i است.
 شروع:

۳-۳- تحلیل مرتبه زمانی

اگر q تعداد برچسب‌ها، $|R|$ تعداد قوانین انجمنی استخراج‌شده از روابط بین برچسب‌ها و $|B|$ حداکثر تعداد برچسب‌ها در سمت تالی قوانین باشد، آن‌گاه تولید گراف ارتباط در روش ARCC1 از مرتبه زمانی $O(q|R||B|)$ است. در روش ARCC2 از آن‌جا که ابتدا باید دقت تک‌برچسب‌ها به روش BR محاسبه شود، مرتبه زمانی اجرا به صورت $O(q|R||B|) + O(BR)$ است. البته باید توجه داشت که مرتبه زمانی مربوط به استخراج قوانین انجمنی از روابط بین برچسب‌ها را باید به هردوی این روش‌ها اضافه کرد.

۴- آزمایش‌ها

در این بخش، به تشریح آزمایش‌های انجام‌شده، نحوه پیاده‌سازی روش پیشنهادی، شرایط آزمایش‌های انجام‌شده، مجموعه داده‌های مورد استفاده و معیارهای ارزیابی، پرداخته شده است.

۴-۱- مجموعه داده‌ها

برای انجام آزمایش‌ها، از ۶ مجموعه داده استفاده شده است. ویژگی‌های این مجموعه داده‌ها در جدول ۱ آمده است. این مجموعه داده‌ها از مجموعه داده‌های رایج و مورد استفاده در پژوهش‌ها هستند [۱۲-۸] و [۲۰-۱۴] و به صورت عمومی قابل دریافت هستند.

جدول ۱: مجموعه داده‌های مورد استفاده در آزمایش‌ها به همراه

مشخصات آن‌ها

مجموعه داده	حوزه	تعداد نمونه	تعداد برچسب	کاردینالیته	چگالی	برچسب مجزا
Yeast	بیولوژی	۲۴۱۷	۱۴	۴/۲۳۷	۰/۳۰۳	۱۹۸
Scene	تصویر	۲۴۰۷	۶	۱/۰۷۴	۰/۱۷۹	۱۵
Birds	صوتی	۶۴۵	۱۹	۱/۰۱۴	۰/۰۵۳	۱۳۳
Music	صوتی	۵۹۲	۶	۱/۸۶۹	۰/۳۱۱	۲۷
Flags	تصویر	۱۹۴	۷	۱/۲۵۲	۰/۴۸۵	۵۴
Emotion	صوتی	۵۹۳	۷	۱/۸۶۹	۰/۳۱۱	۲۷

۴-۲- پیاده‌سازی

برای استخراج قوانین انجمنی مثبت و منفی از چارچوب Weka [۲۷] که رایگان و متن‌باز است، استفاده شده است. مجموعه برچسب‌های تمام نمونه‌ها در قالب مناسب ذخیره شده و برای استخراج قوانین انجمنی به کار رفته است. برای پیاده‌سازی روش پیشنهادی و رده‌بندی چندبرچسبی از چارچوب Meka [۲۸] استفاده شده است که درواقع توسعه یافته چارچوب Weka برای حوزه رده‌بندی چندبرچسبی است.

۴-۳- معیارهای ارزیابی

به دلیل تفاوت‌های فضای کاری در مسائل رده‌بندی چندبرچسبی با مسائل تک‌برچسبی، معیارهای ارزیابی آن‌ها نیز متفاوت است. به طور کلی، معیارهای ارزیابی حوزه رده‌بندی چندبرچسبی به دو دسته معیارهای مبتنی بر نمونه و معیارهای مبتنی بر برچسب تقسیم می‌شوند. در دسته اول، عملکرد مدل روی نمونه‌های آزمون ملاک ارزیابی قرار می‌گیرد و در دسته دوم، عملکرد بر اساس هر برچسب به صورت جداگانه ارزیابی می‌شود. از سوی دیگر، معیارهای ارزیابی افزایشی به معیارهایی گفته می‌شود که بهبود رده‌بندی در بیشینه بودن آن‌هاست و معیارهای ارزیابی کاهش‌ی، معیارهایی هستند که بهبود رده‌بندی در کمینه بودن آن‌هاست.

در ادامه، پنج معیار ارزیابی مورد استفاده در این مقاله را معرفی می‌کنیم که معیارهای ارزیابی اول و دوم، افزایشی هستند (به این معنی که افزایش مقدار آن‌ها نشان‌دهنده بهبود رده‌بندی است) و معیارهای ارزیابی سوم تا پنجم، کاهش‌ی هستند. روابط (۲) تا (۶) معیارهای ارزیابی مورد استفاده در این مقاله را نشان می‌دهند. در این روابط، p نشان‌دهنده تعداد نمونه‌های آزمون است و $h(x_i)$ مجموعه برچسب‌های پیش‌بینی‌شده توسط رده‌بند و $Y(x_i)$ مجموعه برچسب‌های واقعی نمونه x_i را نشان می‌دهند. همچنین در این روابط، $f(x_i, y)$ تابعی است که رتبه برچسب y را در پیش‌بینی رده‌بند برای نمونه x_i برمی‌گرداند.

$$\text{Subset Accuracy: } \frac{1}{p} \sum_{i=1}^p [h(x_i) = Y(x_i)] \quad (2)$$

رابطه (۲) معیار ارزیابی صحت زیرمجموعه یا تطابق کامل^{۱۴} را نشان می‌دهد. این معیار بررسی می‌کند که رده‌بند تمامی برچسب‌ها را صحیح پیش‌بینی کرده‌باشد و حتی اگر یکی از برچسب‌های متعلق به نمونه، توسط رده‌بند پیش‌بینی نشده باشد و یا برچسبی اضافی به نمونه نسبت داده شده‌باشد، این معیار برای نمونه موردنظر برابر با صفر خواهد بود. این معیار، به‌ویژه در مواردی که تعداد برچسب‌ها زیاد باشد بسیار سخت‌گیرانه عمل می‌کند.

