



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

9th
Iranian Joint Congress on Fuzzy and
Intelligent Systems

March 2 - 4 , 2022

Higher Education Complex of Bam , Iran

۱۱ لغایت ۱۳ اسفند ۱۴۰۰

مجمع آموزش عالی بـم

مجموعه مقالات فارسی

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

(CFIS2022)

۱۱ تا ۱۳ اسفندماه ۱۴۰۰

مجمع آموزش عالی بـم

بیستمین کنفرانس سیستم‌های فازی ایران (ICFS2022)

هجدهمین کنفرانس سیستم‌های هوشمند (CIS2022)

پنجمین کنفرانس هوش جمعی و محاسبات تکاملی (CSIEC2022)

مجموعه مقالات فارسی نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

تهیه و تدوین: فاطمه بارانی، معراج عبدی

مجتمع آموزش عالی بم

اردیبهشت ماه ۱۴۰۱



بیاد مرحوم پروفور لطفی زاده

پدر منطق فازی



بیاد مرحوم دکتر بهرام صادقپور کیلده
استاد گروه آمار دانشگاه فردوسی مشهد

(۱۳۴۵-۱۴۰۰)

برنامه زمان بندی نهمین کنگره مشترک سیستم های فازی و هوشمند ایران
(روز اول)

تالار ۳	تالار ۲	تالار ۱	چهارشنبه ۱۱ اسفند ۱۴۰۰
مراسم افتتاحیه کنگره			۰۸:۰۰-۱۰:۰۰
استراحت			۱۰:۰۰-۱۰:۳۰
سخنرانی کلیدی (رئیس نشست: پروفسور ماشاالله ماشینچی) Evolutionary Intelligence in (Big) Data Analytics and Optimization Prof. Amir H. Gandomi, University of Technology Sydney, Australia			۱۰:۳۰-۱۱:۳۰
استراحت و ناهار			۱۱:۳۰-۱۳:۰۰
مجمع عمومی انجمن سیستم های فازی ایران			۱۳:۰۰-۱۴:۰۰
نشست پردازش متن و شبکه های اجتماعی	نشست ویژه هوش مصنوعی در خدمت جامعه انسانی	نشست ویژه بهینه سازی فازی شماره ۱	۱۴:۰۰-۱۶:۰۰
استراحت			۱۶:۰۰-۱۶:۱۵
سخنرانی کلیدی (رئیس نشست: پروفسور سیدناصر حسینی) Overview of new families of fuzzy implications and their applications in fuzzy systems Prof. Michal Baczynski, University of Silesia in Katowice, Katowice, Poland			۱۶:۱۵-۱۷:۱۵
استراحت			۱۷:۱۵-۱۷:۳۰
نشست الگوریتم های تکاملی و هوش جمعی	نشست ویژه ساختارهای جبری منطقی (فازی) شماره ۱	نشست ویژه آنالیز عددی فازی	۱۷:۳۰-۱۹:۳۰
سخنرانی کلیدی (رئیس نشست: پروفسور کامبیز بدیع) The Sound of Health پروفسور محمدرضا اکبرزاده توتونچی، دانشگاه فردوسی مشهد، ایران			۱۹:۳۰-۲۰:۳۰

برنامه زمان بندی نهمین کنگره مشترک سیستم های فازی و هوشمند ایران
(روز دوم)

تالار ۳		تالار ۲		تالار ۱		پنجشنبه ۱۴۰۰ اسفند ۱۲	
نشست آمار و احتمال فازی		نشست سیستمهای هوشمند و کاربردها شماره ۱		نشست ویژه اتوماتا و گرافهای فازی		۰۸:۰۰-۱۰:۰۰ ۰۸:۰۰-۱۰:۳۰	۰۸:۰۰-۱۰:۰۰
استراحت				استراحت			۱۰:۰۰-۱۰:۳۰
سخنرانی کلیدی (رئیس نشست: پروفسور مهدی افتخاری) Role of Intelligent systems in the Interpretation of Deep Neural Networks in a Post-pandemic World Prof. Saman K. Halgamuge, University of Melbourne, Australia							۱۰:۳۰-۱۱:۳۰
استراحت و ناهار							۱۱:۳۰-۱۳:۰۰
مجمع عمومی انجمن سیستمهای هوشمند ایران				نشست ویژه نظریه مفصل و کاربردهای آن		۱۳:۰۰-۱۴:۰۰ ۱۴:۰۰-۱۶:۰۰	۱۳:۰۰-۱۴:۰۰
نشست داده کاوی		نشست کاربردهای بهینه سازی و تحقیق در عملیات فازی					۱۴:۰۰-۱۶:۰۰
استراحت							۱۶:۰۰-۱۶:۱۵
سخنرانی کلیدی (رئیس نشست: پروفسور امیر دانشگر) Data-Driven Fuzzy Modeling Prof. Irina Perfilieva, University of Ostrava, Czech Republic							۱۶:۱۵-۱۷:۱۵
استراحت							۱۷:۱۵-۱۷:۳۰
		نشست یادگیری ماشین		نشست ویژه ساختارهای جبری منطقی (فازی) شماره ۲		۱۷:۳۰-۱۹:۳۰	۱۷:۳۰-۱۹:۳۰
سخنرانی کلیدی (رئیس نشست: پروفسور اسفندیار اسلامی) Some recent extensions of fuzzy integrals applied to the computational brain problem Prof. Humberto Bustince, Universidad Publica de Navarra, Spain							۱۹:۳۰-۲۰:۳۰

برنامه زمان بندی نهمین کنگره مشترک سیستم های فازی و هوشمند ایران

(روز سوم)

تالار ۳	تالار ۲	تالار ۱	جمعه ۱۳ اسفند ۱۴۰۰
	نشست بیوانفورماتیک	نشست ویژه بهینه سازی فازی شماره ۲	۰۸:۰۰-۱۰:۰۰
استراحت			۱۰:۰۰-۱۰:۳۰
سخنرانی کلیدی (رئیس نشست: پروفسور محمدرضا اکبرزاده توتونچی) ادراک بر پایه امکان: نامی بر بازشناسی الگوی اشکال پروفسور کامبیز بدیع، دانشگاه تهران، مرکز تحقیقات مخابرات ایران، ایران			۱۰:۳۰-۱۱:۳۰
استراحت و ناهار			۱۱:۳۰-۱۳:۰۰
نشست پردازش تصویر و بینایی ماشین	نشست سیستمهای فازی و کاربردها	نشست ویژه ساختارهای جبری منطقی (فازی) شماره ۳	۱۳:۰۰-۱۵:۰۰
استراحت			۱۵:۰۰-۱۵:۱۵
سخنرانی کلیدی (رئیس نشست: پروفسور رضا سعادت) مشتق توابع فازی و کاربرد آن در معادلات دیفرانسیل فازی دکتر علیرضا خواستار، دانشگاه تحصیلات تکمیلی علوم پایه، زنجان، ایران			۱۵:۱۵-۱۶:۱۵
استراحت			۱۶:۱۵-۱۶:۳۰
	نشست سیستمهای هوشمند و کاربردها شماره ۲	نشست آنالیز ریاضی و عددی فازی	۱۶:۳۰-۱۸:۳۰
استراحت			۱۸:۳۰-۱۸:۴۵
مراسم اختتامیه کنگره			۱۸:۴۵-۲۰:۴۵

پیام وزیر محترم عتف به نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

خدمت همه شرکت‌کنندگان در نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران، متشکل از بیستمین کنفرانس سیستم‌های فازی، هیجدهمین کنفرانس سیستم‌های هوشمند و پنجمین کنفرانس هوش جمعی و محاسبات تکاملی، مخصوصاً میهمانان گرانقدری که از خارج از کشور در این رویداد علمی شرکت نموده‌اند، عرض سلام و ادب و احترام دارم. امیدوارم همواره در زندگی علمی و در انجام پژوهش‌های ارزنده خود موفق باشید.

به نوبه خود عید سعید مبعث را نیز به همه پژوهشگران مخصوصاً شرکت‌کنندگان در این کنگره تبریک عرض می‌کنم و از خداوند متعال می‌خواهم که ما را از رهروان صدیق پیامبر رحمت، حضرت محمد (ص) قرار دهد. به عنوان وزیر علوم، تحقیقات و فناوری و به نمایندگی از جامعه علمی کشور، از اهتمام مجتمع آموزش عالی بوم و همچنین انجمن‌های علمی سیستم‌های فازی و سیستم‌های هوشمند، برای برگزاری این کنگره عظیم علمی که نشان از اهمیت نقش پژوهش در کشور و نیاز کشور به پژوهش‌های فناورانه است، تشکر و قدردانی می‌کنم. همچنین از انجمن‌های علمی و دانشگاه‌های مختلف که همواره حامیان اصلی رویدادهای علمی هستند نیز تشکر می‌کنم. همانگونه که همه شما مستحضرد ارتقای سطح علمی و توسعه جایگاه علمی ایران در جهان، در جهت رسیدن به قله‌های بلند علمی نیازمند ارزش دادن به تجربیات علمی و تحقیقات اندیشمندان در همه زمینه‌ها است و جامعه‌ای در این زمینه موفق تر عمل می‌کند که به علم و پژوهش همراه با فناوری‌های روز و اثر کاربردی آن در جامعه اهمیت بیشتری دهد. لذا فعالیت‌های پژوهشی مبتنی بر فناوری‌های نوین و مبنی بر نیاز جامعه و در جهت رفع مشکلات جامعه، می‌تواند بسیار مفید باشد و این نشست‌های علمی یکی از جایگاه‌های مهمی است که محققان داخل و خارج از کشور، آخرین دستاوردهای علمی خود را با دیگران به اشتراک گذاشته و با هم‌افزایی و هم‌اندیشی راهکارهای مناسبی را برای گذر از مشکلات جامعه و حل آنها ارائه دهند.

مطمئناً استفاده از فناوری‌های نوین و توسعه پژوهش‌های کاربردی در ارتقای نقش صنعت در جامعه بسیار موثر است و علوم بین رشته‌ای نظیر علوم مهندسی، سیستم‌های هوشمند، سیستم‌های فازی و پردازش تصویر و... می‌توانند برای رسیدن صنعت جامعه به اهداف متعالی خود کمک کند. لذا از همه اندیشمندان، اساتید و دانشجویان عزیز می‌خواهم با تلاش و کوشش خود و با همکاری و همفکری یکدیگر در این زمینه کوشا باشند و با ارائه تحقیقات کاربردی مناسب سبب پیشرفت کشور عزیزمان شوند.

در پایان لازم می‌دانم از همه اساتید، پژوهشگران، دبیران علمی، اجرایی و عوامل برگزاری این کنگره قدردانی نموده و موفقیت همه شما عزیزان را آرزو نمایم.

محمدعلی زلفی گل

وزیر علوم، تحقیقات و فناوری

پیام رئیس کنگره

با نام و یاد خدای بزرگ و مهربان و با عرض سلام و ادب محضر با سعادت اندیشمندان عزیز، میهمانان گرامی و سخنرانان ارجمند، از این که دعوت ما را پذیرفته‌اید و با حضور خود برغنای این همایش افزوده‌اید صمیمانه سپاسگزاری نموده و نسبت به همه شما بزرگواران ادای احترام می‌نمایم.

جا دارد از انجمن سیستم‌های فازی ایران و انجمن سیستم‌های هوشمند ایران که امتیاز میزبانی این کنگره را به مجتمع آموزش عالی بم اهدا نموده و برای برگزاری آن نیز حمایت‌های خود را از نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران دریغ نورزیدند، تشکر و قدردانی نمایم. همچنین از معاونت محترم پژوهشی دانشگاه شهید باهنر کرمان، سایر انجمن‌های علمی و سایر حامیان محترم که ما را در برگزاری این کنگره همراهی نموده‌اند نهایت تشکر را دارم.

از همکاران عزیز و اساتید گرامی خودم که در قالب کمیته‌های علمی و اجرایی تلاش شبانه‌روزی داشتند، همچنین از تمامی دست اندرکاران این کنگره، به ویژه دبیران محترم علمی کنگره جناب آقای دکتر ماشین‌چی و سرکارخانم دکتر فروزش و دبیر اجرایی کنگره جناب آقای دکتر عبدی، که در چند ماه اخیر شاهد تلاش‌ها و نگرانی‌های این عزیزان در امر برگزاری این کنگره بوده‌ام قدردانی می‌کنم.

خداوند را شاکرم که به من توفیق داد تا در این کنگره، در کنار شما عزیزان باشم. از اینکه به واسطه شیوع بیماری کرونا، لذت هم‌صحبتی و دیدار حضوری شما عزیزان از ما سلب شد، بسیار متأسفیم، اما باز بر خودمان می‌بالیم که مجتمع آموزش عالی بم، میزبان فرهیختگان و اندیشمندان بزرگی همچون شما عزیزان است و این فرصت را برای توسعه و پیشرفت علمی غنیمت می‌دانیم. در سال گذشته، دانشمندان و اساتید گرانقدری را از دست دادیم، بر خود لازم می‌دانم یاد و خاطره این اسطوره‌های فقید کشور را گرامی بدارم. زنده یاد پروفسور بهرام صادقیور گیلده، زنده یاد پروفسور محمود لشکری زاده بمی و... که در سال جاری جهان فانی را وداع گفته و به جاودانگی پیوستند. یاد این اساتید فرزانه هرگز از دلها و نامشان هرگز از گستره علم و دانش این مرز و بوم محو نخواهد شد.

به اطلاع شما عزیزان می‌رسانم که مجتمع آموزش عالی بم با قدمتی بیش از سی سال، همواره در جهت اعتلای علم و پژوهش در جامعه کوشا بوده است. دوازدهمین کنفرانس سیستم‌های هوشمند ایران در سال ۱۳۹۲ با همکاری انجمن سیستم‌های هوشمند ایران، مؤسسه مهندسی برق و الکترونیک و گروه مهندسی برق دانشگاه شهید باهنر کرمان با هدف ایجاد زمینه‌ای جهت گردهمایی پژوهشگران و صاحب‌نظران و انتشار یافته‌های نو در زمینه سیستم‌های هوشمند و کاربردهای آن در این دانشگاه برگزار شد. همچنین این دانشگاه، با حمایت انجمن سیستم‌های هوشمند ایران و همکاری دانشگاه شهید باهنر کرمان و دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته کرمان پذیرای محققان و توسعه‌دهندگان در زمینه‌های آکادمیک و صنعتی در اولین و سومین کنفرانس محاسبات تکاملی و هوش جمعی در سال‌های ۱۳۹۴ و ۱۳۹۶ بود.

در سال جاری نیز با حمایت انجمن‌های علمی و دانشگاه‌های استان کرمان، نهمین کنگره سیستم‌های هوشمند و فازی را که شامل بیستمین کنفرانس سیستم‌های فازی، هیجدهمین کنفرانس سیستم‌های هوشمند و پنجمین کنفرانس محاسبات تکاملی و هوش جمعی است، میزبان هستیم. مطمئناً این کنگره، فرصتی مغتنم برای ارائه دیدگاه‌ها و نقطه نظرات خبرگان و اهل فن در حوزه علم فازی و کاربرد آن در علوم مختلف می‌باشد.

مجدداً در پایان لازم می‌دانم، از همه عزیزانی که با حمایت‌های مالی و معنوی خود، مجتمع را در برگزاری هرچه باشکوه‌تر این کنگره، یاری کرده‌اند نهایت تشکر و قدردانی را داشته باشم. همچنین از هیأت رئیسه مجتمع، کمیته اجرایی، کمیته علمی و هیأت داوران، اعضای محترم هیأت علمی و کادر اداری مجتمع و همه عزیزانی که در برنامه‌ریزی و اجرای این همایش نقش داشته‌اند تشکر و

سپاسگزاری می‌کنم. از اساتید ارجمند دانشگاه‌ها و دانشجویان عزیز به خاطر حضور و مشارکت فعال آنها در کنگره و از پژوهشگران، صاحب نظران و بزرگوارانی که با ارائه مقاله یا سخنرانی و با مشارکت در پنل‌ها به بار علمی مباحث مطرح شده می‌افزایند، قدردانی و تشکر می‌نمایم.

امیدوارم که علیرغم محدودیت‌های ناشی از شرایط موجود به دلیل بیماری کووید ۱۹، برگزاری این کنگره مکان مناسبی برای ارائه دستاوردهای علمی پژوهشگران در زمینه‌های گوناگون باشد.

با آرزوی توفیق الهی

محمدعلی نورالهی

رئیس کنگره و رئیس مجتمع آموزش عالی بم

پیام رئیس انجمن سیستم‌های فازی ایران

در ابتدا، برگزاری نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران در مجتمع آموزش عالی بزم را به فال نیک گرفته و از تمامی همکاران ارجمند و عزیزی که در برگزاری این کنگره تلاش نموده و کمک شایانی کرده‌اند، به ویژه همکاران خوبم و مدیران پرتلاش و ریاست محترم مجتمع، تشکر و قدردانی می‌نمایم که علیرغم محدودیت‌های ناشی از بیماری کووید ۱۹ و حتی تنگناهای مالی، اراده‌ای قوی و همتی والا را بکار بستند تا این کنگره در بهترین شکل ممکن برگزار گردد و اما ذکر نکاتی را لازم می‌دانم:

براساس اطلاعات پایگاه علمی اسکوپوس، ج.ا.ایران در ده سال گذشته بعد از کشورهای چین و هند، رتبه سوم در تولید مقالات در حوزه سیستم‌های فازی را کسب نموده است و این در حالی است که آمریکا رتبه چهارم را دارد. ضمن آنکه از نظر جمعیت، کشور ما ۱٪ جمعیت جهان را دارد اما در سال ۲۰۲۱ موفق شده که ۱۳٪ تولید مقالات حوزه سیستم‌های فازی را به خود اختصاص دهد و این در حالی است که تولید مقالات در همه‌ی زمینه‌های علمی کشورمان ۲٪ بوده است. بنابراین نقش نظریه سیستم‌های فازی در تولید علم و مقاله معتبر علمی برای ارتقاء جایگاه علمی کشورمان در منطقه و جهان بسیار موثر بوده است.

نکته قابل توجه آن است که علیرغم رشد ۵۰٪ مقالات علوم انسانی در جهان ظرف ۵ سال گذشته، در کشورمان بیش از ۲۰۰٪ رشد داشته است و همچنین در زمینه بهره‌گیری از نظریه سیستم‌های فازی در علم مواد ۱۰۰٪، کشاورزی و زیست فناوری ۵۰٪ و در داروسازی ۲۰٪ در کشور رشد داشته ایم. لذا اطلاعات فوق نشان می‌دهد که محبوبیت بکارگیری منطق و نظریه سیستم‌های فازی در بین دانشمندان و محققین کشورمان رو به افزایش بوده و در حوزه‌های نوین از جمله سیستم‌های هوشمند، هوش مصنوعی، علوم انسانی، علوم پایه و علوم مهندسی بکار گرفته شده است.

شاید روزی که پروفسور لطفی زاده منطق فازی را برای اولین بار ارائه می‌نمود خودش نیز باور نداشت که این نظریه پس از گذشت حدود شش دهه هنوز یکی از جذاب‌ترین موضوعات تحقیقاتی در دنیا باشد که بسیاری از محققان را مجذوب خود کرده است. آنچه امروزه منطق فازی و سیستم‌های هوشمند را بعد از این همه سال هنوز در جاذبه تحقیقاتی و محبوب علمی و پژوهشی قرار داده است، عجین بودن آن با نیازهای روز جامعه، ذهن بشر و همچنین کاربردهای آن در حوزه‌های مختلف مرتبط با زندگی بشر می‌باشد ولی بایستی اذعان نمود که محققین کشور ما کمتر به جنبه‌های کاربردی آن پرداخته‌اند. بنابراین بایستی به گونه‌ای در نظام علم و فناوری کشور برنامه‌ریزی گردد که موجبات جاذبه‌های پژوهشی-کاربردی و رفع نیازهای کشور براساس سیستم‌های فازی و هوشمند فراهم آید.

براین اساس پیشنهاد می‌شود که طی مطالبه علمی-نخبگانی در این حوزه، یک ستاد علمی به نام

"ستاد علم و فناوری سیستم‌های فازی و هوشمند"

همانند سایر ستادهای علمی ذیل معاونت محترم علم و فناوری ریاست جمهوری ایجاد گردد.

لذا به عنوان رئیس انجمن سیستم‌های فازی ایران از تمامی همکاران و محققین عزیزی که در این حوزه فعالیت دارند، درخواست می‌نمایم که کمک و راهنمایی فرمایند تا بتوانیم شرایط لازم را برای راه اندازی این ستاد فراهم آوریم.

با آرزوی اعتلا و خودکفایی علمی و فناوری کشورمان و تشکر مجدد از دست اندرکاران برگزاری این کنگره مطلب را به پایان می‌برم.

با آرزوی توفیق الهی

محمد مهدی زاهدی

رئیس انجمن سیستم‌های فازی ایران

پیام نائب رئیس انجمن سیستم‌های هوشمند ایران

مدل‌سازی و آنالیز پدیده‌های بسیار پیچیده امروزی، نیازمند شیوه‌های نوین علمی و دقیق بوده و دیگر نمی‌توان تنها به شیوه‌های علمی که صرفاً جویای جواب‌های تحلیلی و صرفاً دقیق هستند بسنده نمود، چون روش‌های تحلیلی در گذشته از حل بسیاری از مسائل عاجز مانده‌اند. محاسبات نرم به روش جدید محاسباتی در ریاضیات، علوم کامپیوتر، یادگیری ماشین، هوش مصنوعی و بسیاری از زمینه‌های کاربردی دیگر مرتبط می‌گردد. این روش‌ها امکان آن را فراهم می‌سازند تا بتوان نقصان روش‌های تحلیلی را با مدل‌سازی و مطالعه پدیده‌های بسیار پیچیده دنبال نموده و با نتایج آن به رفع این نقصان پرداخت.

دانش امروزی در نهایت شایستگی، بیانگر اشتراک روزافزون علوم مختلف می‌باشد. لذا در بحث محاسبات نرم نمی‌توان ارتباط تنگاتنگ ریاضی، آمار، کامپیوتر، مهندسی برق، علوم انسانی و سایر رشته‌های مرتبط را منکر شد. بطوریکه می‌توان گفت که هر چقدر این شاخه‌های علوم و فنون به هم مرتبط می‌گردند، زیبایی و جذابیت آنها بیشتر می‌شود و دیگر نمی‌توان به شیوه سنتی دیوار محکمی بین آنها ایجاد نمود، بلکه باید دیوار ایجاد شده قدیمی را کم‌کم از بین برده و در مرز مشترک رشته‌های علوم و مهندسی گام برداشت و بی‌گمان راه مباحث سیستم‌های فازی و هوشمند نیز این چنین می‌باشد. امروزه شاهد آن هستیم که ریاضیات و منطق فازی در علوم انسانی مانند روانشناسی و جامعه‌شناسی و مدیریت، که بیشتر مفاهیم نادقیق هستند، نیز کاربردهای فراوانی یافته است.

انجمن سیستم‌های هوشمند نیز در این راستا طی سال‌های اخیر با هدف گسترش و پیشبرد و ارتقای علمی سیستم‌های هوشمند و توسعه کیفی نیروهای متخصص و بهبود بخشیدن به امور آموزش و پژوهش در زمینه‌های مربوطه، موفقیت‌های بارزی بدست آورده است که از جمله می‌توان به هم‌افزایی علوم مرتبط با این حوزه از طریق برگزاری کنفرانس‌های متعدد ملی و بین‌المللی، انتشار کتاب‌های جدید در حوزه نظریه نایقینی و ارتباط موثر با پیشگامان و نوآوران این علم در سراسر دنیا اشاره نمود.

در پایان، انجمن سیستم‌های هوشمند ضمن خرسندی و حمایت از برگزاری نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران که شامل سه کنفرانس موفق ادواری: (۱) بیستمین کنفرانس سیستم‌های فازی ایران (ICFS2022)، (۲) هجدهمین کنفرانس سیستم‌های هوشمند (CIS2022) و (۳) پنجمین کنفرانس هوش جمعی و محاسبات تکاملی (CSIEC2022) می‌باشد، آرزومند است برگزاری این کنگره در این دوره و انشاءالله برگزاری حضوری آن پس از سپری شدن پاندمی کرونا در دوره‌های بعدی، مورد توجه بیش از پیش تمامی اساتید و محققین شاخه‌های مرتبط با محورهای مورد اشاره و توسعه مباحث میان‌رشته‌ای قرار گیرد تا بدین وسیله زمینه‌ای مناسب برای رشد، بحث و تبادل علمی فراهم گردد.

با احترام

نائب رئیس انجمن سیستم‌های هوشمند ایران - دکتر بهروز فتحی

اسفند ۱۴۰۰

پیام دبیران علمی کنگره

خداوند بزرگ را سپاسگزاریم که ما را یاری نمود تا نهمین کنگره مشترک سیستم های فازی و هوشمند ایران با همکاری انجمن سیستم های فازی ایران و انجمن سیستم های هوشمند ایران، انجمن بین المللی فازی و بسیاری از انجمن های علمی مرتبط به محورهای کنگره و کوشش همکاران، دانشجویان، کارکنان و حمایت های مسئولین محترم مجتمع آموزشی عالی بم در روز های ۱۱ الی ۱۳ اسفند ماه برگزار نماییم.

امید است که حاصل مساعی و تلاش برگزارکنندگان بتواند رضایت خاطر اساتید و دانشجویان شرکت کننده گرامی را فراهم نماید و موجب گردد تا آنان از منظر لطف و عنایت، کاستی ها و کمبودها را به دیده ای اغماض بنگرند. برنامه ریزی برای برگزاری کنگره از مهرماه سال ۱۳۹۹ آغاز گردید و با توجه به استقبال قابل توجه اساتید، دانشجویان و محققین کشور به دلیل محدودیت زمان ناچار شدیم تعدادی از مقالات که از کیفیت بیشتری برخوردار بودند برای ارائه در کنگره انتخاب کنیم.

متقاضیان شرکت در همایش تعداد حدود ۲۶۰ کد مقاله به دبیرخانه کنفرانس ارسال نمودند که کمیته علمی کنگره پس از پالایش اولیه و داوری های تخصصی دقیق توسط حداقل ۲ داور برای هر مقاله، تعداد ۱۲۰ مقاله به صورت ارائه سخنرانی و تعداد ۲۰ مقاله به صورت پوستر پذیرش نمودند. برای هر سخنرانی تخصصی ۱۵ دقیقه زمان ارائه و ۵ دقیقه زمان پرسش و پاسخ در نظر گرفته شده است، که با توجه به برگزاری مجازی کنگره به صورت فیلم از نویسندگان دریافت گردید.

از ۸ استاد مطرح از داخل و خارج از کشور برای سخنرانی کلیدی دعوت شده است که به صورت مجازی ارائه می شوند. زمان سخنرانی های عمومی ۴۵ دقیقه ارائه همراه با ۱۵ دقیقه پرسش و پاسخ در نظر گرفته شده است. ۵ کارگاه به مدت ۳ تا ۴ ساعت به صورت مجازی برگزار می شود، به علاوه مجمع عمومی دو انجمن سیستم های فازی و سیستم های هوشمند از دیگر برنامه های مجتمع است. برگزاری این کنگره مدیون هم فکری همیاری و همکاری افراد بسیاری به ویژه تلاش های بی شائبه دبیر اجرای این کنگره است همچنین در این جا لازم است تا مراتب قدردانی و امتنان خود را از هیأت رئیسه محترم مجتمع آموزش عالی بم و انجمن سیستم های فازی ایران و انجمن سیستم های هوشمند ایران و سایر حمایت کنندگانی که نام آنها در این کتابچه ذکر گردیده است، ابزار نماییم. در پایان، از اعضای محترم کمیته علمی کنگره بویژه جناب آقای دکتر آرشام برومند سعید و سرکار خانم دکتر مرجان کوچکی رفسنجانی که ما را در داوری مقالات به صورت شبانه روزی همراهی کردند، قدردانی می شود.

با کمال احترام

دکتر ماشاءالله ماشینچی، دکتر فرشته فروزش

دبیران علمی کنگره

پیام دبیر اجرایی کنگره

ستایش و سپاس بی کران پروردگار یکتا را که به ما هستی بخشید و به طریق علم و دانش رهنمونمان شد و به هم‌نشینی رهروان علم و دانش مفتخرمان نمود تا از این اقیانوس بی کران بهره‌مند گردیم.

هر آن کس خدمت جانان به جان کرد به گیتی نام خود را جاودان کرد

مجتمع آموزش عالی بم مفتخر است، پس از برگزاری موفق دوازدهمین کنفرانس سیستم‌های هوشمند ایران در سال ۱۳۹۲، اولین و سومین کنفرانس محاسبات تکاملی و هوش جمعی در سال‌های ۱۳۹۴ و ۱۳۹۶، هم‌اکنون میزبان نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران (CFIS2022) با حمایت انجمن سیستم‌های فازی ایران و انجمن سیستم‌های هوشمند ایران از تاریخ یازدهم تا سیزدهم اسفند ماه ۱۴۰۰ باشد. این کنگره شامل سه کنفرانس موفق ادواری: بیستمین کنفرانس سیستم‌های فازی ایران (ICFS2022)، هجدهمین کنفرانس سیستم‌های هوشمند ایران (CIS2022) و پنجمین کنفرانس هوش جمعی و محاسبات تکاملی (CSIEC2022) می‌باشد.

در کمال احترام و افتخار بدینوسیله مراتب سپاسگزاری خویش از اعضای محترم هیأت علمی، دانشجویان و پژوهشگران حوزه سیستم‌های فازی و هوشمند که با ارائه مقالات وزین و ارزشمند خویش بر غنای این کنگره افزودند، ابراز می‌دارم.

پس از اعلام موافقت و حمایت انجمن سیستم‌های فازی ایران و انجمن سیستم‌های هوشمند ایران و اخذ مجوز در سال ۱۳۹۹، امور اجرایی کنگره از ابتدای فروردین ۱۴۰۰ آغاز گردید. در شرایط مخاطره‌آمیز بیماری کووید ۱۹، بنا بر تصمیم شورای سیاستگذاری، مجازی بودن کنگره در دستور کار کمیته اجرایی قرار گرفت و بر همین اساس کلیه امور از جمله آماده‌سازی سایت و تجهیز آن برای میزبانی دریافت مقالات و سایر امور طبق برنامه زمانبندی انجام شد. ۶۲ عضو کمیته علمی و حدود ۱۰۰ داور با تخصص‌های گوناگون از سراسر کشور و جهان در تبیین سیاست‌های کلی کنگره و اجرای آنها شرکت داشتند. جهت برگزاری هرچه باشکوه‌تر و کاربردی‌تر این کنگره، ۸ سخنرانی کلیدی و ۵ کارگاه تخصصی برای ارائه در کنگره انتخاب گردید.

از ریاست محترم مجتمع به عنوان رییس کنگره، معاونت محترم آموزشی و پژوهشی و همچنین مدیریت محترم امور پژوهشی مجتمع که همواره با حمایت خویش ما را در برگزاری کنگره همراهی کردند قدردانی می‌نمایم. همچنین از حمایت‌های مادی و معنوی معاونت پژوهشی دانشگاه شهید باهنر کرمان و حامیان دیگر کنگره که با دلگرمی حمایت خویش ما را برای اجرای موفق کنگره رهنمون شدند.

جای مسرت است که تجربه‌ای بس گرانقدر در کنار اعضای محترم شورای سیاستگذاری کنگره بالاخص استاد مهربان و بزرگوار، جناب آقای پروفسور ماشالله ماشینچی به عنوان دبیر کمیته علمی داشتیم، که با انتقال بی‌کم و کاست و بی‌منت بسیاری از تجارب سال‌ها حضور خویش در عرصه‌های علمی و پژوهشی، ملی و بین‌المللی به کمیته اجرایی کنگره یاری‌گر ما بودند. همچنین در اینجا فرصتی مغتنم است تا از خانم دکتر فرشته فروزش به عنوان دبیر علمی کنگره، آقای دکتر آرشام برومند سعید و خانم دکتر مرجان کوچکی رفسنجانی که در طول یکسال گذشته همواره یاری‌گر ما در تصمیم‌گیری‌ها و بالاخص مدیریت بخش داوری مقالات بوده‌اند، قدردانی نمایم.

شایان ذکر است، تعداد ۲۶۰ کد مقاله در سایت کنگره دریافت شد و پس از پالایش اولیه مقالات و داوری آنها توسط حداقل دو داور متخصص، ۱۲۰ مقاله به صورت سخنرانی و ۲۰ مقاله به صورت پوستر جهت ارائه برگزیده شد. مقالات پذیرفته شده پس از انجام اصلاحات لازم و ارسال ۵۸ مقاله انگلیسی جهت نمایه IEEE، در قالب ۲۱ نشست تخصصی و ۲ پنل پوستر برنامه‌ریزی گردیده است. با عنایت به گستردگی محل اقامت ارائه کنندگان و پژوهشگران از داخل و خارج از کشور، برای پیشگیری از مشکلات احتمالی قطعی اینترنت و سرعت نامناسب ارائه آنلاین، برای هر سخنرانی ۱۵ دقیقه و برای هر پوستر ۵ دقیقه فیلم ضبط شده از نویسندگان دریافت گردید و ۵ دقیقه زمان پرسش و پاسخ پس از پخش فیلم به هر ارائه اختصاص یافت.

بر خود فرض می‌دانم از همکاران فرزانه، دلسوز و پرتلاش کمیته اجرایی کنگره از جمله خانم مهندس مینا میرحسینی، خانم مهندس فاطمه بارانی، خانم دکتر طیبه عسگری جواران، خانم دکتر محدثه سلیمانپور، خانم مهندس مهدیه مظفری، آقای دکتر سید حسین مهدوی، خانم مهندس آناهیتا امیرشجاعی، خانم دکتر زینب خاتون پورطاهری، خانم مهتا بدرود، آقای مهندس عبدالحمید بحرالعلوم و آقای مهندس حسن حدیدی خانم نفیسه حدیدی و سایر همکاران، اساتید مدعو دانشگاه و دانشجویان گرامی که در برگزاری هر چه باشکوه‌تر این رویداد ارزشمند، همکار و همراه ما بوده‌اند و بنده شاهد تلاش‌های شبانه روزی ایشان بوده‌ام، خاضعانه تقدیر و سپاسگزاری نمایم.

قابل ذکر است، مجوز نمایه ISC با همکاری و تلاش آقای دکتر بیژن امامی‌پور، مدیر محترم آموزشی و پژوهشی مجتمع و مجوز نمایه IEEE، نیز نتیجه زحمات سرکار خانم مهندس فاطمه بارانی بوده است که بدینوسیله قدردان زحمات ایشان هستم.

بر خود وظیفه می‌دانم یاد استاد فقید، مرحوم پروفسور بهرام صادقیور گیلده را به پاس خدمات ارزشمند و ماندگار ایشان گرامی بدارم که اسوه اخلاق و الگوی بنده و بسیاری از دانشجویان ایشان در ره کسب علم و دانش بوده‌اند.

امید است تلاش‌های شبانه‌روزی همکاران کمیته اجرایی موجب رضایت شرکت‌کنندگان محترم و حامیان کنگره را جلب نموده باشد و قدمی هرچند کوتاه جهت اعتلای اهداف علمی و پژوهشی کشورمان برداشته باشیم.

با تقدیم احترام

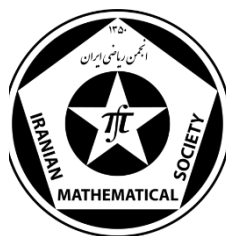
معراج عبدی

دبیر کمیته اجرایی کنگره

حامیان علمی و صنعتی کنگره

۱. مجتمع آموزش عالی بزم
۲. دانشگاه شهید باهنر کرمان
۳. دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته
۴. دانشگاه مازندران
۵. مجتمع آموزش عالی زرنند
۶. انجمن مهندسين برق و الكترونيك بخش ايران (IEEE)
۷. پایگاه استنادی علوم جهان اسلام (ISC)
۸. سیویلیکا (CIVILICA)
۹. انجمن سیستم هوشمند ایران
۱۰. انجمن سیستم‌های فازی ایران
۱۱. انجمن بین‌المللی سیستم‌های فازی (IFSA)
۱۲. انجمن ریاضی ایران
۱۳. انجمن آمار ایران
۱۴. اتحادیه انجمن‌های ایرانی علوم ریاضی
۱۵. انجمن بینایی ماشین و پردازش تصویر ایران
۱۶. انجمن رمز ایران
۱۷. انجمن ایرانی تحقیق در عملیات
۱۸. انجمن مهندسی صنایع ایران
۱۹. قطب علمی رایانش نرم و پردازش هوشمند اطلاعات
۲۰. مرکز منطقه‌ای اطلاع‌رسانی علوم و فناوری
۲۱. کنفرانس یاب
۲۲. شرکت عمران منطقه ویژه اقتصادی ارگ جدید
۲۳. شرکت توزیع نیروی برق جنوب استان کرمان
۲۴. شرکت خدمات مخابراتی ارگ جدید

حامیان علمی و صنعتی کنگره



اعضای کمیته علمی نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

نام و نام خانوادگی عضو کمیته علمی	دانشگاه/سازمان وابسته
دکتر حامد احمد زاده	دانشگاه سیستان و بلوچستان
پروفسور اسفندیار اسلامی	دانشگاه شهید باهنر کرمان
دکتر مهدی افتخاری	دانشگاه شهید باهنر کرمان
پروفسور محمدرضا اکبر زاده توتونچی	دانشگاه فردوسی مشهد
پروفسور توفیق الله‌ویرانلو	دانشگاه باهجه شهیر استانبول
پروفسور محمد امینی	دانشگاه فردوسی مشهد
پروفسور رجبعلی برزویی	دانشگاه شهید بهشتی
دکتر محمود بخشی	دانشگاه بجنورد
پروفسور کامبیز بدیع	دانشگاه تهران - مرکز تحقیقات مخابرات ایران
دکتر آرشام برومند سعید	دانشگاه شهید باهنر کرمان
دکتر میر محسن پدram	دانشگاه خوارزمی
پروفسور سید علی ترابی	دانشگاه تهران
پروفسور محمد تشنه لب	دانشگاه خواجه نصیرالدین طوسی
پروفسور محمد جمشیدی	دانشگاه سن انتونیو تگزاس
دکتر علیرضا خواستار	دانشگاه تحصیلات تکمیلی علوم پایه زنجان
پروفسور امیر دانشگر	دانشگاه صنعتی شریف
پروفسور ولی درهمی	دانشگاه یزد
پروفسور محمد مهدی زاهدی	دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته
دکتر رضا سعادت	دانشگاه علم و صنعت ایران
پروفسور سید محمود طاهری	دانشگاه تهران
دکتر محسن عارفی	دانشگاه بیرجند
پروفسور سعید عباسبندی	دانشگاه بین المللی امام خمینی
پروفسور بهروز فتحی واجارگاه	دانشگاه گیلان
دکتر فرشته فروزش	مجمع آموزش عالی بم
دکتر محمود فضلعلی	دانشگاه شهید بهشتی
دکتر فاطمه کوچکی نژاد	مجله سیستم‌های فازی ایران
دکتر فرشید کی نیا	دانشگاه تحصیلات تکمیلی صنعتی کرمان
پروفسور ماشاالله ماشین چی	دانشگاه شهید باهنر کرمان
دکتر محمدرضا ماشین چی	دانشگاه پیام نور واحد کرمان
دکتر رادکو مسیار	دانشگاه صنعتی اسلواکی
پروفسور حسن میش مست نهی	دانشگاه سیستان و بلوچستان

پروفسور حسین نظام آبادی پور	دانشگاه شهید باهنر کرمان
دکتر محمد علی یعقوبی	دانشگاه شهید باهنر کرمان
پروفسور سید ناصر حسینی	دانشگاه شهید باهنر کرمان، ایران
دکتر مرجان کوچکی رفسنجانی	دانشگاه شهید باهنر کرمان، ایران
پروفسور ولدون لودیک	دانشگاه کلورادو دنور، آمریکا
پروفسور محمدباقر منهج	دانشگاه صنعتی امیر کبیر، ایران
پروفسور جان موردسون	دانشگاه کریتون، آمریکا
پروفسور سوزانا مونتس	دانشگاه اوپیدو، اسپانیا
پروفسور علی وحیدیان کامیاد	دانشگاه فردوسی مشهد، ایران
پروفسور خوزه لوئیس وردگای	دانشگاه دو گرانادا، اسپانیا
پروفسور احسان کایا	دانشگاه فنی یلدوز، ترکیه
دکتر علیرضا عسکرزاده	دانشگاه تحصیلات تکمیلی صنعتی کرمان، ایران
دکتر زشوی شو	دانشگاه سیچوان، چین
دکتر ایوب شیخی	دانشگاه شهید باهنر کرمان، ایران
پروفسور فو گوی شی	موسسه فناوری پکن، چین
دکتر آرش شریفی	دانشگاه آزاد اسلامی - واحد علوم و تحقیقات تهران - ایران
پروفسور الکساندر سوستاک	دانشگاه لتونی، لتونی
پروفسور ناصر ساداتی	دانشگاه صنعتی شریف، ایران
پروفسور دان رالسکو	دانشگاه سین سیناتی، آمریکا
پروفسور آنتونیو دی نولا	دانشگاه سالرنو، ایتالیا
پروفسور برنارد د باتس	دانشگاه گنت، بلژیک
پروفسور گونتر جاگر	دانشگاه علوم کاربردی استرالسوند، آلمان
دکتر عباس پرچمی	دانشگاه شهید باهنر کرمان، ایران
دکتر اکبر پاد	دانشگاه بجنورد، ایران
پروفسور غلامرضا محتشمی برزادران	دانشگاه فردوسی مشهد، ایران
پروفسور میشل باسینسکی	دانشگاه سیلیسیا، لهستان

اعضای کمیته اجرایی نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

سمت اجرایی در کنگره	نام و نام خانوادگی
رئیس کنگره	دکتر محمدعلی نوراللهی
دبیران علمی کنگره	دکتر ماشالله ماشین‌چی
دبیر اجرایی کنگره	دکتر معراج عبدی
دبیر کمیته انفورماتیک	مینا میرحسینی
اعضای کمیته انفورماتیک	عبدالحمید بحر العلوم
دبیر کمیته انتشارات	فاطمه بارانی
اعضای کمیته انتشارات	نقیسه حدیدی حسن حدیدی دکتر بیژن امامی‌پور
دبیر کمیته ثبت‌نام	دکتر طیبه عسکری جواران
اعضای کمیته ثبت‌نام	دکتر زینب خاتون پورطاهری حسن حدیدی
اعضای دبیرخانه	دکتر سید حسین مهدوی مهتا بدرود
دبیر کمیته روابط عمومی	دکتر محمد محمودی‌راد
اعضای کمیته روابط عمومی	مجتبی صادقیان مهدیه فرزانه نیا
دبیر کمیته مالی و پشتیبانی	دکتر حمزه دهقانی
اعضای کمیته مالی و پشتیبانی	حسن عظیمی حسین عظیمی
دبیر کمیته ارتباط با صنعت و برگزاری کارگاه‌ها	دکتر محدثه سلیمان‌پور
اعضای کمیته ارتباط با صنعت و برگزاری کارگاه‌ها	دکتر هادی فرزانه سلیمه قنبری دکتر مریم ضیاءآبادی محمدحسین فروزان‌فر
دبیر کمیته روابط بین‌الملل	آناهیتا امیرشجاعی

اعضای شورای سیاستگذاری نهمین نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

نام و نام خانوادگی	دانشگاه/سازمان وابسته
پروفسور اسفندیار اسلامی	دانشگاه شهید باهنر کرمان
پروفسور محمدرضا اکبرزاده توتونچی	دانشگاه فردوسی مشهد
دکتر آرشام برومند سعید	دانشگاه شهید باهنر کرمان
پروفسور سید محمود طاهری	دانشگاه تهران
پروفسور سعید عباسبندی	دانشگاه بین المللی امام خمینی (ره)
دکتر معراج عبدی	مجتمع آموزش عالی بم
پروفسور بهروز فتحی واجارگاه	دانشگاه گیلان
دکتر فرشته فروزش	مجتمع آموزش عالی بم
پروفسور ماشالله ماشین چی	دانشگاه شهید باهنر کرمان
پروفسور حسین نظام آبادی پور	دانشگاه شهید باهنر کرمان
دکتر محمدعلی نوراللهی	مجتمع آموزش عالی بم

اعضای کمیته داوران نهمین نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

نام و نام خانوادگی	دانشگاه محل خدمت
سپهر ابراهیمی	دانشگاه یزد
خدیدجه ابوالپور	دانشگاه آزاد اسلامی
حامد احمدزاده	دانشگاه سیستان و بلوچستان
صادق اسکندری	دانشگاه گیلان
اسفندیار اسلامی	دانشگاه شهید باهنر کرمان
مهدی افتخاری	دانشگاه شهید باهنر کرمان
مجتبی افتخاری	دانشگاه شهید باهنر کرمان
محبوبه افضلی	دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته
مهدی الله دادی	دانشگاه سیستان و بلوچستان
توفیق الهویرانلو	دانشگاه بین المللی باغچه سهیر، ترکیه
نسبیه امامی	دانشگاه کوثر بجنورد
محمد امینی	دانشگاه فردوسی مشهد
محمد ایزدی خالق آبادی	دانشگاه شهید باهنر کرمان
فاطمه بارانی	مجتمع آموزش عالی بم
محمود بخشی	دانشگاه بجنورد
مجتبی برخوردار	دانشگاه شهید باهنر کرمان
ارشام برومند سعید	دانشگاه شهید باهنر کرمان
اکبر پاد	دانشگاه بجنورد
عباس پرچمی	دانشگاه شهید باهنر کرمان
زینب خاتون پورطاهری	مجتمع آموزش عالی بم
حامد تبریزی	دانشگاه تبریز
سید علی ترابی	دانشگاه تهران
محمد تشنه‌لب	دانشگاه خواجه نصیرالدین طوسی
ناهید جعفری	دانشگاه شهید باهنر کرمان
بهرنگ چابکی	دانشگاه ولی عصر (عج) رفسنجان
حمزه حرف شنو	دانشگاه فرهنگیان کاشان
مهسا حسن شاهی	دانشگاه شهید باهنر کرمان
سوده حسینی	دانشگاه شهید باهنر کرمان
فرهاد حمیدی	دانشگاه سیستان و بلوچستان
محمد حمیدی	دانشگاه پیام نور کاشان
محسن خسروی	دانشگاه شهید باهنر کرمان
عمید خطیبی	دانشگاه آزاد اسلامی
محمدباقر دولتشاهی	دانشگاه لرستان

جعفر رزم آرا	دانشگاه تبریز
اکبر رضایی	دانشگاه پیام نور کرمان
میلاد ریاحی	دانشگاه شهید باهنر کرمان
عظیم ریواز	دانشگاه شهید باهنر کرمان
جعفر زنجانی	دانشگاه فرهنگیان زنجان
رضا سعادت	دانشگاه علم و صنعت ایران
محدثه سلیمان پور مقدم	مجتمع آموزش عالی بم
صادق سلیمانی	دانشگاه کردستان
علی شکیا	دانشگاه ولی عصر(عج) رفسنجان
ایوب شیخی	دانشگاه شهید باهنر کرمان
امیر صباغ ملاحسینی	دانشگاه آزاد اسلامی
سید محمود طاهری	دانشگاه تهران
سعیده ظهیری	مرکز آموزش عالی اقلید
سعید عباسبندی	دانشگاه بین المللی امام خمینی (ره)
امیر عبداللہی	دانشگاه شهید باهنر کرمان
معراج عبدی	مجتمع آموزش عالی بم
علیرضا عرب پور	دانشگاه شهید باهنر کرمان
طیبه عسکری جواران	مجتمع آموزش عالی بم
محمد علایی	دانشگاه ولی عصر(عج) رفسنجان
لعیا علی احمدی پور	دانشگاه شهید باهنر کرمان
آزاده السادات عمرانی زرنندی	دانشگاه شهید باهنر کرمان
حمیده فاطمی دخت	دانشگاه شهید باهنر کرمان
فخرالسادات فانیان	دانشگاه شهید باهنر کرمان
بهروز فتحی واجارگاه	دانشگاه گیلان
حامد فراهانی	دانشگاه دریانوردی و علوم دریایی چابهار
هادی فرزاد	مجتمع آموزش عالی بم
فریده فرساد	مجتمع آموزش عالی بم
فرشته فروزش	مجتمع آموزش عالی بم
محمد قاسم زاده	دانشگاه یزد
شکوفه قربانی	دانشگاه شهید باهنر کرمان
مهدیه قزوینی کر	دانشگاه شهید باهنر کرمان
فرشید کی نیا	دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته
مرجان کوچکی رفسنجانی	دانشگاه شهید باهنر کرمان
فاطمه کوچکی نژاد	مجله سیستمهای فازی ایران
اباذر کیخا	دانشگاه ولایت
علی محمد لطیف	دانشگاه یزد
ماشالله ماشین چی	دانشگاه شهید باهنر کرمان

سمیه معتمد	دانشگاه آزاد اسلامی
ملیحه مغفوری فرسنگی	دانشگاه شهید باهنر کرمان
محمد ملائی امام زاده	دانشگاه شهید باهنر کرمان
نجمه منصوری	دانشگاه شهید باهنر کرمان
مینا میرحسینی	مجمع آموزش عالی بم
حسن میش مست نهی	دانشگاه سیستان و بلوچستان
ابوصالح نادری	مجمع آموزش عالی بم
علی ناصراسدی	مجمع آموزش عالی زرنند
اردوان نجفی	دانشگاه آزاد اسلامی
کمیل نکویی	دانشگاه شهید باهنر کرمان
فهیمة یزدان پناه	دانشگاه ولی عصر(عج) رفسنجان
محمدعلی یعقوبی	دانشگاه شهید باهنر کرمان

سخنرانی‌های کلیدی و کارگاه‌های تخصصی
نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

عنوان سخنرانی کلیدی	زمان
Evolutionary Intelligence in (Big) Data Analytics and Optimization Prof. Amir H. Gandomi, University of Technology Sydney, Australia	چهارشنبه ۱۱ اسفند ۱۴۰۰ ۱۰:۳۰-۱۱:۳۰

Abstract

Evolutionary Computation (EC) has been widely used during the last two decades and has remained a highly-researched topic, especially for complex engineering problems. The EC techniques are a subset of artificial intelligence, but they are slightly different from the classical methods in the sense that the intelligence of EC comes from biological systems or nature in general. The efficiency of EC is due to their significant ability to imitate the best features of nature which have evolved by natural selection over millions of years. The central theme of this presentation is about EC techniques and their application to complex smart cities and infrastructures problems. On this basis, first I will talk about an evolutionary approach called genetic programming for data mining. Applied evolutionary computing will be presented, and then their new advances will be mentioned such as big data mining.

Here, some of my studies on big data mining and modelling using EC and genetic programming, in particular, will be presented. Case studies' topics include forward design, and inverse design. In the second section, the evolutionary optimization algorithms and their key applications in the design optimization of complex and nonlinear engineering systems will be discussed. It will also be explained how such algorithms have been adopted to engineering problems and how their advantages over the classical optimization problems are used in action. Optimization results of large-scale engineering systems and many-objective problems will be presented which show the applicability of EC. Some heuristics will be explained which are adaptable with EC and they can significantly improve the optimization results.

Amir H. Gandomi is a Professor of Data Science and an ARC DECRA Fellow at the Faculty of Engineering & Information Technology, University of Technology Sydney. Prior to joining UTS, Prof. Gandomi was an Assistant Professor at Stevens Institute of Technology, and a distinguished research fellow in BEACON center, Michigan State University. Prof. Gandomi has published over three hundred journal papers and 12 books which collectively have been cited more than 26,000 times (H-index = 74). He has been named as one of the most influential scientific minds and Highly Cited Researcher (top 1% publications and 0.1% researchers) for five consecutive years, 2017 to 2021. He also ranked 17th in GP bibliography among more than 12,000 researchers. He has served as associate editor, editor, and guest editor in several prestigious journals such as AE of IEEE TBD and IEEE IoTJ. His research interests are global optimization and (big) data analytics using machine learning and evolutionary computations in particular.

عنوان سخنرانی کلیدی	زمان
Overview of new families of fuzzy implications and their applications in fuzzy systems Prof. Michal Baczynski, University of Silesia in Katowice, Katowice, Poland	چهارشنبه ۱۱ اسفند ۱۴۰۰ ۱۶:۱۵-۱۷:۱۵

Abstract

The mathematical base of fuzzy systems is a multivalued logic, where among others, we use different generalizations of classical logic connectives. In our talk, we concentrate on analyzing multivalued implications called fuzzy implications in the literature. This class of functions plays a significant role both in theory and applications, as can be seen from their use in, among others, approximate reasoning, fuzzy control, image processing, data analysis and multivalued mathematical logic.

We will review the different classical families of fuzzy implications, focusing primarily on their main properties and possible characterizations. Next, we analyze some new families of fuzzy implication functions; in particular, we present the preference implications and new weighted and mean-implication functions. We will show why these implications can be helpful for engineers and specialists who design and implement fuzzy systems.

Prof. Michał Baczyński

Deputy Dean of the Doctoral School at the University of Silesia in Katowice, Vice President of the Upper Silesian Branch of the Polish Mathematical Society, Faculty of Science and Technology, University of Silesia in Katowice, Bankowa, Katowice, Poland

Prof. Michał Baczyński deals with the mathematical foundations of intelligent systems, especially fuzzy systems. He analyses different methods and systems where multi-valued operators are used. His main object of scientific interest are fuzzy (multi -valued) connectives, in particular fuzzy implication functions, used in many different fields of computer science and computational intelligence.

عنوان سخنرانی کلیدی	زمان
The Sound of Health پروفسور محمدرضا اکبرزاده توتونچی، دانشگاه فردوسی مشهد، ایران	چهارشنبه ۱۱ اسفند ۱۴۰۰ ۱۹:۳۰-۲۰:۳۰

Abstract

The recent Covid19 pandemic gave modern technologies an opportunity to showcase their relevance in serving humanity. The unprecedented availability of inexpensive processing power, wide-scaled network connectivity, and the extraordinary variety of sensing devices enabled many to work remotely, study in virtual classrooms, and buy their necessary everyday goods without leaving home. Nevertheless, the pandemic caused many patients to avoid hospitals and medical facilities when they needed it the most. It also contaminated many patients who did seek help but were healthy otherwise. Telemedicine, a relatively well-established concept long before the pandemic, shows great potential in this regard.

In this talk, I will present the utility of sound, particularly the human voice, as the means for inexpensive and readily available telemedicine technology for remote medical screening. Besides a great utility during pandemics, this means that it could offer increased accessibility to a wide range of underprivileged patients in far and remote areas. However, realistic constraints of low-quality microphones and the narrow bandwidth of transmission lines present significant challenges in this process. The available features are also many, their signal is noisy, and there is considerable uncertainty in mapping these features with the disease. This leads us to the role of intelligent systems for automatic screening, such as fuzzy logic for feature selection and deep neural network structures for signal classification. In particular, I hope to review our recent work in this domain using intelligent feature selection, classification, and testing/training strategy in diagnosing Diseases such as Parkinson's, Covid29, and several other lung diseases.

Mohammad-R. Akbarzadeh-T. (Ph.D. University of New Mexico (UNM), 1998) is currently a professor and director of the Center of Excellence on Soft Computing and Intelligent Information Processing (SCIIP), Electrical Engineering Department, Ferdowsi University of Mashhad. From 1996-2003, he was also with the NASA Center for Autonomous Control Engineering at UNM. In 2006 and 2017, he was also with the Berkeley Initiative on Soft Computing (BISC), UC Berkeley as a visiting scholar. In 2007, he also served as a consulting faculty at the Department of Aerospace and Aeronautic Engineering, Purdue University. Prof. Akbarzadeh is the founding president of the Intelligent Systems Scientific Society of Iran, the founding councilor representing the Iranian Coalition on Soft Computing in IFSA, an IEEE senior member, and the founding faculty councilor of the IEEE student branch until 2008. He has received several awards, including the Caro Lucas Life Achievement Award from the IEEE Iran Section in 2021, the Outstanding Faculty Award from Ferdowsi Academic Foundation in 2021, the National Outstanding Faculty Award in 2019, and the IDB Excellent Leadership Award in 2010. He has also won the FUM university's Outstanding Faculty Award three times. His research interests are bio-inspired computing/optimization, fuzzy logic and control, soft computing, multi-agent systems, complex systems, robotics, cognitive sciences, and medical informatics. He has published over 450 peer-reviewed articles in these and related research fields.

عنوان سخنرانی کلیدی	زمان
Role of Intelligent systems in the Interpretation of Deep Neural Networks in a Post-pandemic World Prof. Saman K. Halgamuge, University of Melbourne, Australia	پنجشنبه ۱۲ اسفند ۱۴۰۰ ۱۰:۳۰-۱۱:۳۰

Abstract

Popular models of AI, in particular machine learning based models have three significant deficiencies: they are mostly manually designed using the experience of AI-experts; they lack human interpretability, i.e., users do not understand the AI architectures either semantically/linguistically or mathematically; and they are unable to dynamically change when new data are acquired. Addressing these deficiencies would provide answers to some of the valid questions about traceability, accountability and the ability to integrate existing knowledge (scientific or linguistically articulated human experience) into the AI model. This keynote addresses these deficiencies in the context of major global problems in the post pandemic world and how intelligent systems can contribute to finding solutions.

Prof. Saman Halgamuge received the B.Sc. Engineering degree in Electronics and Telecommunication from the University of Moratuwa, Sri Lanka, and the Dipl.-Ing and Ph.D. degrees in data engineering from the Technical University of Darmstadt, Germany. He is currently a Professor of the Department of Mechanical Engineering of the School of Electrical Mechanical and Infrastructure Engineering, The University of Melbourne, Australia. He is also an honorary professor of Australian National University. He is a Fellow of IEEE, a distinguished Lecturer of IEEE Computational Intelligence Society and listed as a top 2% most cited researcher for AI and Image Processing in the Stanford database. His research interests are in AI, machine learning including deep learning, optimization, big data analytics and their applications in biomedicine and engineering.

عنوان سخنرانی کلیدی	زمان
Data-Driven Fuzzy Modeling Prof. Irina Perfilieva, University of Ostrava, Czech Republic	پنجشنبه ۱۲ اسفند ۱۴۰۰ ۱۶:۱۵-۱۷:۱۵

Abstract

The talk will focus on efficient data-driven modeling associated with the inverse problem and feature extraction. We show how the theories of manifolds and F-transforms contribute to these delineated areas. The manifold hypothesis states that the shape of observed data is relatively simple and that it lies on a low-dimensional manifold embedded in a higher-dimensional space.

We contribute to the problem of manifold learning. We show that a space whose topological structure is characterized by a fuzzy partition naturally leads to so called Riemannian spaces or Riemannian manifolds. Finally, we show how the discussed notions contribute to the mathematics of deep learning.

Prof. Irina Perfilieva, PhD, dr.h.c., prof.h.c., Centre of Excellence IT4Innovations division of the University of Ostrava, Institute for Research and Applications of Fuzzy Modeling, Ostrava, Czech Republic

Research Activities:

- fuzzy approximation and fuzzy topological spaces;
- fuzzy transforms, image processing, mathematical morphology;
- fuzzy differential equations;
- fuzzy logic, approximate reasoning;
- systems of relation equations and their solvability.

For her long-term scientific achievements she was awarded on the International FLINS 2010. She received the memorial Da Ruan award in 2012. Since 2013, she is an EUSFLAT Honorary Member, and since 2019 she is an IFSA Fellow.

عنوان سخنرانی کلیدی	زمان
Some recent extensions of fuzzy integrals applied to the computational brain problem Prof. Humberto Bustince, Universidad Publica de Navarra, Spain	پنجشنبه ۱۲ اسفند ۱۴۰۰ ۱۹:۳۰-۲۰:۳۰

Abstract

We start revising some extensions of the classical Choquet and Sugeno integrals that have appeared in the literature and which replace sums, products, min or max by more general operators in their definition. The resulting functions are not, in general, aggregation functions, but they belong to the wider class of pre-aggregation functions, for which monotonicity is required only for some (and not every) positive direction. We discuss the utility of these functions in some ensemble techniques in order to fuse data from different channels in order to be able to predict from electroencefalographical signals whether a given subject is thinking on the movement of one hand or another. Experimental result with these new families of functions improves those obtained when the classical Choquet and Sugeno integrals are obtained.

Humberto Bustince Sola is a full professor of Computer Science and Artificial Intelligence at the Public University of Navarra and honorary professor at the University of Nottingham since 2017. He is the main researcher of the Research Group on Artificial Intelligence and Approximate Reasoning, whose research lines are both theoretical (data fusion functions, information and comparison measures, fuzzy sets and their extensions) and applied (Deep learning, image processing, classification, machine learning, data mining, big data or the computational brain). He has led 13 research projects funded by national and regional governments, and two excellence networks on soft computing.

He has been the main researcher in projects with companies and entities such as Caja de Ahorros de Navarra, INCITA, Gamesa Tracasa or the Servicio Navarro de Salud. He has taken part in two international projects. He has authored or coauthored more than 300 works, according to Web of Science, including around 160 in Q1 journals. He was a highly cited researcher among the top 1% most relevant scientists in the world in 2018, according to Clarivate Analytics. He collaborates with first line research groups from countries such as United Kingdom, Belgium, Australia, the Czech Republic, Slovakia, Canada or Brasil. He is editor in chief of the Mathware&Soft Computing online magazine of the European Society of Fuzzy Logic and technologies and of the Axioms journal. Associated editor of the IEEE Transactions on Fuzzy Systems journal and member of the editorial boards of the journals Fuzzy Sets and Systems, Information Fusion, International Journal of Computational Intelligence Systems and Journal of Intelligent & Fuzzy Systems. Moreover, he is a coauthor of a book about averaging functions, and has been the co-editor of several books. He has been in charge of organizing several first level international conferences such as EUROFUSE 2009 and AGOP 2013. He is Senior Member of IEEE y Fellow of the International Fuzzy Systems Association (IFSA). Member of the Basque Academy of Sciences, Arts and Literature, Jakiunde, since 2018. He has advised 11 Ph.D thesis.

He was awarded the Cross of Carlos III the Noble by The Government of Navarra in 2017. He got the National Computer Science Prize José García Santesmases in 2019 and the Scientific Excellence Award of EUSFLAT the same year.

عنوان مخبرانی کلیدی	زمان
ادراک بر پایه امکان: تاملی بر بازشناسی الگوی اشکال پروفسور کامبیز بدیع، دانشگاه تهران، مرکز تحقیقات مخابرات ایران، ایران	جمعه ۱۳ اسفند ۱۴۰۰ ۱۰:۳۰-۱۱:۳۰

چکیده

"امکان" و "احتمال" به عنوان دو موجودیت کلیدی در بیان عدم قطعیت (نایقینی، ناتاشتی)، طی چند دهه اخیر، پا به پای یکدیگر مطرح بوده اند. آنچه در کل ایندو را از هم مجزا می سازد، تاکید امکان بر "وضعیت ساختار موجود" در یک موقعیت است، حال آنکه احتمال بر "کارکردهای تجربه شده" در خصوص یک موقعیت تاکید می ورزد. بدین منوال، شرط لازم برای توجیه کاربست امکان در قبال یک موقعیت، توانایی در نظر گرفتن ساختار برای آن و یا به تعبیری ساختارپذیر بودن است.

زمانی که بحث ساختار به میان می آید، این پرسش مطرح است که اجزای یک الگو و یا یک موقعیت چگونه بر پایه مدل های معین یا به ظهور می نهد، و زمانی که موضوع ادراک بر پایه امکان مطرح است، می توان چشم بر آن داشت که طیفی از امکانات در خصوص این اجزا (که طبیعتی ساختاری دارد) گرد هم می آیند تا به این الگو و یا موقعیت مربوطه که درون آن اجزاء هر کدام برپایه مدل ویژه ای به صحنه آمده اند، هویت بخشند. ما بر این باوریم که ادراک این چنینی بر پایه نوعی ادغام (همجوشی) در قبال مجموعه ای از امکانات رخ می دهد. ادراک را بدین گونه می توان نگاشتی ذهنی از همبستگی مجموعه ای از ساختارهای شکل معنادار بر پایه ادغام اطلاعات مربوط به درجات امکان تعلق به مدلهای مرجع، که منجر به فراخاهی طبقه ای با هویت خاص می گردد، تلقی کرد. به تعبیری می توان ادعا داشت که ادغام وضعیت همبستگی ساختارهای شکل معنادار می تواند زمینه ساز یک دریافت (یا اندریافت) بهین از طبقات ممکن حاکم بر ساختاریابی این الگو باشد.

اساس رهیابی به فرآیند ادغام، که در ادبیات جاری رایانش به آن همجوشی اطلاعات نیز گفته می شود، بدین گونه است که، پس از آنکه برداری برای یک الگو بر پایه اطلاعات مربوط به متوسط درجات امکان (به ازاء جمیع ساختارهای شکل معنادار از منظر هویت طبقات ممکنه الگو) بر ساخته شد، کوشیده می شود تا شباهت اجزاء این بردار با ساختارهای شکل معین در تک تک طبقات ممکنه، بر پایه بکارگیری نوعی تابع شباهت به رایانش درآمده و در گام بعد حاصلضرب مقادیر این شباهت ها به ازاء تک تک این طبقات تعیین گردد. بدیهی است که منطقاً طبقه ای از الگو که میزان این حاصلضرب به ازاء آن بیشترین باشد، می تواند بهترین طبقه ممکنه برای الگوی مربوطه تلقی گردد. بدیهی است که سیاق این چنینی برای ادغام اطلاعات مربوط به درجات امکان، ساده ترین نوع رویکرد به ادغام است که اهمیت وجودی ساختارهای شکل معنادار در هویت بخشی به طبقات الگو را یکسان در نظر می گیرد. معذالک، این انتظار می رود که بتوان از فرمالیزم های متفاوتی در ادغام اطلاعات بهره برد. ادغام یا همجوشی درجات امکان تعلق می تواند در دو سطح صورت گیرد: هویت بخشی به ساختارهای شکل معنادار با هدف تعیین معقول (سنجیده) ترین طبقه ممکنه برای الگوی ورودی و طبقه بندی نهایی الگوی ورودی با هدف تعیین معقول (سنجیده) ترین طبقات ممکنه. به هر باره از تلفیقی از این دو سطح می توان در حالت کلی برای یافتن طبقات معقولی (که به بهترین نحو پاسخگوی الگوی ورودی بوده و امکان را در ابعاد گوناگون قابل قبول سازند) بهره برد.

از جمله الگوها با خواص ساختاری می توان به الگوی اشکال مشتمل بر منحنی های باز و بسته و یا نواحی ویژه، الگوی هستینه های معنایی مشتمل بر ترکیبات واژگانی با معانی / پیامهای ویژه، الگوی عبارات ریاضی / منطقی مشتمل بر نمادهای منطقی، گزاره ها، آرگومان ها، توابع و انواع متغیرها، ... اشاره داشت. اشکال بصورت کلی اجسام (زنده، بی جان، بر ساخته، پدیده ها) کیهانی، انگمانی، سازمانی، فردی (و نمادها) حروف، اعداد، صور قراردادی حاوی پیام (، را در بر می گیرد.

در این سخنرانی، مقوله ادراک بر پایه امکان در خصوص الگوی اشکال با تاکید بر جایگاه همجوشی اطلاعات مورد بحث قرار می گیرد. در وادی اشکال، مثالهایی از حروف (لاتین و عربی) همراه با ساختارهای رابطه ای میان اجزای آنها (از نوع منحنی) که ساختاربخش اجزای شکل معنادار هستند، به بحث گذاشته شده و گفته می شود که چگونه بر پایه درجات امکان متعلق به این اجزای شکل است که نهایتاً طبقات ممکنه برای یک الگو معنا می یابند. منحنی مورد نظر برای یک جزء، منحنی با یک جهت یکسان است که درون آن هیچگونه نقطه عطفی که تغییردهنده جهت باشد، به چشم نمی خورد. دلیل عمده این انتخاب، سهولت در کاربست ساختارهای رابطه ای با هدف ساختار بخشی به ساختارهای شکل معنادار است.

در پایان، ذکر این نکته ضروری است که رویکرد امکانی به ادراک این فرصت را به نحو مناسبی فراهم می سازد تا آن دسته از ساختارهای رابطه ای شکل معنادار که پیام آور هویت کانونی طبقات هستند در فرآیند طبقه بندی الگو مد نظر قرار گرفته و بدین منوال از دخیل ساختن ساختارهای رابطه ای نامربوط که فرآیند طبقه بندی را با هزینه زمانی بالا مواجه می سازد، پرهیز گردد. آنچه به عنوان شهود، در وادی ادراک از آن یاد می شود، به نظر می رسد که به این امر مرتبط گردد.

کامبیز بدیع درجات کارشناسی، کارشناسی ارشد و دکترای خود در زمینه مهندسی برق و الکترونیک را، با تخصص بازشناسی الگو، از انستیتو تکنولوژی توکیو دریافت نموده است. در سال های گذشته، پژوهش ایشان در زمینه هایی از قبیل ماشین یادگیری، مدلسازی شناختی و پردازش و تولید دانش به طور اعم و دانش ورزی تمثیلی، مدل ورزی تجربه و مدل ورزی فرآیند تفسیر به طور ویژه، با تاکید بر آفریدن ایده ها، فنون و محتواهای جدید، متمرکز بوده است.

دکتر بدیع از پژوهشگران فعال در حوزه مطالعات میان گستره ای و تراگستره ای بوده و از انگیزه ای بالا برای کاربست مدل های هوشمند، شناختی و پدیدارشناختی در حوزه های انسانی برخوردار می باشد.

نامبرده، در حال حاضر استاد پژوهشگاه ارتباطات و فناوری اطلاعات، استاد وابسته در دانشکده علوم مهندسی دانشگاه تهران و همزمان سردبیر مجله بین المللی Research Communication Technology & Information و عضو مدعو فرهنگستان علوم در شاخه مهندسی برق و کامپیوتر می باشد.

عنوان سخنرانی کلیدی	زمان
مشتق توابع فازی و کاربرد آن در معادلات دیفرانسیل فازی دکتر علیرضا خواستان، دانشگاه تحصیلات تکمیلی علوم پایه، زنجان، ایران	جمعه ۱۳ اسفند ۱۴۰۰ ۱۵:۱۵-۱۶:۱۵

چکیده

یکی از مفاهیم مهم در ریاضیات فازی، مفهوم مشتق توابع فازی است. در ۳۰ سال گذشته، تعاریف مختلفی برای مشتق فازی معرفی و به کار گرفته شده است. در این سخنرانی، ابتدا تاریخچه مختصری در مورد انواع مشتق توابع فازی شامل "مشتق هوکوهارا"، "مشتق هوکوهارای تعمیم یافته" و "مشتق متریک" بیان شده و برخی ویژگی‌های مهم آنها در مقایسه با یکدیگر مورد بررسی قرار می‌گیرند. سپس در ادامه، کاربرد مشتق فازی در معادلات دیفرانسیل فازی و برخی روش‌های حل آنها ارایه می‌شود.

دکتر علیرضا خواستان

دانشیار دانشگاه تحصیلات تکمیلی علوم پایه زنجان

فارغ التحصیل دوره دکتری از دانشگاه تبریز و دوره پسا دکتری از دانشگاه سانتیاگو دی کامپوستلا، اسپانیا (Santiago de Compostela, Spain)

ادیتور مجله Fuzzy Sets and Systems و مجله فارسی سیستم‌های فازی و کاربردها

کارگاه‌های تخصصی

دو شنبه ۹ اسفند ماه ۱۴۰۰
برنامه‌ریزی ریاضی فازی و خاکستری دکتر سیدهادی ناصری و دکتر داود درویشی، <u>دانشگاه مازندران</u> رئیس نشست: پروفسور بهروز فتحی <u>واجارگاه</u>
آشنایی با گرافهای دانش و کاربرد آنها در سیستم های هوشمند دکتر امین انجم شعاع، <u>دانشگاه ملی گالوی ایرلند</u> رئیس نشست: دکتر مرجان کوچکی <u>رفسنجانی</u>
سه شنبه ۱۰ اسفند ماه ۱۴۰۰
چگونه می‌توانیم یک شبیه‌سازی موفق مونت کارلو در محاسبات داشته باشیم؟ پروفسور بهروز فتحی <u>واجارگاه، دانشگاه گیلان</u> رئیس نشست: دکتر سیدهادی ناصری
فناوری‌های بلاک چین، رمزارزها و قرارداد هوشمند پروفسور حسین نظام آبادی‌پور، <u>دانشگاه شهید باهنر کرمان</u> رئیس نشست: دکتر مهدی افتخاری
دو شنبه ۲۳ اسفند ماه ۱۴۰۰
مجموعه های فازی و منطق فازی ویژه دبیران آموزش و پرورش پروفسور سید محمود طاهری، <u>دانشگاه تهران</u> رئیس نشست: معراج عبدی

فهرست مقالات فارسی

شماره صفحه	عنوان مقاله
مقالات سخنرانی	
۴۰	اتوماتا L - فازی مردد در دسترس و هم-در دسترس
۴۶	ارائه یک روش جدید مبتنی بر ترکیب ویژگی‌های استخراج شده از شبکه‌های عصبی عمیق DenseNet169 و MobileNet و دسته‌بندی‌کننده LightGBM به منظور تشخیص بیماری کرونا از روی تصاویر اشعه ایکس
۵۲	ارائه یک سامانه توصیه‌گر مبتنی بر مدل برای بیماران مبتلا به پرفشاری خون
۵۸	ارائه یک معیار شباهت ترکیبی جهت بهبود سیستم‌های مشارکت جمعی آیت‌محور
۶۴	الگوریتم ژنتیک برای پستچی چینی تحت شرایط نایقینی
۶۹	آموزش شبکه‌های عصبی پیشرو با استفاده از الگوریتم بهینه‌ساز تعادل به‌منظور شناسایی سیستم غیرخطی
۷۵	بررسی تاثیر بکارگیری توابع زیان مختلف بر عملکرد مدل خوشه بندی فازی برای داده های فازی در حضور داده های پرت
۸۱	بهبود ترافیک شهری در شبکه‌های بین خودرویی با استفاده از رویکرد پروتکل وضعیت-اتصال و شبکه‌های عصبی
۸۷	تابعی اکتشافی برای بهبود دقت پیشبینی برنامه‌های جهشیافته آشکار کننده خطا
۹۳	تحلیل الگوی مصرفی مشترکین برق با استفاده از خوشه بندی
۹۸	تحلیل بورس ایران با استفاده از اندیکاتورهای باند بولینگر و MACD
۱۰۴	تشخیص اسکیزوفرنی بر اساس سیگنال الکتروانسفالوگرام با استفاده از یادگیری عمیق
۱۱۰	تشخیص گریه نوزاد از سایر صداهای محیط با استفاده از یادگیری عمیق
۱۱۶	تعیین کارایی در تحلیل پوششی داده فازی با تقریب نزدیکترین بازه
۱۲۱	تغییر الگوریتم GRASP برای خوشه بندی داده ها
۱۲۷	حل مسائل کنترل بهینه دارای عدم قطعیت
۱۳۱	حل مسائل کنترل بهینه کسری بازهای با استفاده از تبدیل لاپلاس
۱۳۶	شبکه‌های عصبی پیچشی برای تشخیص بیماری‌های تنفسی
۱۴۳	قابلیت اعتماد بر اساس α -شک اعداد فازی
۱۴۸	کنترل شبکه‌های تنظیم‌کننده ژن P53 مبتنی بر یادگیری تقویتی عمیق و کاربرد آن در بیماری سرطان

۱۵۴	مروری بر انواع اعداد فازی و دسته بندی آنها با رویکرد رتبه بندی
۱۶۰	مسئله برنامه ریزی خطی فازی نوع-۲ بازه ای با ابهام از نوع وگنس در بردار منابع
۱۶۶	مساله زمانبندی پروژه منابع محدود فازی با چندین تامین کننده
۱۷۲	مسیریابی مؤثر در شبکه های ویژه خودرویی براساس روشهای پشتیبانی از هواپیماهای بدون سرنشین
۱۷۸	یک رویکرد ترکیبی مبتنی بر واژه نامه و یادگیری ماشین برای شناسایی هیجان در نظرهای کاربران فارسی زبان
۱۸۴	یک سیستم بیومتریک ترکیبی هایبرید مبتنی بر تلفیق ویژگیهای چهره، هر دو عنبیه و دو اثرانگشت
۱۹۳	یک مدل ریاضی دو هدفه فازی برای مسأله زمانبندی زنجیره بحرانی با در نظر گرفتن محدودیت بودجه و عدم قطعیت
مقالات پوستر	
۲۰۰	ارائه یک روش تلفیقی مبتنی بر سیستم های چندعاملی در تصاویر ماموگرافی جهت تشخیص زود هنگام بیماری سرطان پستان
۲۰۶	پیاده سازی یادگیری تقویتی عمیق برای کنترل هوشمند مبدل افزایشنده
۲۱۲	تشخیص اخبار جعلی مبتنی بر ترکیبی از ویژگیهای زبانی و صفات گوینده خبر
۲۲۰	تنظیم کننده خودکار چندهدفه آنلاین داده محور برای سیستم کنترل یادگیر تکرارشونده با استفاده از الگوریتم بهینه یابی ازدحام ذرات
۲۲۶	حل مسئله برنامه ریزی خطی کروی فازی
۲۳۲	شناسایی بی درنگ حالات عاطفی چهره با روش یادگیری عمیق
۲۳۷	شناسایی بیماری COVID-19 با استفاده از شبکه عصبی پیچشی
۲۴۳	کاربردی از شبیه سازی مونت کارلو در آزمون کیفیت شیشه خودرو
۲۴۸	گامی فراتر در پیشگویی پیوند: یک مرور سیستماتیک بر پیشگویی پیوند چندلایه
۲۵۶	نقش پایه ای اینترنت اشیا در مدیریت هوشمند فرآیندهای صنعت کشاورزی

مقالات سخنرانی کنگره

اتوماتا L -فازی مردد در دسترس و هم-در دسترس

محمد جواد عاقلی^۱، محمد مهدی زاهدی^۲ و مرضیه شمسی زاده^۳

^۱ بخش ریاضی، دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته کرمان، ایران. mjavadagheli@gmail.com

^۲ بخش ریاضی، دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته کرمان، ایران. zahedi_mm@kgut.ac.ir

^۳ بخش ریاضی، دانشگاه صنعتی خاتم الانبیاء بهبهان، خوزستان، ایران. shamsizadeh.m@gmail.com

چکیده: در نظریه اتوماتا (اتوماتا فازی) موضوع تشخیص زبان از اهمیت خاصی برخوردار است. به ویژه اینکه تحقیقات زیادی صورت گرفته که به روش‌های مختلف اتوماتایی با تعداد حالات کمتر (حتی کمترین حالات ممکن) بیانند که یک زبان مورد نظر را تشخیص دهد. اینک در این مقاله ما همین بحث را برای اتوماتا L -فازی مردد مورد مطالعه قرار می‌دهیم. به عبارت دیگر در این مقاله تلاش شده که از یک اتوماتا L -فازی مردد، به اتوماتا L -فازی مردد در دسترس و هم-در دسترس با تعداد حالت‌های کمتری دست یابیم که زبان اتوماتای اولیه را حفظ کند.

کلمات کلیدی: مجموعه فازی مردد، اتوماتا فازی مردد، مجموعه L -فازی مردد، اتوماتا L -فازی مردد، شبکه، در دسترس، هم-در دسترس، هم-پوشش.

۱ مقدمه

مجموعه حالات هم-در دسترس، بخش در دسترس و بخش هم-در دسترس اتوماتا L -فازی مردد، با تعداد حالات کمتر را تعریف کردیم که زبان اولیه را حفظ کند. سرانجام با تعریف همزمان بخش در دسترس و هم-در دسترس مجموعه حالات، به معرفی بخش هم-پوشش اتوماتا L -فازی مردد، با تعداد حالات کمتر پرداختیم به قسمی که زبان اولیه را حفظ می‌کند.

۲ تعاریف و قضایای اولیه

در این بخش تعاریف و مفاهیم پایه‌ای که در کل مقاله استفاده می‌شوند را ارائه می‌کنیم.

تعریف ۱. [۲] فرض کنید X یک مجموعه و L یک شبکه دلخواه باشد. آنگاه مجموعه L -فازی مردد (*Hesitant*)، روی مجموعه X را با نداشت زیر نمایش می‌دهیم

$$H: X \rightarrow 2^L$$

$$x \mapsto H(x).$$

در سال ۲۰۰۹ مفهوم مجموعه فازی مردد توسط تورا و ناروکاوا معرفی شد [۱۱]. آنها مجموعه فازی مردد را به وسیله ارزش عضویت اعضا به زیرمجموعه دلخواهی از بازه $[0, 1]$ معرفی کردند. سپس در سال ۲۰۱۸ دهمیری، ماشین چی و مسیار با گسترش بازه $[0, 1]$ به شبکه L توانستند مفهوم مجموعه L -فازی مردد را بیان کنند [۲].

وی در سال ۱۹۶۷ اتوماتا فازی را معرفی کرد [۱۲]. دو سال بعد لی و زاده مفهوم اتوماتا فازی حالت متناهی را ارائه کردند [۶]. اتوماتا فازی حالت متناهی، اهمیت کاربردی مهمی دارد [۷]. اخیرا کاستا و بدراگال با استفاده از مفهوم مجموعه‌های فازی مردد توانستند مفهوم اتوماتا فازی مردد را وارد عرصه اتوماتا کنند [۱].

در نظریه اتوماتا، زبان اهمیت خاصی دارد به ویژه آن که تحقیقات زیادی صورت گرفته که به روش‌های مختلف اتوماتای با تعداد حالات کمتر بیابند که یک زبان را تشخیص دهد [۴، ۵، ۸، ۹]. در این مقاله حالت در دسترس و حالت هم-در دسترس را تعریف کردیم. سپس، با تبدیل مجموعه حالات به مجموعه حالات در دسترس و همچنین به

که $H(x)$ درجه عضویت x در H را مشخص می کند و عضو مجموعه L -فازی L دارد.

تعریف ۲. [۲] برای $A, B \in 2^L$ ، تابع $\sqcup : 2^L \times 2^L \rightarrow 2^L$ را به صورت زیر تعریف می کنیم

$$A \sqcup B = \{a \vee b | a \in A \text{ و } b \in B\},$$

و برای تابع $\sqcap : 2^L \times 2^L \rightarrow 2^L$ نیز داریم

$$A \sqcap B = \{a \wedge b | a \in A \text{ و } b \in B\}.$$

تعریف ۳. [۱۲] فرض کنید که X یک مجموعه باشد. آنگاه مجموعه X^* را به صورت زیر در نظر می گیریم

$$X^* = \{u | u = x_1 x_2 \dots x_n, x_i \in X, i \in N\} \cup \{\Lambda\}$$

که در آن Λ عضو خنثی مجموعه X^* است. به هر یک از اعضای مجموعه X^* یک کلمه می گویند و Λ را کلمه تهی می نامند.

تعریف ۴. [۷] ۵-تایی $A = (Q, X, \iota, \delta, \tau)$ را یک اتوماتا فازی جزئی می نامیم که در آن

۱. Q مجموعه متناهی ناتهی است و مجموعه حالات اتوماتا نامیده

می شود،

۲. X مجموعه متناهی ناتهی است و اعضای آن را الفبای اتوماتا می نامیم،

۳. ι یک زیر مجموعه فازی از Q است، یعنی $\iota : Q \rightarrow [0, 1]$ که آن را حالت فازی اولیه می نامیم،

۴. $\delta : Q \times X \rightarrow Q$ یک تابع است که آن را تابع انتقال حالت می نامیم،

۵. τ یک زیر مجموعه فازی از Q است، یعنی $\tau : Q \rightarrow [0, 1]$ که آن را حالت فازی نهایی می نامیم.

۳ اتوماتا L -فازی مردد در دسترس و هم-در

دسترس

تعریف ۵. فرض کنید L یک شبکه دلخواه باشد. یک اتوماتا L -فازی مردد، ۵-تایی $A = (Q, X, \iota, \delta, \tau)$ است که در آن

۱. Q مجموعه متناهی ناتهی است و مجموعه حالات اتوماتا نامیده می شود،

۲. X مجموعه متناهی ناتهی است و اعضای آن را الفبای اتوماتا می نامیم،

۳. $\iota : Q \rightarrow 2^L$ یک مجموعه L -فازی مردد روی Q از حالت های اولیه است،

۴. $\delta : Q \times X \times Q \rightarrow 2^L$ یک مجموعه L -فازی مردد است، که آن را مجموعه L -فازی مردد انتقال حالت می نامیم،

۵. $\tau : Q \rightarrow 2^L$ یک مجموعه L -فازی مردد روی Q از حالت های نهایی است.

برای یک اتوماتا L -فازی مردد، مجموعه L -فازی مردد δ^* به δ توسعه می دهیم و به صورت زیر تعریف می کنیم

$$\delta^*(q, \Lambda, p) = \begin{cases} \{1\} & \text{اگر } q = p \\ \{0\} & \text{صورت این غیر در} \end{cases},$$

و

$$\delta^*(q, ua, p) = \sqcup_{r \in Q} [\delta^*(q, u, r) \sqcap \delta(r, a, p)],$$

جاییکه $a \in X$ و $u \in X^*$ است.

تعریف ۶. فرض کنید $A = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آنگاه مجموعه L -فازی مردد، ξ_A را روی $X^* \times Q$ به صورت زیر تعریف می کنیم:

$$\xi_A : X^* \times Q \rightarrow 2^L$$

$$\xi_A(u, q) = \sqcup \{i(p) \sqcap \delta^*(p, u, q) | p \in Q\}.$$

تعریف ۷. فرض کنید $A = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آنگاه $q \in Q$ را یک حالت در دسترس گوئیم، اگر وجود داشته باشد $u \in X^*$ بطوریکه

$$\xi_A(u, q) \neq \{0\}.$$

بنابراین اگر تمام حالات A در دسترس باشد، آن را یک اتوماتا L -فازی مردد در دسترس می نامیم.

ملاحظه ۱. فرض کنید $A = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آنگاه $q \in Q$ را یک حالت در دسترس است، اگر و فقط اگر وجود داشته باشد $u \in X^*$ و $p \in Q$ بطوریکه

$$\iota(p) \sqcap \delta^*(p, u, q) \neq \{0\}.$$

قضیه ۱. فرض کنید $A = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد و $q \in Q$ یک حالت در دسترس باشد. آنگاه وجود دارد $u \in X^*$ و $p \in Q$ بطوریکه

$$\iota(p) \sqcap \delta^*(p, u, q) \neq \{0\}, |u| < |Q|.$$

تعریف ۸. فرض کنید $A = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آنگاه ι را در دسترس گوئیم، اگر وجود داشته باشد $u \in X^*$ بطوریکه برای هر $q \in Q$ داشته باشیم:

$$\iota(q) \neq \{0\} \iff \xi_A(u, q) \neq \{0\}. \quad (۱)$$

تعریف ۹. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آنگاه ι را بطور قوی در دسترس گوئیم، اگر وجود داشته باشد $u \in X^*$ بطوریکه برای هر $q \in Q$ داشته باشیم:

$$\iota(q) = \xi_{\mathcal{A}}(u, q). \quad (۲)$$

واضح است که اگر ι بطور قوی در دسترس باشد، آنگاه ι در دسترس است.