$$\text{Accuracy}_{\text{exam}}: \frac{1}{p} \sum_{i=1}^p \frac{|h(x_i) \cap Y(x_i)|}{|h(x_i) \cup Y(x_i)|} \quad (3)$$

$\text{Accuracy}_{\text{exam}}$ در رده‌بندی چندبرچسبی، به‌عنوان معیار صحت تعریف شده است. همان‌طور که در رابطه (۳) مشخص است، این تعریف از صحت، متفاوت از تعریف آن در مسائل تک‌برچسبی است و برابر با

نسبت اشتراک برچسب‌های واقعی و برچسب‌های پیش‌بینی‌شده توسط رده‌بند به اجتماع آن‌ها است.

$$\text{One-error: } \frac{1}{p} \sum_{i=1}^p \left[\operatorname{argmax}_{y \in Y(x_i)} f(x_i, y) \right] \notin Y(x_i) \quad (4)$$

در معیار رابطه (۴) ارزیابی بر این اساس صورت می‌گیرد که از بین برچسب‌های پیش‌بینی‌شده توسط رده‌بند، برچسب دارای رتبه برتر در بین برچسب‌های واقعی نمونه باشد.

$$\text{Ranking Loss: } \quad (5)$$

$\frac{1}{p} \sum_{i=1}^p \frac{1}{|Y_i| | \hat{Y}_i |} |\{ (y', y'') | f(x_i, y') \leq f(x_i, y''), (y', y'') \in Y_i \times \hat{Y}_i \}|$ که در آن $\hat{Y}_i$ برچسب‌های غیرمتعلق به نمونه i را نشان می‌دهد. در معیار رابطه (۵)، نکته مهم این است که برای هر نمونه، برچسب‌های غیرمتعلق به آن نمونه از برچسب‌های واقعی نمونه، رتبه کمتری را در رتبه‌بندی پیش‌بینی‌شده از برچسب‌ها توسط رده‌بند داشته‌باشند.

$$\text{Hamming loss: } \frac{1}{p} \sum_{i=1}^p |h(x_i) \Delta Y(x_i)| \quad (6)$$

در معیار رابطه (۶)، Δ به مفهوم تفاوت مقارن بین دو مجموعه $h(x_i)$ و $Y(x_i)$ است. به عبارت دیگر، به‌ازای هر برچسب واقعی از نمونه، اگر رده‌بند آن برچسب را برای نمونه پیش‌بینی نکند، Δ برابر یک و در غیر این صورت برابر صفر است. همچنین اگر رده‌بند برچسبی را به اشتباه به نمونه نسبت دهد، Δ برابر یک و در غیر این صورت برابر صفر است. در بدترین حالت، مقدار این معیار برابر با مجموع تعداد برچسب‌های واقعی و برچسب‌های پیش‌بینی‌شده است که نشان می‌دهد رده‌بند هیچ‌کدام از برچسب‌های واقعی را پیش‌بینی نکرده است.

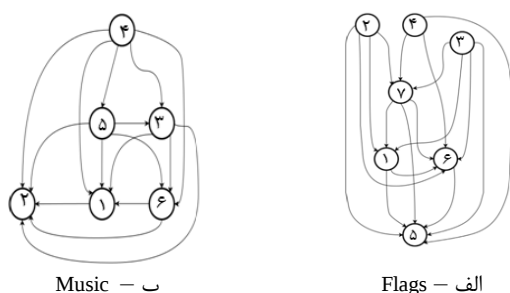
۴-۴- الگوریتم‌های رقیب و شرایط آزمایش

در دو روش پیشنهادی این مقاله، برای استخراج قوانین انجمنی از روش FP-Growth [۲۹] استفاده شده است. ضریب اطمینان در استخراج قوانین ۰/۹ (مقدار پیش‌فرض در WEKA) بوده است. برای ارزیابی، روش‌های پیشنهادی با روش‌های BR, CC, BCC و CT مقایسه شده‌اند. در تمامی روش‌ها از درخت تصمیم j48 با مقادیر پیش‌فرض آن در WEKA به‌عنوان رده‌بند پایه استفاده شده است. ارزیابی بر اساس روش ارزیابی 10-fold Cross Validation صورت گرفته است و نتایج گزارش‌شده در جدول‌ها، میانگین و انحراف معیار نتایج به‌دست‌آمده از ده مرتبه اجرا است. عملکرد روش‌های رقیب فوق بر روی ۶ مجموعه داده Yeast, Scene, Birds, Music, Flags و Emotions در جداول ۲ تا ۶ ثبت شده است.

۵- مشاهدات و تحلیل نتایج

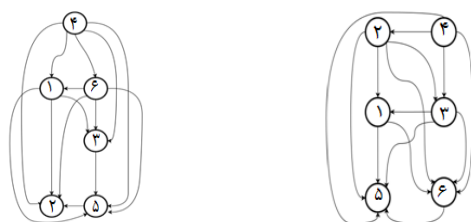
جداول ۲ تا ۶ نتایج آزمایش‌ها روی روش‌های مختلف به تفکیک معیارهای ارزیابی را نشان می‌دهند که با استفاده از آن‌ها نه تنها می‌توان میزان بهبود در نتایج با استفاده از روش‌های پیشنهادی را نسبت به روش CC مشاهده کرد، بلکه می‌توان روش پیشنهادی را با سایر روش‌های رقیب که هر یک به نحوی نتایج CC پایه را بهبود

روش BCC تقریباً برابری می‌کند. صحت روش پیشنهادی در مجموعه داده‌های Emotions و Music نیز بهبود یافته، همچنین نسبت به روش‌های رقیب بهترین عملکرد داشته است.



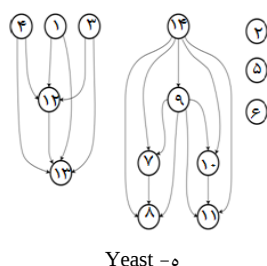
ب - Music

الف - Flags

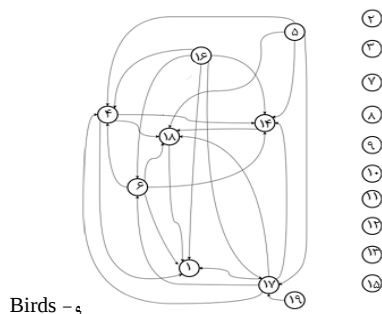


د - Emotions

ج - Scene



ه - Yeast



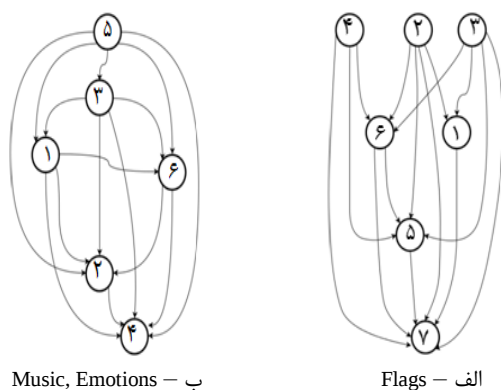
و - Birds

شکل ۴: گراف زنجیره رده‌بند حاصل از روش ARCC2

نتایج جدول ۲ نشان می‌دهد که روش ARCC2 روی مجموعه داده Birds افت داشته است. بالاترین صحت تک‌برچسبی در مورد این مجموعه داده به ترتیب مربوط به برچسب‌های ۱۷، ۶، ۴ و ۱۴ است. با دقت در گراف ترتیب زنجیره در شکل ۴-و، مشاهده می‌شود که باوجود این‌که برچسب ۱۷ بالاترین صحت را داشته و در بالای زنجیره قرار دارد اما متأثر از برچسب‌های ۵، ۱۶ و ۱۹ است که صحتشان از برچسب ۱۷ کمتر است که این مشاهدات، لزوم انتخاب قوانین مناسب و حذف قوانین مضر را آشکار می‌کند. همین تحلیل در خصوص مجموعه داده Emotions نیز صدق می‌کند.

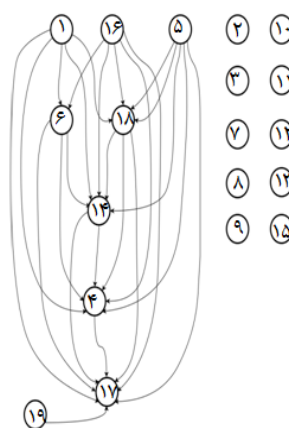
بخشیده‌اند، مقایسه کرد. در مورد هر مجموعه داده، بهترین نتیجه به صورت پررنگ نشان داده شده است.

برای تحلیل هر کدام از معیارهای ارزیابی بر روی ۶ مجموعه داده فوق، گراف ساخته شده بر اساس روش‌های پیشنهادی ARCC1 و ARCC2 در شکل‌های ۳ و ۴ آمده است.

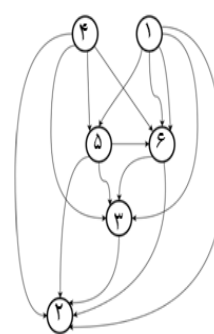


ب - Music, Emotions

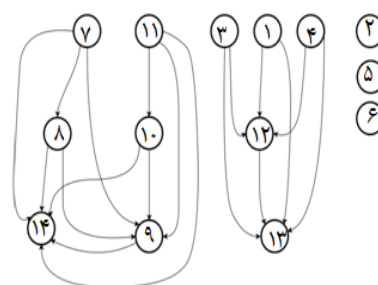
الف - Flags



د - Birds



ج - Scene



ه - Yeast

شکل ۳: گراف زنجیره رده‌بند حاصل از روش ARCC1

جدول ۲ نتایج مربوط به معیار ارزیابی صحت را نشان می‌دهد. مشاهده می‌شود که دو روش پیشنهادی بهترین عملکرد را روی همه مجموعه داده‌ها به استثنای مجموعه داده yeast دارند. همان‌طور که در جدول ۲ مشاهده می‌شود، با استفاده از روش ARCC1 صحت در مجموعه داده Birds نسبت به روش BR و CC اندکی افزایش یافته و با

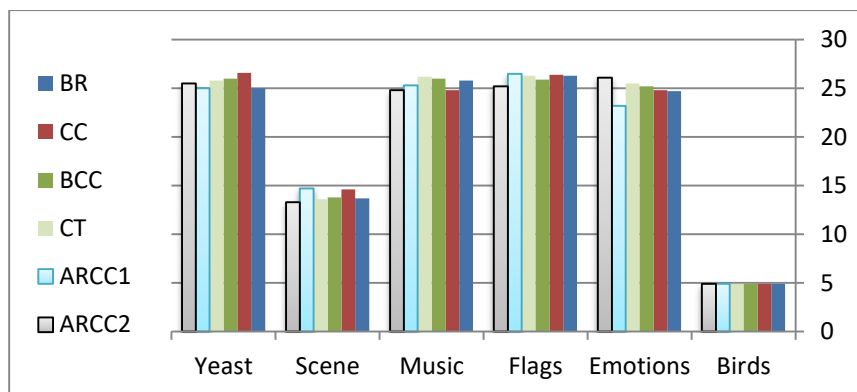
عملکرد آن نسبت به BCC و CT و BR بهتر است. در روش ARCC2، گراف ۴-ج به همان ترتیب صحت تکبرچسبی ساخته شده و از آنجاکه نتایج رده‌بندهای با صحت بالا به رده‌بندهای دیگر داده می‌شود، بهترین صحت کل را به همراه داشته است.