مثال ۱. فرض کنید $L = [0, 1]$ ، $X = \{x\}$ ، $Q = \{q_1, q_2\}$. آنگاه $\iota : Q \rightarrow 2^L$ و $\delta^* : Q \times X^* \times Q \rightarrow 2^L$ را مطابق زیر تعریف می کنیم

$$\begin{aligned} \delta(q_1, x, q_1) &= \{0, \frac{1}{3}\}, & \delta(q_1, x, q_2) &= \{\frac{1}{3}, \frac{1}{2}\}, \\ \delta(q_2, x, q_1) &= \{0\}, & \delta(q_2, x, q_2) &= \{1\}, \\ \iota(q_1) &= \{0, \frac{1}{2}\}, & \iota(q_2) &= \{\frac{1}{2}, 1\}. \end{aligned}$$

ما داریم

$$\begin{aligned} \delta^*(q_1, x^2, q_1) &= \sqcup \left(\delta(q_1, x, q_1) \sqcap \delta(q_1, x, q_1), \delta(q_1, x, q_2) \sqcap \delta(q_2, x, q_1) \right) \\ &= \sqcup \left(\{0, \frac{1}{3}\} \sqcap \{0, \frac{1}{3}\}, \{\frac{1}{3}, \frac{1}{2}\} \sqcap \{0\} \right) = \{0, \frac{1}{3}\} \sqcup \{0\} = \{0, \frac{1}{3}\}, \end{aligned}$$

$$\begin{aligned} \delta^*(q_1, x^2, q_2) &= \sqcup \left(\delta(q_1, x, q_1) \sqcap \delta(q_1, x, q_2), \delta(q_1, x, q_2) \sqcap \delta(q_2, x, q_2) \right) \\ &= \sqcup \left(\{0, \frac{1}{3}\} \sqcap \{\frac{1}{3}, \frac{1}{2}\}, \{\frac{1}{3}, \frac{1}{2}\} \sqcap \{1\} \right) = \{0, \frac{1}{3}\} \sqcup \{\frac{1}{3}, \frac{1}{2}\} = \{\frac{1}{3}, \frac{1}{2}\}, \end{aligned}$$

$$\begin{aligned} \delta^*(q_2, x^2, q_1) &= \sqcup \left(\delta(q_2, x, q_1) \sqcap \delta(q_1, x, q_1), \delta(q_2, x, q_2) \sqcap \delta(q_2, x, q_1) \right) \\ &= \sqcup \left(\{0\} \sqcap \{0, \frac{1}{3}\}, \{1\} \sqcap \{0\} \right) = \{0\}, \end{aligned}$$

$$\begin{aligned} \delta^*(q_2, x^2, q_2) &= \sqcup \left(\delta(q_2, x, q_1) \sqcap \delta(q_1, x, q_2), \delta(q_2, x, q_2) \sqcap \delta(q_2, x, q_2) \right) \\ &= \sqcup \left(\{0\} \sqcap \{\frac{1}{3}, \frac{1}{2}\}, \{1\} \sqcap \{1\} \right) = \{0\} \sqcup \{1\} = \{1\}, \end{aligned}$$

$$\begin{aligned} \xi_{\mathcal{A}}(x^2, q_1) &= \sqcup \left(\iota(q_1) \sqcap \delta^*(q_1, x^2, q_1), \iota(q_2) \sqcap \delta^*(q_2, x^2, q_1) \right) \\ &= \sqcup \left(\{0, \frac{1}{2}\} \sqcap \{0, \frac{1}{3}\}, \{1, \frac{1}{2}\} \sqcap \{0\} \right) = \{0, \frac{1}{3}\} \sqcup \{0\} = \{0, \frac{1}{3}\} \neq \{0\}, \end{aligned}$$

$$\begin{aligned} \xi_{\mathcal{A}}(x^2, q_2) &= \sqcup \left(\iota(q_1) \sqcap \delta^*(q_1, x^2, q_2), \iota(q_2) \sqcap \delta^*(q_2, x^2, q_2) \right) \\ &= \sqcup \left(\{0, \frac{1}{2}\} \sqcap \{\frac{1}{3}, \frac{1}{2}\}, \{1, \frac{1}{2}\} \sqcap \{1\} \right) = \{0, \frac{1}{3}, \frac{1}{2}\} \sqcup \{1, \frac{1}{2}\} = \{1, \frac{1}{2}\} \neq \{0\}. \end{aligned}$$

بنابراین برای هر $q \in Q$ ، $q \neq \{0\}$ و $\xi_{\mathcal{A}}(x^2, q) \neq \{0\}$ است. از این رو می توان گفت برای هر $q \in Q$ ، x^2 در رابطه (۲) صدق می کند. بنابراین ι در دسترس است. اکنون برای هر $x^n \in X^*$ داریم:

$$\delta^*(q_1, x^n, q_1) = \{0, \frac{1}{3}\},$$

$$\delta^*(q_1, x^n, q_2) = \{\frac{1}{3}, \frac{1}{2}\},$$

$$\delta^*(q_2, x^n, q_1) = \{0\},$$

$$\delta^*(q_2, x^n, q_2) = \{1\}.$$

بنابراین

$$\begin{aligned} \xi_{\mathcal{A}}(x^n, q_1) &= \sqcup \left(\iota(q_1) \sqcap (\delta^*(q_1, x^n, q_1), \iota(q_2) \sqcap \delta^*(q_2, x^n, q_1)) \right) \\ &= \{0, \frac{1}{3}\}, \end{aligned}$$

$$\begin{aligned} \xi_{\mathcal{A}}(x^n, q_2) &= \sqcup \left(\iota(q_1) \sqcap (\delta^*(q_1, x^n, q_2), \iota(q_2) \sqcap \delta^*(q_2, x^n, q_2)) \right) \\ &= \{1, \frac{1}{2}\}. \end{aligned}$$

از این رو برای هر $u \in X^*$ داریم $\xi_{\mathcal{A}}(u, q_1) \neq \{0\}$ پس هیچ $u \in X^*$ وجود ندارد که در رابطه (۲) صدق کند در نتیجه ι بطور قوی در دسترس نیست.

قضیه ۲. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد و ι در دسترس باشد. آنگاه وجود دارد $u \in X^*$ بطوریکه $|u| < |Q|$ و برای هر $q \in Q$ داریم:

$$\iota(q) \neq \{0\} \iff \xi_{\mathcal{A}}(u, q) \neq \{0\}. \quad (۳)$$

تعریف ۱۰. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آنگاه $\beta_{\mathcal{A}} : X^* \rightarrow 2^L$ را مجموعه L -فازی مردد زبان $\mathcal{A}$ نامیده و به صورت زیر تعریف می کنیم:

$$\beta_{\mathcal{A}}(u) = \sqcup_{q,p \in Q} (\iota(q) \sqcap \delta^*(q, u, p) \sqcap \tau(p)).$$

تعریف ۱۱. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آنگاه گوئیم $u \in X^*$ توسط $\mathcal{A}$ شناسایی می شود هرگاه $\beta_{\mathcal{A}}(u) \neq \{0\}$. تمام کلمه های که توسط $\mathcal{A}$ شناسایی می شوند را با $R(\mathcal{A})$ نمایش می دهیم و آن را زبان $\mathcal{A}$ می نامیم. بنابراین زبان اتوماتا L -فازی مردد $\mathcal{A}$ برابر است با

$$R(\mathcal{A}) = \{u \in X^* | \beta_{\mathcal{A}}(u) \neq \{0\}\}.$$

تعریف ۱۲. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آن گاه Q' را به صورت

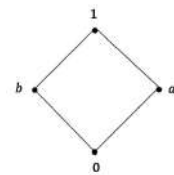
$$Q' = \{q \in Q | \sqcup \{i(p) \sqcap \delta^*(p, u, q) | u \in X^*, p \in Q\} \neq \{0\}\}$$

تعریف می کنیم و آن را بخش در دسترس Q می نامیم. اگر $Q' = Q$ باشد، آنگاه $\mathcal{A}$ را در دسترس می نامیم.

تعریف ۱۳. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد و $\mathcal{A}' = (Q', X, \delta', \iota', \tau')$ آنگاه Q باشد. جاییکه $\iota' = \iota|_{Q'}$, $\delta' = \delta|_{Q' \times X \times Q'}$ می نامیم، جاییکه $\tau' = \tau|_{Q'}$ و $\mathcal{A}'$ یک اتوماتا L -فازی مردد در دسترس است.

مثال ۲. فرض کنید L مشبکه کراندار شکل ۱ و $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد، جاییکه $Q = \{q_1, q_2\}$ و $X = \{x\}$ است. آنگاه $\delta^* : Q \times X^* \times Q \rightarrow 2^L$, $\iota : Q \rightarrow 2^L$ و $\tau : Q \rightarrow 2^L$ را مطابق زیر تعریف می کنیم

$$\begin{aligned} \delta(q_1, x, q_1) &= \{b\}, & \delta(q_1, x, q_2) &= \{a\}, \\ \delta(q_2, x, q_1) &= \{0, b\}, & \delta(q_2, x, q_2) &= \{0, a\}, \\ \iota(q_1) &= \{b\}, & \iota(q_2) &= \{0, b\}, \\ \tau(q_1) &= \{b, 1\}, & \tau(q_2) &= \{a\}. \end{aligned}$$



شکل ۱: مشبکه کراندار L از مثال ۲

بنابراین

$$\begin{aligned} \iota(q_1) \cap \delta(q_1, x, q_2) &= \{b\} \cap \{a\} = \{0\}, \\ \iota(q_2) \cap \delta(q_2, x, q_2) &= \{0, b\} \cap \{0, a\} = \{0\}, \end{aligned}$$

اکنون داریم

$$\delta^*(q_1, x^2, q_2) = \sqcup \left(\delta(q_1, x, q_1) \cap \delta(q_1, x, q_2), \delta(q_1, x, q_2) \cap \delta(q_2, x, q_2) \right) = \sqcup \left(\{b\} \cap \{a\}, \{a\} \cap \{0, a\} \right) = \{0\} \sqcup \{0, a\} = \{0, a\},$$

$$\delta^*(q_2, x^2, q_2) = \sqcup \left(\delta(q_2, x, q_1) \cap \delta(q_1, x, q_2), \delta(q_2, x, q_2) \cap \delta(q_2, x, q_2) \right) = \sqcup \left(\{0, b\} \cap \{a\}, \{0, a\} \cap \{0, a\} \right) = \{0\} \sqcup \{0, a\} = \{0, a\}.$$

و برای هر $x^n \in X^*$ که $n \in \{2, 3, 4, \dots\}$ داریم

$$\begin{aligned} \delta^*(q_1, x^n, q_2) &= \{0, a\}, \\ \delta^*(q_2, x^n, q_2) &= \{0, a\}. \end{aligned}$$

بنابراین برای هر $u \in X^*$ داریم

$$\begin{aligned} i(q_1) \cap \delta^*(q_1, u, q_2) &= \{b\} \cap \{0, a\} = \{0\}, \\ i(q_2) \cap \delta^*(q_2, u, q_2) &= \{0, b\} \cap \{0, a\} = \{0\}. \end{aligned}$$

در نتیجه q_2 یک حالت در دسترس نیست.

اکنون برای q_1 داریم

$$i(q_1) \cap \delta(q_1, x, q_1) = \{b\} \cap \{b\} = \{b\} \neq \{0\},$$

بنابراین q_1 یک حالت در دسترس است. از این رو $Q' = \{q_1\}$ بخش در دسترس Q است. در نتیجه $\mathcal{A}' = (Q', X, \delta', \iota', \tau')$ که در آن

$$\begin{aligned} \iota' : Q' &\rightarrow 2^L, & \tau' : Q' &\rightarrow 2^L \\ \iota'(q_1) &= \{b\}, & \tau'(q_1) &= \{b, 1\} \end{aligned}$$

$$\delta'^* : Q' \times X^* \times Q' \rightarrow 2^L,$$

$$\delta'(q_1, u, q_1) = \{b\},$$

بخش در دسترس اتوماتا $\mathcal{A}$ است، بطوریکه

$$R(\mathcal{A}) = R(\mathcal{A}') = X^*$$

تعریف ۱۴. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آن گاه Q'' را به صورت زیر تعریف می کنیم

$$Q'' = \{q \in Q \mid \sqcup \{\delta^*(q, u, p) \cap \tau(p) \mid u \in X^*, p \in Q\} \neq \{0\}\}$$

و آن را بخش هم-در دسترس Q می نامیم. اگر $Q'' = Q$ باشد، آنگاه $\mathcal{A}$ را هم-در دسترس می نامیم.

ملاحظه ۲. اگر $q \in Q''$ و تنها اگر وجود داشته باشد $p \in Q, u \in X^*$ بطوریکه $\delta^*(q, u, p) \cap \tau(p) \neq \{0\}$.

تعریف ۱۵. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد و $Q'' = (Q'', X, \delta'', \iota'', \tau'')$ آنگاه Q'' بخش هم-در دسترس Q باشد. جاییکه $\delta'' = \delta|_{Q'' \times X \times Q''}$ می نامیم، جاییکه $\tau'' = \tau|_{Q''}$ و $\iota'' = \iota|_{Q''}$ بنابراین $\mathcal{A}''$ یک اتوماتا L -فازی مردد است.

مثال ۳. فرض کنید L مشبکه کراندار شکل ۱ و $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد، جاییکه $Q = \{q_1, q_2\}$ و $X = \{x\}$ است. آنگاه $\delta^* : Q \times X^* \times Q \rightarrow 2^L$, $\iota : Q \rightarrow 2^L$ و $\tau : Q \rightarrow 2^L$ را مطابق زیر تعریف می کنیم

$$\begin{aligned} \delta(q_1, x, q_1) &= \{b\}, & \delta(q_1, x, q_2) &= \{a\}, \\ \delta(q_2, x, q_1) &= \{0, b\}, & \delta(q_2, x, q_2) &= \{0, b\}, \\ \iota(q_1) &= \{b\}, & \iota(q_2) &= \{0, b\}, \\ \tau(q_1) &= \{0, a\}, & \tau(q_2) &= \{b\}. \end{aligned}$$

بنابراین،

$$\delta(q_1, x, q_1) \cap \tau(q_1) = \{b\} \cap \{0, a\} = \{0\},$$

$$\delta(q_1, x, q_2) \cap \tau(q_2) = \{a\} \cap \{b\} = \{0\}.$$

اکنون داریم

$$\begin{aligned} \delta^*(q_1, x^2, q_2) &= \sqcup \left(\delta(q_1, x, q_1) \cap \delta(q_1, x, q_2), \delta(q_1, x, q_2) \cap \right. \\ &\left. \delta(q_2, x, q_2) \right) = \sqcup \left(\{b\} \cap \{a\}, \{a\} \cap \{0, b\} \right) = \{0\} \sqcup \{0\} = \{0\}, \end{aligned}$$

$$\begin{aligned} \delta^*(q_1, x^2, q_1) &= \sqcup \left(\delta(q_1, x, q_1) \cap \delta(q_1, x, q_1), \delta(q_1, x, q_2) \cap \right. \\ &\left. \delta(q_2, x, q_1) \right) = \sqcup \left(\{b\} \cap \{b\}, \{a\} \cap \{0, b\} \right) = \{b\} \sqcup \{0\} = \{b\}. \end{aligned}$$

و برای هر $x^n \in X^*$ که $n \in \{2, 3, 4, \dots\}$ داریم

$$\delta^*(q_1, x^n, q_1) = \{b\},$$

$$\delta^*(q_1, x^n, q_2) = \{0\}.$$

بنابراین، برای هر $u \in X^*$ داریم

$$\delta^*(q_1, u, q_1) \cap \tau(q_1) = \{b\} \cap \{0, a\} = \{0\},$$

$$\delta^*(q_1, u, q_2) \cap \tau(q_2) = \{0\} \cap \{b\} = \{0\},$$

در نتیجه q_1 یک حالت هم-در دسترس نیست. اکنون برای q_2 داریم

$$\delta(q_2, u, q_2) \cap \tau(q_2) = \{0, b\} \cap \{b\} = \{0, b\} \neq \{0\},$$

بنابراین q_2 یک حالت هم-در دسترس است. از این رو $Q'' = \{q_2\}$ بخش هم-در دسترس Q است. در نتیجه $\mathcal{A}'' = (Q'', X, \delta'', i'', \tau'')$ که در آن داریم

$$\iota'' : Q'' \rightarrow 2^L, \quad \tau'' : Q'' \rightarrow 2^L$$

$$\iota''(q_2) = \{0, b\}, \quad \tau''(q_2) = \{b\}$$

$$\delta''^* : Q'' \times X^* \times Q'' \rightarrow 2^L,$$

$$\delta''(q_2, u, q_2) = \{0, b\},$$

بخش هم-در دسترس اتوماتا $\mathcal{A}$ است، بطوریکه

$$R(\mathcal{A}) = R(\mathcal{A}'') = X^*$$

قضیه ۳. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. اگر $\mathcal{A}' = (Q', X, \delta', \iota', \tau')$ و $\mathcal{A}'' = (Q'', X, \delta'', \iota'', \tau'')$ به ترتیب بخش در دسترس و بخش هم-در دسترس از $\mathcal{A}$ باشند، آنگاه داریم

$$1. R(\mathcal{A}) = R(\mathcal{A}')$$

$$2. R(\mathcal{A}) = R(\mathcal{A}'')$$

تعریف ۱۶. فرض کنید $\mathcal{A} = (Q, X, \iota, \delta, \tau)$ یک اتوماتا L -فازی مردد باشد. آنگاه $\mathcal{A}$ را هم-پوشش گوئیم، اگر در دسترس و هم-در دسترس

باشد. همچنین $\mathcal{A}^t = (Q^t, X, \delta^t, i^t, \tau^t)$ را قسمت هم-پوشش $\mathcal{A}$ گوئیم، جایی که

$$\begin{aligned} Q^t &= \{q \in Q \mid \sqcup \{i(p) \cap \delta^*(p, u, q) \mid u \in X^*, p \in Q\} \neq \\ &\{0\}, \sqcup \{\delta^*(q, u, p) \cap \tau(p) \mid u \in X^*, p \in Q\} \neq \{0\}\} \\ \tau^t &= \tau|_{Q^t} \text{ و } i^t = i|_{Q^t}, \delta^t = \delta|_{Q^t \times X \times Q^t} \end{aligned}$$

قضیه ۴. اگر $\mathcal{A}$ یک اتوماتا L -فازی مردد و $\mathcal{A}^t$ قسمت هم-پوشش $\mathcal{A}$ باشد، آنگاه

$$R(\mathcal{A}) = R(\mathcal{A}^t).$$

اکنون نشان می دهیم، اگر $\mathcal{A}$ یک اتوماتا L -فازی مردد و β مجموعه L -فازی مردد زبان $\mathcal{A}$ باشد. آنگاه یک اتوماتا L -فازی مردد در دسترس کامل مانند $\mathcal{A}_{ca}$ وجود دارد، بطوریکه $\beta_{\mathcal{A}} = \beta_{\mathcal{A}_{ca}}$.

تعریف ۱۷. فرض کنید $\mathcal{A}$ یک اتوماتا L -فازی مردد باشد. آنگاه $\mathcal{A}$ را کامل گوئیم، اگر برای هر $q \in Q$ و $a \in X$ وجود داشته باشد $p \in Q$ و $0 \neq \alpha \in L$ بطوریکه $\alpha \in \delta(q, a, p)$.

قضیه ۵. فرض کنید $\mathcal{A}$ یک اتوماتا L -فازی مردد غیر کامل و $\beta_{\mathcal{A}}$ مجموعه L -فازی مردد زبان $\mathcal{A}$ باشد. آنگاه یک اتوماتا L -فازی مردد کامل مانند $\mathcal{A}_c$ وجود دارد، بطوریکه $\beta_{\mathcal{A}} = \beta_{\mathcal{A}_c}$.

قضیه ۶. فرض کنید $\mathcal{A}$ یک اتوماتا L -فازی مردد و $\beta_{\mathcal{A}}$ مجموعه L -فازی مردد زبان $\mathcal{A}$ باشد. آنگاه یک اتوماتا L -فازی مردد در دسترس کامل مانند $\mathcal{A}_{ca}$ وجود دارد بطوریکه $\beta_{\mathcal{A}} = \beta_{\mathcal{A}_{ca}}$.

ملاحظه ۳. ما تلاش می کنیم در یک مقاله دیگر موضوع بهینه سازی (مینیمال کردن) اتوماتای L -فازی مردد را که زبان آن را حفظ می کند، مورد مطالعه و بررسی قرار دهیم.

مراجع

- [1] VS. Costa, B. Bedregal, On typical hesitant fuzzy automata. Soft Computing, 2020. 24(12):8725–8736.
- [2] A. H. Dehmiry, M. Mashinchi, R. Mesiar, Hesitant L-Fuzzy Sets, International journal of intelligent systems, 33 (2018), 1027-1042.
- [3] J. M. Howie, Automata and languages. Oxford University Press, Inc., 1992.
- [4] R. Jakubows, Application of formal language and fuzzy automata in designing, in: J. Madey (ed.): Int. Conf. on Inf. Processing IFIP-INFOPOL-76 (1977).
- [5] E.T. Lee, Fuzzy languages and their relation to automata (1972), Ph.D. Thesis, Dept. of Elect. Eng. Comp. Sci., Univ. of California, Berkeley, CA.

- [6] E. T. Lee, L. A. Zadeh, Note on fuzzy languages, *Information Sciences*, 1 (1969), 421-434.
- [7] D.S. Malik and J.N. Mordeson, *Fuzzy Automata and Languages: Theory and Applications*, Chapman Hall, CRC Boca Raton, London, New York, Washington DC, 2002.
- [8] E.S. Santos: Fuzzy automata and languages, *Inform. Sci.* 10 (1976) 193-197.
- [9] M.G. Thomason: Finite fuzzy automata, regular fuzzy languages and pattern recognition, *Pattern Recogn.* 5 (1973) 383-390.
- [10] V. Torra, Hesitant fuzzy sets, *International Journal of Intelligent Systems* 25.6 (2010): 529-539.
- [11] V. Torra, Y. Narukawa, On hesitant fuzzy sets and decision, 2009 IEEE International Conference on Fuzzy Systems. IEEE, 2009.
- [12] W. G. Wee, On generalization of adaptive algorithm and application of the fuzzy sets concept to pattern classification, Ph.D. Thesis, Purdue University, Lafayette, IN, 1967.
- [13] L. A. Zadeh, Fuzzy sets, *Information and Control*, 8 (1965,) 338-353



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بهم

ارائه یک روش جدید مبتنی بر ترکیب ویژگی‌های استخراج شده از شبکه‌های عصبی عمیق DenseNet169 و MobileNet و دسته‌بندی کننده LightGBM به منظور تشخیص بیماری کرونا از روی تصاویر اشعه ایکس

حمید نصیری^۱، غزل خیرالدین^۲، مرتضی دری‌گیو^۳

^۱ دانشجوی دکتری، دانشکده مهندسی کامپیوتر، دانشگاه صنعتی امیرکبیر، h.nasiri@aut.ac.ir

^۲ دانشجوی کارشناسی، دانشکده مهندسی برق و کامپیوتر، دانشگاه سمنان، gh.kheyroddin@semnan.ac.ir

^۳ استادیار، دانشکده مهندسی برق و کامپیوتر، دانشگاه سمنان، dorrigiv@semnan.ac.ir

چکیده: بیماری کووید-۱۹ اولین بار در شهر ووهان چین یافت شد و به سرعت به مناطق مختلف جهان راه پیدا کرد. پس از شیوع بیماری کرونا در سراسر دنیا، پژوهشگران بسیاری شروع به پیدا کردن راهی برای تشخیص بیماری کرونا از روی تصاویر اشعه ایکس قفسه سینه مراجعان کردند. از آنجایی که تشخیص به موقع این بیماری می‌تواند مراحل درمان را اثربخش کند، وجود روشی که بتواند تشخیص بیماری را سرعت ببخشد بسیار حیاتی می‌باشد. در این مقاله روش جدیدی برای این مسئله پیشنهاد شده است که علاوه بر دقت بالاتر از روش‌های پیشین، سرعت بسیار بالایی نیز دارد. در روش پیشنهادی از ترکیب شبکه‌های عصبی عمیق DenseNet169 و MobileNet به منظور استخراج ویژگی‌های تصاویر اشعه ایکس قفسه سینه بیماران استفاده شده است؛ سپس با الگوریتم انتخاب ویژگی تک متغیره مهم‌ترین این ویژگی‌ها انتخاب شده و به عنوان ورودی به الگوریتم LightGBM داده شده است تا عمل دسته‌بندی را انجام دهد. به منظور ارزیابی کارایی روش پیشنهادی از مجموعه داده ChestX-ray8 که شامل ۱۱۲۵ تصویر اشعه ایکس قفسه سینه بیماران می‌باشد، استفاده شده است. روش پیشنهادی در مسئله دو کلاسه (کرونا، سالم) دقت ۹۸/۵۴ درصد و در مسئله سه کلاسه (کرونا، سالم، ذات‌الریه) دقت ۹۱/۱۱ درصد را بدست آورده است.

کلمات کلیدی: ویروس کرونا، شبکه DenseNet169، شبکه MobileNet، الگوریتم LightGBM، انتخاب ویژگی تک متغیره

۱- مقدمه

کووید-۱۹ در صورت عدم رسیدگی و مراقب از بیمار، موجب درگیری ریه و مرگ انسان می‌شود. تا ۱۸ شهریور ۱۴۰۰، تعداد ۲۲۳,۷۲۰,۲۴۶ نفر به این ویروس مبتلا شده‌اند و ۴,۶۱۴,۲۵۰ نفر بر اثر این ویروس جان باختند. بهترین روش تشخیص ویروس کرونا آزمایش سنجش خلط^۱ می‌باشد ولی به علت زمان‌بر بودن این روش، از تصاویر CT اسکن قفسه سینه بیمار استفاده می‌شود که در مراکز دارای دستگاه‌های رادیوگرافی قابل تهیه است [2]. در این پژوهش از تصاویر اشعه ایکس قفسه سینه بیمار استفاده خواهد شد.

در سال ۱۹۶۵ کرونا ویروس‌ها که عضوی از خانواده ویروس‌ها هستند کشف شدند [1]. این ویروس‌ها از ویروس سرماخوردگی تا کووید-۱۹ را شامل می‌شوند. تاکنون هفت کرونا ویروس منتقل شده به انسان کشف شده است که آخرین نوع آن‌ها، کرونا ویروس سندروم حاد تنفسی^۲، در سال ۲۰۱۹ در شهر ووهان چین یافت شد و توسط سازمان بهداشت جهانی (WHO) کووید-۱۹ نامیده شد. از علائم این ویروس می‌توان به سرفه، نفس تنگی، گلو درد، تب و بدن درد اشاره کرد.

^۱ RT-PCR

۳-۱- شبکه‌های عصبی عمیق

تعداد داده‌های محدود باعث روی آوردن به یادگیری انتقالی برای به دست آوردن نتایج قابل قبول می‌شود. در این پژوهش از مدل‌های از پیش آموزش دیده DenseNet169 و MobileNet برای استخراج ویژگی‌های تصاویر استفاده شده است.

شبکه DenseNet169 [11] شبکه‌ای از گروه بزرگ DenseNet^۸ است که از قبل بر روی مجموعه داده ImageNet آموزش داده شده‌اند. ورودی این شبکه یک تصویر واحد می‌باشد. این شبکه، با اتصال هر لایه فعلی به لایه‌های قبلی شبکه، امکان فشرده‌تر شدن شبکه را فراهم می‌سازد که باعث کاهش تعداد کانال‌ها در شبکه و پیچیدگی محاسباتی می‌شود. حجم شبکه DenseNet169 برابر با ۵۷ مگابایت می‌باشد و توانسته است روی مجموعه داده ImageNet دقت ۹۳/۲ درصد را به دست آورد.

شبکه عصبی عمیق MobileNet [12] توسط محققان گوگل جهت معرفی یک شبکه عصبی عمیق سبک پیاده‌سازی شد. در شبکه عصبی عمیق موبایل نت یک نوع کانولوشن جدید به نام کانولوشن تفکیک‌پذیر بر حسب کانال^۹ معرفی شد که شامل دو فیلتر کانولوشن عمیق و کانولوشن نقطه‌ای می‌باشد. فیلتر کانولوشن عمیق یک کانولوشن واحد را روی هر کانال ورودی انجام می‌دهد و فیلتر کانولوشن نقطه‌ای خروجی کانولوشن عمیق را به صورت خطی با کانولوشن ۱×۱ ترکیب می‌کند. این شبکه دارای ۱۶ مگابایت حجم می‌باشد و توانسته است روی مجموعه داده ImageNet دقت ۸۹/۵ درصد را کسب کند.

۳-۲- انتخاب ویژگی تک متغیره

در این الگوریتم برای پیدا کردن رابطه بین متغیرها و میزان وابستگی هر متغیر غیرمستقل از آزمون^{۱۰} C^۲ استفاده می‌شود. این آزمون برای تعیین رابطه بین دو متغیر استفاده می‌شود. بدین صورت می‌توان میزان استقلال دو متغیر را از یکدیگر تشخیص داد. پس با کمک گرفتن از این آزمون می‌توان ویژگی‌هایی که اثربخشی بیشتری روی تولید کلاس‌ها دارند را به دست آورد.

۳-۳- الگوریتم LightGBM

الگوریتم LightGBM اولین بار توسط بخشی از تیم تحقیق و توسعه شرکت مایکروسافت به وجود آمد [13]. LightGBM الگوریتمی مبتنی بر درخت تصمیم^{۱۱} است که به دلیل تولید درخت‌های پیچیده‌تر دقت بالاتری نسبت به سایر الگوریتم‌های مبتنی بر تقویت درخت^{۱۱} دارد. برخلاف تولید درخت‌های پیچیده، به

روش پیشنهادی در این مقاله از شبکه عصبی عمیق MobileNet و DenseNet169 برای استخراج ویژگی‌های تصاویر اشعه ایکس قفسه سینه بیماران استفاده می‌کند و پس از ترکیب ویژگی‌های استخراج شده این دو شبکه با یکدیگر، مهم‌ترین آن‌ها انتخاب شده و به عنوان ورودی به الگوریتم دسته‌بندی کننده LightGBM^{۱۲} داده می‌شوند تا این الگوریتم برای دسته‌بندی تصاویر از آن‌ها استفاده نماید.

۲- پژوهش‌های پیشین

از تورک^۲ و همکارانش [3] در سال ۲۰۲۰ برای پژوهش خود در این زمینه از شبکه DarkNet استفاده کردند. مدل آموزش دیده آن‌ها از ۱۷ لایه کانولوشن تشکیل شده‌است. آن‌ها موفق شدند در دسته‌بندی سه کلاس به دقت ۸۷/۰۲ و در دسته‌بندی دو کلاس به دقت ۹۸/۰۸ درصد دست پیدا کنند. آن‌ها برای پژوهش خود از مجموعه داده‌ای متشکل از ۱۱۲۵ تصویر که ۵۰۰ تصویر ریه افراد سالم، ۵۰۰ تصویر ریه افراد با بیماری ذات‌الریه و ۱۲۵ تصویر افراد با بیماری کووید-۱۹ استفاده کردند. عباس^۳ و همکارانش [4] از معماری کانولوشن DeTraC برای دسته‌بندی دو کلاس استفاده کردند و به دقت ۹۳/۱ درصد رسیدند. آن‌ها از مجموعه داده‌ای متشکل از تصاویر قفسه سینه ۱۲۶ نفر مبتلا به کرونا و ۸۰ نفر سالم برای پژوهش خود استفاده کردند. مینایی و همکارانش [5] از چهار مدل با معماری کانولوشن برای مسئله دو کلاس استفاده کردند. این مدل‌ها شامل ResNet18، SqueezeNet، DenseNet121 و ResNet50 بودند.

وانگ^۴ و همکارانش [6] در پژوهش خود از شبکه با معماری کانولوشن COVID-Net استفاده کردند و دقت ۹۳/۳ درصد را گزارش کردند. همدان^۵ و همکارانش [7] در تحقیقات خود پس از بررسی شبکه‌های مختلف، با استفاده از شبکه‌های VGG19 و DenseNet201 به دقت ۹۰ درصد دست پیدا کردند. نارین^۶ و همکارانش [8] با آزمایش پنج شبکه عصبی از پیش آموزش داده شده روی سه مجموعه داده مختلف، شبکه عصبی ResNet50 را انتخاب کردند، آن‌ها با استفاده از این شبکه عصبی به میانگین دقت ۹۸ درصد در مسئله دو کلاس دست پیدا کردند. بارسوگان^۷ و همکارانش [9] ویژگی‌های تصاویر را از چهار مجموعه داده مختلف استخراج و برای عمل دسته‌بندی از ماشین بردار پشتیبان استفاده کردند؛ نتیجه پژوهش آن‌ها دست‌یافتن به دقت ۹۹/۶۸ در مسئله دو کلاس بود. نصیری و حسنی [10] از شبکه عصبی عمیق DenseNet169 به منظور استخراج ویژگی تصاویر و از الگوریتم XGBoost برای دسته‌بندی آن‌ها استفاده کردند. آن‌ها در مسئله دو کلاس به میانگین دقت ۹۸/۲۴ درصد و در مسئله سه کلاس به دقت ۸۹/۷۰ درصد رسیدند.

۳- پیش‌نیازها

در روش پیشنهادی از شبکه‌های عصبی عمیق DenseNet169 و MobileNet، الگوریتم انتخاب ویژگی تک متغیره و الگوریتم LightGBM استفاده شده است که در این بخش به معرفی هر یک از آن‌ها می‌پردازیم.

^۸ Barstugan

^۹ Densely connected convolutional neural networks

^{۱۰} Depth-wise separable convolution

^{۱۱} Decision Tree

^{۱۲} Tree boosting

^۱ Light Gradient Boosting Machine

^۲ Ozturk

^۳ Abbas

^۴ Wang

^۵ Hemdan

^۶ Narin

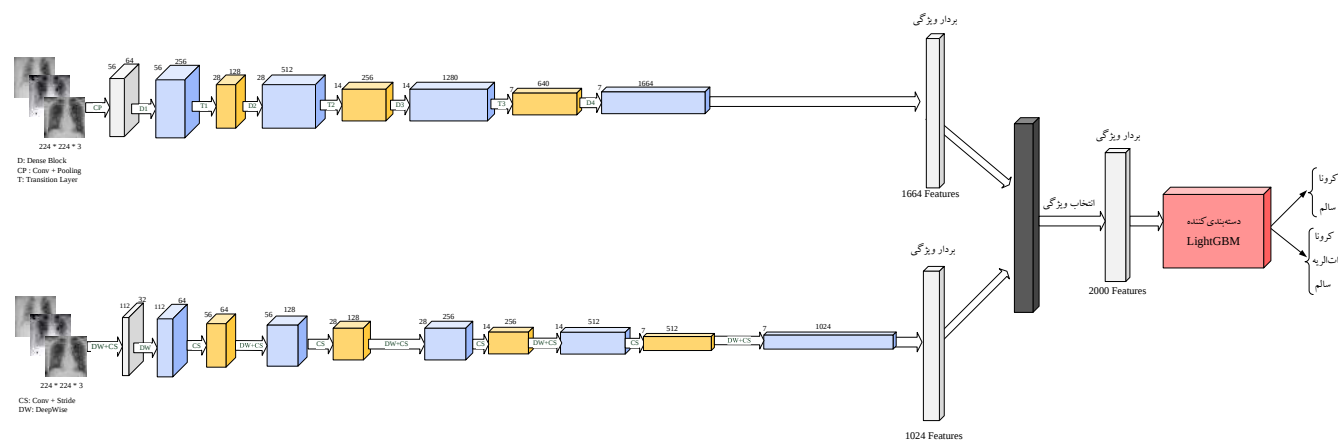
دلیل استفاده از الگوریتم‌های GOSS^۱ و EFB^۲ این الگوریتم بسیار سریع است و حافظه بسیار کمی مصرف می‌کند. همچنین به دلیل وجود پارامترهای مناسب، با تغییر این پارامترها می‌توان از بیش‌برازش^۳ مدل جلوگیری کرد [13].

$$\bar{V}_j(d) = \frac{1}{n} \frac{\sum_{x_i \in A_j} (\hat{a}_{x_i} g_i + \frac{1-a}{b} \hat{a}_{x_i} g_i)^2}{nl(d)} + \frac{(\hat{a}_{x_i} g_i + \frac{1-a}{b} \hat{a}_{x_i} g_i)^2}{nr(d)} \quad (1)$$
$$\begin{aligned} A_l &= \{x_i \hat{=} A : x_{i,j} \leq d\}, A_r = \{x_i \hat{=} A : x_{i,j} > d\} \\ B_l &= \{x_i \hat{=} B : x_{i,j} \leq d\}, B_r = \{x_i \hat{=} B : x_{i,j} > d\} \end{aligned} \quad (2)$$

۴- روش پیشنهادی

۵- آزمایشات ارزیابی

مرحله اول: در این مرحله، دوازده شبکه عصبی از پیش آموزش داده شده مورد ارزیابی قرار گرفتند و با استفاده از هر کدام از آن‌ها ویژگی‌های تصاویر استخراج شد و به ازای هر شبکه دسته‌بندی‌کننده آموزش داده شد و در نهایت تمامی نتایج با یکدیگر مقایسه شدند. با توجه به دقت دسته‌بندی‌کننده LightGBM، شبکه‌هایی که بهترین ویژگی‌ها را استخراج کرده بودند تشخیص داده شدند. شبکه‌های عصبی از پیش آموزش داده شده DenseNet169 و MobileNet به عنوان استخراج‌کننده‌های ویژگی در این پژوهش انتخاب شدند. به منظور استخراج ویژگی توسط این شبکه‌های عصبی، ابتدا اندازه تمامی تصاویر به 224×224 تغییر



[†] Exclusive Function Bundling



شکل ۲: نمونه تصاویر موجود در مجموعه داده

همانطور که در جدول (۲) قابل مشاهده است، در مسئله سه کلاس به معیارهای صحت، امتیاز F_1 و دقت روش پیشنهادی عملکرد بهتری نسبت به سایر روش‌ها داشته است. در معیار تشخیص روش ارائه شده توسط نصیری و حسنی بهترین عملکرد را داشته و در معیار حساسیت روش پیشنهادی و روش ارائه شده توسط نصیری و حسنی بهترین عملکرد را داشته‌اند. در مسئله دو کلاس به توجه به نتایج جدول (۳) در معیار حساسیت روش پیشنهادی این مقاله بهترین عملکرد را داشته است و به میانگین ۹۹/۶۰ درصد رسیده است. در این معیار روش ارائه شده توسط از ترک و همکارانش در جایگاه دوم قرار گرفته است. در معیار تشخیص، الگوریتم ارائه شده توسط نصیری و حسنی از سایر روش‌ها عملکرد بهتری داشته است و به میانگین ۹۹/۷۸ درصد رسیده است. اگرچه در معیار صحت نتایج بدست آمده توسط سه مقاله به یکدیگر نزدیک است ولی روش ارائه شده توسط نصیری و حسنی بهترین عملکرد را در این معیار داشته است. در معیارهای امتیاز F_1 و دقت نیز بهترین عملکرد مربوط به روش پیشنهادی مقاله حاضر بوده است. در این دو معیار الگوریتم ارائه شده توسط نصیری و حسنی نسبت به روش ارائه شده توسط از ترک و همکارانش عملکرد بهتری داشته است. لازم به ذکر است که روش پیشنهادی به علت استفاده از دسته‌بندی‌کننده LightGBM و عدم آموزش شبکه عصبی عمیق، به نسبت روش ارائه شده توسط از ترک و همکارانش که یک شبکه عصبی عمیق را آموزش می‌دهد، پیچیدگی محاسباتی کمتری دارد و مراحل استخراج ویژگی، انتخاب ویژگی و دسته‌بندی با سرعت بسیار بالایی قابل انجام است. در شکل‌های (۳) و (۴) ماتریس درهم‌ریختگی برای مسائل سه و دو کلاس قابل مشاهده است. جدول (۴) مقایسه بین نتایج روش پیشنهادی با دیگر پژوهش‌ها را نشان می‌دهد.

۶- نتیجه‌گیری

در این پژوهش یک مدل مبتنی بر یادگیری عمیق برای تشخیص بیماری کرونا از روی تصاویر اشعه ایکس قفسه سینه بیماران ارائه شده است. روش پیشنهادی با استفاده از شبکه‌های عصبی عمیق DenseNet169 و MobileNet و ویژگی‌های تصاویر را استخراج و با استفاده از روش انتخاب ویژگی تک متغیره تعدادی از مهم‌ترین ویژگی‌ها را انتخاب می‌کند و سپس با استفاده از الگوریتم LightGBM تصاویر را دسته‌بندی می‌نماید. این روش به دقت ۹۸/۵۴ درصد در دسته‌بندی دو کلاس (سالم و مبتلا به کرونا) و ۹۱/۱۱ درصد در دسته‌بندی سه کلاس (سالم، مبتلا به کرونا و ذات‌الریه) دست یافته است. نتایج ارزیابی روش پیشنهادی نشان می‌دهد که این روش به دلیل استفاده از الگوریتم LightGBM، سرعت و دقت بالاتری نسبت به پژوهش‌هایی که در گذشته در این زمینه انجام شده است، دارد. امید است تا روش ارائه شده در این مقاله کمکی به کادر زحمت‌کش درمان کند و سرعت تشخیص و مهار بیماری کرونا را بهبود بخشد.

۷- دسترسی به پیاده‌سازی مقاله

پیاده‌سازی انجام شده برای روش پیشنهادی این مقاله به همراه مجموعه داده مورد استفاده جهت ارزیابی روش پیشنهادی از طریق گیت‌هاب^۵ در دسترس می‌باشد.

پیدا کرد. پس از آن مقادیر پیکسل‌های تصویر نرمال شد و به عنوان ورودی به شبکه عصبی عمیق داده شد. پس از اتمام عملیات، شبکه DenseNet169 تعداد ۱۶۶۴ ویژگی و شبکه MobileNet تعداد ۱۰۲۴ ویژگی را استخراج کردند.

مرحله دوم: به منظور بالا بردن دقت مدل و قدرت تعمیم آن، ویژگی‌های استخراج شده توسط دو شبکه، با یکدیگر ترکیب می‌شوند. در این مرحله بردار ویژگی با ۲۶۸۸ ویژگی بدست می‌آید. از آنجایی که ویژگی‌های غیر ضروری، کارایی مدل را پایین می‌آورند و همچنین برای بالا بردن سرعت آموزش دسته‌بندی‌کننده، باید تعداد ویژگی‌ها تا حد امکان کاهش پیدا کند. برای انتخاب ویژگی از روش انتخاب ویژگی تک متغیره استفاده شد که از آزمون χ^2 به منظور انتخاب K بهترین ویژگی استفاده می‌کند. این آزمون میزان مستقل بودن هر ویژگی نسبت به کلاس‌ها را محاسبه می‌نماید. در این مقاله مقدار K با آزمون و خطا ۲۰۰۰ انتخاب شد.

مرحله سوم: ویژگی‌های انتخاب‌شده به عنوان ورودی به الگوریتم LightGBM داده می‌شوند و عمل دسته‌بندی انجام می‌شود. علت انتخاب این الگوریتم سرعت بالای آن در آموزش و دقت بالای آن است. در این پژوهش، پارامترها با آزمون و خطا انتخاب شده‌اند. در جدول (۱) پارامترهای دسته‌بندی‌کننده LightGBM قابل مشاهده است.

جدول ۱: پارامترهای دسته‌بندی‌کننده LightGBM

پارامتر	مقدار
نرخ یادگیری	۰/۲۴
حداکثر تعداد برگ‌ها	۱۰۵
تعداد تکرار الگوریتم	۲۵۰
حداکثر عمق درخت	۷
تعداد درخت‌ها	۳۰۰
حداقل تعداد داده مورد نیاز در برگ	۴۰

روش پیشنهادی در مسئله سه کلاس به دقت ۹۱/۱۱ درصد را بدست آورده است. همچنین میانگین دقت در آزمایش اعتبارسنجی متقابل ۵ بخشی^۱ برای مسئله دو کلاس ۹۸/۵۴ درصد محاسبه شده است. نتایج بدست آمده برای مسائل سه کلاس و دو کلاس به ترتیب در جداول (۲) و (۳) نشان داده شده است. نتایج بدست آمده با نتایج مقاله از ترک و همکارانش [3] و مقاله نصیری و حسنی [10] مقایسه شده است. به منظور مقایسه بهتر علاوه بر دقت، معیارهای حساسیت^۲، صحت^۳، تشخیص^۴ و امتیاز F_1 نیز گزارش شده‌اند.