با مقایسه میانگین رتبه الگوریتم‌های رقیب در جدول ۲ مشاهده می‌شود که الگوریتم پیشنهادی ARCC1 بهترین میانگین رتبه را به دست آورده است و الگوریتم‌های ARCC2 و CC میانگین رتبه یکسانی دارند. با مقایسه میانگین رتبه الگوریتم‌ها با استفاده از آزمون فریدمان، مقادیر $F = 1/0.2$ و $x^2_F = 5/0.9$ با مقدار بحرانی $2/0.9$ در سطح بحرانی $0/1$ به دست می‌آید و بنابراین نمی‌توان فرضیه پوچ را رد کرد که این بدین معنی است که تفاوت بین رده‌بندهای رقیب با توجه به این معیار از نظر آماری بااهمیت نیست.

در مورد مجموعه داده Scene، روش پیشنهادی ARCC2 با در نظر گرفتن معیار صحت بهتر از سایر روش‌های رقیب عمل کرده است. این در حالی است که ARCC1 نسبت به روش پایه CC دچار افت صحت شده است. همان‌طور که دیده می‌شود این افت صحت نسبت به CC در مورد روش‌های رقیب BCC و CT نیز وجود دارد. با بررسی مجموعه داده Scene برای فهمیدن دلیل این مشاهده، روشن می‌شود که این مجموعه داده هرچند جزء مجموعه داده‌های رایج در حوزه چندبرچسبی است، ولی فقط حدود ۷ درصد از کل ۲۴۰۷ نمونه آن بیش از یک برچسب دارند و سایر نمونه‌های آن تکبرچسبی هستند. این مسئله باعث شده است که روش پیشنهادی ARCC1 که به دنبال یافتن زنجیره و ترتیب از روی قوانین است، به دلیل ناکافی بودن قوانین باکیفیت، نسبت به روش پایه CC افت کند، ولی بالاین حال،

جدول ۲: مقایسه دو روش ARCC با BR، CC، BCC، CT بر اساس معیار صحت

مجموعه داده	BR	CC	BCC	CT	ARCC1	ARCC2
Birds	$57/3 \pm 4/1$	$56/3 \pm 4/5$	$57/5 \pm 4/3$	$57/1 \pm 3/9$	$57/5 \pm 4/3$	$57/4 \pm 5/7$
Emotions	$46/2 \pm 3/7$	$48/3 \pm 4/7$	$46/7 \pm 4/0$	$46/2 \pm 7/1$	$50/9 \pm 5/5$	$45/8 \pm 4/4$
Flags	$60/1 \pm 6/8$	$59/5 \pm 8/2$	$60/2 \pm 7/3$	$59/0 \pm 8/1$	$59/2 \pm 8/0$	$62/3 \pm 5/1$
Music	$43/0 \pm 5/1$	$48/3 \pm 3/8$	$45/1 \pm 4/6$	$45/2 \pm 3/3$	$47/9 \pm 4/7$	$48/7 \pm 4/8$
Scene	$53/5 \pm 3/4$	$58/2 \pm 2/3$	$54/8 \pm 3/1$	$54/7 \pm 2/2$	$56/2 \pm 3/2$	$61/1 \pm 2/0$
Yeast	$44/0 \pm 2/0$	$42/9 \pm 2/3$	$41/9 \pm 1/3$	$42/1 \pm 1/9$	$42/5 \pm 2/0$	$41/6 \pm 2/8$
میانگین رتبه	۴/۰۸	۳	۳/۴۱	۴/۷۵	۲/۷۵	۳



نمودار ۱: مقایسه دو روش پیشنهادی ARCC با BR، CC، BCC، CT بر اساس معیار صحت

جدول ۳: مقایسه دو روش ARCC با BR، CC، BCC، CT بر اساس معیار صحت سخت گیرانه

مجموعه داده	BR	CC	BCC	CT	ARCC1	ARCC2
Birds	$48/5 \pm 3/6$	$47/3 \pm 4/3$	$49/0 \pm 4/5$	$48/4 \pm 3/6$	$48/5 \pm 3/6$	$47/8 \pm 7/3$
Emotions	$18/4 \pm 5/3$	$24/3 \pm 5/8$	$18/7 \pm 4/9$	$19/7 \pm 7/4$	$27/8 \pm 5/4$	$21/1 \pm 5/1$
Flags	$15/5 \pm 9/7$	$23/7 \pm 12/4$	$25/2 \pm 9/4$	$18/5 \pm 11/2$	$21/2 \pm 10/1$	$17/0 \pm 7/7$
Music	$17/7 \pm 6/4$	$23/6 \pm 3/7$	$20/0 \pm 5/0$	$20/4 \pm 3/4$	$23/6 \pm 4/1$	$25/0 \pm 5/5$
Scene	$42/7 \pm 3/5$	$52/7 \pm 3/2$	$45/8 \pm 2/9$	$44/8 \pm 2/8$	$48/3 \pm 2/9$	$55/0 \pm 2/7$
Yeast	$6/8 \pm 1/7$	$14/2 \pm 1/9$	$5/1 \pm 3/1$	$7/2 \pm 2/1$	$7/9 \pm 1/8$	$7/7 \pm 0/9$
میانگین رتبه	۵/۲۵	۲/۵۸	۳/۶۷	۴/۱۷	۲/۳۳	۳

رتبه الگوریتمها با استفاده از آزمون فریدمان، مقادیر $x^2_F = ۲/۹۳$ و $F_F = ۰/۵۴$ با مقدار بحرانی $۲/۰۹$ در سطح بحرانی $۰/۱$ به دست می‌آید و بنابراین نمی‌توان فرضیه پوچ را رد کرد که این بدین معنی است که تفاوت بین رده‌بندی‌های رقیب از نظر آماری بااهمیت نیست.