جدول ۲: مقایسه روش پیشنهادی با سایر مقالات (مسئله سه کلاس)

روش	حساسیت (%)	تشخیص	صحت	امتیاز F_1	دقت
روش پیشنهادی	۹۵/۲۰	۹۷/۱۶	۹۵/۱۲	۹۵/۶	۹۱/۱۱
از ترک و همکارانش [3]	۸۸/۱۷	۹۳/۶۶	۹۰/۹۷	۸۹/۴۴	۸۹/۳۳

^۱ 5 Fold Cross Validation

^۲ Sensitivity

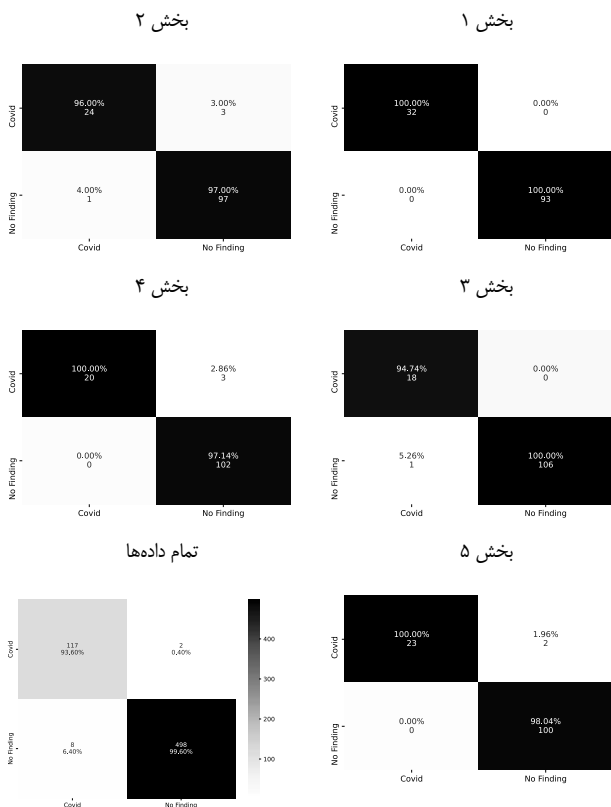
^۳ Precision

^{*} Specificity

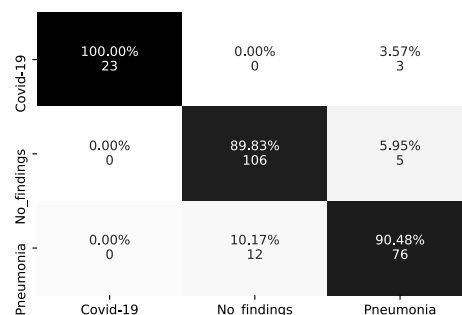
^۵ <https://github.com/gha7all/Covid19-Automated-Detection>

جدول ۳: مقایسه روش پیشنهادی با سایر مقالات (مسئله دو کلاسه)

معیار ارزیابی (%)	روش	بخش ۱	بخش ۲	بخش ۳	بخش ۴	بخش ۵	میانگین
حساسیت	روش پیشنهادی	۱۰۰	۹۸/۹۷	۹۹/۰۶	۱۰۰	۱۰۰	۹۹/۶۰
	ازترک و همکارانش	۱۰۰	۹۶/۴۲	۹۰/۴۷	۹۳/۷۵	۹۳/۱۸	۹۵/۱۳
	نصیری و حسنی	۹۵/۲۰	۹۵/۴۰	۹۶/۷۰	۸۱/۴۰	۹۱/۴۰	۹۲/۰۸
تشخیص	روش پیشنهادی	۱۰۰	۸۸/۸۸	۱۰۰	۸۶/۹۵	۹۲/۰۰	۹۳/۵۶
	ازترک و همکارانش	۱۰۰	۹۶/۴۲	۹۰/۴۸	۹۳/۷۵	۹۳/۱۸	۹۵/۳۰
	نصیری و حسنی	۱۰۰	۱۰۰	۱۰۰	۸۹/۹۰	۱۰۰	۹۹/۷۸
صحت	روش پیشنهادی	۱۰۰	۹۷/۰۰	۱۰۰	۹۷/۱۴	۹۸/۰۳	۹۸/۴۳
	ازترک و همکارانش	۱۰۰	۹۴/۵۲	۹۷/۰۶	۹۹/۰۲	۹۹/۴۸	۹۸/۰۳
	نصیری و حسنی	۹۹/۵۰	۹۹/۵۰	۹۹/۴۰	۹۵/۳۰	۹۹/۰۲	۹۸/۵۴
امتیاز F1	روش پیشنهادی	۱۰۰	۹۷/۹۷	۹۷/۵۳	۹۸/۵۵	۹۹/۰۰	۹۹/۰۱
	ازترک و همکارانش	۱۰۰	۹۸/۵۲	۹۳/۷۹	۹۵/۹۳	۹۵/۶۲	۹۶/۵۱
	نصیری و حسنی	۹۸/۵۰	۹۸/۵۰	۹۸/۲۰	۹۲/۵۰	۹۷/۳۰	۹۷/۰۰
دقت	روش پیشنهادی	۱۰۰	۹۶/۸۰	۹۹/۲۰	۹۷/۶۰	۹۸/۴۰	۹۸/۴۰
	ازترک و همکارانش	۱۰۰	۹۷/۶۰	۹۶/۸۰	۹۷/۶۰	۹۷/۶۰	۹۸/۰۸
	نصیری و حسنی	۹۹/۲۰	۹۹/۲۰	۹۹/۲۰	۹۵/۲۰	۹۸/۴۰	۹۸/۲۴



شکل ۴: ماتریس درهم‌ریختگی (مسئله دو کلاسه)



شکل ۳: ماتریس درهم‌ریختگی (مسئله سه کلاسه)

جدول ۴: مقایسه نتایج روش پیشنهادی با مقالات مشابه

نام محقق	تعداد تصاویر	روش مورد استفاده	دقت (%)
آپستولوپولوس [15]	۱۴۲۸	VGG19	۹۳/۴۸
وانگ [6]	۱۳۶۴۵	COVID-Net	۹۲/۴۰
همدان [7]	۵۰	COVIDX-Net	۹۰/۰۰
ستی [16]	۵۰	ResNet50+SVM	۹۵/۳۸
نارین [8]	۱۰۰	Deep CNN ResNet50	۹۸/۰۰
ازتورک [3]	۶۲۵	DarkCovidNet	۹۸/۰۸
	۱۱۲۵		۸۹/۳۳
نصیری [10]	۶۲۵	DenseNet169+XGBoost	۹۸/۲۳
	۱۱۲۵		۸۹/۷۰
پژوهش فعلی	۶۲۵	DenseNet169+MobileNet + LightGBM	۹۸/۵۴
	۱۱۲۵		۹۱/۱۱

مراجع

- [1] E. Mahase, "Covid-19: Coronavirus was first described in The BMJ in 1965," *BMJ*, vol. 369, no. April, p. m1547, 2020, doi: 10.1136/bmj.m1547.
- [2] T. Singhal, "A Review of Coronavirus Disease-2019 (COVID-19)," *Indian J. Pediatr.*, vol. 87, no. 4, pp. 281–

“Detection of coronavirus disease (COVID-19) based on deep features and support vector machine,” *Int. J. Math. Eng. Manag. Sci.*, vol. 5, no. 4, pp. 643–651, 2020.

- [3] T. Ozturk, M. Talo, E. A. Yildirim, U. B. Baloglu, O. Yildirim, and U. R. Acharya, “Automated detection of COVID-19 cases using deep neural networks with X-ray images,” *Comput. Biol. Med.*, vol. 121, 2020, doi: 10.1016/j.compbiomed.2020.103792.
- [4] A. Abbas, M. M. Abdelsamea, and M. M. Gaber, “Classification of COVID-19 in chest X-ray images using DeTraC deep convolutional neural network,” *Appl. Intell.*, vol. 51, no. 2, pp. 854–864, 2021, doi: 10.1007/s10489-020-01829-7.
- [5] S. Minaee, R. Kafieh, M. Sonka, S. Yazdani, and G. J. Soufi, “Deep-covid: Predicting covid-19 from chest x-ray images using deep transfer learning,” *Med. Image Anal.*, vol. 65, p. 101794, 2020.
- [6] H. Gunraj, L. Wang, and A. Wong, “COVIDNet-CT: A Tailored Deep Convolutional Neural Network Design for Detection of COVID-19 Cases From Chest CT Images,” *Front. Med.*, vol. 7, pp. 1–12, 2020, doi: 10.3389/fmed.2020.608525.
- [7] E. E.-D. Hemdan, M. A. Shouman, and M. E. Karar, “Covidx-net: A framework of deep learning classifiers to diagnose covid-19 in x-ray images,” *arXiv Prepr. arXiv:2003.11055*, 2020.
- [8] A. Narin, C. Kaya, and Z. Pamuk, “Automatic detection of coronavirus disease (covid-19) using x-ray images and deep convolutional neural networks,” *Pattern Anal. Appl.*, pp. 1–14, 2021.
- [9] M. Barstugan, U. Ozkaya, and S. Ozturk, “Coronavirus (COVID-19) Classification using CT Images by Machine Learning Methods,” *arXiv Prepr. arXiv:2003.09424*, 2020.
- [10] H. Nasiri and S. Hasani, “Automated detection of COVID-19 cases from chest X-ray images using deep neural network and XGBoost,” *arXiv Prepr. arXiv:2109.02428*, 2021, [Online]. Available: <http://arxiv.org/abs/2109.02428>.
- [11] G. Huang, Z. Liu, L. Van Der Maaten, and K. Q. Weinberger, “Densely connected convolutional networks,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 4700–4708.
- [12] A. G. Howard *et al.*, “Mobilenets: Efficient convolutional neural networks for mobile vision applications,” *arXiv Prepr. arXiv:1704.04861*, 2017.
- [13] G. Ke *et al.*, “Lightgbm: A highly efficient gradient boosting decision tree,” *Adv. Neural Inf. Process. Syst.*, vol. 30, pp. 3146–3154, 2017.
- [14] X. Wang, Y. Peng, L. Lu, Z. Lu, M. Bagheri, and R. M. Summers, “Chestx-ray8: Hospital-scale chest x-ray database and benchmarks on weakly-supervised classification and localization of common thorax diseases,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 2097–2106.
- [15] I. D. Apostolopoulos and T. A. Mpesiana, “Covid-19: automatic detection from x-ray images utilizing transfer learning with convolutional neural networks,” *Phys. Eng. Sci. Med.*, vol. 43, no. 2, pp. 635–640, 2020.
- [16] P. K. Sethy, S. K. Behera, P. K. Ratha, and P. Biswas,



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بوم

ارائه یک سامانه توصیه‌گر مبتنی بر مدل برای بیماران مبتلا به پرفشاری خون

عارفه ولی‌اللهی^۱، مهرداد کارگری^۲، سید میثم علوی^۳

^۱دپارتمان مهندسی فناوری اطلاعات، دانشگاه تربیت مدرس، تهران، ایران، arefeh.valiollahi@modares.ac.ir

^۲دانشیار، گروه مهندسی فناوری اطلاعات، دانشگاه تربیت مدرس، تهران، ایران، m_kargari@modares.ac.ir

^۳دپارتمان مهندسی فناوری اطلاعات، دانشگاه تربیت مدرس، تهران، ایران، meysamalavi@modares.ac.ir

چکیده: پرفشاری خون یک بیماری مزمن است که به عنوان اصلی‌ترین عامل خطر برای بیماری‌های قلبی عروقی شناخته می‌شود. شواهد پزشکی نشان می‌دهد که تشخیص زودهنگام پرفشاری خون، اصلاح عادات‌های زندگی و کنترل دقیق آن می‌تواند روند پیشروی بیماری را کند کرده و پیامدهای بعدی آن را کاهش دهد. در سال‌های اخیر سامانه‌های توصیه‌گر به موازات پیشرفت‌های مشهود در فناوری‌های اطلاعات و هوش مصنوعی به طور قابل توجهی توسعه یافته‌اند که یکی از کاربردهای آن در حوزه‌های مختلف علوم پزشکی است و از آن جمله می‌توان به سامانه‌های توصیه‌گر جهت پیشگیری، کنترل و درمان بیماری‌ها اشاره کرد. در همین راستا در سال‌های قبل تلاش‌هایی در خصوص طراحی سامانه توصیه‌گر در زمینه کنترل و درمان پرفشاری خون در قالب توصیه‌گرهای دارویی یا توصیه‌گرهای سبک زندگی خاص یک بیمار انجام شده است. با توجه به اهمیت موضوع در این مقاله با استفاده از رویکرد پالایش مشارکتی به عنوان یکی از انواع سامانه‌های توصیه‌گر و نیز الگوریتم خوشه‌بندی **K-means** یک مدل توصیه‌گر برای بیماری پرفشاری خون ارائه شده است. نتایج این پژوهش نشان می‌دهد که مدل توصیه‌گر ارائه شده، از عملکرد قابل قبولی در مقایسه با نظرات خبرگان این حوزه برخوردار است.

کلمات کلیدی: سامانه توصیه‌گر، پرفشاری خون، خوشه‌بندی، یادگیری ماشین

۱- مقدمه

نسخه که می‌تواند ناشی از عوارض جانبی دارو، تحمل‌پذیری ضعیف بیمار نسبت به داروها یا عدم اثربخشی داروها باشد، با روش‌های درمانی فعلی به خوبی کنترل نمی‌شود [10]. علاوه بر این، تصمیم‌گیری پزشکی که معمولاً بر اساس برخی دستورالعمل‌ها انجام می‌شود، ممکن است منجر به سوگیری‌های نامطلوب، خطاهای انسانی و هزینه‌های بالای پزشکی شده و کیفیت خدمات ارائه‌شده به بیماران را تحت تأثیر قرار دهد [11]. در نظر گرفتن ویژگی‌های فردی بیماران در کنار این دستورالعمل‌ها می‌تواند نقش مهمی در کنترل و درمان پرفشاری خون داشته باشد. به طور کلی تلاش در جهت بهبود روش‌های تشخیص، کنترل و درمان بیماری پرفشاری خون از اهمیت بالایی برخوردار است و سامانه‌های توصیه‌گر مبتنی بر یادگیری ماشین می‌توانند در این مسیر تأثیر بسزایی داشته باشند.

پرفشاری خون یک بیماری مزمن غیر واگیر [1] جدی است که اصلی‌ترین عامل خطر برای بیماری‌هایی نظیر مشکلات قلبی عروقی، سکته مغزی، نارسایی‌های مزمن کلیوی و دیابت است [4]–[2]. این بیماری در حال تبدیل شدن به یک مشکل جدی سلامت در جهان است [5] به طوری که طبق آمار سازمان بهداشت جهانی، فشارخون بالا عامل ۷٫۵ میلیون مرگ (۱۲٫۸ درصد از کل مرگ‌ومیر جهانی) است [6]. همچنین طبق آمارهای این سازمان از هر ۴ بزرگسال ۱ نفر دچار بیماری پرفشاری خون است که پیش‌بینی می‌شود تا سال ۲۰۲۵ بیش از ۱٫۵ میلیارد نفر در سراسر جهان به این بیماری مبتلا شوند [7]. به همین دلیل پژوهش‌های زیادی با هدف پیشگیری، تشخیص زودهنگام، کنترل و درمان این بیماری صورت گرفته است [8]، [9]. این بیماری به دلیل موانعی مانند پایداری ضعیف بیماران به

۱-۱- سامانه‌های توصیه‌گر

سامانه توصیه‌گر ابزاری است که به کاربران کمک می‌کند مواردی را انتخاب کنند که مطابق علایق و ترجیحات آن‌ها باشد. به‌طور کلی سامانه‌های توصیه‌گر در چهار دسته مبتنی بر محتوا، پالایش مشارکتی، مبتنی بر دانش و ترکیبی طبقه‌بندی می‌شوند. در سامانه توصیه‌گر مبتنی بر محتوا، توصیه‌های یک کاربر بر اساس محتوا انجام می‌گیرد. محتوا می‌تواند شامل دسته‌بندی‌های محصولات، عادات، علایق یا اولویت‌های کاربران باشد. در حالت کلی این سامانه‌ها نیازمند اطلاعات دقیق درباره محصولات و کاربران هستند [12]. دسته پالایش مشارکتی یکی از روش‌های مرسوم در توصیه‌دهی است که به بررسی ارتباط کاربران می‌پردازد. یکی از مزیت‌های آن عدم نیاز به آنالیز و تفسیر محصول در توصیه‌دهی است. این سامانه بر این مبنا بنا شده که کاربرانی که به یک محصول علاقه دارند ممکن است به محصول دیگر نیز علاقه داشته باشند [13]. سامانه‌های پالایش مشارکتی را می‌توان به دو زیردسته مبتنی بر مدل و مبتنی بر حافظه تقسیم کرد. در رویکرد مبتنی بر مدل، برای تولید توصیه‌ها مدلی بر اساس رتبه‌بندی کاربر ایجاد می‌شود. رویکرد دوم برای ارائه توصیه‌ها عملیات جستجوی کاربران مشابه را در کل داده انجام می‌دهد [13]. سامانه‌های توصیه‌گر مبتنی بر دانش از اطلاعات کاربر و محصول برای توصیه‌دهی استفاده می‌کنند. در این سامانه‌ها کاربران با بیان علایق خود معیارهای جستجویشان را بیان می‌کنند و سامانه با بررسی این اطلاعات و تطبیق آن‌ها با اطلاعاتی که از قبل در پایگاه‌های داده ذخیره شده، نتایج توصیه‌دهی را بیان می‌کند. از آنجاکه این روش وابسته به رتبه‌دهی کاربر نیست مشکل ارزیابی اولیه و پراکندگی را ندارد و می‌تواند مکملی برای روش‌های دیگر باشد. دسته‌ی آخر سامانه‌های توصیه‌گر ترکیبی هستند که باهدف دستیابی به کارایی بالاتر و کاهش محدودیت‌های ناشی از نقاط ضعف روش‌های قبل، آن‌ها را بر اساس الگوریتم‌های مختلفی باهم ترکیب می‌کند [14].

سامانه‌های توصیه‌گر در بسترهایی مانند خرده‌فروشی آنلاین، خدمات پخش و شبکه‌های اجتماعی مورد استفاده قرار می‌گیرند تا روند انتخاب محصول را برای کاربران تسهیل کنند [16]. در سال‌های اخیر، این سامانه‌ها به‌طور گسترده‌ای در حوزه مراقبت‌های بهداشتی با عنوان «سامانه‌های توصیه‌گر سلامت»^۱ برای حمایت بهتر از تصمیمات پزشکی مورد استفاده قرار گرفته [17] که می‌توانند ضمن ایجاد تجربه بهتری برای بیماران، وضعیت سلامتی آن‌ها را ارتقا داده و آن‌ها را تشویق کنند تا از شیوه زندگی سالم‌تری پیروی کنند. علاوه بر این، به‌عنوان دستیار تشخیص^۲ [18] در پیش‌بینی یا درمان بیماری به متخصصان مراقبت‌های بهداشتی کمک می‌کنند. مهم‌ترین هدف سامانه‌های توصیه‌گر سلامت این است که ضمن اطمینان از صحت، قابلیت اطمینان و حفظ حریم خصوصی اطلاعات بیمار، توصیه‌های لازم را در زمان مناسب ارائه کند. علاوه بر این، انتظار می‌رود این سامانه‌ها هزینه فرایند تصمیم‌گیری مربوط به مراقبت‌های بهداشتی را نیز به حداقل برسانند [17].

در همین راستا در سال‌های اخیر طراحی و استفاده از سامانه‌های توصیه‌گر در حوزه پیشگیری و درمان بیماری پرفشاری خون مورد توجه برخی

محققین قرار گرفته است [19]. با توجه به اهمیت موضوع در این مقاله یک سامانه توصیه‌گر پزشکی برای بیماری پرفشاری خون ارائه شده است که نتایج نشان می‌دهد در مقایسه با تحقیقات پیشین از عملکرد قابل قبولی برخوردار است.

ادامه این مقاله به شرح ذیل سازمان‌دهی می‌شود: در بخش دوم ادبیات موضوع به‌طور خلاصه بررسی می‌شود. در بخش سوم سامانه پیشنهادی معرفی شده است. در بخش چهارم به ارزیابی نتایج پرداخته می‌شود. بخش پنجم مربوط به نتیجه‌گیری و پیشنهاد کارهای آتی است.

۲- مرور ادبیات

یادگیری ماشین در اواخر دهه ۱۹۷۰ آغاز شد. حضور علوم کامپیوتر در دهه ۱۹۹۰ همراه با ایده شبکه عصبی، ماشین بردار پشتیبان^۳ و روش‌های ترکیبی پررنگ شد [20] و در اواخر دهه ۱۹۹۰ سامانه‌های توصیه‌گر برای اولین بار تعریف شدند [13]. در زمینه بیوتکنولوژی، هوش مصنوعی از اواخر دهه ۱۹۹۰ نقش مهمی ایفا کرده است و از روش‌های دسته‌بندی برای تشخیص و پیش‌بینی بسیاری از بیماری‌ها استفاده شده است [21]. تشخیص بیماری‌های قلبی عروقی مبتنی بر یادگیری ماشین در سال ۱۹۹۲ آغاز شد که تا سال ۲۰۰۸ مقالات کمی در این حوزه منتشر می‌شد (حدود ۱ مقاله در سال) اما پس از آن جهش فزاینده‌ای در تعداد مقالات به وجود آمد [22].

در مطالعه [9]، یک سامانه توصیه‌گر ترکیبی کارآمد با رویکرد دسته‌بندی چند کلاسه برای بیماری‌های قلبی عروقی ارائه شده است. پژوهش انجام شده نشان داد این سامانه برای بیماران که معمولاً دسترسی به متخصص قلب و عروق ندارند، عملکرد خوبی دارد. این روش پیشنهادی می‌تواند به متخصصین قلب و عروق جوان و تازه‌کار در تصمیم‌گیری سریع پزشکی کمک کند. پژوهش انجام شده توسط ما ستقیم^۴ و همکاران در راستای ایجاد یک سامانه توصیه‌گر برای بیماران قلبی عروقی ارائه شده که بر مبنای پالایش مشارکتی و استفاده از خوشه‌بندی و زیر خوشه‌بندی در دو مرحله است. نتایج این پژوهش نشان داده است که روش ارائه‌شده با کاهش دامنه جستجو برای رکورد جدید، زمان مورد نیاز برای ارائه توصیه‌ها را هم کاهش می‌دهد [15]. در [22] یک سامانه توصیه‌گر با استفاده از تکنیک‌های تبدیل فوری و یادگیری ماشین برای پیش‌بینی بیماری‌های مزمن قلبی ارائه شده است. نتایج این مطالعه نشان می‌دهد که این مدل برای ارائه توصیه‌های مطمئن و دقیق پیش‌بینی احتمال وقایع مختلف در سلامت افراد ایجاد شده است. هدف از این مطالعه ارائه بینش‌های شخصی، به‌موقع و کاربردی در رابطه با مراقبت‌های بهداشتی بود. در تحقیقی دیگر، مدلی برای ردیابی نتایج درمان دارویی در بیماران مبتلا به پرفشاری خون و یافتن مؤثرترین دارو به‌عنوان نسخه بیمار ارائه شد [24]. این پژوهش، شخصی‌سازی را بر این اساس انجام می‌دهد که آنچه برای بیماران مشابه خوب عمل کرده، برای بیمار جدید هم قابل توصیه است. در [18] سامانه‌ای برای توصیه فعالیت‌های بدنی برای بیماران مبتلا به

^۱ Health Recommender Systems - HRS

^۲ Diagnostic Assistants

^۳ Support vector machine (SVM)

^۴ Mustaqeem

می‌شود. در این مدل از معیار شباهت کسینوسی^۶ استفاده شده است. این معیار تشابه اقلام را به عنوان بردار در نظر می‌گیرد و با اندازه‌گیری زاویه بین آن‌ها، همبستگی را محاسبه می‌کند. شباهت حداکثر زمانی است که دو بردار باهم زاویه صفر درجه داشته باشند، درحالی‌که دو بردار با زاویه ۹۰ درجه هیچ شباهتی ندارند [۸]. رابطه ۱ نحوه محاسبه شباهت کسینوسی را نشان می‌دهد [15].

$$sim(a, b) = \frac{(a) \cdot (b)}{\|a\| \|b\|} = \frac{\sum_{i=1}^n a_i b_i}{\sqrt{\sum_{i=1}^n a_i^2} \sqrt{\sum_{i=1}^n b_i^2}} \quad (1)$$

همچنین برای اعتبارسنجی مدل از روش اعتبارسنجی متقابل^۷ با $k=10$ استفاده شده است. شکل (۱) فرایند کلی اجرای روش را نشان می‌دهد.



شکل ۱: فرایند اجرای روش

۴- ارزیابی

معیارهای ارزیابی استفاده‌شده در این مدل عبارتند از: MAE^۸، Precision، Recall، F1-Score و MAE روشی برای مقایسه مقدار پیش‌بینی‌شده با نتیجه واقعی است و رابطه ۲ نحوه محاسبه آن را نشان

پرفشاری خون ایجاد شد. همچنین در [25] یک مدل توصیه‌گر بر اساس خوشه‌بندی فازی و مدل معنایی ارائه شد که نتایج آن حاکی از بهبود سرعت و دقت توصیه منابع بود. در پژوهشی دیگر، یک سامانه توصیه‌گر غذایی مبتنی بر بیماری پیشنهاد شده که بر ایجاد یک رژیم غذایی شخصی سازی شده با توجه به نیازهای تغذیه‌ای، بیماری‌ها و ترجیحات آن‌ها تمرکز دارد [14].

علی‌رغم اینکه در مطالعات گذشته تلاش شده که مسئله توصیه‌گری برای بیماری پرفشاری خون از زاویه‌های مختلف بررسی شود، وجود سامانه‌ای که بتواند بطور هم‌زمان توصیه‌های دارویی و شخصی را در بر داشته باشد تاکنون ارائه نشده است. همچنین اکثر مطالعات انجام‌شده، با چالش مقیاس‌پذیری و پراکندگی در توصیه‌گری مواجه هستند که ابزارهای یادگیری ماشین می‌توانند در این زمینه کمک‌کننده باشند. با توجه به نقاط ضعف پژوهش‌های پیشین در این مقاله یک سامانه توصیه‌گر مبتنی بر مدل و با استفاده از ابزارهای یادگیری ماشین ارائه شده است.

۳- روش

هدف این مطالعه ارائه درمان بهتر و دقیق‌تر برای بیماران مبتلا به فشارخون بالا بر اساس ویژگی‌های موجود است. به همین منظور، با استفاده از رویکرد مبتنی بر مدل که زیر شاخه‌ای از دسته پالایش مشارکتی است، روشی برای توصیه به بیماران دارای پرفشاری خون پیشنهاد می‌شود.

۳-۱- خوشه‌بندی

مجموعه داده استفاده شده در این تحقیق شامل ۲۳۰ رکورد جمع‌آوری شده از بیماران مبتلا به پرفشاری خون است. پس از پیش‌پردازش داده‌ها عملیات خوشه‌بندی انجام شد. در روش پیشنهادی، خوشه‌بندی با استفاده از الگوریتم K-means و با $k = 8$ انجام شد که تعداد خوشه‌های مناسب با روش آرنج^۵ تعیین شده است. پس از خوشه‌بندی داده‌ها در ۸ گروه، توصیه‌های متناسب با هر خوشه مشخص شد. این مرحله با استفاده از سه روش مختلف برای یافتن این توصیه‌ها صورت گرفت، که در نهایت شامل سه خروجی متفاوت بود. جدول (۱) روش‌های مختلف انتخاب توصیه‌ها را نشان می‌دهد.

جدول ۱: روش‌های مختلف انتخاب توصیه‌های هر خوشه

مدل	نحوه انتخاب توصیه‌های هر خوشه
مدل ۱	انتخاب ۱۰ توصیه برتر برای هر خوشه
مدل ۲	انتخاب ۱۰ درصد اول توصیه‌های هر خوشه
مدل ۳	انتخاب توصیه‌هایی با اشتراک حداقل ۱۰ درصد اعضای هر خوشه

۳-۲- محاسبه شباهت

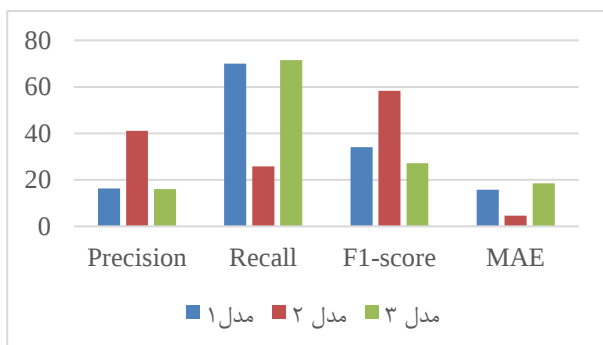
در روش پیشنهادی برای بیمار جدید ابتدا مدل، فردی با بیشترین میزان تشابه با بیمار جدید را مشخص کرده و خوشه‌ای که بیمار مذکور به آن تعلق دارد تعیین می‌شود سپس توصیه‌های مربوط به آن خوشه به بیمار جدید ارائه

^۵ Elbow method

^۶ Cosine similarity

^۷ 10- fold Cross validation

^۸ Mean Absolute Error



شکل ۲: مقایسه نتایج سه مدل

در جدول (۳) دو نمونه از خروجی های مدل توصیه گر نشان داده شده است. در این جدول قسمت توصیه های موجود برای بیمار شامل مجموعه توصیه هایی که پزشک در مواجه با بیمار دارد و بخش توصیه های پیشنهاد شده، مجموعه توصیه هایی است که سامانه توصیه گر پیشنهادی در مورد همان بیمار پیشنهاد کرده است. همانطور که مشخص است سامانه توصیه گر ارائه شده در مقایسه با توصیه های پزشک عملکرد قابل قبولی از خود نشان داده است. پیش تر نیز بیان شد رویکرد پیشنهادی در این پژوهش علاوه بر توصیه های شخصی دربرگیرنده توصیه های دارویی نیز می باشد. که این توصیه ها مربوط به میزان مصرف داروهاست و در این پژوهش این میزان با دو مقدار کم (Low) و زیاد (High) در نظر گرفته شده است. مثلاً توصیه برای بیمار با شماره شناسه ۱۵۹ مصرف کم داروهای مسدودکننده گیرنده آنژیوتانسین^۹ (ABR) می باشد که سامانه توصیه گر پیشنهادی نیز آن را به درستی پیشنهاد کرده است. همچنین برخی توصیه ها از نوع شخصی است که سامانه آن را بر حسب ویژگی ها و حال بیمار به او پیشنهاد می کند. بطور مثال در مورد بیمار با شماره شناسه ۱۹۰ سامانه توصیه گر پیشنهادی به او « لزوم چک کردن مرتب فشار خون در منزل » را پیشنهاد کرده است.

جدول ۳: نمونه خروجی

عنوان	نمونه ۱	نمونه ۲
شناسه بیمار	۱۵۹	۱۹۰
توصیه های موجود برای بیمار	ARB-LOW, Beta blocker-LOW, Diuretics-LOW, Statin-LOW	ARB-LOW, Beta blocker-LOW, Diuretics-LOW, Statin-LOW
توصیه های پیشنهاد شده	ARB-LOW, Beta blocker-LOW, Statin-LOW, ACEI-LOW, Diuretics-LOW, CCU بستری	ARB-LOW, Beta blocker-LOW, Statin-LOW, Calcium channel blocker-LOW, Diuretics-LOW, CCU بستری
	نوبت آنژیوگرافی	چک کردن مرتب فشار خون در منزل
Precision	۶۶.۶۶٪	۵۰٪
Recall	۱۰۰٪	۱۰۰٪
F1-Score	۷۹.۹۹٪	۶۶.۶۶٪
MAE	۵.۱۷	۶.۸۹

می دهد که در آن N تعداد داده های آزمایشی، P نتایج پیش بینی شده و A نتایج واقعی است [26].

$$MAE = \frac{1}{N} \sum_{i=1}^N \|P_i - A_i\| \quad (2)$$

معیارهای Precision و Recall برای ارزیابی کیفیت توصیه های ارائه شده استفاده می شوند. Precision عبارت است از نسبت تعداد توصیه های ارائه شده صحیح به بیمار به تعداد کل توصیه های ارائه شده که محاسبه آن از رابطه ۳ انجام می شود [15].

$$Precision = \frac{\text{Correct recommendations}}{\text{Total recommendations}} \quad (3)$$

recall عبارت است از نسبت تعداد توصیه های ارائه شده ای که بیمار آن ها را دارد به تعداد کل توصیه های موجود برای بیمار و از طریق رابطه ۴ محاسبه می شود [15].

$$Recall = \frac{\text{Correct recommendations}}{\text{Relevant recommendations}} \quad (2)$$

در ادامه نتایج به دست آمده از مدل با در نظر گرفتن سه حالت انتخاب توصیه ها در جدول (۲) ارائه شده است.

جدول ۲: نتایج ارزیابی مدل

معیار ارزیابی	مدل ۱	مدل ۲	مدل ۳
Precision	۱۶.۲۹	۴۱.۰۸	۱۵.۹۵
Recall	۷۰.۰۵	۲۵.۸۱	۷۱.۵۳
F1-Score	۳۴.۰۱	۵۸.۳۵	۲۷.۱۴
MAE	۱۵.۷۱	۴.۵۶	۱۸.۵۳

مقادیر جدول (۲) نشان می دهد که بسته به روش استفاده شده در انتخاب توصیه های هر خوشه، مدل کارایی متفاوتی خواهد داشت. مدل دوم عملکرد بهتری به لحاظ نتیجه معیار F1-score و MAE دارد. این در حالی است که مقدار Recall آن از دو مدل دیگر کمتر است. از طرفی مدل های ۱ و ۳ با Recall بالای ۷۰ درصد توانسته اند میزان قابل قبولی از توصیه های موجود برای بیماران را پیش بینی کنند اما با توجه به کاهش نسبت توصیه های پیش بینی شده به توصیه های موجود برای بیمار، مقدار F1-score کمتری دارند. به عبارتی می توان چنین نتیجه گرفت که برای بهبود معیار Precision، با کاهش معیار Recall مواجه می شویم و بالعکس. درواقع اگر هدف کاهش اختلاف مقادیر پیش بینی شده با توصیه های موجود باشد معیار Precision اهمیت بیشتری پیدا می کند و اگر هدف افزایش احتمال حضور توصیه های بیمار در مقادیر پیش بینی شده باشد معیار Recall دارای اهمیت است. شکل (۲) مقایسه عملکرد سه مدل را نشان می دهد.

^۹ Angiotensin Receptor Blockers (ARB)

در این پژوهش با استفاده از یادگیری ماشین سامانه توصیه‌گری برای بیماران دارای پرفشاری خون به منظور کاهش خطاهای احتمالی، کاهش هزینه‌های درمانی و تسریع روند درمان ارائه شد. به همین منظور با استفاده از روش خوشه‌بندی k-means داده‌ها به ۸ خوشه تقسیم شدند. سپس روش اعتبارسنجی متقابل با $k=10$ برای محاسبه معیارهای ارزیابی به کار گرفته شد. در پژوهش‌های پیشین توصیه‌گرهای از گروه‌های دارویی یا شخصی ارائه شده‌اند و سامانه‌ای که در برگیرنده هر دو باشد وجود نداشته است بنابراین در این پژوهش تلاش شد تا هر دو دسته‌بندی از توصیه‌ها ارائه شود. علاوه بر خطای پایین نشان‌دهنده این است که توصیه‌های پیش‌بینی شده تطابق قابل قبولی با موارد موجود برای بیماران دارند. با توجه به اینکه عوامل بسیاری در تغییر فشارخون افراد تأثیر دارد، به عنوان کارهای آتی می‌توان با افزودن ویژگی‌های دیگر مانند نتایج آزمایشگاهی بیماران، توصیه‌های دقیق‌تری ارائه کرد.

مراجع

- Advances in Science, Engineering and Robotics Technology (ICASERT), pp. 1-6. IEEE, 2019.
- [9] Jabeen, Fouzia, Muazzam Maqsood, Mustansar Ali Ghazanfar, Farhan Aadil, Salabat Khan, Muhammad Fahad Khan, and Irfan Mehmood. "An IoT based efficient hybrid recommender system for cardiovascular disease." *Peer-to-Peer Networking and Applications* 12, no. 5 (2019): 1263-1276.
- [10] Jalali, Niloofer. "Temporal CBR for Personalizing Hypertension Treatment." PhD diss., Northeastern University, 2016.
- [11] Lafta, Raid Luaibi. "An intelligent recommender system based on short-term disease risk prediction for patients with chronic diseases in a telehealth environment." PhD diss., University of Southern Queensland, 2018.
- [12] Nasiri, Mahdi, Behrouz Minaei, and Amir Kiani. "Dynamic recommendation: Disease prediction and prevention using recommender system." *International Journal of Basic Science in Medicine* 1, no. 1 (2016): 13-17.
- [13] Portugal, Ivens, Paulo Alencar, and Donald Cowan. "The use of machine learning algorithms in recommender systems: A systematic review." *Expert Systems with Applications* 97 (2018): 205-227.
- [14] Kuzelewska, Urszula. "Clustering algorithms in hybrid recommender system on movielens data." *Studies in logic, grammar and rhetoric* 37, no. 1 (2014): 125-139.
- [15] Pawar, Raksha, Shazia Lardkhan, Stuti Jani, and Krishna Lakhi. "NutriCure: A Disease-Based Food Recommender System."
- [16] Mustaqeem, Anam, Syed Muhammad Anwar, and Muhammad Majid. "A modular cluster based collaborative recommender system for cardiac patients." *Artificial intelligence in medicine* 102 (2020): 101761.
- [17] Tran, Thi Ngoc Trang, Alexander Felfernig, Christoph Trattner, and Andreas Holzinger. "Recommender systems in the healthcare domain: state-of-the-art and research issues." *Journal of Intelligent Information Systems* 57, no. 1 (2021): 171-201.
- [18] Sharma, Deepika, Gagangeet Singh Aujla, and Rohit Bajaj. "Deep neuro-fuzzy approach for risk and severity prediction using recommendation systems in connected health care." *Transactions on Emerging Telecommunications Technologies* 32, no. 7 (2021): e4159.
- [19] Ferretto, Luciano Rodrigo, Ericles Andrei Bellei, Daiana Biduski, Luiz Carlos Pereira Bin, Mirella Moura Moro, Cristiano Roberto Cervi, and Ana Carolina Bertoletti De Marchi. "A physical activity recommender system for patients with arterial hypertension." *IEEE Access* 8 (2020): 61656-61664.
- [20] Amaratunga, Dhammika, Javier Cabrera, Davit Sargsyan, John B. Kostis, Stavros Zinonos, and William J. Kostis. "Uses and opportunities for machine learning in hypertension research." *International Journal of Cardiology Hypertension* 5 (2020): 100027.
- [21] Zaman, Sabina, and Rizoan Toufiq. "Comparative Study of Back Propagation Neural Network and Support Vector Machine to Classify Hypertension Gene Sequences from Relative Codon Frequency." In *2018 International Conference on Computer, Communication, Chemical, Material and Electronic Engineering (IC4ME2)*, pp. 1-5. IEEE, 2018.
- [22] Alizadehsani, Roohallah, Moloud Abdar, Mohamad Roshanzamir, Abbas Khosravi, Parham M. Kebria, Fahime Khozeimeh, Saeid Nahavandi, Nizal Sarrafzadegan, and U. Rajendra Acharya. "Machine learning-based coronary
- [1] Soh, Desmond Chuang Kiat, E. Y. K. Ng, V. Jahmunah, Shu Lih Oh, Ru San Tan, and U. Rajendra Acharya. "Automated diagnostic tool for hypertension using convolutional neural network." *Computers in Biology and Medicine* 126 (2020): 103999.
- [2] Islam, Md Mazharul, and Rittika Shamsuddin. "Machine learning to promote health management through lifestyle changes for hypertension patients." *Array* 12 (2021): 100090.
- [3] Martinez-Ríos, Erick, Luis Montesinos, Mariel Alfaro-Ponce, and Leandro Pecchia. "A review of machine learning in hypertension detection and blood pressure estimation based on clinical and physiological data." *Biomedical Signal Processing and Control* 68 (2021): 102813.
- [4] Ambika, M., G. Raghuraman, L. SaiRamesh, and A. Ayyasamy. "Intelligence-based decision support system for diagnosing the incidence of hypertensive type." *Journal of Intelligent & Fuzzy Systems* 38, no. 2 (2020): 1811-1825.
- [5] Sookrah, Romeshwar, Jaysree Devesh Dhowtal, and Soukashmee Devi Nagowah. "A DASH diet recommendation system for hypertensive patients using machine learning." In *2019 7th International Conference on Information and Communication Technology (ICoICT)*, pp. 1-6. IEEE, 2019.
- [6] Ghaderi Niri, Amin, Nacer Farajzadeh, and Elahe Baybordi. "A Case Study of the Impact of Parental Diseases on the Probability of Hypertension Using Data Mining Techniques." *Journal of Health and Biomedical Informatics* 7, no. 4 (2021): 354-367. [in Persian]
- [7] Alzubi, Raid, Naeem Ramzan, Hadeel Alzoubi, and Stamos Katsigiannis. "SNPs-based hypertension disease detection via machine learning techniques." In *2018 24th International Conference on Automation and Computing (ICAC)*, pp. 1-6. IEEE, 2018.
- [8] Habib, Sumaya, Maisha Binte Moin, Sujana Aziz, Kalyan Banik, and Hossain Arif. "Heart failure risk prediction and medicine recommendation using exploratory data analysis." In *2019 1st International Conference on*

- artery disease diagnosis: A comprehensive review." *Computers in biology and medicine* 111 (2019): 103346.
- [23] Narayan, Subhashini, and E. Sathiyamoorthy. "A novel recommender system based on FFT with machine learning for predicting and identifying heart diseases." *Neural Computing and Applications* 31, no. 1 (2019): 93-102.
 - [24] Jamshidi, Soheil, Mohamad Ali Torkamani, Jynelle Mellen, Malhar Jhaveri, Penny Pan, James Chung, and Hakan Kardes. "A Hybrid Health Journey Recommender System Using Electronic Medical Records." In *HealthRecSys@ RecSys*, pp. 57-62. 2018.
 - [25] Jalali, Niloofer, Stephen Agboola, Kamal Jethwani, Ibrahim Zeid, and Sagar Kamarthi. "Temporal Case-Based Reasoning for Personalized Hypertensive Treatment." In *ASME International Mechanical Engineering Congress and Exposition*, vol. 50688, p. V014T07A018. American Society of Mechanical Engineers, 2016.
 - [26] Yang, Ye, and Yunhua Zhang. "Collaborative filtering recommendation model based on fuzzy clustering algorithm." In *AIP Conference Proceedings*, vol. 1967, no. 1, p. 040050. AIP Publishing LLC, 2018.



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

ارائه یک معیار شباهت ترکیبی جهت بهبود سیستم‌های مشارکت جمعی آیت‌محور

پردیس پورسیستانی^۱، مسعود سعید^۲، حسین نظام آبادی پور^۳

^۱ دانشجوی دکترا، بخش مهندسی کامپیوتر، دانشگاه آزاد اسلامی کرمان، کرمان، poursistani64@gmail.com
^۲ استادیار، دانشکده فنی مهندسی، بخش مهندسی کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، msaeedmz@uk.ac.ir
^۳ استاد، دانشکده فنی مهندسی، بخش مهندسی برق، دانشگاه شهید باهنر کرمان، کرمان، nezam@uk.ac.ir

چکیده: سیستم‌های مشارکت جمعی حافظه‌محور، از متداول‌ترین روش‌های مورد استفاده در پیاده‌سازی سیستم‌های توصیه‌گر هستند که دارای دو مدل آیت‌محور و کاربرمحور می‌باشند. در هر دو مدل، با استفاده از یک معیار شباهت همسایگان شناسایی شده و از این همسایگان جهت پیش‌بینی رای استفاده می‌شود. با توجه به نقش بسیار مهمی که معیار شباهت در کیفیت پیش‌بینی رای ایفا می‌کند، در این مقاله یک معیار شباهت ترکیبی بر اساس معیارهای شباهت کسینوسی و جاکارد ارائه شده است. نتایج آزمایش‌ها و مقایسه با معیارهای شباهت متداول بیانگر کیفیت قابل قبول معیار شباهت پیشنهادی می‌باشد.

کلمات کلیدی: سیستم‌های توصیه‌گر، فیلترینگ مشارکت جمعی، معیار شباهت، حافظه‌محور، کاربرمحور، آیت‌محور.

ابزارهای تحلیلی متفاوتی استفاده می‌کنند. به طور کلی روش‌هایی که در سیستم‌های توصیه‌گر استفاده می‌شوند بر اساس نوع منابع داده و نحوه تحلیل این داده‌ها، به سه دسته کلی (الف) فیلترینگ مشارکت جمعی^۱، (ب) محتوای محور^۲ و (ج) ترکیبی^۳ تقسیم می‌شوند [۱].

الگوریتم‌های پایه فیلترینگ مشارکت جمعی، به طور معمول یک ماتریس رای را که معرف نظرات کاربران نسبت به آیت‌ها است را به عنوان تنها داده ورودی دریافت کرده و پیش‌بینی رای یا پیشنهادها را بر اساس علائق کاربران مشابه تولید و ارائه می‌کنند. در الگوریتم‌های محتوا محور، پیشنهادها بر اساس محتوای آیت‌ها شکل گرفته و آیت‌هایی که از لحاظ محتوا مشابه آیت‌هایی هستند که قبلاً به آن‌ها ابراز علاقه شده است، پیشنهاد می‌شوند. این روش‌ها

۱- مقدمه

با پیشرفت و توسعه اینترنت و سیستم‌های تجارت الکترونیک، حجم اطلاعات زیادی در اختیار کاربران قرار گرفته است که استفاده از این حجم از اطلاعات به راحتی امکان‌پذیر نمی‌باشد. به همین دلیل، امروزه استفاده از سیستم‌هایی که قادر به ارائه پیشنهادهایی متناسب با نیاز کاربر باشند امری لازم و ضروری است. سیستم‌های توصیه‌گر به دنبال پیش‌بینی رای یا اولویتی هستند که یک کاربر به یک آیت می‌دهد. به عبارت دیگر، هدف این سیستم‌ها شناسایی سلیقه‌های کاربران و پیشنهاد آیت‌های مناسبی است که کاربران از وجود آن‌ها بی‌اطلاع هستند. سیستم‌های توصیه‌گر شخصی‌سازی شده به کاربر کمک می‌کنند تا بر مشکل حجم بالای اطلاعات فائق آمده و به او در یافتن آیت‌های مورد نظرش یاری رسانند.