جدول ۵ نتایج به‌دست‌آمده بر اساس معیار one error را نشان می‌دهد. همان‌طور که در بخش ۴-۳ بیان شد، این معیار به‌ازای هر نمونه بررسی می‌کند که آیا برجسب دارای رتبه بالاتر در پیش‌بینی رده‌بند، در بین برجسب‌های واقعی نمونه قرار دارد یا خیر؟ و اگر نتیجه منفی باشد، مقدار error یک واحد افزایش پیدا می‌کند. نتایج نشان می‌دهند روش‌های پیشنهادی ARCC1 و ARCC2 از نظر میانگین رتبه در این معیار به ترتیب رتبه‌های اول و دوم را کسب نموده‌اند که نشان‌دهنده برتری روش‌های پیشنهادی بر اساس این معیار در بین الگوریتم‌های رقیب است. با مقایسه میانگین رتبه الگوریتمها با استفاده از آزمون فریدمان، مقادیر $x^2_F = ۵/۷۴$ و $F_F = ۱/۱۸$ با مقدار بحرانی $۲/۰۹$ در سطح بحرانی $۰/۱$ به دست می‌آید و بنابراین نمی‌توان فرضیه پوچ را رد کرد که این بدین معنی است که تفاوت بین رده‌بندی‌های رقیب با در نظر گرفتن معیار one error از نظر آماری بااهمیت نیست.

تحلیل Rank loss: این معیار به ازای هر نمونه دو لیست می‌سازد. لیستی از برجسب‌های واقعی و لیست دیگر شامل برجسب‌هایی است که متعلق به نمونه نیست. این معیار کنترل می‌کند برجسب‌های متعلق به یک نمونه در رتبه‌بندی خروجی توسط رده‌بند، رتبه بالاتری نسبت به برجسب‌های غیرمتعلق به آن داشته باشند.

جدول ۶ نتایج به‌دست‌آمده برای روش‌های مختلف را بر مبنای معیار rank loss نشان می‌دهد. مشاهده می‌شود که روش‌های پیشنهادی ARCC1 و ARCC2 به‌ترتیب بهترین و سومین میانگین رتبه را در بین الگوریتم‌های مورد مقایسه به‌دست آورده‌اند. با مقایسه میانگین رتبه الگوریتمها با استفاده از آزمون فریدمان، مقادیر به‌دست‌آمده به‌صورت $x^2_F = ۴/۵۷$ و $F_F = ۰/۸۹$ با مقدار بحرانی $۲/۰۹$ در سطح بحرانی $۰/۱$ است و بنابراین نمی‌توان فرضیه پوچ را رد کرد که این بدین معنی است که تفاوت بین رده‌بندی‌های رقیب با در نظر گرفتن معیار rank loss از نظر آماری بااهمیت نیست.

در جدول ۳ نتایج مربوط به صحت سخت‌گیرانه آمده‌است. مشاهده می‌شود که الگوریتم پیشنهادی ARCC1 بهترین میانگین رتبه را در بین الگوریتم‌های رقیب به‌دست آورده است و الگوریتم‌های CC و ARCC2 به ترتیب رتبه‌های دوم و سوم را کسب کرده‌اند. با مقایسه میانگین رتبه الگوریتمها با استفاده از آزمون فریدمان، مقادیر به‌دست‌آمده برابر $x^2_F = ۱۰/۲۶$ و $F_F = ۲/۵۹$ است که با توجه به مقدار بحرانی $۲/۰۹$ در سطح بحرانی $۰/۱$ ، نمی‌توان فرضیه پوچ را رد کرد که این بدین معنی است که بین عملکرد رده‌بندها تفاوت معناداری از نظر آماری وجود دارد. با اعمال آزمون ثانویه McNemar، با توجه به اینکه میزان تفاوت بحرانی در سطح $۰/۱$ برابر $۲/۷۹$ است، می‌توان نتیجه گرفت که روش پیشنهادی ARCC1 با در نظر گرفتن معیار صحت سخت‌گیرانه بسیار بهتر از روش BR است.

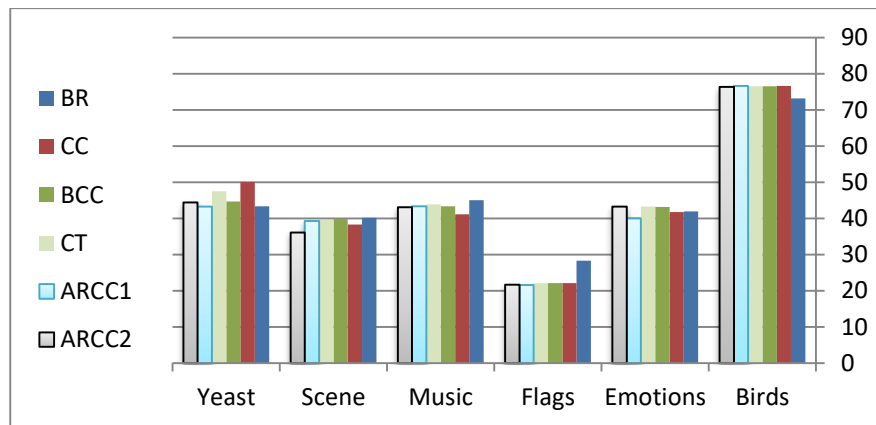
نتایج جدول ۳ نشان می‌دهد که دو روش ARCC در معیار صحت سخت‌گیرانه در مورد مجموعه داده Flags افت داشته است. در این مجموعه داده صحت تک‌برجسبی برجسب‌های ۵ و ۷ از سایر برجسبها بیشتر است. از طرفی با توجه به ترتیب زنجیره در شکل‌های ۳-الف و ۴-الف، خروجی سایر رده‌بندها به رده‌بند برجسب‌های ۵ و ۷ می‌رسد و بنابراین در ترتیب‌های پیشنهادی، خطای سایر رده‌بندها به این رده‌بند سرایت می‌کند و در نتیجه، صحت کل کاهش می‌یابد. روش پیشنهادی ARCC1 در مورد مجموعه داده Emotions بهترین عملکرد را داشته است.

جدول ۴ نتایج حاصل از الگوریتم‌های رقیب را روی داده‌های مورد آزمایش بر اساس معیار Hamming loss نشان می‌دهد. این معیار در مورد هر نمونه، با یافتن هر اختلافی بین برجسب‌های واقعی و برجسب‌های پیش‌بینی شده، مقدار loss را اضافه می‌کند. مقادیر کم برای این معیار، نتایج بهتر را نشان می‌دهند. روش پیشنهادی ARCC1 معیار را در دو مجموعه داده و ARCC2 در سه مجموعه داده بهبود داده به‌طوری که بهترین عملکرد داشته‌اند. در خصوص مجموعه داده Birds، تفاوتی بین نتایج روش‌های رقیب مشاهده نمی‌شود.

نتایج جدول ۴ نشان می‌دهد که در این معیار از نظر میانگین رتبه، روش پیشنهادی ARCC2 رتبه اول، روش BR رتبه دوم و روش پیشنهادی ARCC1 رتبه سوم را کسب کرده‌اند. با مقایسه میانگین

جدول ۴: مقایسه دو روش ARCC با BR، CC، BCC، CT بر اساس معیار Hamming loss

مجموعه داده	BR	CC	BCC	CT	ARCC1	ARCC2
Birds	$۴/۹ \pm ۰/۶$	$۴/۹ \pm ۰/۵$	$۴/۹ \pm ۰/۶$	$۴/۹ \pm ۰/۶$	$۴/۹ \pm ۰/۶$	$۴/۹ \pm ۰/۸$
Emotions	$۲۴/۷ \pm ۲/۶$	$۲۴/۸ \pm ۳/۲$	$۲۵/۲ \pm ۱/۹$	$۲۵/۵ \pm ۴/۰$	$۲۳/۲ \pm ۳/۶$	$۲۶/۱ \pm ۱/۸$
Flags	$۲۶/۳ \pm ۵/۲$	$۲۶/۴ \pm ۵/۲$	$۲۵/۹ \pm ۴/۹$	$۲۶/۳ \pm ۵/۲$	$۲۶/۵ \pm ۵/۷$	$۲۵/۲ \pm ۳/۹$
Music	$۲۵/۸ \pm ۲/۷$	$۲۴/۸ \pm ۱/۷$	$۲۶/۰ \pm ۲/۵$	$۲۶/۲ \pm ۱/۹$	$۲۵/۳ \pm ۲/۶$	$۲۴/۸ \pm ۱/۶$
Scene	$۱۳/۷ \pm ۰/۸$	$۱۴/۶ \pm ۰/۸$	$۱۳/۸ \pm ۱/۰$	$۱۳/۶ \pm ۰/۷$	$۱۴/۷ \pm ۱/۱$	$۱۳/۳ \pm ۰/۹$
Yeast	$۲۵/۰ \pm ۱/۰$	$۲۶/۶ \pm ۱/۰$	$۲۶/۰ \pm ۴/۰$	$۲۵/۸ \pm ۱/۰$	$۲۵/۰ \pm ۰/۸$	$۲۵/۵ \pm ۰/۵$
میانگین رتبه	۲/۹۱	۴	۳/۹۱	۴	۳/۵	۲/۶۶



نمودار ۲: مقایسه دو روش پیشنهادی ARCC با BR, CC, BCC, CT بر اساس معیار Hamming loss

جدول ۵: مقایسه دو روش ARCC با BR, CC, BCC, CT بر اساس معیار One Error

مجموعه داده	BR	CC	BCC	CT	ARCC2	ARCC2
Birds	73.2 ± 4.7	73.2 ± 4.7	73.2 ± 4.7	73.2 ± 4.7	73.2 ± 4.7	73.2 ± 4.7
Emotions	42.0 ± 6.3	41.8 ± 6.3	42.2 ± 5.8	42.3 ± 8.9	40.0 ± 7.5	43.3 ± 6.4
Flags	28.3 ± 11.8	22.1 ± 11.0	22.1 ± 11.0	22.1 ± 11.0	21.6 ± 11.2	21.7 ± 9.5
Music	45.1 ± 7.9	41.2 ± 5.6	43.4 ± 7.6	43.9 ± 4.7	43.4 ± 5.9	43.1 ± 5.3
Scene	40.2 ± 2.8	38.3 ± 2.3	39.9 ± 3.5	39.6 ± 2.7	39.3 ± 3.5	36.1 ± 2.2
Yeast	43.4 ± 3.9	50.1 ± 3.2	44.7 ± 2.2	47.5 ± 2.8	43.2 ± 3.4	44.4 ± 2.9
میانگین رتبه	4	3.41	4	4.5	2.5	2.58