از زمان مطرح شدن سیستم‌های توصیه‌گر، پژوهشگران به صورت پیوسته رویکردهای جدیدی را برای ساخت آن‌ها پیشنهاد کرده‌اند که از منابع داده و

^۳ Hybrid filtering

^۱ Collaborative Filtering

^۲ Content-based

با توجه به الگوریتم‌های استفاده شده به دو دسته حافظه محور^۴ و مدل محور^۵ تقسیم می‌شوند.

الگوریتم‌های مدل محور از مجموعه‌ای از رای‌ها برای یافتن الگوی رای‌دهی استفاده می‌کنند و با استفاده از ماتریس رای‌ها و روش‌های یادگیری ماشین، به صورت برون خط یک مدل را آموزش می‌دهند و در زمان اجرا از این مدل پیش‌ساخته جهت پیش‌بینی رای بهره می‌برند [۲]. سپس بر اساس مدل استخراج شده، رای آیتم مورد نظر را پیش‌بینی می‌کنند.

الگوریتم‌های حافظه‌محور از کلیه آیتم‌هایی که قبلاً توسط کاربر به آنها رای داده شده است، جهت پیش‌بینی رای استفاده می‌کنند. این روش‌ها، به علت نگهداری رای‌ها در حافظه و استفاده مستقیم از آن جهت پیش‌بینی رای، حافظه محور نامیده می‌شوند [۱]. در این روش‌ها با استفاده از یک معیار شباهت، کاربران یا آیتم‌های مشابه کاربر فعال یا آیتم فعال محاسبه شده و با استفاده از رای‌های مشترک، پیش‌بینی رای برای آیتم‌های جدید صورت می‌گیرد. این الگوریتم‌ها خودشان به دو دسته کاربر محور^۶ و آیتم محور^۷ تقسیم می‌شوند. روش‌های فیلترینگ مشارکت جمعی کاربر محور، ابتدا کاربران مشابه با کاربر مورد نظر را یافته و سپس با استفاده از رای‌های آن‌ها، پیش‌بینی رای یا پیشنهادها را ارائه می‌دهند. روش‌های فیلترینگ مشارکت جمعی آیتم محور، ابتدا آیتم‌های مشابه با آیتم مورد نظر را پیدا می‌کنند و سپس بر اساس آن‌ها به کاربران پیشنهاد می‌دهند. بنابراین، روش‌های حافظه‌محور برای یافتن کاربران یا آیتم‌های مشابه وابسته به یک معیار شباهت هستند که در این راستا از معیارهای شباهت مرسوم مانند کسینوسی و ضرائب همبستگی پیرسون استفاده می‌کنند.

مسائل مطرح در حوزه سیستم‌های توصیه‌گر را می‌توان به دو صورت کلی مسئله پیش‌بینی رای^۸ و مسئله پیشنهاد N-بالاترین^۹ مدل‌سازی نمود [۳]. در رویکرد اول هدف پیش‌بینی ارزش رای یک کاربر به آیتمی است که هنوز به آن رای نداده است و منظور از N-بالاترین آیتم‌ها، لیستی حاوی N آیتم می‌باشد که به احتمال زیاد کاربر به آنها علاقه‌مند است. این لیست مرتب بوده و آیتم‌هایی که علاقه کاربر به آنها بیشتر باشد در ابتدای لیست قرار می‌گیرند. یکی از مهم‌ترین چالش‌های سیستم‌های توصیه‌گر، پراکندگی داده‌ها است. بیشتر خانه‌های ماتریس رای نامعلوم و بالطبع ماتریس رای کاربران به آیتم‌ها دارای پراکندگی بالایی است. بنابراین، در رویکردهای آیتم‌محور تعداد کاربران مشترکی که جهت محاسبه شباهت استفاده می‌شوند، اندک بوده و منجر به کاهش کیفیت پیش‌بینی‌ها و ارائه پیشنهادها می‌شود. راه حل‌های متنوعی جهت بهبود کیفیت پاسخ سیستم‌های توصیه‌گر در شرایط پراکندگی داده‌ها ارائه شده است. یکی از این رویکردها استفاده از مجموعه فازی جهت مدل‌سازی رفتار کاربر و مشخص کردن میزان علاقه او به آیتم‌ها است [۹، ۱۷].

در رویکرد دیگری، معیارهای شباهت ترکیبی جدیدی ارائه شده است که در شرایط پراکندگی داده‌ها، کارایی بهتری را ارائه می‌دهند. به عنوان نمونه، بابادیل و همکاران او [۴] از طریق ترکیب وزن‌دار شش معیار شباهت با یکدیگر، معیار شباهت جدیدی را ارائه کردند. در روش پیشنهادی آن‌ها، وزن‌ها با استفاده

از یک شبکه عصبی، یادگیری می‌شد. این روش دارای پیچیدگی محاسباتی و زمانی بالایی بود. در تلاشی دیگر، آهین [۵] از سه عامل شباهت، تقریب اثرگذاری و محبوبیت جهت ارائه یک معیار شباهت استفاده کرد. در این روش، مقادیر رای‌ها و نسبت آیتم‌های دارای رای مشترک لحاظ نمی‌شود. همچنین، ترجیحات کلی کاربران در این معیار شباهت در نظر گرفته نمی‌شود و به راحتی با دیگر روش‌ها قابل ترکیب نیست.

پاترا و همکاران او [۶] یک معیار شباهت مبتنی بر جاکارد و ضرائب باتاچاریا^{۱۰} برای غلبه بر مشکل پراکندگی در شرایطی که تعداد آیتم‌های رای مشترک اندک بود ارائه نمودند. در تحقیق دیگری، ایوب و همکاران او [۱۸]، معیار شباهت جاکارد اصلاح شده را ارائه کردند که از محاسبه نسبت مقادیر رای‌ها به تعداد آیتم‌های رای مشترک حاصل می‌شود. در تحقیق دیگری، پولاتیدیس و همکاران او [۷] معیار شباهتی چند سطحی ارائه نمودند. در این معیار، هرگاه کاربران مورد مقایسه به تعداد معینی آیتم مشترک رای نداده باشند یا معیار شباهت پیرسون آنها به حد نصاب مشخصی نرسیده باشد، از معیارهای شباهت مختلفی برای کاربران استفاده می‌شود. در تحقیق دیگری توسط فنگ و همکاران او [۸]، برای غلبه بر مشکل پراکندگی داده‌ها از ترکیب سه عامل اثرگذار استفاده شده است. عامل اول استفاده از معیار شباهت کسینوسی و پر کردن ماتریس رای با میانگین رای کاربر در شرایط پراکندگی داده‌ها، عامل دوم به منظور کاهش اثر کاربران با آیتم مشترک کمتر است و عامل سوم جهت لحاظ کردن ترجیحات مختلف کاربران در رای‌دهی می‌باشد. گوا و همکاران او [۹]، با استفاده از معیارهای شباهت مبتنی بر گوگل^{۱۱} و KL فازی شده^{۱۲}، یک معیار آیتم‌محور ترکیبی ارائه نمودند. در معیار شباهت پیشنهادی، از معیار شباهت کسینوسی جهت محاسبه شباهت بین آیتم‌ها بر اساس ارزش رای‌ها و از معیار شباهت جاکارد جهت اعمال اطلاعات ساختاری استفاده شده است.

در همه این روش‌ها، سعی بر این بوده است که از کلیه اطلاعات در دسترس جهت بهبود مشکل پراکندگی داده‌ها استفاده شود. معیار شباهت کسینوسی مبتنی بر مثلث، رویکرد دیگری است که توسط عامر و همکاران او [۱۶]، ارائه شده است. این معیار شباهت، از ترکیب معیار کسینوسی و فاصله اقلیدسی پیشنهاد شده است. در این معیار از نسبت مساحت مثلث پایه به کل مساحت مثلث به عنوان یک عامل تقویت‌کننده برای فاصله اقلیدسی استفاده شده و ضمن سادگی و کارایی بالا، نقاط ضعف آن را کاهش داده است.

در راستای مشکلات ذکر شده، در این مقاله یک معیار شباهت ترکیبی جهت استفاده در مسائل پیش‌بینی رای با رویکرد فیلترینگ مشارکت جمعی آیتم محور ارائه شده است. این ترکیب منجر به در دسترس بودن اطلاعات بیشتر برای پیش‌بینی رای کاربران به آیتم‌های نامعلوم و بهبود کیفیت پاسخ‌های تولید شده خواهد شد و در نهایت منجر به افزایش دقت پیش‌بینی می‌گردد.

این مقاله مشتمل بر ۵ بخش است. در بخش دوم مفاهیم اولیه شرح داده می‌شود. در بخش سوم به شرح معیار شباهت پیشنهادی و نحوه بکارگیری آن می‌پردازیم. در ادامه در بخش چهارم نحوه ارزیابی، مجموعه داده مورد استفاده

^۹ Top-N Recommendation Problem

^{۱۰} Bhattacharyya

^{۱۱} adjusted Google

^{۱۲} intuitionistic fuzzy set (IFS) based Kullback-Leibler (KL)

^۴ Memory-based collaborative filtering

^۵ Model-based collaborative filtering

^۶ User-based

^۷ Item-based

^۸ Rating Prediction Problem

و نیز نتایج پیاده سازی مدل را بررسی کرده و در پایان به جمع‌بندی مطالب مطرح شده می‌پردازیم.

۲- مفاهیم پایه

۲-۱- تعاریف و نمادها

در مسائل مطرح شده در حوزه سیستم‌های توصیه‌گر، دارای تعدادی کاربر هستیم که نظرات خود را نسبت به تعدادی آیت‌م ابراز کرده‌اند. از نماد u برای نمایش مجموعه کاربران و I برای نمایش مجموعه آیت‌م‌ها استفاده می‌شود. نظر هر کاربر به یک آیت‌م را $r_{u,i}$ می‌نامیم و آن را به صورت $r_{u,i}$ نمایش می‌دهیم که معرف ارزش رای داده شده توسط کاربر u به آیت‌م i می‌باشد. حروف u و v برای نمایش کاربران و حروف i و j را برای نمایش آیت‌م‌ها به کار می‌بریم. مجموعه همه رای‌ها در یک ماتریس پراکنده R با ابعاد $m \times n$ تحت عنوان ماتریس امتیازات کاربران - آیت‌م‌ها^{۱۳}، ماتریس رای، شناخته می‌شود.

از نماد U_i برای نمایش زیرمجموعه‌ای از کاربران که به آیت‌م i رای داده‌اند و همچنین از نماد $U_{i,j}$ برای نشان دادن زیرمجموعه‌ای از کاربران که به هر دو آیت‌م i و j رای داده‌اند استفاده می‌شود. نماد U_{i+j} نیز معرف کاربرانی است که حداقل به یکی از آیت‌م‌های i یا j ، رای داده‌اند و $\bar{r}_{*,i}$ نشان‌گر میانگین رای‌هایی است که به آیت‌م i داده شده است. تعداد اعضای مجموعه S را با $|S|$ نشان می‌دهند و در نهایت رای پیش‌بینی شده کاربر u به آیت‌م i را با نماد $p_{u,i}$ مشخص می‌کنیم.

۲-۲- فیلترینگ مشارکت جمعی آیت‌م‌محور

در سیستم‌های توصیه‌گر فیلترینگ مشارکت جمعی فرض بر این است که اگر کاربرانی در گذشته دارای علایق یکسانی بوده‌اند، در آینده نیز دارای علایق مشترکی خواهند بود. در این سیستم‌ها، رای‌های کاربران به آیت‌م‌ها به صورت یک ماتریس رای ذخیره می‌شود و از آن جهت یافتن کاربران با علایق مشابه استفاده می‌شود. برای پیش‌بینی رای کاربر فعال u به آیت‌م فعال i ، آیت‌م‌های مشابه آیت‌م فعال، یک گروه همسایگی را تشکیل داده و بر اساس رای‌هایی که کاربر فعال به آیت‌م‌های همسایه داده است، رای احتمالی کاربر فعال به آن آیت‌م پیش‌بینی می‌شود. این پیش‌بینی یک مقدار $p_{u,i}$ را تولید می‌کند که در واقع نمره‌ای است که کاربر u به آیت‌م i می‌دهد (رابطه ۱). عبارت $N_i(u; k)$ معرف تعداد k آیت‌م مانند i است که دارای بیشترین شباهت $sim(i, j)$ به آیت‌م فعال i هستند و توسط کاربر فعال u به آنها رای داده شده است.

$$p_{u,i} = \frac{\sum_{j \in N_i(u; k)} r_{u,j} \times sim(i, j)}{\sum_{j \in N_i(u; k)} sim(i, j)} \quad (1)$$

۲-۳- معیارهای شباهت مبتنی بر رای

روش‌های معمول محاسبه شباهت مانند ضرائب همبستگی پیرسون و معیار کسینوسی فقط از ماتریس رای برای محاسبه شباهت بین آیت‌م‌ها استفاده

می‌کنند. به عنوان نمونه معیار شباهت کسینوسی تعدیل‌شده^{۱۴} با استفاده از رابطه زیر محاسبه می‌شود:

$$sim^{AC}(i, j) = \frac{\sum_{u \in U_{i,j}} (r_{u,i} - \bar{r}_{u,*}) \times (r_{u,j} - \bar{r}_{u,*})}{\sqrt{\sum_{u \in U_i} (r_{u,i} - \bar{r}_{u,*})^2} \sqrt{\sum_{u \in U_j} (r_{u,j} - \bar{r}_{u,*})^2}} \quad (2)$$

به واسطه برخی مشکلات، معیارهای شباهت مبتنی بر رای، به تنهایی نمی‌توانند شباهت بین آیت‌م‌ها را به خوبی ارائه دهند. از جمله آنها می‌توان به موارد زیر اشاره کرد [۱۰]:

- در محاسبه شباهت آیت‌م‌ها، به نظرات همه آیت‌م‌های همسایه فارغ از تعداد رای‌های دخیل در محاسبه شباهت، وزن یکسانی داده می‌شود.
- آیت‌م‌هایی که مورد علاقه بیشتر کاربران هستند، در محاسبه اغلب شباهت‌های بین آیت‌م‌ها مورد استفاده قرار می‌گیرند. در حالیکه، این آیت‌م‌ها نماینده خوبی برای اندازه‌گیری شباهت نیستند.

۲-۴- معیارهای شباهت مبتنی بر ساختار

معیارهای شباهت مبتنی بر ساختار، صرفاً از رای کاربران مشترک استفاده نکرده و سعی می‌کنند از کلیه رای‌هایی که به یک آیت‌م داده شده است جهت محاسبه شباهت استفاده نمایند. به این ترتیب رای پیش‌بینی‌شده، فقط به سمت کاربران مشترک متمایل نشده و کلیه رای‌های دریافتی توسط یک آیت‌م در نظر گرفته می‌شود [۱۰]. از جمله این معیارها می‌توان به معیار شباهت جاکارد اشاره کرد که نسبتی است از تعداد کاربرانی که به هر دو آیت‌م رای داده‌اند به تعداد کاربرانی که به هر یک از دو آیت‌م رای داده‌اند (رابطه ۳).

$$sim^{JD}(i, j) = \frac{|U_{i,j}|}{|U_{i+j}|} \quad (3)$$

بنابراین، با استفاده از معیارهای شباهت ساختاری می‌توان آیت‌م‌های مرتبط را که کاربران به آنها رای داده‌اند مشخص نمود. ولی بدون در نظر گرفتن رای کاربر به آیت‌م نمی‌توان علاقه کاربر به هر دو آیت‌م را تعیین نمود. بنابراین با ترکیب کردن معیارهای شباهت ساختاری و معیارهای مبتنی بر رای، می‌توان از کلیه اطلاعات در دسترس جهت پیش‌بینی رای استفاده نمود و بر برخی مشکلات آن‌ها غلبه کرد.

۳- مدل پیشنهادی

مدل پیشنهادی این تحقیق، یک روش آیت‌م‌محور می‌باشد. با توجه به توضیحات بخش قبل، معیار شباهتی که برای محاسبه شباهت آیت‌م‌ها مورد استفاده قرار می‌گیرد، نقشی اساسی در دقت و کارایی این مدل ایفا می‌کند. در ادامه، معیار شباهت پیشنهادی ارائه خواهد شد.

۳-۱- معیار شباهت پیشنهادی

از آنجایی که اغلب معیارهای شباهت از جمله معیار شباهت کسینوسی متأثر از تعداد کاربران دارای رای مشترک هستند، در شرایط پراکندگی داده‌ها کیفیت نتایج آنها کاهش پیدا می‌کند. در نتیجه، در معیار شباهت پیشنهادی از معیار

رای درست $r_{u,i}$ محاسبه شده و سپس برای کلیه پیش‌بینی‌های موجود در مجموعه داده آزمون میانگین گرفته می‌شود.

۴-۳- الگوریتم‌های مقایسه

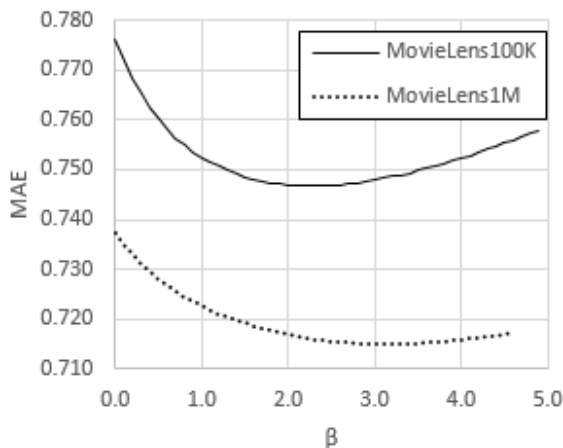
برای ارزیابی معیار شباهت پیشنهادی الگوریتم فیلترینگ مشارکت جمعی آیت‌محور با استفاده از معیارهای شباهت ضریب همبستگی پیرسون^{۱۹}، ضریب همبستگی پیرسون مقید^{۲۰}، معیار سالتون/اچ‌ای^{۲۱}، معیار کسینوسی^{۲۲} و معیار شباهت جاکارد^{۲۳} انتخاب شده‌اند.

۴-۴- روش ارزیابی

به منظور ارزیابی معیار پیشنهادی، از روش ارزیابی متقابل با استفاده از k بخش^{۲۴} استفاده شده است که مجموعه داده‌ها را به صورت تصادفی به ۵ بخش مساوی تقسیم کرده و در هر مرحله ۴ بخش را به عنوان مجموعه داده آموزشی و ۱ بخش باقیمانده را برای تست در نظر گرفتیم و در نهایت از نتایج حاصل از آن‌ها میانگین گرفته شده است.

۴-۵- تنظیم پارامترها

با توجه به اهمیتی که پارامتر β در کیفیت عملکرد الگوریتم پیشنهادی دارد، در این بخش به نحوه انتخاب آن پرداخته می‌شود. به منظور یافتن ارزش بهینه پارامتر β ، مقدار آن را از ۰ تا ۵ با فاصله ۰٫۱ تغییر داده و برای مجموعه داده‌های MovieLens100K و MovieLens1M، ارزش میانگین خطای مطلق محاسبه می‌شود (شکل ۱).



شکل ۱. تغییرات MAE برای مقادیر مختلف β

با توجه به شکل ۱، مشاهده می‌شود که برای مجموعه داده MovieLens100K با افزایش پارامتر β تا مقدار 2.3 میزان خطای مدل کاهش یافته و بعد از آن سیر صعودی دارد و برای مجموعه داده

شباهت جاکارد استفاده شده است تا هم با استفاده از اطلاعات بیشتر کیفیت نتایج بهتر شود و هم در شرایط اندک بودن تعداد رای‌های مشترک، شباهت میان آیت‌ها قابل محاسبه باشد.

بنابراین، معیار شباهت پیشنهادی با استفاده از ترکیب معیارهای شباهت کسینوسی تعدیل شده و جاکارد شکل گرفته است که در آن از پارامتر β جهت تعدیل معیار شباهت جاکارد استفاده شده است. نحوه محاسبه شباهت بین آیت‌های i و j ، طبق رابطه ۳ صورت می‌گیرد.

$$sim^{ACJD}(i,j) = sim^{JD}(i,j)^{\beta} \times sim^{AC}(i,j) \quad (4)$$

بعد از اینکه با استفاده از رابطه ۴ شباهت بین آیت‌ها محاسبه شد، تعداد k آیت مشابه‌تر به آیت فعال انتخاب شده و رای کاربر فعال u به آیت فعال i با استفاده از رابطه ۱ محاسبه می‌شود.

۴-۴- ارزیابی نتایج

۴-۱- مجموعه داده‌ها

به منظور تحقیق در حوزه سیستم‌های توصیه‌گر، مجموعه داده‌های استاندارد گوناگونی وجود دارد که در مقالات مختلف مورد استفاده قرار گرفته است. برای داده‌های تجربی مورد استفاده در این تحقیق از دو نسخه از مجموعه داده مووی لنز^{۱۵} [۱۱] استفاده شده است. مجموعه داده اول MovieLens100K^{۱۶} است که دارای ۱۰۰,۰۰۰ رای و مشتمل بر ۹۴۳ کاربر است که به ۱,۶۸۲ فیلم رای داده‌اند. مجموعه داده دوم MovieLens1M^{۱۷} است که دارای ۱,۰۰۰,۰۰۰ رای و مشتمل بر ۶۰۴۰ کاربر است که به ۳۹۵۲ فیلم رای داده‌اند. در هر دو مجموعه داده‌ها رای‌ها اعدادی از ۱ (بد) تا ۵ (عالی) می‌باشند که رای‌های ۱ و ۲ معرف نگاه منفی کاربران، ۳ معرف بی‌نظر بودن و ۴ و ۵ معرف دید مثبت کاربران به آیت‌ها می‌باشد.

۴-۲- معیارهای ارزیابی

در این تحقیق، با هدف ارزیابی کارایی معیار شباهت پیشنهادی، از معیار متداول میانگین خطای مطلق^{۱۸} استفاده شده است [۱۲]. میانگین خطای مطلق، یکی از متداول‌ترین معیارهای بررسی خطا است که از محاسبه اختلاف رای محاسبه‌شده با رای واقعی بدست می‌آید. کم شدن این معیار معرف افزایش دقت پیش‌بینی الگوریتم توصیه‌گر است و نحوه محاسبه آن طبق رابطه ۵ می‌باشد.

$$MAE = \frac{\sum_{u \in U} \sum_{i \in Test_u} |p_{u,i} - r_{u,i}|}{\sum_{u \in U} |Test_u|} \quad (5)$$

که در این رابطه $Test_u$ معرف مجموعه آزمون کاربر u می‌باشد. در این معیار، برای هر پیش‌بینی رای، قدرمطلق اختلاف رای پیش‌بینی شده $p_{u,i}$ با

- ^{۲۱} Salton/Ochiai similarity
- ^{۲۲} Adjusted cosine (AC)
- ^{۲۳} Jaccard
- ^{۲۴} K-fold cross validation

- ^{۱۵} MovieLens
- ^{۱۶} MovieLens 100K
- ^{۱۷} MovieLens 1M
- ^{۱۸} Mean Absolute Error (MAE)
- ^{۱۹} Pearson Correlation Coefficient (PCC)
- ^{۲۰} Constrained Pearson Correlation Coefficient (CPCC)

دو آیتم رای داده‌اند، منجر به بهبود معیار شباهت پیشنهادی شده است. از سوی دیگر، معیار شباهت جاکارد و سالتون صرفاً از اطلاعات ساختاری استفاده کرده و مقدار رای کاربر را در محاسبه شباهت لحاظ نمی‌کند. بنابراین، اعمال ارزش رای‌های مشترک نیز باعث برتری معیار شباهت پیشنهادی شده است. همان‌طور که از نتایج مشخص است در هر دو مجموعه داده مورد آزمایش، مقادیر میانگین خطای مطلق از الگوی یکسانی پیروی کرده و معیار شباهت پیشنهادی در مقایسه با معیارهای شباهت حافظه‌محور متداول از لحاظ دقت پیش‌بینی بهتر عمل نموده است. همچنین، با توجه به اینکه معیار شباهت پیشنهادی با استفاده از ترکیب معیارهای شباهت کسینوسی تعدیل شده و جاکارد شکل گرفته است، در مقایسه با این معیارهای شباهت پایه، از پیچیدگی زمانی قابل قبولی برخوردار است.

۵- جمع بندی

در این مقاله یک معیار شباهت ترکیبی جهت استفاده در سیستم‌های مشارکت جمعی حافظه‌محور ارائه شد که از لحاظ دقت پیش‌بینی در مقایسه با روش‌های مشابه بهتر عمل کرده است. لازم به ذکر است که افزایش اندکی در دقت پیش‌بینی مدل گرچه در مسئله پیش‌بینی رای ممکن است ناچیز لحاظ شود؛ اما در مسائل پیشنهاد آیت‌ها به کاربران اثرات قابل توجهی خواهد داشت. برای کارهای آتی می‌توان روش پیشنهادی را در شرایط پراکندگی داده‌ها بررسی نموده و با پر کردن ماتریس رای با مقادیر مجازی، مسئله خالی بودن ماتریس رای و در نتیجه مشکل پراکندگی داده‌ها را تا حدودی بهبود بخشید.

مراجع

- [1] Adomavicius, G. and Tuzhilin, A., "Toward the next generation of recommender systems: A survey of the state-of-the-art and possible extensions". IEEE transactions on knowledge and data engineering, 17(6), pp.734-749, 2005.
- [2] Ning, X., Desrosiers, C. and Karypis, G., "A comprehensive survey of neighborhood-based recommendation methods". Recommender systems handbook, pp.37-76, 2015.
- [3] Charu, C.A., "Recommender Systems": The Textbook, 2016.
- [4] Bobadilla, J., Ortega, F., Hernando, A. and Bernal, J., "A collaborative filtering approach to mitigate the new user cold start problem". Knowledge-based systems, 26, pp.225-238, 2012.
- [5] Ahn, H.J., "A new similarity measure for collaborative filtering to alleviate the new user cold-starting problem". Information sciences, 178(1), pp.37-51, 2008.
- [6] Patra, B.K., Launonen, R., Ollikainen, V. and Nandi, S., "A new similarity measure using Bhattacharyya coefficient for collaborative filtering in sparse data. Knowledge-Based Systems", 82, pp.163-177, 2015.
- [7] Polatidis, N. and Georgiadis, C.K., "A multi-level collaborative filtering method that improves recommendations". Expert Systems with Applications, 48, pp.100-110, 2016.
- [8] Feng, J., Feng, X., Zhang, N. and Peng, J., "An improved collaborative filtering method based on similarity". PloS one, 13(9), p.e0204003, 2018.
- [9] Guo, J., Deng, J. and Wang, Y., "An intuitionistic fuzzy set based hybrid similarity model for recommender system". Expert Systems with Applications, 135, pp.153-163, 2019.

MovieLens1M بهترین حالت برای معیار پیشنهادی در $\beta = 3.1$ اتفاق می‌افتد. بنابراین در آزمایش‌های بعدی به ترتیب از ارزش‌های $\beta = 2.3$ و $\beta = 3.1$ برای مجموعه داده‌های MovieLens100K و MovieLens1M استفاده می‌شود. لازم به ذکر است که ارزش‌های متفاوت این پارامتر برای مجموعه داده‌های مختلف، طبیعی بوده و متاثر از اختلاف توزیع رای‌ها و میزان پراکندگی داده‌ها در آن‌ها می‌باشد. تنظیمات پارامترهای الگوریتم‌های انتخاب شده در جدول ۱ نشان داده شده است. لازم به ذکر است که جهت محاسبه شباهت، از کلیه همسایگانی که میزان شباهت آنها با آیتم مورد نظر عددی مثبت بوده است، جهت محاسبه مقدار رای استفاده شده است.

جدول ۱: پارامترهای استفاده شده در معیارهای شباهت مقایسه

معیار شباهت	پارامترها
ضریب همبستگی پیرسون [۱۴]	$k=\infty, \theta=0$
ضریب همبستگی پیرسون مقید [۱۶]	$k=\infty, \theta=0$
جاکارد [۱۳]	$k=\infty, \theta=0$
سالتون/اچای [۱۰]	$k=\infty, \theta=0$
کسینوسی تعدیل شده [۱۵]	$k=\infty, \theta=0$
روش پیشنهادی (MovieLens100K)	$k=\infty, \theta=0, \beta=2.3$
روش پیشنهادی (MovieLens1M)	$k=\infty, \theta=0, \beta=3.1$

۴-۶- نتایج و تحلیل

در این بخش به ارائه نتایج حاصل از پیاده‌سازی و مقایسه روش پیشنهادی می‌پردازیم. نتایج حاصل از پیاده‌سازی معیار شباهت پیشنهادی و معیارهای شباهت مورد مقایسه در جدول ۲ آورده شده است. بر حسب میانگین خطای مطلق، شاهد برتری معیار شباهت پیشنهادی هستیم. برای مجموعه داده MovieLens1M، معیار شباهت پیشنهادی از معیار شباهت کسینوسی تعدیل شده ۳٪، معیار شباهت جاکارد ۱۰٪، معیار شباهت ضریب همبستگی پیرسون ۱۰٪، معیار شباهت ضریب همبستگی پیرسون مقید ۱۱٪ و معیار سالتون/اچای ۱۱٪ بهتر عمل کرده است. برای مجموعه داده MovieLens100K نیز شاهد هستیم که از معیار شباهت کسینوسی تعدیل شده ۴٪، معیار شباهت جاکارد ۸٪ و معیار شباهت ضریب همبستگی پیرسون ۹٪، معیار سالتون/اچای ۹٪ و معیار شباهت ضریب همبستگی پیرسون مقید ۱۰٪ بهتر عمل کرده است.

جدول ۲: نتایج حاصل از پیاده‌سازی معیارهای شباهت

معیار شباهت	میانگین خطای مطلق MovieLens100K	میانگین خطای مطلق MovieLens1M
ضریب همبستگی پیرسون مقید [۱۶]	۰/۸۳۰	۰/۸۰۳
سالتون/اچای [۱۰]	۰/۸۲۰	۰/۸۰۵
ضریب همبستگی پیرسون [۱۴]	۰/۸۲۵	۰/۷۹۷
جاکارد [۱۳]	۰/۸۱۲	۰/۷۹۵
کسینوسی تعدیل شده [۱۵]	۰/۷۷۶	۰/۷۳۷
روش پیشنهادی	۰/۷۴۷	۰/۷۱۵

معیارهای شباهت پیرسون و کسینوسی بر اساس کاربران مشترک محاسبه می‌شوند. از آنجاییکه تعداد کاربران مشترک در محاسبات این شباهت لحاظ نمی‌شود، افزودن نسبت کاربران مشترک به کلیه کاربرانی که حداقل به یکی از

- [10] Wang, D., Yih, Y. and Ventresca, M., "Improving neighbor-based collaborative filtering by using a hybrid similarity measurement". *Expert Systems with Applications*, 160, p.113651, 2020.
- [11] Herlocker, J. L., Konstan, J. A., Borchers, A., & Riedl, J. "An algorithmic framework for performing collaborative filtering". In *Proceedings of the 22nd annual international ACM SIGIR conference on Research and development in information retrieval*(pp.230–237). ACM, 1999.
- [12] Hyndman, R.J. and Koehler, A.B., "Another look at measures of forecast accuracy". *International journal of forecasting*, 22(4), pp.679-688, 2006.
- [13] Koutrika, G., Bercovitz, B., & Garcia-Molina, H. lexRecs:" expressing and combining flexible recommendations". . In *Proceedings of the 2009 ACM SIGMOD International Conference on Management of data*, 745-758 , 2009.
- [14] Resnick, P., Iacovou, N., Suchak, M., Bergstrom, P., & Riedl, J. GroupLens: "an open architecture for collaborative filtering of netnews". In *Proceedings of ACM Conference on Computer Supported Cooperative Work*, 175 - 186, 1994.
- [15] Breese, J. S., Heckerman, D., & Kadie, C. "Empirical analysis of predictive algorithms for collaborative filtering". In *Proceedings of the Fourteenth conference on Uncertainty in artificial intelligence*, 43-52, 1998.
- [16] Amer, A.A., Abdalla, H.I. and Nguyen, L., "Enhancing recommendation systems performance using highly-effective similarity measures". *Knowledge-Based Systems*, 217, p.106842, 2021.
- [17] Saeed, M. and Mansoori, E.G., "A novel fuzzy-based similarity measure for collaborative filtering to alleviate the sparsity problem". *Iranian Journal of Fuzzy Systems*, 14(5), pp.1-18, 2017.
- [18] Ayub, M., Ghazanfar, M.A., Khan, T. and Saleem, A., "An Effective Model for Jaccard Coefficient to Increase the Performance of Collaborative Filtering". *Arabian Journal for Science & Engineering (Springer Science & Business Media BV)*, 45(12), 2020.

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بوم

الگوریتم ژنتیک برای پستچی چینی تحت شرایط نایقینی

سمیرا سامانی‌فر^۱، حسن میش‌مست‌نهی^۲، حامد احمدزاده^۳

^۱ دانشکده ریاضی، دانشگاه سیستان و بلوچستان، samanifar71@gmail.com

^۲ دانشکده ریاضی، دانشگاه سیستان و بلوچستان، hmnehi@hamoon.usb.ac.ir

^۳ دانشکده ریاضی، دانشگاه سیستان و بلوچستان، ahmadzade.h.63@gmail.com

چکیده: مسئله پستچی چینی از پرکاربردترین مسائل در دنیای واقعی است. مسائل پستچی با عبور از تمام یال‌ها برای ایجاد کوتاه‌ترین مسیر با کمترین هزینه، امکان دستیابی به نقاط مختلف و بازگشت دوباره به نقطه شروع را فراهم می‌آورند. مدل‌سازی چنین برنامه‌های کاربردی در دنیای واقعی نیاز به در نظر گرفتن برخی عوامل نامشخص دارد. این مقاله به بررسی مسئله پستچی چینی در چارچوب نظریه نایقینی می‌پردازد. هدف مسئله پستچی چینی حداقل کردن هزینه‌ها در شرایط نایقینی است. سپس مسئله نایقینی را تبدیل به مسئله قطعی می‌نماییم و با استفاده از الگوریتم ژنتیک حل می‌کنیم. الگوریتم ژنتیک اغلب جواب‌های تقریبی خوبی را برای انواع مسائل مختلف می‌دهد. مثالی از مسئله پستچی چینی تحت شرایط نایقینی آورده شده است، سپس مسئله با استفاده از الگوریتم ژنتیک حل می‌شود.

کلمات کلیدی: مسئله پستچی چینی، متغیر نایقینی، نظریه نایقینی، الگوریتم ژنتیک

۱- مقدمه

مسئله پستچی یکی از مهم‌ترین مسائل بهینه‌سازی ترکیبی با طیف گسترده-ای از برنامه‌های کاربردی در دنیای واقعی است. همان‌طور که می‌دانیم، انگیزه پستچی ناشی از کار یک پستچی است که حداقل یک‌بار خیابان منطقه خود را بپوشاند و سپس در همان جایی که مسیر خود را آغاز کرده است، پایان یابد. مسئله پستچی تلاش می‌کند، تا بهترین مسیر را برای یک پستچی جهت توزیع نامه دریافتی از اداره پست و بازگشت به اداره پست با کمترین فاصله پیاده‌روی در هنگام توزیع نامه تعیین کند. مسائل پستچی می‌تواند در تمام مسائل بزرگ شامل مسیریابی روزنامه، تحویل آب، تمیز کردن جاده‌ها در روزهای برفی بر اساس اولویت و ... مورد استفاده قرار گیرند.

عواملی مثل آب و هوا، ترافیک و یا زمان هزینه شده برای پستچی در جاده نامشخص است. به عنوان مثال هزینه سفر بین دو شهر را در نظر می‌گیریم که به عوامل مختلفی مانند مالیات عوارض، هزینه کار، هزینه نگهداری خودرو و ... بستگی دارد که این عوامل ماهیت پویایی دارد. مسائل نایقینی را می‌توانیم به مسائل قطعی تبدیل نمود و پس از تبدیل به یک مسئله قطعی با استفاده از الگوریتم ژنتیک حل می‌شود. الگوریتم ژنتیک به تولید تصادفی جمعیت اولیه می‌پردازد. با استفاده از عملگرهای مختلف فرزندان جدید به وجود می‌آیند و در نهایت با ترکیب فرزندان و والدین نسل جدیدی به وجود می‌آورد. جمعیت تا جایی تکامل می‌یابد که حل‌های مناسب‌تر یا کروموزوم‌های مناسب‌تر جایگزین شود [1].

۲- پیش‌نیازهای مورد نیاز

در این بخش برخی مفاهیم اولیه نظریه نایقینی و الگوریتم ژنتیک مورد نیاز را ارائه می‌دهیم.

۲-۱- نظریه نایقینی

در این قسمت برخی از خصوصیات نظریه نایقینی را بیان می‌کنیم. توزیع نایقینی $\Phi(x)$:

که اگر گراف مربوط به مسیر پستچی، اولیری باشد از هر یال کافی است یک‌بار عبور نماییم، در صورتی که گراف غیراولیری باشد، بعضی از یال‌ها دو بار تکرار می‌شوند. با این شرایط پستچی مایل است کوتاه‌ترین راه را برای تکمیل مسیر تعیین شده خود در خیابان‌ها جستجو کند.

مسئله پستچی چینی قطعی به فراوانی مورد بررسی قرار گرفته است. با این حال به دلیل کمبود داده‌ها، عدم موفقیت و ... مسائل نایقینی ظاهر می‌شود،

اگر $\xi \in Z(a, b, c)$ متغیر زیگزاگ نایقینی باشد، در این صورت امید ریاضی ξ بصورت زیر است:

$$E[\xi] = \frac{a + 2b + c}{4} \quad (11)$$

۲-۲- الگوریتم ژنتیک

الگوریتم ژنتیک اغلب جواب‌های تقریبی خوبی را برای انواع مسائل انجام می‌دهد. روند کلی الگوریتم ژنتیک بصورت زیر است [3]:

مرحله اول: جمعیت اولیه را به صورت تصادفی تولید می‌کنیم.
مرحله دوم: تابع برازندگی برای هر کروموزوم در جمعیت اولیه را ایجاد می‌کنیم. در الگوریتم ژنتیک متعارف افراد با رشته‌های باینری با طول ثابت نشان داده می‌شوند. تابع برازندگی بصورت $\frac{f_i}{f}$ تعریف می‌شود، به گونه‌ای که f_i تابع مرتبط با رشته i می‌باشد.

مرحله سوم: با تکرار مراحل زیر می‌توانیم جمعیت جدید را تولید کنیم:

(۱) انتخاب: دو کروموزوم والدین را از یک جمعیت، با توجه به تابع برازندگی آن‌ها انتخاب می‌کنیم. هر چه تابع برازندگی بیشتر باشد، شانس انتخاب آن‌ها بیشتر است.

(۲) ترکیب: برای عبور از والدین و تشکیل فرزندان جدید به کار می‌رود. می‌توان فرزندی تولید کرد که مزایای هر دو والدین را دارد. نوترکیبی به طور معمول شامل عناصر تصادفی است یا با انتخاب قسمت‌هایی از والدین تلاقی ایجاد نمود. الگوریتم ژنتیک متعارف از عملگر ترکیب نقطه‌ای استفاده می‌کند، این بدان معنا است که فرزند ۱، قسمت اول رشته خود را از والد ۱ و قسمت دوم را از والد ۲ دریافت می‌کند و برای فرزند دوم برعکس می‌باشد. نقطه برش بین دو قسمت به طور تصادفی انتخاب می‌شود.

(۳) جهش: این عملگر افراد جدید را از قدیمی ایجاد می‌کند. بنابراین به حفظ تنوع جمعیت کمک می‌کنند. جهش یک عملگر تغییر است، که یک فرد را به عنوان ورودی دریافت می‌کند و یک فرد جدید را از آن ایجاد می‌نماید.

(۴) فرزندان جدید که با استفاده از عملگرهای انتخاب، ترکیب و جهش تولید شده‌اند را در جمعیت جدید قرار دهید.

مرحله چهارم: برای ادامه الگوریتم از جمعیت جدید تولید شده، استفاده نمایید. مرحله پنجم: اگر شرط نهایی برآورده شد، متوقف شوید و بهترین جواب در جمعیت فعلی را به عنوان جواب نهایی قرار دهید.

مرحله ششم: در صورتی که شرط نهایی برآورده نشد، به مرحله دوم بازگردید و دوباره الگوریتم را تکرار کنید.

۳- مدل مسئله پستی چینی

$G = (V_G, E_G)$ گراف غیرجهت‌داری می‌باشد که V_G مجموعه رئوس و E_G مجموعه یال‌ها می‌باشد. هر یال e_{ij} نشان دهنده یک خیابان است. در مدل بیان شده در این مقاله، پستی می‌خواهد کمترین هزینه را داشته باشد و از شهری شروع کند و دوباره به شهر مبدأ بازگردد.

در صورتی که داشته باشیم، متغیر ξ متغیر نایقینی روی فضای نایقینی $(\Gamma, \mathcal{L}, \mathcal{M})$ باشد، توزیع نایقینی بصورت زیر می‌باشد [2]:

$$\Phi(x) = \mathcal{M}\{\xi \leq x\} \quad (1)$$

به طور مثال اگر ξ متغیر نایقینی خطی باشد، توزیع آن بصورت زیر است:

$$\Phi(x) = \begin{cases} 0, & x \leq a \\ \frac{x-a}{b-a}, & a \leq x \leq b \\ 1, & x \geq b \end{cases} \quad (2)$$

متغیر نایقینی ξ را زیگزاگ نامیده می‌شود، هرگاه توزیع آن بصورت زیر باشد:

$$\Phi(x) = \begin{cases} 0, & x \leq a \\ \frac{x-a}{2(b-a)}, & a \leq x \leq b \\ \frac{x+c-2b}{2(c-b)}, & b \leq x \leq c \\ 1, & x \geq c \end{cases} \quad (3)$$

متغیر نایقینی ξ را نرمال می‌گویند، هرگاه توزیع آن بصورت زیر باشد:

$$\Phi(x) = \left(1 + \exp\left(\frac{\pi((e-x))}{\sqrt{3}\sigma}\right) \right)^{-1}, x \in \mathbb{R} \quad (4)$$

معکوس توزیع نایقینی Φ^{-1} :

Φ^{-1} معکوس توزیع نایقینی می‌باشد اگر و تنها اگر

$$\mathcal{M}\{\xi \leq \Phi^{-1}(\alpha)\} = \alpha \quad \forall \alpha \in [0,1] \quad (5)$$

به طور مثال توزیع نایقینی معکوس برای متغیر نایقین خطی $\mathcal{L}(a, b)$ بصورت زیر است:

$$\Phi^{-1}(x) = (1-\alpha)a + \alpha b \quad (6)$$

توزیع نایقینی معکوس متغیر نایقین زیگزاگ $Z(a, b, c)$ مطابق فرمول زیر است:

$$\Phi^{-1}(x) = \begin{cases} (1-2\alpha)a + 2\alpha b, & \alpha < .5 \\ (2-2\alpha)b + (2\alpha-1)c, & \alpha \geq .5 \end{cases} \quad (7)$$

و توزیع نایقینی معکوس برای متغیر نایقین نرمال $\mathcal{N}(e, \sigma)$ بصورت زیر است:

$$\Phi^{-1}(x) = e + \frac{\sigma\sqrt{3}}{\pi} \ln \frac{\alpha}{1-\alpha} \quad (8)$$

امید ریاضی متغیرهای نایقین:

اگر توزیع نایقینی منظم باشد، امید ریاضی آن بصورت زیر تعریف می‌شود:

$$E[\xi] = \int_0^1 \Phi^{-1}(\alpha) d\alpha \quad (9)$$

اکنون می‌خواهیم مثال‌هایی از $E[\xi]$ داشته باشیم:

اگر $\xi \in \mathcal{L}(a, b)$ یک متغیر خطی نایقینی باشد، در این صورت امید ریاضی ξ بصورت زیر است:

$$E[\xi] = \frac{a+b}{2} \quad (10)$$

۱-۳- مدل نایقینی مسئله پستیچی

ξ_{ij} هزینه سفر یال e_{ij} باشد و یک متغیر نایقینی است. در این صورت مدل ریاضی مسئله پستیچی بصورت زیر بیان می شود [4].

$$\text{Min } Z = \sum_{i=1}^n \sum_{j=1}^n x_{ij} \xi_{ij} \quad (12)$$

Subject to

$$\sum_{j=1}^n x_{ij} - \sum_{j=1}^n x_{ji} = 0, \quad i = 1, 2, \dots, n$$

$$x_{ij} + x_{ji} \geq 1, \quad e_{ij} \in E_G$$

$$x_{ij} \in \{0, 1\}$$

در این مدل متغیر x_{ij} مقدار یک را اتخاذ می کند، اگر یال e_{ij} در مسیر پستیچی باشد و x_{ij} مقدار صفر در نظر گرفته می شود، اگر این یال در مسیر پستیچی قرار نگرفته باشد. قید اول مسئله، ایجاد تور شبکه را تضمین می کند. قید دوم، نشان دهنده آن است که از هر یال حداقل یکبار عبور کنیم.