جدول ۶: مقایسه دو روش ARCC با BR, CC, BCC, CT بر اساس معیار Rank loss

مجموعه داده	BR	CC	BCC	CT	ARCC1	ARCC2
Birds	15.7 ± 2.6	18.5 ± 2.5	18.1 ± 2.9	18.4 ± 2.6	18.2 ± 2.8	18.6 ± 3.9
Emotions	29.2 ± 2.7	28.5 ± 4.5	28.4 ± 3.3	29.4 ± 6.9	26.6 ± 4.5	30.2 ± 4.1
Flags	27.8 ± 7.1	27.7 ± 5.3	28.5 ± 6.0	29.4 ± 5.2	29.4 ± 4.5	29.0 ± 6.9
Music	30.1 ± 4.5	28.4 ± 3.0	30.1 ± 4.1	30.6 ± 3.9	28.3 ± 3.2	27.4 ± 3.2
Scene	24.9 ± 2.5	22.8 ± 2.0	22.1 ± 2.7	22.1 ± 2.0	22.1 ± 2.0	19.9 ± 2.0
Yeast	30.7 ± 1.2	29.6 ± 1.6	29.0 ± 1.1	29.1 ± 1.7	28.1 ± 1.5	29.0 ± 1.4
میانگین رتبه	3.91	3.66	2.83	4.58	2.58	3.41

۶- نتیجه گیری و پیشنهادهایی برای آینده

بسیار عمومی است، به نحوی که نه تنها لزوماً همبند نیست، بلکه الزامی در درخت بودن آن نیز وجود ندارد.

نتایج حاصل از آزمایش‌های گسترده ما بر روی مجموعه داده‌های متنوع و پرکاربرد در حوزه رده‌بندی چندبرچسبی نشان دادند که نوع قوانین استخراج‌شده و تعداد آن‌ها در ایجاد گراف ارتباط نقش مؤثری دارد. هرچه تعداد قوانین و برچسب‌های مؤثر در سمت تالی قوانین بیشتر باشد، گراف شلوغ‌تر و کامل‌تر می‌شود و در نتیجه میزان بهبود در معیارهای ارزیابی بیشتر خواهد بود.

نتایج حاصل از مقایسه روش‌های پیشنهادی با روش‌های موجود روی مجموعه داده‌های رایج در این مقاله، شایستگی روش‌های پیشنهادی نسبت به روش‌های موجود را تایید می‌کنند. البته نتایج مشاهده‌شده بسیار وابسته به معیارهای ارزیابی مورد استفاده هستند،

در این مقاله، دو روش جدید برای تعیین ترتیب رده‌بندها برای اجرای الگوریتم زنجیره رده‌بندها (CC) ارائه شد که بر اساس آن‌ها، برخلاف روش پایه الگوریتم CC که در آن ترتیب رده‌بندها در زنجیره به صورت تصادفی تعیین می‌شود، روابط بین برچسب‌ها در مسئله چندبرچسبی ترتیب رده‌بندها را تعیین می‌کنند. در روش‌های پیشنهادی، از قوانین انجمنی استخراج‌شده از فضای برچسب‌ها که نشان‌دهنده روابط بین برچسب‌ها هستند، برای ایجاد گراف ارتباط استفاده می‌شود که نهایتاً مبنای تعیین ترتیب رده‌بندها قرار می‌گیرد. در الگوریتم پایه CC، گراف ارتباط در حقیقت به صورت یک خط (توالی یا زنجیر) است که نوع ساده درخت همبند است. در حالی که در روش‌های پیشنهادی، گراف ارتباط