۲-۳- تبدیل مسئله پستیچی نایقینی به مسئله قطعی

همان طور که از اطلاعات نظریه نایقینی می دانیم، مقدار متوسط متغیر نایقینی به معنای اندازه گیری نایقینی است و اندازه متغیر نایقینی را نشان می دهد. از ویژگی خطی بودن امید ریاضی می توان استفاده نمود و مسئله را بصورت زیر نوشت:

$$\text{Min } Z = \sum_{i=1}^n \sum_{j=1}^n x_{ij} E[\xi_{ij}] \quad (13)$$

Subject to

$$\sum_{j=1}^n x_{ij} - \sum_{j=1}^n x_{ji} = 0, \quad i = 1, 2, \dots, n$$

$$x_{ij} + x_{ji} \geq 1, \quad e_{ij} \in E_G$$

$$x_{ij} \in \{0, 1\}$$

۴- الگوریتم ژنتیک برای پستیچی

همان طور که گفته شد مسئله پستیچی یافتن کوتاه ترین مسیر در گراف $G(V, E)$ است که از هر یال عبور کند و به نقطه شروع بازگردد. به عنوان مثال در گراف های غیراولیری پستیچی باید از بعضی از یال ها دو بار عبور نماید، تا بتواند به نقطه شروع بازگردد. گراف های غیراولیری را می توان به حداقل هزینه مسیر پستیچی تبدیل نمود و با الگوریتم ژنتیک حل نمود. به طور کلی اگر گراف اولیری داده شده باشد، جواب ما یک جواب منحصر به فرد است [5].

۱-۴- نوع کروموزوم مسئله پستیچی

نوع کروموزوم در مسئله پستیچی را به عنوان یک رشته صحیح در نظر بگیرید، که هر عنصر آن تعداد دفعات عبور پستیچی از هر یال می باشد.

طول یک رشته تعداد یال های در گراف G می باشد. بنابراین جواب ها به صورت رشته ای از اعداد صحیح نشان داده شده اند. تکامل از جمعیت کاملاً تصادفی شروع می شود و نسل ها ادامه می یابد [6].

ما می توانیم یال ها را از ۱ تا t شماره گذاری کنیم. در این صورت هر مسیر را می توانیم بصورت زیر نمایش دهیم:

$$(e_1, e_2, \dots, e_t) \quad (14)$$

به طوری که هر عنصر به عنوان ژن می باشد، عنصرها بصورت زیر تعریف می شود:

e_t : تعداد دفعاتی که پستیچی از یال i عبور می کند. که می تواند یک عدد صحیح غیرمنفی باشد.

الگوریتم ژنتیک جمعیت اولیه را به صورت تصادفی انتخاب می کند.

۲-۴- تابع برازندگی مسئله پستیچی

اگر G گرافی غیراولیری باشد، جواب مطلوب این است که وزن کل کمان ها را به حداقل برسانیم. بنابراین تابع برازندگی بصورت زیر می باشد:

$$f = \begin{cases} \frac{1}{\text{مجموع وزن یالها}}, & \text{اگر مسیر پستیچی باشد} \\ 0, & \text{در غیر این صورت} \end{cases} \quad (15)$$

۳-۴- ترکیب

در گام بعدی دو رشته از والدین را به طور تصادفی از جمعیت اولیه انتخاب می کنیم و با استفاده از عملگر ترکیب، فرزندان جدید را ایجاد می شوند. که می توان از نمونه معمول آن که جابجایی بیت ها دو والد می باشد، استفاده نمود. گام اول: دو کروموزوم والدین بصورت تصادفی انتخاب می کنیم. سپس دو زیر رشته به طور تصادفی انتخاب می شود، که به شیوه زیر انجام می شود:

جدول ۱: کروموزوم های والد

Parent1	۱	۷	۲	۳	۴	۶	۵	۸
Parent2	۴	۶	۳	۵	۷	۱	۸	۲

گام دوم: جابجایی هر کدام از زیر رشته ها می باشد.

جدول ۲: جابجایی زیر رشته ها

۳	۵	۷
۲	۳	۴

گام سوم: نگاشت ژن ها در هر زیر رشته را تعیین می کنیم و بر اساس آن فرزندان را می نویسیم.

$$2 \leftrightarrow 3 \leftrightarrow 5, 4 \leftrightarrow 7$$

گام چهارم: کروموزوم ها را بروزرسانی کنید و با اطلاعات داده شده در گام سوم جابجایی ژن ها را انجام دهید.

جدول ۳: تولید فرزندان

offspring1	۱	۴	۳	۵	۷	۶	۲	۸
------------	---	---	---	---	---	---	---	---

offspring2	۷	۶	۲	۳	۴	۱	۸	۵
------------	---	---	---	---	---	---	---	---

۴-۴- جهش

عملگر جهش که برای جلوگیری از دست رفتن اطلاعات با تبادل اطلاعات درون کروموزوم ایجاد می‌شود. می‌توانیم از جهش وارونگی که معمولاً مرسوم است، استفاده نماییم. این کار با انتخاب یک موقعیت درون یک کروموزوم بصورت تصادفی انجام می‌شود و سپس رشته فرعی بین دو موقعیت را معکوس می‌کند.

جدول ۴: ایجاد فرزند با عملگر جهش

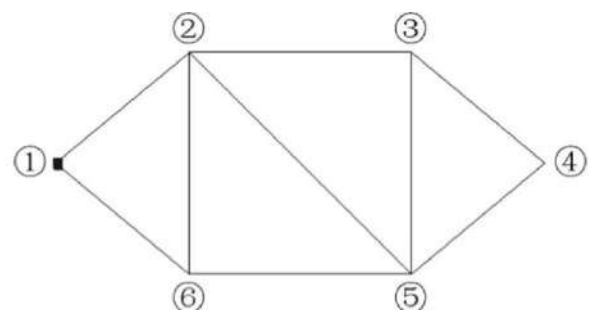
parent	۴	۶	۳	۵	۷	۱	۸	۲
offspring	۴	۶	۷	۵	۳	۱	۸	۲

۴-۵- انتخاب

هنگام اجرای الگوریتم ژنتیک نحوه انتخاب کروموزوم‌ها برای نسل بعدی بسیار مهم است، که باید تعادل بین بهره‌برداری و اکتشاف را در فضای جستجو فراهم کند. در اینجا ما یکی از بهترین روش‌های متناسب به نام چرخ رولت را انتخاب می‌کنیم. اگر بهترین کروموزوم برای نسل بعد انتخاب نشود، یکی از کروموزوم‌ها را به طور تصادفی انتخاب کرده و بهترین کروموزوم را جایگزین آن می‌کنیم.

۵- نتایج عددی

در این بخش مثالی برای نشان دادن روش‌های ذکر شده آورده‌ایم. وظیفه پستیچی آن است که از راس v_1 شروع کند و سپس به همین راس برگردد. پستیچی به گونه‌ای مسیر را طی می‌کند که از هر یال حداقل یک‌بار عبور نماید و کمترین هزینه را صرف کند. اکنون گراف غیرجهت‌دار زیر را در نظر بگیرید [7]:



شکل ۱: گراف غیرجهت‌دار G

اکنون اگر جدول هزینه‌ها بصورت زیر باشد:

جدول ۵: هزینه عبور از هر یال

هزینه سفر نایقینی	یال
$Z(48,50,52.7)$	e_{12}
$Z(48.9,52.4,53.2)$	e_{23}
$N(71,8.2)$	e_{25}
$N(42.6,7.2)$	e_{35}
$L(55,58)$	e_{56}
$L(52,54)$	e_{16}
$N(68.2,8.2)$	e_{26}
$N(70.7,10.8)$	e_{34}
$L(67,70)$	e_{45}

امید ریاضی متغیرهای نایقینی جدول بالا را محاسبه می‌کنیم، به این صورت مسئله پستیچی نایقینی تبدیل به یک مسئله پستیچی قطعی می‌شود. هر یک از توزیع‌های زیگزاگ، خطی و نرمال با استفاده از برنامه نویسی Matlab حل شده است. ابتدا به تعداد ۱۰۰۰ عدد از توزیع یکنواخت گرفته شده است. با محاسبه این اعداد در توزیع نایقینی معکوس و در نهایت میانگین گرفتن به محاسبه امید ریاضی توزیع‌های زیگزاگ، خطی و نرمال می‌پردازیم.

جدول ۶: امید ریاضی متغیرهای نایقینی

$E[\xi_{ij}]$	یال
50.1569	e_{12}
51.7104	e_{23}
70.9512	e_{25}
42.5898	e_{35}
56.4983	e_{56}
53.0067	e_{16}
68.1182	e_{26}
70.8255	e_{34}
68.5009	e_{45}

تمامی برنامه‌های این مقاله با استفاده از Matlab نوشته شده است. در این مسئله احتمال عملگر ترکیب را یک و احتمال عملگر جهش را برابر ۰.۵ در نظر می‌گیریم. مسیر بهینه این مسئله بصورت

$$v_1 \rightarrow v_2 \rightarrow v_6 \rightarrow v_5 \rightarrow v_2 \rightarrow v_3 \rightarrow v_4 \rightarrow$$

$$v_5 \rightarrow v_3 \rightarrow v_5 \rightarrow v_6 \rightarrow v_1$$

یال‌های v_{35} و v_{56} دو بار تکرار شده‌اند.

بقیه یال‌ها یک‌بار تکرار شده‌اند.

هزینه کل سفر ۶۳۱،۴۴۶ می‌باشد.

۶- نتیجه

در این مقاله، پستی می‌خواهد نامه‌های خود را از دفتر پست تحویل بگیرد و تمام خیابان‌های موجود را حداقل یک‌بار طی کند و سپس به دفتر پست برگردد. به دلیل وجود شرایط نامشخص در جاده‌ها، هزینه‌ها بصورت متغیر نایقینی در نظر گرفته شده است. در این مقاله از برنامه‌ریزی نایقینی برای حل مسئله پستی در شرایط نایقینی استفاده شده است. سپس از الگوریتم ژنتیک استفاده شده است که با یک تکنیک جستجو چند جهته به یک جواب تقریباً مطلوب همگرا می‌شود.

مراجع

- [1] Sokmen, O. C., Emec, S., Yilmaz, M., & Akkaya, G. *An overview of chinese postman problem*. Paper presented at the 3rd International Conference on Advanced Engineering Technologies, 2019.
- [2] Zhang, B., & Peng, J. *Uncertain programming model for Chinese postman problem with uncertain weights*. *Industrial Engineering and Management Systems*, 11(1), 18-25, 2012. <https://doi.org/10.7232/iems.2012.11.1.018>.
- [3] Jun, B. H., Kang, M. J., & Han, C. G. *A Genetic Algorithm for the Chinese Postman Problem on the Mixed Networks*. *Journal of the Korea Society of Computer and Information*, 10(1), 181-188, 2005.
- [4] Majumder, S., Kar, S., & Pal, T. *Uncertain multi-objective Chinese postman problem*. *Soft Computing*, 23(22), 11557-11572, 2019. <https://doi.org/10.1007/s00500-018-03697-3>.
- [5] Keresztury, B. *Genetic algorithms and the Traveling Salesman Problem*. Paper presented at the Proceedings of the first International Conference on Genetic Algorithms and their Applications, 2017. <https://doi.org/10.2174/2213275912666190617155828>
- [6] Masuyama, H., Ichimori, T., & Sasama, T. *On ability of orthogonal genetic algorithms for the mixed chinese postman problem*. Paper presented at the ICSoft (1), 2006.
- [7] Xu, Q., Zhu, Y., & Yan, H. *Uncertain Random Programming Models for Chinese Postman Problem*. In *LISS2019*(pp.651-664):Springer, 2020. http://doi.org/10.1007/978-981-15-5682-1_47.



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

آموزش شبکه‌های عصبی پیشرو با استفاده از الگوریتم بهینه‌ساز تعادل به منظور شناسایی سیستم غیرخطی

سید سینا محمدی^۱، محمد ملایی امامزاده^{۲*}، مجتبی برخوردار یزدی^۳

^۱ دانشگاه شهید باهنر، بخش مهندسی برق، کرمان، Seyedsinamohammadi@gmail.com

^۲ دانشگاه شهید باهنر، بخش مهندسی برق، کرمان، Molaie@uk.ac.ir

^۳ دانشگاه شهید باهنر، بخش مهندسی برق، کرمان، Barkhordari@uk.ac.ir

چکیده: اخیراً استفاده از شبکه‌های عصبی مصنوعی (ANN) در حوزه‌های تحقیقاتی متعددی رواج یافته است و کاربردهای بسیاری برای آن پیشنهاد و استفاده شده است. یکی از مهم‌ترین موضوعات در شبکه‌های عصبی، بحث یادگیری آن‌ها می‌باشد لذا انتخاب و اعمال یک روش آموزشی مناسب برای رسیدن به بهترین جواب، جزء مسائل ضروری می‌باشد. دو روش معمول برای آموزش شبکه‌های عصبی وجود دارد: روش پس‌انتشار (BP) که به دلیل ماهیت خود در بسیاری از مسائل، به سرعت به جواب‌های محلی همگرا می‌شود و روش استفاده از الگوریتم‌های هوش جمعی که با توجه به نوع عملکرد و استفاده از هوش جمعی می‌توانند از همگرایی به جواب‌های محلی جلوگیری کرده و در زمان مناسب به جواب‌های بهتری دست پیدا کنند. جدیدترین تحقیقات انجام شده در زمینه الگوریتم‌های هوش جمعی منجر به معرفی روشی به نام الگوریتم بهینه‌ساز تعادل (EO) شده است که می‌تواند یک مسئله بهینه‌سازی را در زمان کوتاه و با جواب مناسبی حل نماید. هدف اصلی این مقاله، استفاده از این الگوریتم بهینه‌ساز تعادل در بحث آموزش شبکه‌های عصبی می‌باشد که برای شناسایی یک رابطه غیرخطی مورد استفاده قرار گرفته است. برای بررسی عملکرد روش پیشنهادی، نتایج حاصل از این روش با سه روش شناخته شده دیگر: روش پس‌انتشار خطا، الگوریتم بهینه‌ساز ازدحام ذرات (PSO) و الگوریتم بهینه‌ساز سالپ (SSO)، مورد مقایسه و ارزیابی قرار گرفته و مزیت‌های روش پیشنهادی در یادگیری شبکه عصبی بیان شده است.

کلمات کلیدی: شبکه عصبی پیشرو، الگوریتم بهینه‌ساز تعادل، آموزش شبکه عصبی، شناسایی سیستم.

پس‌انتشار خطا، اولین روش برای آموزش یک شبکه عصبی بوده که همچنان به صورت رایج در این زمینه مورد استفاده قرار می‌گیرد.

روش‌های هوش جمعی به طور معمول الهام گرفته شده از طبیعت، فیزیک، رفتار حیوانات و قوانین حاکم مابین آن‌ها می‌باشند. الگوریتم PSO یک روش بر پایه جمعیت است که با استفاده از عملکرد ذرات به عنوان یک گروه، قابلیت حرکت در فضای جستجو را دارا می‌باشد و در آن موقعیت نقطه بهینه با هماهنگی بین ذرات و یک عضو برتر به عنوان رهبر، پیدا می‌شود [۶]. پاسخ نهایی الگوریتم PSO تحت تأثیر دو رفتار فردی و اجتماعی است که عملکردهای فردی یک عضو اشاره به قابلیت جستجو و عملکردهای اجتماعی در یک گروه اشاره به بهره‌برداری دارد. الگوریتم PSO قابلیت کنترل بر روی هر دو مؤلفه فردی و اجتماعی را دارد که کمک می‌کند با انتخاب مناسب هر کدام با توجه به مسئله به جواب مناسب‌تری دست پیدا کنیم. این الگوریتم به عنوان یکی از برترین روش‌های آموزشی شبکه عصبی شناخته شده است که جهت مقایسه با روش پیشنهادی مورد بررسی نیز قرار گرفته است [۷-۱۱].

۱- مقدمه

یک شبکه عصبی مصنوعی، الهام گرفته شده از شبکه عصبی بیولوژیکی می‌باشد که قادر است همانند شبکه‌های عصبی موجود در مغز موجودات، با تلاش، شکست و تکرار دوباره آن وظایف را در طول گذشت زمان یاد بگیرد [۱]. این شبکه‌های عصبی مصنوعی کاربردهای خود را در طیف گسترده‌ای از تحقیقات عملی پیدا کرده‌اند که می‌توان به استفاده از شبکه‌های عصبی در بحث پزشکی و پیش‌بینی بیماری قلبی [۲]، تشخیص مبتلا بودن بیماران به کووید ۱۹ [۳]، پیش‌بینی وضعیت سهام در بورس [۴] و تنظیم نمودن پارامترهای یک کنترل کننده PID [۵] اشاره کرد.

در شبکه‌های عصبی، استفاده از الگوریتم یادگیری مناسب اثر بسیار قابل توجهی بر عملکرد کلی سیستم دارد که باید به درستی انتخاب شود. الگوریتم یادگیری وظیفه پیدا نمودن مقادیر مناسب وزن‌ها و بایاس‌ها را بر عهده دارد. روش

الگوریتم SSO همانند PSO از تحت تأثیر رفتارهای فردی و اجتماعی می باشد [۱۲]. در مقایسه با الگوریتم PSO، الگوریتم SSO دارای ساختار ریاضی ساده تری می باشد و با توجه به اینکه در این الگوریتم پارامتری برای تنظیم وجود ندارد، کنترل کمتری بر روی حل مسئله را دارا می باشد. طبقه حرکت ذرات به این گونه می باشد که همانند یک زنجیره، هر عضوی، عضو بعدی که دارای عملکرد بهتری هست را دنبال می کند که به آن زنجیره سالپ گویند. این نوع حرکت به صورت قابل توجهی باعث افزایش قابلیت بهره وری و افزایش همگرایی به نقطه بهینه خواهد شد ولی همچنین موجب کاهش وسعت جستجوی ذرات برای پیدا نمودن نقطه بهینه می شود که در نهایت ممکن است تنها به یک نقطه بهینه محلی به سرعت همگرا شود. این الگوریتم به دلیل جدید بودن خود نسبت به روش های بهینه سازی دیگر، با روش پیشنهادی مورد مقایسه قرار گرفته است [۱۳].

داشتن تنها یک مؤلفه جهت تنظیم نرخ به روزرسانی در روش پس انتشار خطا، معمولاً به دلیل سادگی آن در اعمال، به عنوان یک ویژگی خوب و نقطه قوت بیان می گردد اما این سادگی می تواند به دلیل عدم وجود انعطاف پذیری در حل مسائل به یک نقطه ضعیف نیز تبدیل گردد. از طرف دیگر، داشتن پارامترهای بیشتر جهت تنظیم در الگوریتم های هوش جمعی، به ما قابلیت کنترل بیشتر و انعطاف پذیری بالاتری را برای حل مسئله و رسیدن به جواب بهتری خواهد داد که همین امر باعث پیچیده تر شدن الگوریتم جهت راه اندازی و انتخاب پارامترهای مناسب می شود. روش پس انتشار خطا نیز به علت بودن متداول ترین و شناخته ترین الگوریتم یادگیری شبکه عصبی مورد بررسی با سایر روش ها قرار خواهد گرفت [۱۴].

در این مقاله، الگوریتم بهینه ساز تعادل (EO) جهت آموزش یک شبکه عصبی پیشرو برای درک بهتر ویژگی های این الگوریتم در بحث یادگیری شبکه عصبی مورد بررسی قرار گرفته است. دلیل انتخاب این الگوریتم از سوی ما جهت آموزش شبکه عصبی، داشتن تعادل مناسبی میان دو قابلیت کاوش و بهره وری در عملکرد آن می باشد که می تواند احتمال رسیدن به پاسخ بهینه سراسری را افزایش دهد. در این الگوریتم، برخلاف الگوریتم های PSO و SSO که دارای یک عضو برتر به عنوان رهبر هستند، الگوریتم EO دارای پنج رهبر می باشد که ذرات در هر بار تکرار یکی از آن ها را به صورت تصادفی جهت همگرایی انتخاب می کنند. در این الگوریتم همچنان یک نرخ تولید (generation rate) نیز به خوبی تعریف شده است که هم باعث افزایش قابلیت کاوش و هم قابلیت بهره وری خواهد شد [۱۵].

این مقاله به صورت روبرو ارائه شده است. در بخش ۲، ما الگوریتم EO و نحوه عملکرد آن را شرح خواهیم داد. در بخش ۳، روش آموزش شبکه عصبی با استفاده از الگوریتم EO بیان خواهد شد. در بخش ۴، نتایج شبیه سازی مورد بحث قرار خواهد گرفت. بخش ۵، خلاصه ای از نتایج ارائه خواهد گردید.

۲- الگوریتم بهینه ساز تعادل

در الگوریتم EO، در ابتدای شروع عمل بهینه سازی، تمامی ذرات به صورت تصادفی در فضای جستجو پخش خواهند شد و موقعیت هر ذره به عنوان یک جواب برای حل مسئله در نظر گرفته می شود. یک بردار به نام بردار کاندیدهای تعادل از چهار عضو برتر تشکیل می شود. تمامی ذرات برای رسیدن به جواب بهینه باید به نقطه تعادل خود در حالت نهایی دست پیدا نمایند. در طول هر بار تکرار، هر ذره موقعیت خود را به توجه یکی از عضوهای داخل بردار کاندیدهای تعادل به صورت تصادفی انتخاب نموده و به روزرسانی می نماید که این امر باعث

افزایش قابلیت کاوش و جستجوی هر ذره می شود. موقعیت اولیه هر ذره که به صورت تصادفی در شروع الگوریتم در فضای جستجو پخش می شود به صورت رابطه (۱) می باشد.

$$C_i^{initial} = C_{min} + rand_i(C_{max} - C_{min}), i = 1, 2, \dots, n \quad (1)$$

C_{ith} بیانگر موقعیت اولیه ذره i ام، C_{min} و C_{max} معرف محدوده پایینی و بالایی فضای جستجو، $rand_i$ یک بردار با مقادیر تصادفی مابین $[0, 1]$ تعداد ذرات تعریف شده در الگوریتم می باشد.

پس از آنکه هر ذره در ابتدا شروع الگوریتم در موقعیت های تصادفی خود قرار گرفتند، موقعیت هر ذره به عنوان یک جواب مورد ارزیابی قرار می گیرد که بعد مشخص شدن شایستگی هر ذره، می توان بردار کاندیدهای تعادل را با توجه به چهار عضو برتر تعیین نمود. در ابتدای الگوریتم، تنها بردار کاندیدهای تعادل الگو حرکات ذرات دیگر را مشخص می کند اما در ادامه هر ذره به صورت تصادفی یکی از کاندیدهای تعادل را با توجه به بردار (۲)، برای همگرایی انتخاب می کند.

$$C_{eq, pool} = \{C_{eq(1)}, C_{eq(2)}, C_{eq(3)}, C_{eq(4)}, C_{eq(ave)}\} \quad (2)$$

C_{eq} بیانگر میانگین موقعیت چهار عضو برتر دیگر در بردار کاندیدهای تعادل است که این عضو میانگین، به صورت مستقیم باعث تأثیر گذاری در قابلیت بهره وری و چهار عضو دیگر باعث افزایش قابلیت کاوش خواهند شد.

$$C_{new} = C_{eq} + (C_{old} - C_{eq})F + \frac{G}{\lambda}(1 - F) \quad (3)$$

رابطه (۳)، بیانگر روش و قانون به روزرسانی ذرات در این الگوریتم می باشد. همان طور که مشاهده می شود، قانون به روزرسانی تشکیل شده از سه قسمت مختلف می باشد. قسمت اول بیانگر یکی از اعضای برتر در بردار کاندیدهای تعادل (۲) است. قسمت دوم به عنوان یک مکانیسم جستجوی مستقیم عمل می کند که برگرفته شده از اختلاف میان موقعیت کنونی ذره و نقطه تعادل است. در قسمت سوم، نرخ جهش نقش یک اصلاح کننده و همچنین یک کاوشگر را ایفا می نماید. تعادل قابل قبولی بین قابلیت های کاوش و بهره وری با استفاده از رابطه نمایی F امکان پذیر خواهد بود.

$$F = e^{-\lambda(t-t_0)} \quad (4)$$

لاندا (λ)، یک بردار با مقادیر تصادفی مابین $[0, 1]$ می باشد و زمان های t و t_0 می توانند با توجه به روابط ۵ و ۶ به دست آورده شوند.

$$t = \left(1 - \frac{Iter}{Max_{iter}}\right)^{\left(a_2 \frac{Iter}{Max_{iter}}\right)} \quad (5)$$

$$t_0 = \frac{1}{\lambda} \ln(-a_1 \text{sign}(r - 0.5)[1 - e^{-\lambda t}] + t) \quad (6)$$

Max_{iter} نشان دهنده تعداد تمام تکرارهای تعریف شده برای الگوریتم می باشد و $Iter$ نشان دهنده تکراری هست که در حال اجرا می باشد. A_2 یک مقدار ثابت است که توسط ما پیش از اجرای الگوریتم تنظیم می شود و جهت کاهش سرعت همگرایی و کنترل قابلیت بهره وری به کار برده شده است و همچنین a_2 نیز به همین صورت جهت کنترل قابلیت کاوش مورد استفاده قرار می گیرد و r نیز یک بردار تشکیل شده از مقادیر تصادفی در بازه $[0, 1]$ هست که از تصادفی بودن حرکت هر ذره اطمینان حاصل می کند. نرخ جهش G طبق رابطه ۷ تعریف می گردد.

$$G = G_0 F \quad (7)$$

به صورتی که :

$$G_0 = G_{CP}(C_{eq} - \lambda C) \quad (8)$$

$$G_{CP} = \begin{cases} 0.5r_1 & r_2 \geq 0.5 \\ 0 & r_2 < 0.5 \end{cases} \quad (9)$$

r_1 و r_2 اعداد تصادفی با مقادیری در بازه $[0, 1]$ می باشد.

۳- آموزش شبکه عصبی پیشرو با استفاده از الگوریتم بهینه ساز تعادل (EO-FANN)

شبکه های عصبی را می توان در بسیار از زمینه های علمی از جمله تشخیص تقلب های مالی، ارزیابی ریسک، مهندسی، تشخیص بیماری، صنعت، آموزش و غیره استفاده نمود. همان طور که در بخش اول اشاره شد، روش های متعددی جهت آموزش و یادگیری شبکه های عصبی تاکنون پیشنهاد شده است. هر روشی به دلیل عملکرد متفاوت و داشتن ویژگی های متفاوت ذاتی، دارای مزیت ها و معایبی خواهد بود که کاربرد هر کدام را در حوزه های مختلف پررنگ تر خواهد کرد. همگرایی سریع الگوریتم SSO، می تواند در یادگیری های آنلاین که زمان همگرایی برای ما نسبت به جوابی عالی ضروری تر هست مناسب باشد [۱۶]. اما همین طور الگوریتم های تکاملی، قابلیت فرار از بهینه های محلی را دارند اما در عوض دارای سرعت همگرایی پایینی می باشند که این ویژگی می تواند در بحث یادگیری آنلاین و جایی که جواب بهتر برای ما از زمان رسیدن به آن جواب، مهم تر هست استفاده قرار گیرد. آموزش یک شبکه عصبی که دارای تعادل در میان قابلیت کاوش و بهره وری آن وجود داشته باشد، می تواند باعث افزایش عملکرد شبکه شود.

شبکه های عصبی پیشرو تشکیل شده از گره ها و وزن های مربوط به آن هستند که با انتخاب یک روش آموزشی مناسب جهت یادگیری شبکه، می توانند هر داده پیچیده و غیر خطی را به خوبی شناسایی نمایند. در فرآیند یادگیری، وزن ها و بایاس ها دو پارامتری هستند که یک الگوریتم آموزشی باید بتواند مقدار مناسب آن ها به دست آورده. یک شبکه عصبی پیشرو می تواند با توجه به مسئله دارای ساختاری متفاوتی در تعداد لایه ها و گره ها باشد. یک شبکه عصبی پیشرو به صورت معمول دارای سه لایه به نام های، لایه ورودی، لایه پنهان و لایه خروجی است.

در هنگام پروسه یادگیری، الگوریتم آموزشی بایاس ها و وزن ها را تا هنگامی که خطای تعریف شده به حداقل برسد، تغییر و تنظیم می نمایند. با توجه به نوع مسئله روش های آموزشی می توانند به صورت آنلاین یا آفلاین مورد استفاده قرار گیرند. در هنگام آموزش آفلاین، معمولاً رسیدن به بهترین جواب نسبت به زمان یادگیری دارای اولویت می باشد.

میانگین خطای مربع (MSE) میان خروجی شبکه (f_i) و مقدار هدف (y_i) طبق رابطه ۱۰ بعنوان تابع ارزیابی محاسبه می گردد.

$$MSE = \frac{1}{N} \sum_{i=1}^N (f_i - y_i)^2 \quad (10)$$

که در آن، N تعداد داده، f_i مقدار بازگردانده شده توسط شبکه و y_i مقدار هدف برای داده i می باشد.

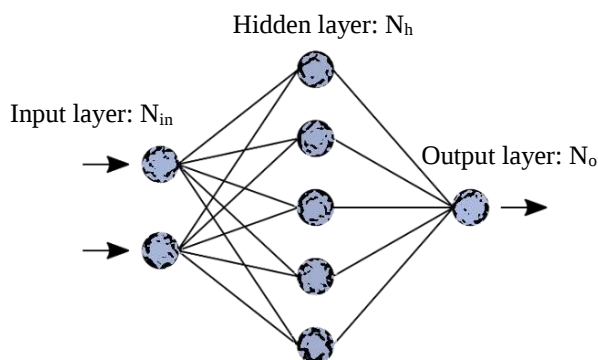
یادگیری شبکه عصبی آموزش دیده توسط الگوریتم بهینه ساز تعادل را می توان به صورت زیر خلاصه نمود:

گام اول. پارامترهای $a1$ و $a2$ را در رابطه ۵ و ۶ تنظیم نماید.

گام دوم. موقعیت اولیه هر عضو را به صورت تصادفی با توجه به رابطه ۱ تنظیم نماید.

گام سوم. موقعیت هر عضو را بعنوان یک راه حل در نظر گرفته و با توجه به شکل ۱، ابعاد هر عضو برابر است با رابطه ۱۱:

$$size \{X\} = \{N_{in} \times N_h + N_h + N_h \times N_o + N_o\} \quad (11)$$



شکل ۱: ساختار پایه یک شبکه عصبی پیشرو

۴- شبیه سازی

شبکه های عصبی قادر به تقریب و شناسایی هر گونه تابع پیوسته غیر خطی هستند. سه مجموعه داده مختلف بر اساس یک تابع غیر خطی همان طور که در رابطه های ۱۲، ۱۳ و ۱۴ نشان داده است، برای شناسایی استفاده شده است.

Data A : $x_A = Inputs = linspace(-1,1,100)$

$$y_A = Targets = \frac{1 + \sin(\frac{\pi}{4} x_A^2)}{2} \quad (12)$$

Data B : $x_B = Inputs = \cos x_A$

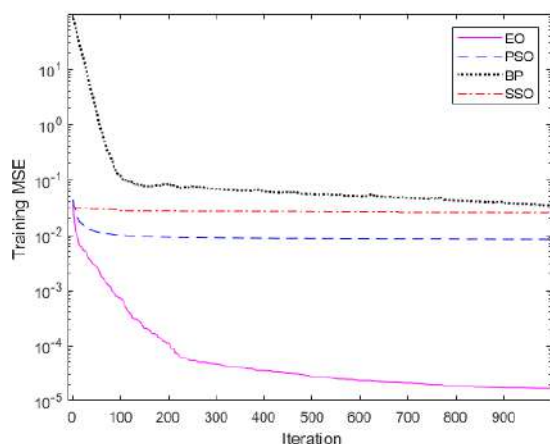
$$y_B = Targets = \frac{1 + \sin(\frac{\pi}{4} x_B^2)}{2} \quad (13)$$

Data C : $x_C = Inputs = rand(100,1)$

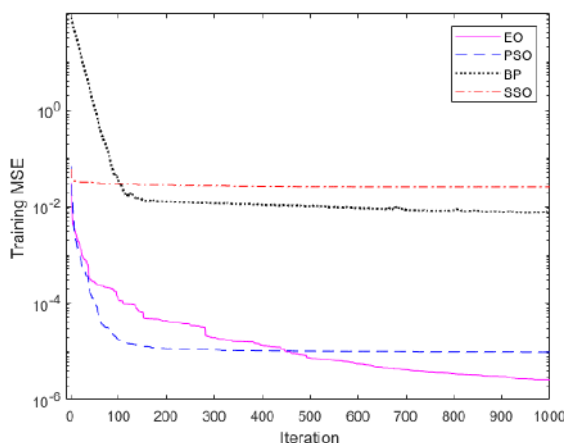
$$y_C = Targets = \frac{1 + \sin(\frac{\pi}{4} x_C^2)}{2} \quad (14)$$

برای ارزیابی روش پیشنهادی، شبکه آموزش دیده با الگوریتم پیشنهادی، با سه الگوریتم رایج دیگر بانام های الگوریتم PSO (یکی از قویترین روش های بهینه سازی)، SSO (یکی از جدیدترین روش های جستجو) و پس از آن شار خطا (متداول ترین روش یادگیری شبکه عصبی) مورد مقایسه قرار می گیرد. الگوریتم های یادگیری با دو شرط توقف متفاوت مورد آزمایش و بررسی قرار گرفته اند: یکی بر اساس تعداد تکرارها و دیگری بر اساس رسیدن خطای یادگیری به یک سطح آستانه قابل قبول. در الگوریتم های آموزشی به دلیل آنکه موقعیت اولیه هر ذره در هر بار اجرا به صورت تصادفی انتخاب می شود، از آنجایی که این مقدار اولیه تصادفی می تواند در نتیجه نهایی آن تأثیر زیادی بگذارد، بنابراین شبیه سازی های انجام شده به تعداد ۲۰ بار برای هر الگوریتم تکرار شده است تا ارزیابی مناسب تری انجام گردد.

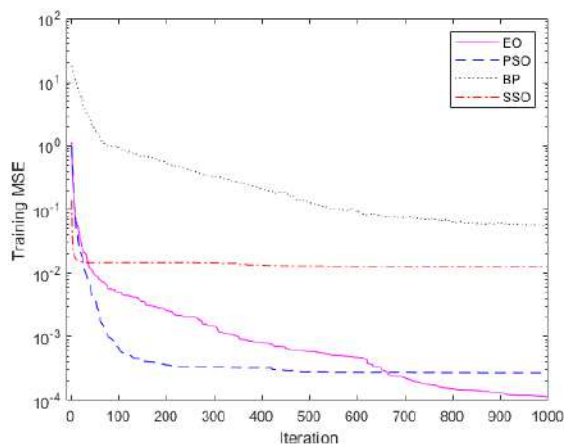
در این مقاله، ساختار شبکه عصبی به صورت سه لایه می باشد و تعداد نرون های لایه پنهان با توجه به مسئله، متغیر خواهد بود. تابع فعال سازی در لایه پنهان و لایه خروجی به ترتیب توابع tansig و purelin انتخاب شده اند. شبیه سازی ها با استفاده از یک کامپیوتر با ۱۶ گیگابایت رم و پردازنده Core i5 7400 در محیط ویندوز ۱۰ انجام شده است.



شکل ۲: خطای آموزش برای داده A بر حسب تکرار یادگیری



شکل ۳: خطای آموزش برای داده B بر حسب تکرار یادگیری



شکل ۴: خطای آموزش برای داده C بر حسب تکرار یادگیری

با توجه به شکل‌های ۲، ۳ و ۴ مشاهده می‌شود که الگوریتم بهینه‌ساز تعادل در ابتدای آموزش، به یک نقطه بهینه محلی همگرا نمی‌شود و نسبت به بقیه روش‌ها سرعت همگرایی اولیه آن کم بوده ولی در ادامه با جستجوی بهتر جواب در فضای جستجو تقریباً به بهترین جواب خواهد رسید. در ادامه، شرط دوم برای توقف الگوریتم‌ها، رسیدن خطای مورد به یک سطح آستانه مطلوب (thr) تعریف شده است که شبیه‌سازی‌ها برای دو سطح مختلف از این حد آستانه (thr=0.001 و thr=0.0001) اجرا شده‌اند تا تجزیه و تحلیل مناسبی بتوان انجام داد. نتایج به دست آمده در جداول ۵-۱۰ آورده

در انجام شبیه‌سازی، مقادیر پارامترهای هر الگوریتم به صورت زیر انتخاب شده‌اند. در الگوریتم بهینه‌ساز تعادل a_1 و a_2 در رابطه ۵ و ۶ به صورت زیر: $a_1=2$, $a_2=1$

و در الگوریتم PSO مقادیر زیر:

inertia weight=1, inertia weight damping ratio = 0.99, personal and global learning coefficients = 1.5, 2.0

انتخاب شده‌اند. همچنین در روش پس‌انتشار خطا، نرخ به‌روزرسانی برای داده‌های A و B مقدار ۰.۰۰۲ و برای داده C مقدار ۰.۰۰۰۲ در نظر گرفته شده است. همچنین تعداد نرون‌ها در لایه پنهان و همچنین فضای جستجو هر عضو در جدول ۱ برای هر داده آورده شده است.

جدول ۱: مشخصات ساختار اولیه شبکه عصبی و الگوریتم آموزشی

Dataset	Parameter		
	Number of hidden neurons	Dimension of particles	Search space of particles
A	5	16	[-1,1]
B	10	31	[-2,2]
C	100	301	[-2,2]

تعداد اجرا برای الگوریتم‌های بهینه‌ساز تعادل، PSO و SSO به مقدار ۱۰۰۰ تکرار در نظر گرفته شده است و در روش پس‌انتشار خطا این تکرار به تعداد ۴۵۰۰ افزایش داده شده است تا از نظر زمانی و کارایی با الگوریتم‌های دیگری قابل مقایسه باشد. نتایج این شبیه‌سازی‌ها در جدول‌های ۲ تا ۴ آمده است.

جدول ۲: عملکرد الگوریتم‌های آموزشی برای داده A

Training algorithms	Cost function: Mean Square Error		
	Worst	Average	Best
EO (1000 itr)	4.14E-05	1.67E-05	7.83E-07
PSO (1000 itr)	4.79E-03	8.28E-03	3.14E-04
SSO (1000 itr)	2.56E-02	2.25E-02	1.86E-02
BP (45000 itr)	1.41E-02	1.08E-03	2.63E-05

جدول ۳: عملکرد الگوریتم‌های آموزشی برای داده B

Training algorithms	Cost function: Mean Square Error		
	Worst	Average	Best
EO (1000 itr)	1.15E-05	2.25E-06	7.61E-09
PSO (1000 itr)	8.80E-05	9.54E-06	2.29E-07
SSO (1000 itr)	5.47E-02	2.53E-02	1.36E-06
BP (45000 itr)	3.29E-01	2.53E-03	2.97E-04

جدول ۴: عملکرد الگوریتم‌های آموزشی برای داده C

Training algorithms	Cost function: Mean Square Error		
	Worst	Average	Best
EO (1000 itr)	3.84E-04	1.12E-04	8.42E-07
PSO (1000 itr)	1.85E-03	2.73E-04	1.48E-05
SSO (1000 itr)	1.62E-01	1.28E-02	8.98E-05
BP (45000 itr)	7.10E-03	3.06E-03	1.20E-03

همان‌طور که در جداول ۲، ۳ و ۴ دیده می‌شود، الگوریتم بهینه‌ساز تعادل توانسته است تا به جواب بهتری (خطای آموزش کم‌تری) در مقایسه با سایر الگوریتم‌های یادگیری برسد.

همچنین نتایج این شبیه‌سازی‌ها در شکل‌های ۲ تا ۴ نیز رسم شده است؛ که در این شکل‌ها، خطای رسم شده در حقیقت میانگین خطای ۲۰ بار اجرا برای هر الگوریتم می‌باشد که می‌تواند مشخص‌کننده میزان کارایی و عملکرد این الگوریتم‌ها در یادگیری شبکه عصبی باشد.

در ادامه برای ارزیابی این روش‌ها در حالتی که دقت یادگیری بسیار بالایی نیاز است، حد آستانه خطا را به $\text{thr}=0.0001$ کاهش داده‌ایم و دوباره شبیه‌سازی‌ها را تکرار کرده‌ایم و نتایج در جدول‌های ۸، ۹ و ۱۰ آمده است.

جدول ۸: میانگین عملکرد الگوریتم‌های آموزشی برای رسیدن به حد آستانه $\text{thr}=0.0001$ برای داده A در ۲۰ بار تکرار

Training algorithms	Computational time (sec)	Iteration number	success rate of algorithms
EO	32.57	158.4	100%
PSO	-	-	0%
SSO	-	-	0%
BP	+483	46654.5	10%

در نتایج جدول ۸ مشابه نتایج جدول ۵، روش EO دارای بهترین جواب بوده و سایر روش‌ها از جمله PSO نتیجه قابل قبولی نداشته‌اند.

جدول ۹: میانگین عملکرد الگوریتم‌های آموزشی برای رسیدن به حد آستانه $\text{thr}=0.0001$ برای داده B در ۲۰ بار تکرار

Training algorithms	Computational time (sec)	Iteration number	success rate of algorithms
EO	14.22	70.3	100%
PSO	+222	+1000	70%
SSO	+475.5	+1000	5%
BP	-	-	0%

در جدول ۹ که دقت یادگیری بیشتری نسبت به جدول ۶ نیاز است، مشاهده می‌شود که روش PSO برتری خود را از دست داده است و فقط روش EO دارای موفقیت ۱۰۰٪ در رسیدن به $\text{thr}=0.0001$ می‌باشد.

جدول ۱۰: میانگین عملکرد الگوریتم‌های آموزشی برای رسیدن به حد آستانه $\text{thr}=0.0001$ برای داده C در ۲۰ بار تکرار

Training algorithms	Computational time (sec)	Iteration number	success rate of algorithms
EO	191.54	783.5	90%
PSO	+240.03	+1000	55%
SSO	+498.3	+1000	5%
BP	-	-	0%

در نتایج جدول ۱۰ برای داده C هم مشاهده می‌شود که در حالتی که دقت یادگیری بالایی نیاز باشد، روش EO نسبت به سایر روش‌ها برتری فراوانی دارد.

به‌طور کلی در نتایج جدول‌های ۵ تا ۷ مشاهده شد که اگر دقت یادگیری را بسیار بالا در نظر بگیریم و به یک جواب بهینه نه لزوماً بهترین جواب راضی شویم، الگوریتم PSO ممکن است نسبت به الگوریتم بهینه‌ساز تعادل بتواند به این نقطه بهینه محلی (سطح آستانه $\text{thr}=0.0001$) سریع‌تر همگرا شود اما در عوض اگر مشابه جدول‌های ۸ تا ۱۰، دقت یادگیری را بسیار بالا و دقیق در نظر بگیریم که معادل بهینه سراسری باشد، در این صورت دیگر روش PSO برتری خود را از دست داده و روش EO دارای بهترین نتایج در رسیدن به نقطه بهینه سراسری (سطح آستانه $\text{thr}=0.0001$) می‌باشد.

در این بخش، شبکه عصبی پیشرو با الگوریتم‌های مختلفی آموزش دیده شد که نتایج به‌دست آمده از شبیه‌سازی‌ها نشان‌دهنده عدم همگرایی الگوریتم‌های PSO، SSO و BP به بهینه سراسری را می‌باشد. الگوریتم بهینه‌ساز تعادل توانسته است تا با توجه به ویژگی‌های منحصر به فردی که دارد، بتواند به بهترین راه‌حل‌ها در نقاط بهینه سراسری دست پیدا نماید.

شده‌اند. در این آزمایش‌ها، ستون “success rate of algorithms” بیان‌گر آن می‌باشد که از کل ۲۰ بار اجرا، چند بار (به درصد) الگوریتم موفق به رسیدن به آستانه خطای مورد قبول شده است. جهت ثبت زمان محاسباتی، حداکثر زمان برای رسیدن به حد آستانه برای هر بار اجرا مقدار ۵۰۰ ثانیه فرض شده است و در صورتی که الگوریتم در این زمان قادر به رسیدن حد آستانه تعیین شده نباشد، زمان اجرای آن مرحله ۵۰۰+ ثانیه در نظر گرفته شده است.

جدول ۵: میانگین عملکرد الگوریتم‌های آموزشی برای رسیدن به حد آستانه $\text{thr}=0.001$ برای داده A در ۲۰ بار تکرار

Training algorithms	Computational time (sec)	Iteration number	success rate of algorithms
EO	17.84	90.4	100%
PSO	+403.6	+1000	20%
SSO	-	-	0%
BP	197.32	48366.3	100%

در جدول ۵ که مشاهده می‌شود روش EO دارای بهترین جواب می‌باشد و روش PSO دارای موفقیت ۲۰٪ (۴ بار موفق و ۱۶ بار ناموفق) در رسیدن به سطح آستانه می‌باشد. روش BP با تکرار و زمان بیشتری به جواب رسیده است.

جدول ۶: میانگین عملکرد الگوریتم‌های آموزشی برای رسیدن به حد آستانه $\text{thr}=0.001$ برای داده B در ۲۰ بار تکرار

Training algorithms	Computational time (sec)	Iteration number	success rate of algorithms
EO	3.70	19	100%
PSO	2.38	15.4	100%
SSO	+467	+1000	25%
BP	-	-	0%

با توجه به جدول ۶ هر دو روش EO و PSO دارای موفقیت ۱۰۰٪ در رسیدن به سطح آستانه $\text{thr}=0.001$ می‌باشند ولی روش PSO دارای نتایج بهتری است. روش SSO دارای موفقیت ۲۵٪ (۵ بار موفق و ۱۵ بار ناموفق) در رسیدن به سطح آستانه مطلوب می‌باشد. روش BP در این حالت نتوانسته است به این حد آستانه برسد.

جدول ۷: میانگین عملکرد الگوریتم‌های آموزشی برای رسیدن به حد آستانه $\text{thr}=0.001$ برای داده C در ۲۰ بار تکرار

Training algorithms	Computational time (sec)	Iteration number	success rate of algorithms
EO	59.12	279.5	100%
PSO	+41.53	+179.3	95%
SSO	+483	+1000	20%
BP	-	-	0%

با توجه به نتایج جدول ۷، روش EO دارای بهترین نتایج و درصد موفقیت ۱۰۰٪ می‌باشد. روش PSO در صورت رسیدن به سطح آستانه، زمان کمتری نسبت به EO نیاز داشته است ولی فقط در ۹۵٪ به سطح آستانه رسیده است و در ۵٪ موارد به سطح آستانه مطلوب نرسیده است.

با توجه به جداول ۵، ۶ و ۷ می‌توان مشاهده نمود که الگوریتم بهینه‌ساز تعادل هر چند که از لحاظ زمانی و سرعت رسیدن به این جواب (حد آستانه $\text{thr}=0.001$) همیشه بهترین نتایج را نسبت به PSO نداشته است ولی توانسته است تا در تمامی موارد شبیه‌سازی با موفقیت به حد آستانه تعریف شده برسد. به دلیل عدم انعطاف‌پذیری لازم دو روش BP و SSO، این دو روش نتوانسته‌اند به جواب قابل قبولی دست پیدا کنند و در عوض در همان مراحل اولیه به یک نقطه بهینه محلی همگرا شده‌اند.

5- نتیجه گیری

در این مقاله، الگوریتم بهینه‌ساز تعادل برای آموزش یک شبکه عصبی پیش‌رو مورد استفاده قرار گرفت تا ویژگی‌های این الگوریتم در بحث یادگیری مورد بررسی قرار گیرند. برای ارزیابی روش پیشنهادی، سه روش PSO، SSO و BP نیز برای یادگیری شبکه عصبی مورد استفاده قرار گرفتند. در نتایج به دست آمده مشاهده شد که علی‌رغم سرعت بالای همگرایی الگوریتم SSO در زمان اولیه اجرا و به دلیل پایین بودن تنوع جمعیت و نوع حرکات زنجیره‌وار ذرات، این الگوریتم در حل مسائل ذکر شده نتوانست تا به جواب مناسبی دست پیدا نماید. الگوریتم پس‌انتشار خطا (BP)، به علت عدم انعطاف‌پذیری در حل مسئله، به یک نقطه بهینه محلی همگرا شده و قادر به یافتن راه‌حل‌های بهتر نبوده است. الگوریتم PSO تنها دارای یک عضو برتر به عنوان رهبر در میان کل مجموعه ذرات می‌باشد که این عامل می‌تواند سرعت همگرایی ذرات را افزایش دهد، اما در مقابل نیز می‌تواند باعث همگرایی به یک جواب بهینه محلی شده و ذرات نقاط بهینه دیگر را در هنگام کاوش، از دست دهند. در الگوریتم بهینه‌ساز تعادل با استفاده از پنج عضو پیش‌رو برای راهنمایی ذرات دیگر باعث و سبب تر شدن فضای کاوش و جستجو شده است و از همگرایی سریع ذرات به یک جواب بهینه محلی جلوگیری کرده است. همین امر باعث افزایش عملکرد و انعطاف‌پذیری بالاتر این الگوریتم در حل مسائل مختلف خواهد شد.

با توجه به نتایج به دست آمده از شبیه‌سازی‌های انجام شده، می‌توان چنین نتیجه گرفت که روش‌های مبتنی بر هوش جمعی مانند PSO و EO نسبت به روش متداول BP دارای کارایی بهتری در یادگیری شبکه‌های عصبی هستند. همچنین روش PSO به دلیل استفاده از یک عضو رهبر، با وجود سرعت بیشتر در ابتدای مراحل یادگیری، در ادامه ممکن است که در نقاط بهینه محلی گیر افتاده و از رسیدن به جواب بهینه جا بماند. ولی روش EO به دلیل استفاده از چند عضو پیش‌رو، هرچند که در مراحل اولیه دارای بهترین جواب نیست، اما در ادامه مراحل یادگیری توانسته به نتایج بهتری نسبت به سایر روش‌ها دست پیدا کند. این مزیت روش EO بیشتر خودش را در حالتی که دقت یادگیری بالاتری مدنظر است مانند $\text{thr}=0.0001$ خودش را نشان می‌دهد. لذا این مقاله روش EO را به عنوان یک روش مناسب برای یادگیری دقیق شبکه‌های عصبی پیشنهاد می‌کند.

مراجع

- [6] Kennedy, J., & Eberhart, R., "Particle swarm optimization", International Conference on Neural Networks, vol. 4, pp. 1942–1948, 1995.
- [7] F. G. E. Alfassio Grimaldi, "PSO as an effective learning algorithm for neural network applications," 3rd International Conference on Computational Electromagnetics and Its Applications, 2004.
- [8] J. Zhou, Z. Duan, Y. Li, J. Deng, and D. Yu, "PSO-based neural network optimization and its utilization in a boring machine," Journal of Materials Processing Technology, vol. 178, no. 1-3, pp. 19–23, 2006.
- [9] S. Chatterjee, S. Hore, N. Dey, S. Chakraborty, and A. S. Ashour, "Dengue fever classification using Gene Expression Data: A PSO based Artificial Neural Network Approach," Advances in Intelligent Systems and Computing, pp. 331–341, 2017.
- [10] M. Shariati, M. S. Mafipour, et al., "Application of a hybrid artificial neural network-particle swarm optimization (Ann-PSO) model in behavior prediction of channel shear connectors embedded in normal and high-strength concrete," Applied Sciences, vol. 9, no. 24, p. 5534, 2019.
- [11] B. Wang, H. Moayedi, H. Nguyen, L. K. Foong, and A. S. Rashid, "Feasibility of a novel predictive technique based on artificial neural network optimized with particle swarm optimization estimating pullout bearing capacity of helical piles," Engineering with Computers, vol. 36, no. 4, pp. 1315–1324, 2019.
- [12] Mirjalili, S.A., Gandomi, A.H., Mirjalili, S.Z., Saremi, S., Faris, H., & Mirjalili, S.M., "Salp Swarm Algorithm: A bio-inspired optimizer for engineering design problems", Advances in Engineering Software, vol. 114, pp. 163–191, 2017.
- [13] L. Abualigah, M. Shehab, M. Alshinwan, and H. Alabool, "SALP SWARM ALGORITHM: A comprehensive survey," Neural Computing and Applications, vol. 32, no. 15, pp. 11195–11215, 2019.
- [14] A. T. C. Goh, "Back-propagation neural networks for Modeling Complex Systems," Artificial Intelligence in Engineering, vol. 9, no. 3, pp. 143–151, 1995.
- [15] Faramarzi, A., Heidarinejad, M., Stephens, B., & Mirjalili, S., "Equilibrium optimizer: A novel optimization algorithm", Knowledge-Based Systems, vol. 191, article 105190, 2020.
- [16] Labbaf Khaniki, M. A., Behzad Hadi, M., & Manthouri, M., "Feedback Error Learning Controller based on RMSprop and Salp Swarm Algorithm for Automatic Voltage Regulator System", 10th International Conference on Computer and Knowledge Engineering, pp. 425–430, 2020.
- [1] Hornik, Kurt, et al., "Multilayer Feedforward Networks Are Universal Approximators", Neural Networks, vol. 2, pp. 359–366, 1989.
- [2] Thanga Selvi, R., Muthulakshmi, I., "An optimal artificial neural network based big data application for heart disease diagnosis and classification model", J Ambient Intell Human Comput, vol. 12, pp. 6129–6139, 2021.
- [3] Abbas, A., Abdelsamea, M.M. & Gaber, M.M., "Classification of COVID-19 in chest X-ray images using DeTraC deep convolutional neural network", Appl Intell, vol. 51, pp. 854–864, 2021.
- [4] Zhang, D., and Sha L., "The Application Research of Neural Network and BP Algorithm in Stock Price Pattern Classification and Prediction", Future Generation Computer Systems, vol. 115, pp. 872–879, 2021.
- [5] Zeng, Guo-Q., et al., "Adaptive Population Extremal Optimization-Based PID Neural Network for Multivariable Nonlinear Control Systems", Swarm and Evolutionary Computation, vol. 44, pp. 320–334, 2019.

بررسی تاثیر بکارگیری توابع زیان مختلف بر عملکرد مدل خوشه‌بندی فازی برای داده‌های فازی در حضور داده‌های پرت

الهام اسکندری^{۱*} و علیرضا خواستار^۲

^۱ دانشجوی دکتری، دانشکده ریاضی، دانشگاه تحصیلات تکمیلی علوم پایه زنجان، زنجان، eskandari.e@iasbs.ac.ir

^۲ دانشیار، دانشکده ریاضی، دانشگاه تحصیلات تکمیلی علوم پایه زنجان، زنجان، khastan@iasbs.ac.ir

چکیده: در این مقاله، استفاده از توابع زیان لگاریتم کسینوس هذلولوی و چندکی در مدل خوشه‌بندی فازی برای داده‌های فازی را پیشنهاد و تاثیر استفاده از توابع زیان مربع، هابر، خطی، سیگموئیدی، لگاریتمی، لگاریتم کسینوس هذلولوی و چندکی بر عملکرد مدل در حضور داده‌های پرت را در شبیه‌سازی صورت پذیرفته بررسی کرده‌ایم. مجموعه داده‌های مصنوعی و واقعی مورد استفاده، از نظر تعداد ویژگی‌ها (۲، ۳ و ۹)، تعداد کلاس‌ها (۳، ۴ و ۷) و تعداد و پخش داده‌های پرت، دارای تنوع مناسبی هستند. نتایج شبیه‌سازی مویید آن است که توابع زیان لگاریتمی و لگاریتم کسینوس هذلولوی، نسبت به وجود داده‌های پرت مقاوم هستند.

کلمات کلیدی: تابع زیان، خوشه‌بندی فازی مقاوم، داده‌های فازی $L-R$ ، داده‌های پرت، تابع زیان لگاریتم کسینوس هذلولوی، تابع زیان چندکی، تابع زیان هابر.

۱ مقدمه

توسط بزدک [۱] تعمیم داده شد، شناخته شده‌ترین و کاربردی‌ترین روش خوشه‌بندی فازی است [۲]. ساتو^۴ و ساتو^۵ [۳] نخستین کسانی بودند که نقش عدم قطعیت در داده‌ها را توأم با نقش عدم قطعیت در تخصیص داده‌ها به خوشه‌ها، مورد توجه قرار داده و در سال ۱۹۹۵ یک مدل خوشه‌بندی فازی برای داده‌های فازی ارائه کردند. برخی پژوهشگران توجه خود را به توسعه روش‌های مقاوم^۶ در مقابل داده‌های پرت معطوف کرده‌اند؛ بعضی از مهم‌ترین رویکردها عبارت‌اند از انتخاب واسط^۷ داده‌ها به عنوان سرخوشه^۸ به جای انتخاب میانگین آن‌ها، استفاده از یک معیار عدم تشابه مقاوم، در نظر گرفتن یک خوشه‌ی

خوشه‌بندی به فرایند گروه‌بندی داده‌ها در خوشه‌هایی که اعضای آن‌ها دارای بیش‌ترین تشابه درون‌خوشه‌ای و کم‌ترین تشابه میان‌خوشه‌ای باشند، اطلاق می‌شود. لحاظ کردن عدم قطعیت در تخصیص داده‌ها به خوشه‌ها، منجر به پیدایش الگوریتم‌های خوشه‌بندی فازی شده است که در آن‌ها، برخلاف الگوریتم‌های خوشه‌بندی قطعی^۱، هر داده فقط به یک خوشه تعلق نداشته و داده‌ها با یک درجه عضویت خاص به هر خوشه اختصاص داده می‌شوند. روش خوشه‌بندی C -میانگین که در سال ۱۹۷۴ به طور مستقل توسط دان^۲ و بزدک^۳ معرفی و در سال ۱۹۸۱

^۴Sato

^۵Sato

^۶Robust

^۷Medoid

^۸Prototype

^۱Crisp

^۲Dunn

^۳Bezdek

*نویسنده‌ی مسئول

اضافی به نام خوشه‌ی نويز، حذف کسری از داده‌های پرت، وزن‌دهی کم به تاثیر داده‌های پرت و احتمالاتی.^۹

رویکردی دیگر، استفاده از توابع زیان مقاوم است که با مقاوم بودن معیار فاصله مرتبط است. دو تابع زیان بسیار معمول، توابع زیان مربع^{۱۰} یعنی $\mathcal{L}_{\text{SQR}}(e) = e^2$ و خطی^{۱۱} یعنی $\mathcal{L}_{\text{LIN}}(e) = |e|$ هستند. اگر مقدار خطا کمتر از ۱ باشد، زیان مربع، کوچک و اگر بیشتر از ۱ باشد، زیان مربع بزرگ می‌شود. در صورت وجود داده‌های پرت در مجموعه داده که موجب افزایش خطا می‌شوند؛ تابع زیان مربع، زیان زیادی را محاسبه می‌کند. تابع زیان خطی، نسبت به وجود داده‌های پرت مقاوم‌تر است؛ اما اشکال بزرگ این تابع، بزرگ بودن مقدار گرادیان آن برای مقدارهای کوچک خطاست. هدف ما از کارکردن بر روی این موضوع، آزمودن توابع زیان مختلف در مدل خوشه‌بندی فازی برای داده‌های فازی و یافتن توابع زیانی است که کارایی مدل در حضور داده‌های پرت را بهبود دهند. در این مقاله، استفاده از تابع زیان مقاوم لگاریتم کسینوس هذلولوی^{۱۲} و تابع زیان چندکی^{۱۳} پیشنهاد شده است. این دو تابع زیان خاص به دلیل دارا بودن ویژگی‌های جالبی انتخاب شده‌اند که در بخش ۴ به آن خواهیم پرداخت.

این مقاله شامل شش بخش است. در بخش ۲ توضیحاتی مقدماتی ایراد می‌شود. در بخش ۳ به تاریخچه‌ی موضوع مورد مطالعه اشاره می‌شود. ویژگی‌های توابع زیان پیشنهادی، در بخش ۴ تبیین می‌شود. در بخش ۵ شبیه‌سازی صورت گرفته معرفی و نتایج ارائه می‌شود. بخش ۶ نیز به تفسیر نتایج دست‌یافته اختصاص دارد.

۲ پیش‌نیازها

مجموعه داده‌ی فازی را می‌توان در یک ماتریس $N \times J$ ذخیره کرد که N ، تعداد داده‌ها و J ، تعداد ویژگی‌های فازی است.

تعریف ۱ (عدد فازی $L-R$). فرض کنید $L, R: [0, 1] \rightarrow [0, 1]$ دو تابع پیوسته و نزولی باشند که $L(0) = R(0) = 1$ و $L(1) = R(1) = 0$. مجموعه‌ی فازی $x_{ij}: \mathbb{R} \rightarrow [0, 1]$ که

$$x_{ij}(u_{ij}) = \begin{cases} L\left(\frac{c_{1ij} - u_{ij}}{l_{ij}}\right), & u_{ij} \leq c_{1ij}, \\ 1, & c_{1ij} \leq u_{ij} \leq c_{2ij}, \\ R\left(\frac{u_{ij} - c_{2ij}}{r_{ij}}\right), & u_{ij} \geq c_{2ij}, \end{cases}$$

نشان‌دهنده‌ی زامین ویژگی فازی $L-R$ است که در زامین داده مشاهده شده است که به صورت چهارتایی $x_{ij} = (c_{1ij}, c_{2ij}, l_{ij}, r_{ij})$ نمایش داده می‌شود. مولفه‌های c_{1ij} و c_{2ij} به ترتیب، مراکز

^۹Possibilistic

^{۱۰}Squared

^{۱۱}Linear

^{۱۲}Logarithm of hyperbolic cosine

^{۱۳}Quantile

چپ و راست و l_{ij} و r_{ij} به ترتیب، گستره‌های چپ و راست عدد فازی x_{ij} نامیده می‌شوند.

در حالت خاص، وقتی توابع L و R خطی‌اند، اعداد فازی ذوزنقه‌ای به دست می‌آیند. چنان‌چه علاوه بر شرط خطی بودن توابع L و R داشته باشیم $c_{1ij} = c_{2ij}$ ، آن‌گاه x_{ij} را عدد فازی مثلثی می‌نامیم.

۳ مروری بر کارهای پیشین

در سال ۱۹۹۶، یانگ^{۱۴} و کو^{۱۵} [۴]، تعریف [۱] فاصله‌ی

$$d^2(x_i, x_{i'}) = (c_i - c_{i'})^2 + ((c_i - \lambda l_i) - (c_{i'} - \lambda l_{i'}))^2 + ((c_i + \rho r_i) - (c_{i'} + \rho r_{i'}))^2, \quad (1)$$

را بین دو عدد فازی $L-R$ تک‌متغیره پیشنهاد کردند که در آن $\lambda = \int_0^1 L^{-1}(\omega) d\omega$ و $\rho = \int_0^1 R^{-1}(\omega) d\omega$ پارامترهایی هستند که شکل تابع عضویت را تعیین می‌کنند و برای هر مقدار λ و ρ یک تابع عضویت خاص داریم؛ به عنوان مثال، وقتی $\lambda = \rho = \frac{2}{3}$ ، تابع عضویت از نوع سهموی، وقتی $\lambda = \rho = \frac{1}{2}$ ، از نوع مثلثی و وقتی $\lambda = \rho = \frac{1}{3}$ ، از نوع ریشه‌ی دوم خواهد بود [۲]. مدل پیشنهادی یانگ و کو به شدت تحت تاثیر داده‌های پرت است و هنگامی که مجموعه داده، شامل تعداد زیادی کلاس نامتوازن باشد، نتایج خوشه‌بندی ضعیفی حاصل می‌شود. طی سه دهه‌ی اخیر، مدل‌های خوشه‌بندی فازی مقاومی از جمله مدل‌های یانگ و لیو^{۱۶} [۵]، بوتکیویچ^{۱۷} [۶]، هونگ^{۱۸} و یانگ [۷]، زرنندی^{۱۹} و رضایی^{۲۰} [۸] و دورسو^{۲۱} و جیوانی^{۲۲} [۹] برای داده‌های فازی ارائه شدند.

در سال ۲۰۲۰، دورسو و لسکی^{۲۳} [۱۰] با در نظر گرفتن روش افراز حول واسطه‌ها در یک چارچوب فازی، یک مدل خوشه‌بندی فازی مقاوم برای داده‌های فازی را بر اساس ترکیبی از توابع زیان و عملگر OWA (ordered weighted averaging) یاگر^{۲۴} ارائه کردند. آن‌ها با در نظر گرفتن نسخه‌ی چندمتغیره معیار فاصله‌ی پیشنهاد شده توسط یانگ و کو (۱) و توابع زیان مربع، خطی، هابر^{۲۵} یعنی

$$\mathcal{L}_{\text{HUB}}(e) = \begin{cases} \frac{e^2}{\delta^2}, & |e| \leq \delta, \\ |e|, & |e| > \delta, \end{cases}$$

^{۱۴}Yang

^{۱۵}Ko

^{۱۶}Liu

^{۱۷}Butkiewicz

^{۱۸}Hung

^{۱۹}Zarandi

^{۲۰}Razae

^{۲۱}Durso

^{۲۲}Giovanni

^{۲۳}Leski

^{۲۴}Yager

^{۲۵}Huber

۲-۴ تابع زیان چندکی

تابع زیان چندکی یعنی

$$\mathcal{L}_{\text{Quantile}}(e) = \begin{cases} \gamma e, & e \geq 0, \\ (\gamma - 1)e, & e < 0, \end{cases}$$

که $\gamma \in [0, 1]$ ، تعمیم تابع زیان خطی است (به ازای $\gamma = 0.5$). با افزایش مقدار γ ، شیب این تابع برای خطاهای منفی افزایش یافته و برای خطاهای مثبت کاهش می‌یابد.

۵ شبیه‌سازی

مدل (۲) با $\beta_i = 1$ و $m = 1.5$ و در نظر گرفتن توابع زیان مربع، خطی، هابر با $\delta = 5$ ، سیگموییدی با $\alpha = \beta = 2$ ، لگاریتمی، لگاریتم کسینوس هذلولوی و چندکی با $\gamma = 0.25$ در MATLAB R2021a پیاده‌سازی شده‌است. همه‌ی آزمایش‌ها بر روی یک ماشین (سیستم عامل: macOS، ظرفیت حافظه RAM: 8GB، پردازنده: 1.1GHz dual-core Intel Core i3) انجام شده‌است.

۱-۵ مجموعه داده‌های تولید شده

برای تولید مجموعه داده‌های فازی مثلثی، سه سناریوی مرکز پرت مثل شکل ۱ا، گستره‌های پرت مثل شکل ۱ب و مراکز و گستره‌های پرت مثل شکل ۱ج در نظر گرفته شده‌است [۹]. هر مجموعه داده در $\mathbb{R}^2$ شامل ۱۲۰ داده در ۳ کلاس دارای هم‌پوشانی با تعداد اعضای برابر و یک گروه ۲۰-عضوی از داده‌های پرت است. جدول ۱، سازوکار تولید مجموعه داده‌ها را بیان می‌کند.

جدول ۱: سازوکار تولید مجموعه داده‌های فازی مثلثی طبق سناریوهای مرکز پرت، گستره‌های پرت و مرکز و گستره‌های پرت.

مجموعه داده	موقعیت	ویژگی	مرکز	گستره‌ها
کلاس ۱	چپ	۱ و ۲	$U[1, 2]$	$U[0, 1]$
کلاس ۲	راست بالا	۱ و ۲	$U[2, 3]$	$U[0, 1]$
کلاس ۳	راست پایین	۱ و ۲	$U[2, 3]$	$U[0, 1]$
داده‌های پرت		۱ و ۲	$N(5, 2)$	$U[0, 1]$
کلاس ۱	درونی	۱ و ۲	$U[1, 2]$	$U[0, 1]$
کلاس ۲	میانی	۱ و ۲	$U[1, 2]$	$U[1, 2]$
کلاس ۳	بیرونی	۱ و ۲	$U[1, 2]$	$U[2, 3]$
داده‌های پرت		۱ و ۲	$U[1, 2]$	$N(5, 2)$
کلاس ۱	چپ	۱ و ۲	$U[1, 2]$	$U[0, 1]$
کلاس ۲	راست بالا	۱ و ۲	$U[2, 3]$	$U[1, 2]$
کلاس ۳	راست پایین	۱ و ۲	$U[2, 3]$	$U[2, 3]$
داده‌های پرت		۱ و ۲	$N(5, 2)$	$N(5, 2)$

که $\delta > 0$ ، سیگموییدی^{۲۶} یعنی $\mathcal{L}_{\text{SIG}}(e) = \frac{1}{1 + \exp(-\alpha(|e| - \beta))}$ و لگاریتمی^{۲۷} یعنی $\mathcal{L}_{\text{LOG}}(e) = \log(1 + e^2)$ معیار عدم تشابه

$$\mathcal{D}(\mathbf{x}_i, \mathbf{h}_g) = \mathcal{L}(\mathbf{c}_{1i} - \mathbf{c}_{1g}) + \mathcal{L}(\mathbf{c}_{2i} - \mathbf{c}_{2g}) +$$

$$\mathcal{L}[(\mathbf{c}_{1i} - \lambda \mathbf{l}_i) - (\mathbf{c}_{1g} - \lambda \mathbf{l}_g)] +$$

$$\mathcal{L}[(\mathbf{c}_{2i} + \rho \mathbf{r}_i) - (\mathbf{c}_{2g} + \rho \mathbf{r}_g)],$$

را بین داده i ام و سرخوشه‌ی g ام پیشنهاد و مدل

$$J(\mathbf{U}, \mathbf{H}) = \sum_{g=1}^c \sum_{i=1}^N \beta_i (u_{gi})^m \mathcal{D}(\mathbf{x}_i, \mathbf{h}_g), \quad (2)$$

را ارایه کردند که در آن u_{gi} درجه عضویت داده i ام در خوشه‌ی g ام، $m \in [1, \infty)$ پارامتر فازی کننده و $\beta_i \in [0, 1]$ یک پارامتر مناسب است که تاثیر داده‌های پرت را کاهش می‌دهد.

تابع زیان هابر، برای مقدارهای کوچک $|e|$ ، مربع و برای مقدارهای بزرگ آن، خطی است. در نتیجه نسبت به تابع زیان مربع، کمتر تحت تاثیر داده‌های پرت است. همچنین، بر خلاف تابع زیان خطی، مشتق‌پذیر بوده و کمینه‌سازی آن شدنی است. بنابراین تابع زیان هابر از مزایای هر دو تابع زیان مربع و خطی بهره‌مند و فاقد معایب آن‌هاست. وجه تمایز کار ما با دوسو و لسکی آن است که ما تاثیر استفاده از توابع زیان مختلف را با وزن‌دهی یکنواخت به داده‌ها ($\beta_i = 1$) مورد مطالعه قرار می‌دهیم.

۴ توابع زیان پیشنهادی

در این بخش، ویژگی‌های دو تابع زیان پیشنهاد شده برای بکارگیری در مدل خوشه‌بندی فازی برای داده‌های فازی را تبیین می‌کنیم.

۱-۴ تابع زیان لگاریتم کسینوس هذلولوی

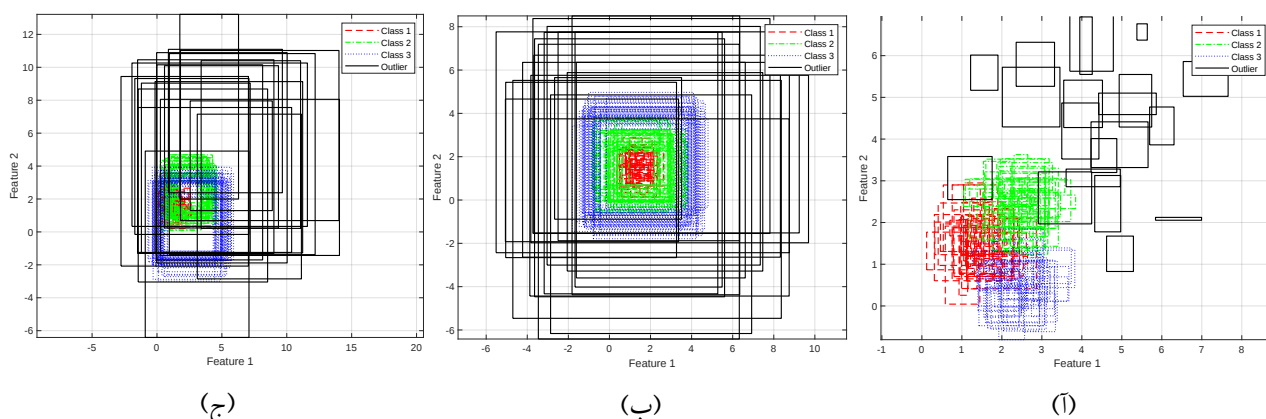
تابع زیان لگاریتم کسینوس هذلولوی یعنی

$$\mathcal{L}_{\text{Log-cosh}}(e) = \log(\cosh(e)),$$

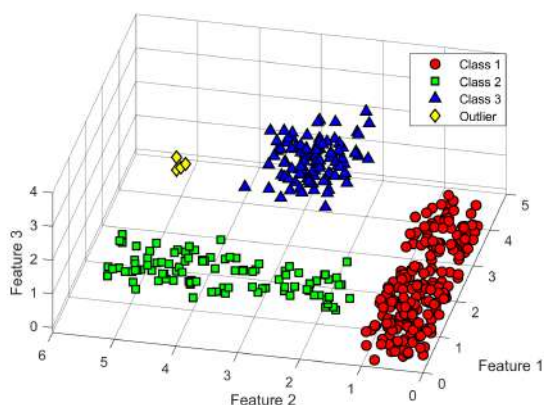
برای مقدارهای کوچک e ، تقریباً برابر $\frac{e^2}{2}$ و برای مقدارهای بزرگ آن، تقریباً برابر $|e| - \log(2)$ است. در نتیجه تحت تاثیر داده‌های پرت قرار نمی‌گیرد. وجود مشتق دوم، مزیت دیگر این تابع زیان است که در تابع زیان هابر از آن بی‌بهره بودیم. بنابراین با استفاده از تابع زیان لگاریتم کسینوس هذلولوی، ضمن برخورداری از مزایای تابع زیان هابر، استفاده از روش‌هایی که برای بهینه‌سازی نیاز به مشتق دوم دارند نیز میسر است.

²⁶Sigmoidal

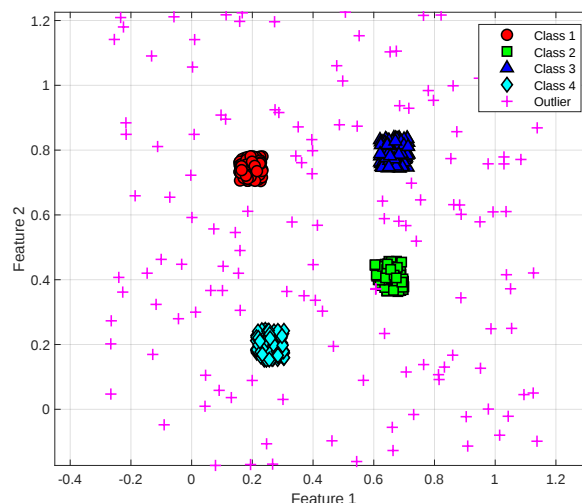
²⁷Logarithmic



شکل ۱: نمونه‌هایی از مجموعه داده‌های فازی مثلثی تولید شده طبق سناریوهای (آ) مرکز پرت، (ب) گستره‌های پرت و (ج) مرکز و گستره‌های پرت.



شکل ۳: مجموعه داده‌ی مصنوعی Lsun3D.



شکل ۲: مجموعه داده‌ی مصنوعی Zelnik4.

با برچسب دو، ۱۷ نمونه با برچسب سه، ۱۳ نمونه با برچسب پنج، ۹ نمونه با برچسب شش و ۲۹ نمونه با برچسب هفت. هیچ نمونه‌ای برای کلاس چهار در مجموعه داده وجود ندارد و کلاس شش پرت است. برای فازی‌سازی داده‌ها، هر سه مجموعه داده را به عنوان مرکز داده‌های فازی مثلثی در نظر گرفته‌ایم و فرض کرده‌ایم که گستره‌های این داده‌ها نیز از توزیع آماری $\mathcal{N}(1, 0.3)$ پیروی می‌کنند.

۳-۵ نتایج

از شاخص رند فازی ^{۲۸}frand [۱۴] و شاخص هولرمایر ^{۲۹}HUL [۱۵] که معیارهای اندازه‌گیری کلاس‌بندی غلط ^{۳۰} هستند، استفاده شده است و میانگین این شاخص‌ها و بازه اطمینان آن‌ها برای ۳۰ بار تکرار آزمایش در جدول ۲ گزارش شده است. هرچه مقدار به دست آمده برای این شاخص‌ها به یک نزدیک‌تر باشد، نتیجه‌ی خوشه‌بندی، رضایت‌بخش‌تر خواهد بود.

۲-۵ مجموعه داده‌های معیار

مجموعه داده‌ی مصنوعی Zelnik4 [۱۱] در $\mathbb{R}^2$ که در شکل ۲ نشان داده شده است، شامل ۶۲۲ نمونه در ۴ کلاس و یک گروه ۱۳۸-عضوی از داده‌های پرت است؛ تعداد ۱۱۱ نمونه با برچسب یک، ۱۱۴ نمونه با برچسب دو، ۱۵۰ نمونه با برچسب سه و ۱۰۹ نمونه با برچسب چهار.

مجموعه داده‌ی مصنوعی Lsun3D [۱۲، صفحه ۱۰۹] در $\mathbb{R}^3$ که در شکل ۳ نشان داده شده است، شامل ۴۰۴ نمونه در ۳ کلاس جدا از هم با اشکال هندسی متفاوت و یک گروه کوچک ۴-عضوی از داده‌های پرت است؛ تعداد ۲۰۰ نمونه با برچسب یک، ۱۰۰ نمونه با برچسب دو و ۱۰۰ نمونه با برچسب سه.

مجموعه داده‌ی واقعی Glass [۱۳] در $\mathbb{R}^9$ که برای تجسم، آن را با روش LDA (linear discriminant analysis) کاهش بعد داده و در شکل ۴ نمایش داده‌ایم، شامل ۲۱۴ نمونه در ۷ کلاس است. تعداد نمونه‌ها در هر کلاس نامتوازن است؛ تعداد ۷۰ نمونه با برچسب یک، ۷۶ نمونه

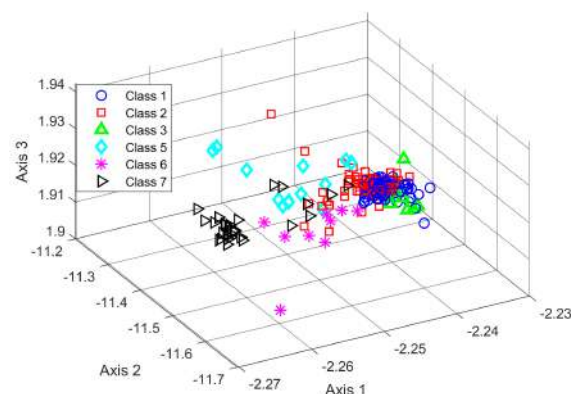
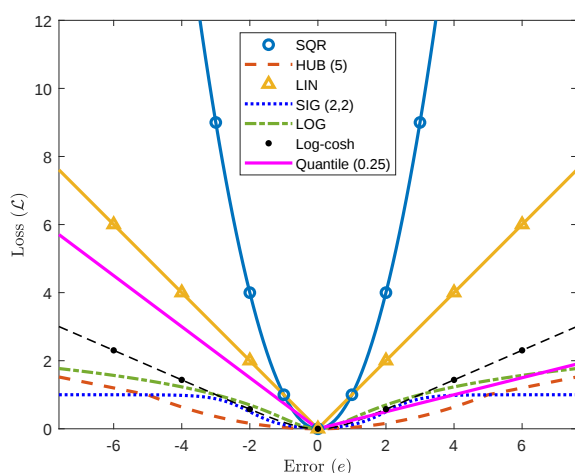
²⁸Fuzzy rand

²⁹Hullermeier

³⁰Misclassification

جدول ۲: نتایج بکارگیری توابع زیان مختلف در مدل خوشه‌بندی فازی برای داده‌های فازی.

میانگین	Glass		Lsun3D		Zelnik4		مرکز و گستره‌ها		گستره‌ها		مرکز		مجموعه داده	
	95% CI	میانگین	95% CI	میانگین	95% CI	میانگین	95% CI	میانگین	95% CI	میانگین	95% CI	میانگین	شاخص	تابع زیان
۰/۷۳ ۰/۵۸	[۰/۶۲, ۰/۶۴]	۰/۶۳	[۰/۷۳, ۰/۸۲]	۰/۷۸	[۰/۸۵, ۰/۹۳]	۰/۸۹	[۰/۶۴, ۰/۷۷]	۰/۷۰	[۰/۶۰, ۰/۶۸]	۰/۶۴	[۰/۶۸, ۰/۷۷]	۰/۷۲	frand	SQR
	[۰/۳۳, ۰/۳۷]	۰/۳۵	[۰/۶۱, ۰/۷۳]	۰/۶۷	[۰/۷۴, ۰/۸۶]	۰/۸۰	[۰/۴۷, ۰/۶۷]	۰/۵۷	[۰/۴۰, ۰/۵۱]	۰/۴۵	[۰/۵۵, ۰/۶۷]	۰/۶۱	HUL	
۰/۷۳ ۰/۵۹	[۰/۶۲, ۰/۶۴]	۰/۶۳	[۰/۶۹, ۰/۷۶]	۰/۷۳	[۰/۸۴, ۰/۹۲]	۰/۸۸	[۰/۷۰, ۰/۸۱]	۰/۷۵	[۰/۶۱, ۰/۶۷]	۰/۶۴	[۰/۷۳, ۰/۷۸]	۰/۷۶	frand	HUB
	[۰/۳۳, ۰/۳۷]	۰/۳۵	[۰/۵۷, ۰/۶۶]	۰/۶۱	[۰/۷۴, ۰/۸۵]	۰/۸۰	[۰/۵۸, ۰/۷۳]	۰/۶۵	[۰/۴۲, ۰/۵۱]	۰/۴۶	[۰/۶۲, ۰/۶۹]	۰/۶۵	HUL	
۰/۶۹ ۰/۴۷	[۰/۵۷, ۰/۵۹]	۰/۵۸	[۰/۶۷, ۰/۷۴]	۰/۷۰	[۰/۷۲, ۰/۷۹]	۰/۷۵	[۰/۷۱, ۰/۸۱]	۰/۷۶	[۰/۵۵, ۰/۵۸]	۰/۵۶	[۰/۷۴, ۰/۸۳]	۰/۷۸	frand	LIN
	[۰/۱۸, ۰/۲۲]	۰/۲۰	[۰/۵۰, ۰/۵۹]	۰/۵۴	[۰/۵۰, ۰/۶۰]	۰/۵۵	[۰/۵۵, ۰/۶۸]	۰/۶۱	[۰/۲۵, ۰/۳۰]	۰/۲۷	[۰/۵۹, ۰/۶۹]	۰/۶۴	HUL	
۰/۶۶ ۰/۴۰	[۰/۵۵, ۰/۵۸]	۰/۵۷	[۰/۷۳, ۰/۸۱]	۰/۷۷	[۰/۵۷, ۰/۶۰]	۰/۵۸	[۰/۷۰, ۰/۸۱]	۰/۷۵	[۰/۵۲, ۰/۵۶]	۰/۵۴	[۰/۶۷, ۰/۷۷]	۰/۷۲	frand	SIG
	[۰/۱۴, ۰/۱۹]	۰/۱۶	[۰/۶۰, ۰/۷۲]	۰/۶۶	[۰/۱۵, ۰/۱۷]	۰/۱۶	[۰/۵۶, ۰/۷۰]	۰/۶۳	[۰/۲۰, ۰/۲۶]	۰/۲۳	[۰/۴۸, ۰/۶۵]	۰/۵۷	HUL	
۰/۷۷ ۰/۶۲	[۰/۶۰, ۰/۶۲]	۰/۶۱	[۰/۷۲, ۰/۸۱]	۰/۷۷	[۰/۸۶, ۰/۹۳]	۰/۹۰	[۰/۷۸, ۰/۸۸]	۰/۸۳	[۰/۶۳, ۰/۶۷]	۰/۶۵	[۰/۸۰, ۰/۹۰]	۰/۸۵	frand	LOG
	[۰/۲۸, ۰/۳۲]	۰/۳۰	[۰/۶۰, ۰/۷۱]	۰/۶۵	[۰/۷۵, ۰/۸۷]	۰/۸۱	[۰/۶۸, ۰/۸۱]	۰/۷۵	[۰/۴۱, ۰/۴۷]	۰/۴۴	[۰/۷۰, ۰/۸۴]	۰/۷۷	HUL	
۰/۷۵ ۰/۶۰	[۰/۶۱, ۰/۶۳]	۰/۶۲	[۰/۷۱, ۰/۸۰]	۰/۷۶	[۰/۸۵, ۰/۹۳]	۰/۸۹	[۰/۶۸, ۰/۸۲]	۰/۷۵	[۰/۶۸, ۰/۷۲]	۰/۷۰	[۰/۷۳, ۰/۸۶]	۰/۸۰	frand	Log-cosh
	[۰/۲۹, ۰/۳۳]	۰/۳۱	[۰/۵۷, ۰/۶۸]	۰/۶۳	[۰/۷۵, ۰/۸۷]	۰/۸۱	[۰/۵۴, ۰/۷۴]	۰/۶۴	[۰/۵۰, ۰/۵۷]	۰/۵۴	[۰/۶۲, ۰/۷۸]	۰/۷۰	HUL	
۰/۷۰ ۰/۴۹	[۰/۵۷, ۰/۵۹]	۰/۵۸	[۰/۶۵, ۰/۷۳]	۰/۶۹	[۰/۷۳, ۰/۸۰]	۰/۷۶	[۰/۷۶, ۰/۸۵]	۰/۸۰	[۰/۵۷, ۰/۵۹]	۰/۵۸	[۰/۷۵, ۰/۸۳]	۰/۷۹	frand	Quantile
	[۰/۱۸, ۰/۲۱]	۰/۲۰	[۰/۴۶, ۰/۵۸]	۰/۵۲	[۰/۵۰, ۰/۶۲]	۰/۵۶	[۰/۶۲, ۰/۷۵]	۰/۶۸	[۰/۲۷, ۰/۳۲]	۰/۳۰	[۰/۶۳, ۰/۷۳]	۰/۶۸	HUL	



شکل ۴: مجموعه داده‌ی واقعی Glass که برای مصورسازی، ابعاد آن را از ۹ به ۳ کاهش داده‌ایم.

شکل ۵: مقایسه‌ی نمودارهای مربوط به توابع زیان مربع، هابر، خطی، سیگموئیدی، لگاریتمی، لگاریتم کسینوس هذلولوی و چندکی.

داشت.

نتایج شبیه‌سازی نیز موید آن است که توابع زیان خطی و چندکی، به دلیل شیب زیاد آن‌ها در نزدیکی $e = 0$ ، عملکرد ضعیفی دارند. تابع زیان مربع برای $e > 1$ ، زیان زیادی را محاسبه می‌کند و به همین دلیل، نسبت به وجود داده‌های پرت مقاوم نیست. توابع زیان لگاریتمی و لگاریتم کسینوس هذلولوی بهترین عملکردها و تابع زیان سیگموئیدی، ضعیف‌ترین عملکرد را دارا هستند.

مراجع

- [1] J. Bezdek. *Pattern Recognition with Fuzzy Objective Function Algorithms*. Plenum Press, New York, 1981.

۶ نتیجه‌گیری

بکارگیری توابع زیان لگاریتم کسینوس هذلولوی و چندکی در مدل خوشه‌بندی فازی برای داده‌های فازی، پیشنهاد و تاثیر استفاده از هریک از این دو تابع و نیز توابع زیان مربع، هابر، خطی، سیگموئیدی و لگاریتمی، در حضور داده‌های پرت بررسی شده است.