- [10] K. Dembczyński, W. Cheng, and E. Hüllermeier, "Bayes Optimal Multilabel Classification via Probabilistic Classifier Chains", 27th International Conference on Machine Learning, Haifa, 2010.
- [11] J. Read, L. Martino, P. M. Olmos, and D. Luengo, "SCALABLE MULTI-OUTPUT LABEL PREDICTION: FROM CLASSIFIER CHAINS TO CLASSIFIER TRELLISES", arXiv: 1501.04870v1 [stat.ML], 2015.
- [12] T. Kajdanowicz and P. Kazienko, "Heuristic Classifier Chains for Multi-label Classification", Springer-Verlag Berlin Heidelberg FQAS, LNAI. 8132, pp. 555–566, 2013.
- [13] E. A. Cherman, C. M. Maria, and M. Jean, "Multi-label Problem Transformation Methods: a Case Study", CLEI ELECTRONIC JOURNAL, Vol. 14, pp. 4–14, 2011.
- [14] J. Read, B. Pfahringer, and G. Holmes, "Multi-label Classification using Ensembles of Pruned Sets", 8th IEEE International Conference on Data Mining, 2008.
- [15] M. Zhang and Z. Zhi-Hua, "ML-KNN: A lazy learning approach to multi-label learning", ELSEVIER, Pattern Recognition Vol. 40, pp. 2038 – 2048, 2007.
- [16] C. Vens, J. Struyf, L. Schietgat, S. Džeroski, and H. Blockeel, "Decision trees for hierarchical multi-label classification", Springer, Mach Learn Vol. 73, pp. 185–214, 2008.
- [17] J. Read and F. Perez-Cruz, "Deep Learning for Multi-label Classification", arXiv:1502.05988v1 [cs.LG], 2014.
- [18] L. E. Mencía and J. Fürnkranz, "Pairwise Learning of Multilabel Classifications with Perceptrons", IEEE 978-1-4244-1821-3/08, 2008.
- [19] G. Tsoumakas, I. Katakis, and I. Vlahavas, "Random k-Labelsets for Multilabel Classification", IEEE Transactions on Knowledge and Data Engineering Vol. 23, NO. 7, 2011.
- [20] Sh. Wang, J. Wang, Z. Wang, Q. and Ji, "Enhancing multi-label classification by modeling dependencies among labels", Elsevier, Pattern Recognition Vol. 47, pp. 3405–3413, 2014.
- [21] F. Charte, A. Rivera, J.M. Jesus, and F. Herrera, "Improving Multi-label classifiers via Label Reduction with Association Rules", Springer-Verlag Berlin Heidelberg HAIS part II, LNCS 7209, pp. 188–199, 2012.
- [22] R. Alazaidah, F. Thabtah, and Q. Al-Raaidehi, "A Multi-Label Classification Approach Based on Correlations Among Labels", International Journal of Advanced Computer Science and Applications Vol. 6, No. 2, pp. 52–59, 2015.
- [23] K. Patel, N. Kapadia, and M. Parikh, "Discover Multi-label Classification using Association Rule Mining", International journal of Advance Engineering and Research Development, Vol. 1, Issue. 1, ISSN 2348-4470, 2014.
- [24] Y. Rong, Q. Yanpeng and D. Ansheng, *Multi-label Fuzzy Similarity-Based Nearest-Neighbour Classification Using Association Rule*, International Conference on Intelligent Data Engineering and Automated Learning. Springer International Publishing, LNCS 9937, pp. 542–551, 2016.
- به نحوی که در مورد معیار صحت سخت گیرانه، برتری روش پیشنهادی ARCC1 نسبت به روش BR از نظر آماری در سطح بحرانی ۰/۱ تأیید شد. همچنین مشاهدات ما نشان دادند که متوازن بودن مجموعه داده و افزایش تعداد نمونه‌های چندبرچسبی در مجموعه داده به افزایش بهبود در نتایج خواهد انجامید.
- پیشنهادهای این مقاله برای پژوهش‌های آینده به صورت زیر است:
- انتخاب هوشمندانه قوانین بهینه از بین قوانین استخراج شده.
 - حذف قوانین مضر استخراج شده از مجموعه داده‌های با تعداد برچسب زیاد و کاردینالیتی برچسب کم.
 - بررسی میزان تأثیر پارامترهای استخراج قوانین در نتایج.
 - انتخاب برچسب‌های اولیه مناسب برای کاهش انتشار خطا.
- ### مراجع
- [۱] محمدعلی زارع چاهوکی و حمیدرضا محمدی، "بهینه‌سازی هسته‌های چندگانه در ماشین بردار پشتیبان جفتی برای کاهش شکاف معنایی تشخیص صفحات فریب‌آمیز"، مجله مهندسی برق دانشگاه تبریز، دوره ۴۶، شماره ۴، زمستان ۱۳۹۵، صفحه ۱۴۵–۱۳۵
- [۲] ندا خانبانی، امیرمسعود افتخاری مقدم، "ارائه یک روش تشخیص زبان علامت مبتنی بر رویکرد MLRF فازی با استفاده از اطلاعات عمق تصویر"، مجله مهندسی برق دانشگاه تبریز، انتشار آنلاین از تاریخ ۱۳ اسفند ۱۳۹۵
- [3] J. Read, B. Pfahringer, G. Holmes, and E. Frank, "Classifier chains for multi-label classification", Springer, Mach Learn Vol 85, pp 333–359, 2011.
- [4] M. Zhang and Z. Zhi-Hua, "A Review on Multi-Label Learning Algorithms", IEEE Transactions on Knowledge and Data Engineering, Vol. 26, Issue 8, 2013.
- [5] W. Bi, J. Kwok, "Multi-Label Classification on Tree- and DAG-Structured Hierarchies", Appearing in Proceedings of the 28th International Conference on Machine Learning, Bellevue, WA, USA, 2011.
- [6] J. Han, J. Pei, and M. Kamber, *Data Mining concepts and techniques*, Elsevier third Edition Book, pp. 243–278, 2012.
- [7] R. Agrawal and R. Srikant, *Fast Algorithms for mining Association Rules in Large Databases*, 20th International Conference on Very Large Data Bases, pp. 478–499, 1998.
- [8] J. Read, B. Pfahringer, G. Holmes, and E. Frank, "Classifier Chains for Multi-label Classification", Springer-Verlag Berlin Heidelberg, LNAI. 5782, pp. 254–269, 2009.
- [9] L. S. Enrique, C. Bielza, E. F. Morales, P. Hernandez-Leal, J. H. Zaragoza, and P. Larrañaga, "Multi-label classification with Bayesian network-based chain classifiers", ELSEVIER, Pattern Recognition Letters Vol. 41, pp. 14–22, 2014.

- [27] Waikato Environment for Knowledge Analysis, Version 3.6.10, 2013, www.cs.waikato.ac.nz/~ml/weka (دسترسی در ۱۳۹۴/۱۱/۵)
- [28] A Multi-label Extension to WEKA, Version 1.7.7, 2015, www.meka.sourceforge.net (دسترسی در ۱۳۹۴/۱۱/۵)
- [29] J. Han, Pei J. and Yin Y, "Mining frequent patterns without candidate generation", Proceeding of the 2000 ACM-SIGMID International Conference on Management of Data, pp. 1-12, 2000.
- [25] F. Benites and E. Sapozhnikova, "Improving Multi-label Classification by Means of Cross-Ontology Association Rules." Bioinformatics, vol. 2, no. 9, 2015.
- [۲۶] خلیل غفوری‌پور و زهرا میرزامؤمن، "بهبود رده‌بندی چندبرچسبی بر اساس استخراج قوانین انجمنی مثبت و منفی از روی برچسب‌ها"، بیست و یکمین کنفرانس ملی سالانه انجمن کامپیوتر ایران، پژوهشگاه دانش‌های بنیادی، تهران، ۱۸-۲۰ اسفند ۹۴.

زیرنویس‌ها

-
- ۱ Binary Relevance
 - ۲ Classifier Chains
 - ۳ Hierarchical Multilabel Classification
 - ۴ Ensembles of Classifier Chains
 - ۵ Pairwise
 - ۶ Label Powerset
 - ۷ Pruned Sets
 - ۸ Multi Label k Nearest Neighbor
 - ۹ Bayesian Classifier chain
 - ۱۰ Tree Naïve Bayesian Classifier Chain
 - ۱۱ Classifier Trellises
 - ۱۲ Incomplete Label
 - ۱۳ Association Rule Classifier Chains
 - ۱۴ Subset Accuracy