شکل ۵ نمودار مربوط به هریک از این توابع زیان را با یکدیگر مقایسه می‌کند. به طور کلی، آن دسته از توابع زیان که برای مقادیرهای نزدیک به صفر خطا، زیان کمی را محاسبه کنند و با افزایش خطا و دور شدن از نقطه‌ی $e = 0$ ، شیب تابع با آهنگ ملایمی افزایش یابد؛ و از سوی دیگر، در صورت وجود داده‌های پرت و افزایش چشمگیر مقدار خطا، زیان مربوطه را با بزرگ‌نمایی خطا محاسبه نکنند، کارایی بهتری خواهند

- [2] P. D’Urso, “Fuzzy clustering of fuzzy data,” *Advances in Fuzzy Clustering and Its Applications*, pp.155–192, 2007.
- [3] M. Sato and Y. Sato, “Fuzzy clustering model for fuzzy data,” *Proceedings of 1995 IEEE International Conference on Fuzzy Systems*, vol.4, pp.2123–2128, 1995.
- [4] M. Yang and C. Ko, “On a class of fuzzy c-numbers clustering procedures for fuzzy data,” *Fuzzy Sets and Systems*, vol.84, no.1, pp.49–60, 1996.
- [5] M. Yang and H. Liu, “Fuzzy clustering procedures for conical fuzzy vector data,” *Fuzzy Sets and Systems*, vol.106, no.2, pp.189–200, 1999.
- [6] B. Butkiewicz, “Robust fuzzy clustering with fuzzy data,” *Advances in Web Intelligence, Third International Atlantic Web Intelligence Conference, AWIC 2005, Lecture Notes in Computer Science, Springer, Berlin, Heidelberg*, vol.3528, pp.76–82, 2005.
- [7] W. Hung and M. Yang, “Fuzzy clustering on LR -type fuzzy numbers with an application in Taiwanese tea evaluation,” *Fuzzy Sets and Systems*, vol.150, no.3, pp.561–577, 2005.
- [8] M. Zarandi and Z. Razaee, “A fuzzy clustering model for fuzzy data with outliers,” *International Journal of Fuzzy System Applications (IJFSA)*, vol.1, no.2, pp.29–42, 2010.
- [9] P. D’Urso and L. De Giovanni, “Robust clustering of imprecise data,” *Chemometrics and Intelligent Laboratory Systems*, vol.136, pp.58–80, 2014.
- [10] P. D’Urso and J. Leski, “Fuzzy clustering of fuzzy data based on robust loss functions and ordered weighted averaging,” *Fuzzy Sets and Systems*, vol.389, pp.1–28, 2020.
- [11] L. Zelnik-Manor and P. Perona, “Self-tuning spectral clustering,” vol.17, 2004.
- [12] M. Thrun. *Projection-Based Clustering through Self-Organization and Swarm Intelligence*. Springer Vieweg, Wiesbaden, 2018.
- [13] D. Dua and C. Graff, “UCI machine learning repository,” 2017.
- [14] R. Campello, “A fuzzy extension of the Rand index and other related indexes for clustering and classification assessment,” *Pattern Recognition Letters*, vol.28, no.7, pp.833–841, 2007.
- [15] E. Hullermeier, M. Rifqi, S. Henzgen, and R. Senge, “Comparing fuzzy partitions: a generalization of the Rand index and related measures,” *IEEE Transactions on Fuzzy Systems*, vol.20, no.3, pp.546–556, 2012.

بهبود ترافیک شهری در شبکه‌های بین خودرویی با استفاده از رویکرد پروتکل وضعیت-اتصال و شبکه‌های عصبی

مسعود مؤمنی کلاگری^۱، مرجان کوچکی رفسنجانی^۲، حمیده فاطمی دخت^۳

^۱ دانشجوی کارشناسی ارشد، بخش علوم کامپیوتر، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، momeni@math.uk.ac.ir

^۲ دانشیار، بخش علوم کامپیوتر، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، kuchaki@uk.ac.ir

^۳ دکتری ریاضی کاربردی، بخش علوم کامپیوتر، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، h.fatemidokht@math.uk.ac.ir

چکیده: شبکه‌های بین خودرویی اخیراً به عنوان یک حوزه پژوهشی و کاربردی نمایان شده‌اند که کاربردهای بسیار زیادی در حوزه‌های مختلف از جمله مدیریت ترافیک و مسیر حرکت خودرو دارند. یکی از چالش‌ها در شبکه‌های بین خودرویی، تغییر مکرر در حرکت وسایل نقلیه می‌باشد که باعث ایجاد تغییرات پویا در اتصال ارتباطی و توپولوژی شبکه می‌شود. از این رو در سناریو زمان واقعی به دلیل این تغییرات، شناسایی تراکم ترافیک با شکست مواجه می‌شود. در این مقاله با بکارگیری شبکه‌های عصبی در شبکه‌های بین خودرویی و با الهام گرفتن از پروتکل مسیریابی وضعیت-اتصال به بررسی و پیش بینی وضعیت ترافیکی مسیرها پرداخته‌ایم. بمنظور بررسی کارایی روش پیشنهادی، از نرم‌افزارهای سومو و پایتون استفاده شده است. نتایج حاکی از آن است که با بروزرسانی اطلاعاتی از وضعیت ترافیکی مسیرها، روش پیشنهادی رهیافت مناسبی جهت پیش بینی ترافیک جاده‌ای می‌باشد.

کلمات کلیدی: شبکه‌های بین خودرویی، شبکه عصبی، پروتکل وضعیت-اتصال، ترافیک شهری، دقت پیش‌بینی.

سیستم‌های هوشمند حمل و نقل (ITS)^۶ هستند که به کنترل ترافیک و جلوگیری از تصادف و کاربردهای دیگر اختصاص یافته‌اند. خودروها، اطلاعات ترافیکی و امنیتی را با تجهیزات کنار جاده‌ای و خودروهای دیگر مبادله می‌کنند تا سفری مطمئن داشته باشند.

در تاریخچه مسیریابی وسایل نقلیه معمولاً فرض بر این است که هر وسیله نقلیه از ابتدا تا انتهای مسیر توسط یک راننده هدایت شود. حمل و نقل یکی از مسائل اصلی در جامعه است و مسیریابی وسایل نقلیه یکی از موضوعات مورد مطالعه در سال‌های اخیر بوده است. بطور معمول، تحقیقات در این زمینه بمنظور پیدا کردن مسیر برای وسیله نقلیه است که این مورد همان مسئله مسیریابی خودرو (VRP)^۷ نام دارد [2]. مسئله مسیریابی خودرو یک مسئله بهینه‌سازی ترکیبی^۸ و برنامه‌ریزی گسسته^۹ است که هدف آن

۱- مقدمه

شبکه‌های بین خودرویی (VANETs)^۱ نوعی از شبکه‌های سیار موردی هستند با این تفاوت که ارتباط میان وسایل نقلیه نزدیک به هم و همینطور بین وسایل نقلیه و تجهیزات ثابت که اغلب در کنار جاده‌ها نصب می‌شوند را فراهم می‌سازد. معماری این شبکه تشکیل شده از خودرو، واحد پردازنده (OBU)^۲، گذرگاه واحدهای کنار جاده (RSUs-G)^۳، مرکز کنترل حمل و نقل (TCC)^۴ و مرکز ابری (CC)^۵ [1]. شبکه‌های بین خودرویی مبنای

^۱ Vehicular Ad hoc Networks (VANETs)

^۲ On-Board Unit (OBU)

^۳ Road Side Units Gateway (RSUs-G)

^۴ Transportation Control Center (TCC)

^۵ Cloud Center (CC)

^۶ Intelligent Transportation Systems (ITS)

^۷ Vehicle Routing Problem (VRP)

^۸ Combinational optimizer

^۹ Discrete programming

سرویس‌دهی به مشتریان با استفاده از ناوگانی از وسایل نقلیه و همچنین کمینه کردن هزینه کل مسیر از مبدأ به مقصد است. مسئله مسیریابی خودرو برگرفته از مسئله فروشنده دوره‌گرد (TSP)^{۱۰} یا مسئله مسیریابی قوس (ARP)^{۱۱} می‌باشد؛ به طوری که در آن‌ها می‌توان راهی بهینه برای ساخت مسیر پیدا نمود به شیوه‌ای که مسیر، کوتاه‌ترین مسیر ممکن باشد.

یادگیری ماشین (ML)^{۱۲} فرآیندی است که به طور خودکار از مطالعه یا تجربه یاد می‌گیرد و بدون اینکه به صراحت برنامه‌ریزی شده باشد، عمل می‌کند. در فرآیند پردازش، یادگیری ماشین کارآمدتر، قابل اطمینان‌تر و مقرون به صرفه‌تر شده است. یادگیری ماشین، مدل‌ها را با تجزیه و تحلیل داده‌های پیچیده‌تر به طور خودکار، سریع‌تر و دقیق‌تر تولید می‌کند و عمدتاً به یادگیری نظارت شده^{۱۳}، یادگیری بدون نظارت^{۱۴}، یادگیری نیمه نظارت شده^{۱۵} و یادگیری تقویتی^{۱۶} طبقه‌بندی می‌شود. قدرت یادگیری ماشین در توانایی آن‌ها برای ارائه راه‌حل‌های تعمیم‌یافته از طریق یک معماری که می‌تواند عملکرد آن را بهبود بخشد، نهفته است. به دلیل ماهیت میان‌رشته‌ای، آن نقش محوری در زمینه‌های مختلفی از جمله مهندسی، پزشکی و محاسبات دارد. عملکرد یک سیستم در استفاده از یادگیری ماشین بهبود می‌یابد و مداخله انسان یا برنامه‌ریزی مجدد سیستم نیز محدود می‌گردد [3].

یکی از موضوعات اساسی و کاربردی در تردهای درون‌شهری رفع معضل ترافیک می‌باشد. یکی از عوامل اصلی ترافیک، عدم وجود اطلاعات درست نسبت به حال و آینده از وضعیت جاده‌ها برای راننده است. اگر بتوان میزان ارسال اطلاعات مناسب به هر راننده را به منظور بالا بردن آگاهی محیطی (مسیرهای منتهی به مقصد) در زمان حال و آینده افزایش داد، آنگاه راننده می‌تواند در تشخیص انتخاب مسیر مناسب بهتر عمل کند. در نتیجه بار ترافیکی مسیرهای موجود، تحت الشعاع قرار نمی‌گیرد.

توپولوژی پویا و متراکم شبکه‌های ویژه خودرویی، که از تحرک زیاد و تراکم بالای گره‌ها حاصل می‌شود، چالش‌های بسیاری را برای برقراری ارتباط و مسیریابی ایجاد کرده است. از این رو در بدست آوردن تراکم ترافیک در سناریو زمان واقعی به دلیل تغییر سریع در زمان تراکم ترافیک با شکست مواجه می‌شود. بنابراین لازم است نگرشی ایجاد شود که به طور مؤثر تراکم ترافیک را تحت نظر داشته باشد یا اطلاعاتی بیشتر نسبت به مسیر مورد تقاضا را فراهم کند که موجب بهبود ترافیک درون شهری خواهد شد.

به عبارتی دیگر، از آنجایی که شبکه‌های بین خودرویی اطلاعات کنونی مسیرها را در اختیار خودرو قرار می‌دهند، تصمیم‌گیری برای تعیین وضعیت یک مسیر برای لحظاتی بعد مشکل است. از این رو ما نیاز به یک پیش‌بینی از وضعیت مسیر داریم که بتوان با استفاده از آن یک تصمیمی مناسب اتخاذ کرد. بر این اساس با توجه به حجم بالای داده‌ها در شبکه‌های بین خودرویی، بکارگیری شبکه‌های عصبی جهت پیش‌بینی با دقت مناسب و تصمیم‌گیری،

امری لازم و مناسب است. با بکارگیری شبکه‌های بین خودرویی در کنار شبکه‌های عصبی، ضریب اطمینان نسبت به پیش‌بینی وضعیت ترافیکی مسیرها، بالاتر می‌رود.

در ادامه این مقاله، در بخش ۲ شرح مختصری از موارد مورد نیاز جهت تفهیم بهتر موضوع ارائه شده است. در بخش ۳ پژوهش‌های صورت گرفته در راستای موضوع مقاله به طور مختصر توضیح داده شده است. در بخش ۴ روش پیشنهادی ارائه شده است. در بخش ۵ روند شبیه‌سازی شهر هوشمند و محیط شبیه‌سازی به همراه جزئیات مربوط به شبیه‌سازی روش پیشنهادی بیان شده است؛ در ادامه‌ی این بخش ارزیابی نتایج آورده شده به‌طوری‌که روش پیشنهادی با دو روش دیگر مقایسه گردیده است و بالاخره، در بخش ۶، نتیجه‌گیری مقاله بیان شده است.

۲- پیش‌نیازها

با توجه به گسترش شهرنشینی و ترافیک‌های زیاد، فرآیند پیش‌بینی ترافیک، دارای اهمیت بالایی است. این امکان با استفاده از اطلاعات دریافتی در شبکه‌های بین خودرویی و به‌کارگیری آن در شبکه‌های عصبی جهت آگاهی از وضعیت مسیرهای منتهی از مبدأ به مقصد، اتفاق می‌افتد. با توجه به اینکه روش اتخاذ شده مبتنی بر طبقه‌بندی می‌باشد، جهت پیش‌بینی و ارائه وضعیت جاده‌ها با یک درصد مناسب، از تئوری شبکه‌های عصبی در یادگیری ماشین استفاده شده است.

شبکه‌های عصبی مصنوعی [4] ساختاری قدرتمند برای فرآیندهای کنترلی و طبقه‌بندی داده‌ها دارند. وضعیت یک شبکه عصبی مصنوعی مرحله به مرحله بر اساس اطلاعات ورودی که دریافت می‌کند تغییر کرده و هر چه میزان این اطلاعات بیشتر شود شبکه خطای خود را بیشتر کاهش می‌دهد.

یکی از مشهورترین انواع شبکه‌های عصبی، شبکه عصبی پرسپترون چندلایه (MLP)^{۱۷} است که به‌طور موفقیت‌آمیزی در بسیاری از کاربردها از جمله طبقه‌بندی داده، مورد استفاده قرار گرفته است. هنگام کار با شبکه‌های عصبی MLP با دو مسئله رو به رو هستیم: انتخاب معماری مناسب و انتخاب الگوریتم آموزشی مناسب. معماری مناسب به معنی انتخاب بهینه تعداد لایه‌ها، تعداد نورون‌ها در هر لایه و نوع تابع تحریک هر نورون می‌باشد. معماری بهینه شبکه‌های عصبی مبتنی بر مجموعه داده‌ها و ویژگی‌های آن‌ها است.

پروتکل مسیریابی وضعیت-اتصال^{۱۸} از دو جزء کلی مسیریاب^{۱۹} و جدول مسیریابی^{۲۰} تشکیل می‌شود [5]:

مسیریاب با دریافت مبدأ و مقصد و مراجعه به جدول مسیریابی، با توجه به اطلاعات ثبت شده، یک مسیر را انتخاب کرده و به سمت مقصد حرکت می‌کند. جدول مسیریابی، پایگاه داده‌ای است که مسیرها مانند نقشه در آن ذخیره می‌شود. یعنی اطلاعات هر مسیر با توجه به شاخص‌های در نظر گرفته شده در این جدول قرار دارند. جدول مسیریابی به شکل یک فایل در اختیار مسیریاب قرار می‌گیرد.

^{۱۰} Travelling Salesman Problem (TSP)

^{۱۱} Arc Routing Problem (ARP)

^{۱۲} Machine Learning (ML)

^{۱۳} Supervised learning

^{۱۴} Unsupervised learning

^{۱۵} Semi-supervised learning

^{۱۶} Reinforcement learning

^{۱۷} Multi-Layer Perceptron (MLP)

^{۱۸} Link-State

^{۱۹} Router

^{۲۰} Routing table

۳- کارهای مرتبط

تائو^{۲۱} و همکارانش [6]، یک رویکرد یادگیری عمیق مبتنی بر تأخیر برای پیش‌بینی حجم ترافیک شهری ارائه کرده‌اند، تأخیر در سفر یکی دیگر از مشکلاتی است که به طور گسترده نادیده گرفته شده است، اما می‌تواند به طور قابل توجهی بر نتیجه پیش‌بینی تأثیر بگذارد. به طور خاص، وسایل نقلیه برای جابجایی از مکانی به مکان دیگر به زمان نیاز دارند و به این زمان تأخیر سفر^{۲۲} می‌گویند. برای افزایش بیشتر عملکرد پیش‌بینی در سناریوی شهری، آنها یک چارچوب یادگیری عمیق مبتنی بر تأخیر را برای بهبود دقت پیش‌بینی کوتاه مدت جریان ترافیک پیشنهاد می‌کنند، که در آن تأخیر سفر به شکل یک ماتریس وزن‌دار به عنوان یک ورودی چند متغیره انباشته^{۲۳} شبکه عصبی بازگشتی (RNN)^{۲۴} ثبت می‌شود. ورودی چند متغیره باعث می‌شود این رویکرد توانایی استخراج قوی‌تری برای تشخیص روابط فضایی^{۲۵} داشته باشد و ساختار انباشته^{۲۶} منجر به فرآیند یادگیری الگوی دقیق‌تری می‌شود. نتایج نشان می‌دهد که رویکرد آنها دقیق و قابل اعتماد است.

خان^{۲۷} و همکارانش [7] یک پروتکل مسیریابی مبتنی بر ترافیک را در دو بخش، ۱) اطلاعات ترافیکی وسایل نقلیه در زمان واقعی برای انتخاب مسیر، امکان محاسبه درجه اتصال پیش‌بینی شده در بخش‌های مختلف و ۲) یک روش انتقال جدید بر اساس اطلاعات جغرافیایی بسته‌ها را از گره مبدأ به مقصد، ارائه می‌دهند. معیار جدید، چگالی‌های خودرو^{۲۸} که از طریق پارامترهایی چون شناسه گره‌های مبدأ و مقصد، موقعیت بعدی گره، سرگروه‌های با موقعیت جاری، میزان اتصال مورد انتظار، بسته کشف مسیر، گره میانی و بخشی که گره در آن قرار دارد محاسبه خواهد شد را در نظر گرفته و اتصال در هر بخش تخمین زده شده است و در نتیجه اتصال گره‌ها و نسبت تحویل داده برای ارسال بسته‌ها محاسبه شده است. نتایج شبیه‌سازی نشان می‌دهد که این روش از پروتکل‌های مسیریابی رقیب مؤثرتر است.

خاتری^{۲۹} و همکارانش [8] یک مطالعه موردی با توجه موضوعات مربوط به ایمنی، ارتباطات و مسائل مربوط به ترافیک در شبکه‌های بین خودرویی و اجرای آنها در امکان سنجی و بررسی چگونگی غلبه بر این مسائل را مورد بررسی قرار داده‌اند. آنها روش‌هایی را که از الگوریتم‌های یادگیری ماشین برای بالا بردن کارایی شبکه‌های بین خودرویی استفاده کرده بودند را بررسی کردند و همچنین گیلانی^{۳۰} و همکارانش [9]، بای^{۳۱} و همکارانش [10] نیز مطالعات مروری بر روی این موضوعات داشتند.

سو^{۳۲} و همکارانش [11] الگوریتمی بر اساس هسته K-نزدیکترین همسایه^{۳۳} برای پیش‌بینی وضعیت‌های ترافیکی جاده‌ای در سری‌های زمانی ارائه دادند. حالت‌های ترافیکی جاده‌ای یک الگوی منظمی در سری‌های زمانی دارد. بدین ترتیب از طریق تطبیق مؤثر بین وضعیت‌های ترافیک جاده‌ای قبلی و فعلی، می‌توان وضعیت ترافیکی جاده‌ای آینده را به طور مؤثر پیش‌بینی کرد. پیش‌بینی حالات ترافیکی جاده‌ای بر اساس الگوریتم ارائه شده، انجام گرفته است. بر این اساس، عملکرد این روش با شبکه‌های عصبی (صرفاً ابزاری جهت پیش‌بینی) و مدل سری زمانی ARIMA^{۳۴} ارزیابی و نتایج حاکی از برتری این روش نسبت به دو روش گفته شده در تعیین وضعیت ترافیکی جاده‌ای، می‌باشد. از این رویکرد در شهر پکن استفاده شده است که نتایج نهایی آنها نشان داد که وضعیت ترافیکی جاده‌ای رویکرد ارائه شده می‌تواند به سطح بالایی از دقت برسد.

دو^{۳۵} و همکارانش [12] بر روی یک روش تخمین ترافیک کارآمد، جهت فراهم آوردن اطلاعات ترافیکی به خودروها و مرکز نظارت ترافیک متمرکز شده‌اند. داده‌های ترافیک توسط VANET جمع‌آوری می‌شوند تا از هزینه زیاد استقرار و نگهداری ردیاب ترافیک در جاده جلوگیری شود. از آنجا که وسایل نقلیه کاوشگر^{۳۶} ممکن است تمام جاده‌ها را پوشش ندهند، بسیاری از عناصر در محیط ماتریس ترافیک خالی هستند. در این مقاله، به جهت پوشش داده‌های گم شده الگوریتمی ارائه شده که در آن تداوم موقتی و حدود داده‌ها را مدنظر قرار می‌دهد. بنابراین عناصر خالی در ماتریس ترافیک را برآورد کرده و پر می‌کند. دو مقاله‌ای آخر ذکر شده در این بخش با روش پیشنهادی ارائه شده در بخش بعد، به جهت پیش‌بینی وضعیت ترافیک مقایسه می‌گردند.

۴- روش پیشنهادی

در این مقاله، به جهت رفع چالش ترافیک، روشی ارائه داده‌ایم که از شبکه‌های عصبی در شبکه‌های بین خودرویی استفاده می‌شود. ترکیب این دو با هم باعث می‌گردد پیش‌بینی وضعیت ترافیکی مسیرها دقت بالاتری داشته باشد. رویکرد این روش برگرفته از پروتکل مسیریابی وضعیت-اتصال می‌باشد که برای پیش‌بینی وضعیت ترافیکی، این پروتکل را الگو قرار داده و معادل‌سازی می‌کنیم.

بدین منظور جدا از سه فاکتور ترافیک، سرعت و فاصله، فاکتور مربوط به عرض خیابان نیز در نظر گرفته می‌شود. اطلاعات مربوط به وضعیت ترافیک، سرعت خودرو، فاصله تا مقصد و عرض خیابان برای مسیرهای مختلف به وسیله شبکه‌های بین خودرویی تعیین می‌شود که می‌توان آن را همان جدول مسیریابی در نظر گرفت. بنابراین پایگاه داده‌ای^{۳۷} به وسیله شبکه‌های بین خودرویی تشکیل می‌شود. اطلاعات موجود در این جدول مسیریابی، در اختیار شبکه‌های عصبی قرار می‌گیرد. اگر بخواهیم معادل شبکه‌های عصبی در

^{۲۱} Tao

^{۲۲} Travel delay

^{۲۳} Multivariate input stacked

^{۲۴} Recurrent Neural Network (RNN)

^{۲۵} Spatial relationships capture

^{۲۶} The stacked structure

^{۲۷} Khan

^{۲۸} Vehicle densities

^{۲۹} Khatri

^{۳۰} Gillani

^{۳۱} Bai

^{۳۲} Xu

^{۳۳} Kernel-KNN

^{۳۴} AutoRegressive Integrated Moving Average (ARIMA)

^{۳۵} Du

^{۳۶} Probe vehicle

^{۳۷} Database

یک مبدأ به نام میدان امام و یک مقصد به نام میدان امام حسین انتخاب کرده و مسیرها را مشخص می‌کنیم. در این مقاله ۷ مسیر را مورد بررسی داده‌ایم (شکل ۱ و جدول ۲).



شکل ۱: شماتیک مسیرها

جدول ۲: تعیین نقاط هفت مسیر

مسیر	نقاط
مسیر شماره ۱	ABCDEFGHGI
مسیر شماره ۲	ABCNEFGHI
مسیر شماره ۳	ABCNFGHI
مسیر شماره ۴	ABCDLMHI
مسیر شماره ۵	ABKLMHI
مسیر شماره ۶	AJMHI
مسیر شماره ۷	AJKLMHI

همان‌طور که گفته شد از رویکرد شبکه‌های عصبی به عنوان مسیریاب، برای تشخیص کلاس وضعیت امور ترافیکی جاده‌ای مسیرها که عبارتند از: عدم وجود ترافیک و وجود ترافیک، کمک گرفته شده است.

متغیرهای تعیین شده برای ایجاد مدل پیش‌بینی به دو دسته متغیر هدف و متغیرهای پیشگو (مستقل) دسته‌بندی شده‌اند. متغیر هدف، وجود و عدم وجود ترافیک و سایر متغیرها شامل سرعت، حجم خودروها در مسیر، مسافت مسیر، عرض مسیر و زمان سفر به عنوان متغیر پیشگو مورد استفاده قرار گرفته است.

برای ساخت مدل شبکه‌های عصبی، داده‌ها به دو بخش آموزش و آزمایش تقسیم شده‌اند. به منظور آموزش شبکه در حالت یادگیری نظارت شده معمولاً از الگوریتم پس انتشار خطا^{۳۹} استفاده می‌شود.

شبکه‌های عصبی معمولاً زمان زیادی برای آموزش نیاز دارند و برای کار کردن با آن‌ها باید به مواردی از جمله معماری مطلوب شبکه که به صورت تجربی انتخاب می‌شود و همچنین حدس مهندسی در خصوص تنظیمات پارامترهای مدل نظیر نرخ یادگیری، نرخ منظم‌ساز و ... توجه شود.

از آنجایی که تفسیر خروجی‌های مربوط به وزن‌های شبکه‌های عصبی امری دشوار می‌باشد از طریق روش اتصال وزن‌های لایه‌ها، اطلاعات وزن‌ها را بدست آورده‌ایم. پارامترهای پیش‌فرض برای شبکه عصبی در جدول ۳ ارائه شده‌اند.

پروتکل‌های وضعیت-اتصال را بیان کنیم، می‌توان آن را به عنوان مسیریاب در نظر گرفت. با وجود اطلاعات هر مسیر در جدول مسیریابی، شبکه‌های عصبی به پیش‌بینی وضعیت ترافیکی مسیرها می‌پردازند. الگوریتم مربوط به این روش در ادامه آورده شده است.

الگوریتم پیشنهادی

Input:

- Routing table, X ; the target value Y (i.e class label);
- Static factor table (S)
- Number of rout (R)

Output: the forecasting output ($\hat{Y}$).

Data separation:

- Routing table (X) is separated as X_{tr} and X_{tx}

Learning:

- Learning by doing MLP (W)
- Combining VANET and ML methods part (router):
- Determine the status vector of training features with apply Link State methods for each routers (W) (SX_{tr})
- Determine the status vector of testing features with apply Link State methods for each routers (W) (SX_{tx})

Loop:

- **For** $i \in \{1, \dots, R\}$ **do**
Absolute error:
Calculate the distance D_i between $SX_{tr}[i]$ and $SX_{tx}[i]$
End for

Sort the distance D from smallest to largest value;

Set the forecasting the class label ($\hat{Y}$) based on frequent class of processed index for determine of rout.

Update: Append $\hat{Y}$ to routing table

return $\hat{Y}$

۵- شبیه‌سازی و ارزیابی نتایج

از آنجایی که در حال حاضر هیچ شهر هوشمندی مبتنی بر شبکه‌های بین خودرویی در ایران وجود ندارد، ابتدا به شبیه‌سازی شهر هوشمند می‌پردازیم و با استفاده از این شبیه‌سازی یک پایگاه اطلاعاتی (داده) از وضعیت امور جاده‌ای مربوط به شهرستان ساری واقع در استان مازندران، ایجاد می‌کنیم. سپس براساس این پایگاه اطلاعاتی یک پیش‌بینی نسبت به وضعیت ترافیکی جاده‌ها، با رویکرد گفته شده ارائه می‌شود.

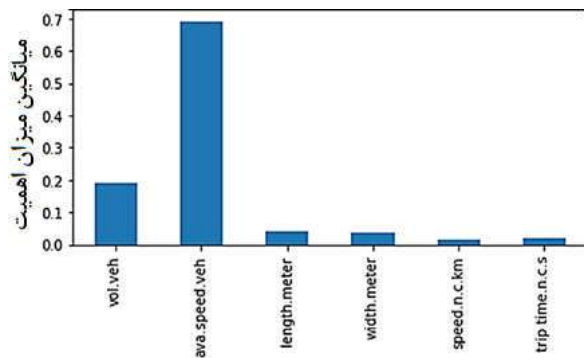
شبیه‌سازی در محیط نرم‌افزاری سومو^{۳۸} صورت گرفته و وضعیت تردهای لحظه‌ای بر پایه شبکه‌های بین خودرویی خروجی گرفته شده است. همچنین تمام مدل‌های ساخته شده در محیط برنامه‌نویسی پایتون صورت گرفته است. در جدول ۱ پارامترهای شبیه‌سازی ارائه شده‌اند.

جدول ۱: پارامترهای شبیه‌سازی

پارامتر	مقدار
کمترین عرض جغرافیایی	۳۶/۵۴۸۸
بیشترین عرض جغرافیایی	۳۶/۵۶۹۶
کمترین طول جغرافیایی	۵۳/۰۳۴
بیشترین طول جغرافیایی	۵۳/۰۷۷۶
تعداد خودرو	۸۶۴۰۰
زمان شبیه‌سازی	۲۴ ساعت
کل داده‌ی خام	۲۶۲۴۷۱۲۲
داده‌ی مورد استفاده (برای هر مسیر)	۴۶۳۲۶۷

^{۳۹} Backpropagation of error

^{۳۸} Sumo



شکل ۲: اهمیت ویژگی‌های استفاده شده به جهت پیش‌بینی کلاس وضعیت

۶- نتیجه‌گیری

با توجه به نتایج ارائه شده، می‌توان گفت صرفاً با ارائه وضعیت فعلی وضعیت ترافیکی توسط شبکه‌های بین خودرویی، پیش‌بینی ترافیکی و تصمیم‌گیری، عملی با ریسک و خطا می‌باشد؛ چرا که امکان تغییر وضعیت ترافیکی مسیرها وجود دارد. با اتخاذ یک تصمیم و استراتژی اشتباه، می‌توان تبعاتی مانند افزایش ترافیک، مصرف سوخت بیش از انتظار، افزایش زمان سفر و حتی استهلاک خودرویی را مخصوصاً در شهرهای بزرگ، انتظار داشت. در این مقاله برای پیش‌بینی وضعیت ترافیکی مسیرها از ترکیب کردن شبکه‌های عصبی و شبکه‌های بین خودرویی با الهام گرفتن از پروتکل وضعیت-اتصال، استفاده شده است. نتایج شبیه‌سازی نشان می‌دهد روش پیشنهادی با استفاده از شبکه عصبی در مقایسه با دو روش بررسی شده دیگر، برای مجموعه داده‌های آزمایشی، دقت پیش‌بینی نسبت به وضعیت ترافیکی را افزایش می‌دهد.

مراجع

- [1] Senouci O., Aliouat Z., Harous S., "MCA-V2I: A Multi-hop Clustering Approach over Vehicle-to-Internet Communication for Improving VANETs performances", Future Generation Computer Systems, Vol. 96, pp. 309-323, 2019.
- [2] Domínguez-Martín B., Rodríguez-Martín I., Salazar-González J.-J., "The Driver and Vehicle Routing Problem", Computers & Operations Research, Vol. 92, pp. 56-64, 2018.
- [3] Kumar D.P., Amgoth T., Annavarapu C.S.R., "Machine Learning Algorithms for Wireless Sensor Networks: A Survey", Information Fusion, Vol. 49, pp. 1-25, 2019.
- [4] Yegnanarayana B., Artificial Neural Networks. PHI Learning Pvt. Ltd., 2009.
- [5] Yi J., Adnane A., David S., Parrein B., "Multipath Optimized Link State Routing for Mobile Ad Hoc Networks", Ad Hoc Networks, Vol. 9, No. 1, pp. 28-47, 2011.
- [6] Tao Y., Sun P., Boukerche A., "A Delay-Based Deep Learning Approach for Urban Traffic Volume Prediction", in IEEE International Conference on Communications (ICC), 2020.
- [7] Khan S., Alam M., Fränzle M., Müllner N., Chen Y., "A Traffic Aware Segment-based Routing Protocol for VANETs in Urban Scenarios", Computers & Electrical Engineering, Vol. 68, pp. 447-462, 2018.

جدول ۳: مشخصات شبکه عصبی MLP

۱	تعداد لایه پنهان
۵۰	تعداد نورون
خطی	تابع فعال‌سازی
۰/۰۱	نرخ یادگیری
۰/۰۰۰۱	نرخ تنظیم

تعداد کل خودروهای تردد داده شده در این شبیه‌سازی ۸۶۴۰۰ می‌باشد. تعداد رکوردهای ثبت شده از حجم ترافیک در زمان‌های مختلف برای هر هفت مسیر ۶۶۱۸۱ و تعداد کل رکوردها در مجموع ۴۶۳۲۶۷ می‌باشد. به جهت مدل‌سازی و بررسی دقت، ۹۰ درصد داده‌ها برای آموزش و ۱۰ درصد برای آزمایش در نظر گرفته شده است. نتایج دقت پیش‌بینی و دقت تشخیص اندازه سرعت در هر مسیر براساس توابع زیان‌های معروف جذر میانگین مربعات خطا (RMSD) و میانگین قدر مطلق خطاها (MAE) بررسی شده که این معیارها به صورت زیر تعریف می‌شوند:

$$RMSD = \frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2 \quad (۱)$$

$$MAE = \frac{1}{N} \sum_{i=1}^N |y_i - \hat{y}_i| \quad (۲)$$

که در آنها y و $\hat{y}$ به ترتیب کلاس وضعیت و پیش‌بینی کلاس وضعیت می‌باشند. معیار MAE به نوعی نشان دهنده متوسط اندازه خطاها و معیار RMSD بیانگر پراکندگی خطاها حول صفر می‌باشند، بنابراین هرچه مقدار این دو معیار کمتر باشد بیانگر دقت بیشتر پیش‌بینی‌ها است. معمولاً مقدار MAE کوچکتر از مقدار RMSD است.

جدول ۴: مقایسه نتایج در تعیین کلاس وضعیت ترافیکی مسیرها

روش‌ها	متوسط دقت پیش‌بینی	انحراف از معیار
دو و همکارانش [12]	۰/۹۴۵۶۳۵۷۸۴	۰/۰۰۵۰۱۵۴۳۲
سو و همکارانش [11]	۰/۹۷۳۷۹۷۹۲۷۴	۰/۰۰۴۵۷۱۴۲
روش پیشنهادی	۰/۹۷۵۹۵۳۳۳۱	۰/۰۰۴۳۴۱۹۷۲

با توجه به جدول ۴، روش پیشنهادی با استفاده از شبکه عصبی برای مجموعه داده‌های آزمایشی دارای دقت بالاتری نسبت به دو روش دیگر هست. بنابراین می‌توان نتیجه گرفت برای پیش‌بینی کلاس درست از وضعیت هر مسیر با استفاده از متغیرهای پیشگویانه‌ی ذکر شده، بکارگیری شبکه‌های عصبی مصنوعی در شبکه‌های بین خودرویی می‌تواند رهیافت مناسبی باشد. همچنین اهمیت ویژگی‌های استفاده شده در این مقاله جهت پیش‌بینی در شکل ۲ قابل مشاهده است، که مشخصاً سهم سرعت و حجم خودروها بیشتر از بقیه‌ی فاکتورها می‌باشد.

^{۴۰} Root Mean Square Deviation (RMSD)

^{۴۱} Mean Absolute Error (MAE)

- [8] Khatri S., Vachhani H., Shah S., Bhatia J., Chaturvedi M., Tanwar S., Kumar N., *"Machine Learning Models and Techniques for VANET based Traffic Management: Implementation Issues And Challenges"*, Peer-to-Peer Networking and Applications, Vol. 14, No. 3, pp. 1778-1805, 2021.
- [9] Gillani M., Niaz H.A., Tayyab M., *"Role of Machine Learning in WSN and VANETs"*, International Journal of Electrical and Computer Engineering Research, Vol. 1, No. 1, pp. 15-20, 2021.
- [10] Bai R., Chen X., Chen Z.-L., Cui T., Gong S., He W., Jiang X., Jin H., Jin J., Kendall G., *"Analytics and Machine Learning in Vehicle Routing Research"*, arXiv preprint arXiv:2102.10012, 2021.
- [11] Xu D., Wang Y., Peng P., Beilun S., Deng Z., Guo H., *"Real-Time Road Traffic State Prediction Based on Kernel-KNN"*, Transportmetrica A: Transport Science, Vol. 16, No. 1, pp. 104-118, 2020.
- [12] Du R., Chen C., Yang B., Guan X., *"VANET based Traffic Estimation: A Matrix Completion Approach"* in IEEE Global Communications Conference (GLOBECOM), 2013.

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران ۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

تابعی اکتشافی برای بهبود دقت پیش‌بینی برنامه‌های جهش‌یافته آشکار کننده خطا

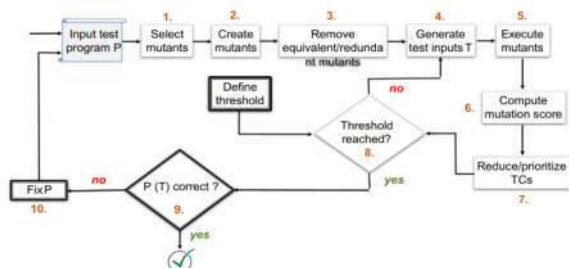
طه رستمی^۱، سعید جلیلی^۲

^۱ گروه مهندسی کامپیوتر، دانشگاه تربیت مدرس، تهران، taha.rostami@modares.ac.ir

^۲ گروه مهندسی کامپیوتر، دانشگاه تربیت مدرس، تهران، sjalili@modares.ac.ir

چکیده: آزمون جهش روشی قدرتمند است که در آزمون نرم‌افزار استفاده می‌شود. تعداد زیاد برنامه‌های جهش‌یافته، مقیاس پذیری و تعداد زیاد برنامه‌های جهش‌یافته غیر مرتبط با خطای واقعی اعتبار آن را تهدید می‌کند. برای مواجهه شدن با تهدیدهای گفته شده، اخیراً رویکردی مبتنی بر یادگیری ماشین برای شناسایی برنامه‌های جهش‌یافته آشکار کننده خطا پیشنهاد شده است. در این پژوهش، یک تابع اکتشافی برای کاهش داده‌های مجموعه آموزش پیشنهاد شده است که در ترکیب با سایر روش‌ها توانسته **AUC** را روی دو مجموعه داده **Codeflaws** و **CoREBench** حدود ۲٪ بهبود دهد.

کلمات کلیدی: آزمون جهش، آزمون نرم‌افزار، برنامه‌های جهش‌یافته آشکار کننده خطا، یادگیری ماشین



شکل ۱: فرآیند آزمون جهش مدرن [1]

سپس با ردگیری نتایج اجرای بدست آمده، امتیاز جهش یعنی تعداد برنامه‌های جهش‌یافته کشف شده تقسیم بر تعداد کل برنامه‌های جهش‌یافته محاسبه می‌شود. اگر امتیاز جهش به آستانه از پیش تعیین شده برسد، می‌توان مجموعه آزمون را مجموعه آزمون باکیفیت مناسبی در نظر گرفت. اما در غیر این صورت نمی‌توان مجموعه آزمون را مجموعه با کیفیتی دانست و آزمون‌گر باید مجموعه آزمون را بهبود دهد.

در آزمون جهش تعداد بسیار زیادی برنامه‌جهش‌یافته تولید می‌شود [2] که اکثریت آن‌ها ربطی به خطای واقعی ندارند (یعنی کشف آن‌ها خطاهای موجود در برنامه را آشکار نمی‌کند) [3,4]. تعداد زیاد برنامه‌های جهش‌یافته، آزمون جهش را از نظر محاسباتی بسیار پرهزینه می‌کند و مقیاس پذیری آن را

۱- مقدمه

آزمون جهش، روشی قدرتمند است که در آزمون نرم‌افزار برای فعالیتهای گوناگون از جمله راهنمایی برای تولید آزمون و ارزیابی کیفیت مجموعه آزمون استفاده می‌شود و تقریباً تمام معیارهای دیگر کیفیت مجموعه آزمون را شامل می‌شود [1].

شکل ۱ فرآیند آزمون جهش مدرن را نشان می‌دهد. مطابق این شکل فرآیند کلی این آزمون به این صورت است که ابتدا یک برنامه به عنوان ورودی داده می‌شود. سپس به کمک یک ابزار آزمون جهش، تعداد زیادی برنامه جهش‌یافته (یعنی برنامه‌هایی که از نظر گرامری تفاوتی جزئی با برنامه اصلی دارند) انتخاب و تولید می‌شود. در مرحله بعدی آزمون‌گر سعی می‌کند مجموعه آزمونی طراحی کند که بتواند تمام برنامه‌های جهش‌یافته را کشف کند. به عبارت دیگر، اگر رفتار برنامه اصلی و برنامه جهش‌یافته وقتی که آن‌ها را علیه یک مورد آزمون اجرا می‌کنیم متفاوت باشد، می‌گوییم برنامه جهش‌یافته کشف شده است، در غیر این صورت می‌گوییم برنامه جهش‌یافته نجات پیدا کرده است. بعد از طراحی مجموعه آزمون، برنامه اصلی و تمام برنامه‌های جهش‌یافته علیه یک به یک موارد آزمون موجود در مجموعه آزمون اجرا می‌شوند.

تهدید می‌کند [5,6]. از طرف دیگر وجود تعداد زیادی برنامه جهش‌یافته غیر مرتبط با خطای واقعی، اعتبار امتیاز جهش را تهدید می‌کند [7]. به همین دلیل پژوهشگران روش‌های گوناگونی مثل جهش انتخابی [8]، انتخاب تصادفی برنامه‌های جهش‌یافته [9] و غیره را برای کاهش تعداد برنامه‌های جهش‌یافته پیشنهاد داده‌اند.

از آنجایی که در بین مجموعه برنامه‌های جهش‌یافته، تعداد زیادی برنامه جهش‌یافته کم ارزش وجود دارد، روش‌های پیشنهادی گذشته که سعی می‌کنند زیر مجموعه‌ای از برنامه‌های جهش‌یافته را انتخاب کنند که بتواند همه برنامه‌های جهش‌یافته را نمایندگی کند تعداد زیادی برنامه‌های جهش‌یافته کم ارزش را هم در مجموعه انتخابی قرار می‌دهند.

اخیراً چکم و همکاران [10] نشان دادند که اگر به جای انتخاب زیرمجموعه‌ای که بخواهد کل برنامه‌های جهش‌یافته را نمایندگی کند، فقط برنامه‌های جهش‌یافته آشکار کننده خطا را شناسایی و انتخاب کنیم، با هزینه کمتر خطاهای بیشتری نسبت به سایر روش‌های گذشته آشکار خواهد شد.

برنامه‌های جهش‌یافته آشکار کننده خطا، برنامه‌هایی کشف شدنی هستند که کشف شدنشان منجر به نمایان شدن خطا در برنامه اصلی می‌شود و از آنجایی که فقط حدود ۲٪ از برنامه‌های جهش‌یافته، آشکار کننده خطا هستند شناسایی چنین برنامه‌های جهش‌یافته کار چالش برانگیزی است. در این پژوهش روشی ارائه شده است که دقت AUC روش چکم و همکاران [10] را بهبود داده است.

سازمان دهی مقاله به این صورت است: در بخش ۲، پژوهش‌های مرتبط بررسی شده است. در بخش ۳، روش پیشنهادی معرفی می‌شود. سپس در بخش ۴، نتایج ارزیابی‌ها و مقایسه روش پیشنهادی با روش چکم و همکاران [10] ارائه می‌شود. در بخش ۵، تهدیدهایی معرفی می‌شود که اعتبار روش پیشنهادی را به چالش می‌کشند. در نهایت در بخش ۶ جمع‌بندی و نتیجه‌گیری آورده شده است.

۲- تاریخچه پژوهش

از آنجایی که مساله انتخاب برنامه‌های جهش‌یافته آشکار کننده خطا اخیراً معرفی شده است فقط چکم و همکاران [10] به آن پرداخته‌اند. با این حال برای آشنایی بهتر با مساله انتخاب برنامه‌های جهش‌یافته آشکار کننده خطا و درک اهمیت و تفاوت آن با روش‌های سنتی، در این بخش علاوه بر روش چکم و همکاران [10]، بعضی از مهم‌ترین روش‌های سنتی انتخاب زیرمجموعه‌ای از برنامه‌های جهش‌یافته نیز بررسی شده است.

نمونه برداری تصادفی از برنامه‌های جهش‌یافته، ساده ترین روش کاهش تعداد برنامه‌های جهش‌یافته است که به طور جالبی مطابق پژوهش‌های گذشته [11,12] موثر بوده است.

روش دیگر محدود کردن تعداد عملگرهای جهش است که در نتیجه آن تعداد کمتری برنامه جهش‌یافته تولید می‌شود. محدود کردن عملگرهای جهش به فقط ۵ عملگر جهش مرتبط با تطابق رابطه‌ای، منطقی، حسابی، تک عملوندی و مقدار مطلق توسط آفت و همکاران [13] پیشنهاد شد که محبوب ترین مجموعه عملگرهای جهش محسوب می‌شوند و در اکثر ابزارهای مدرن آزمون جهش نیز وجود دارند [1].

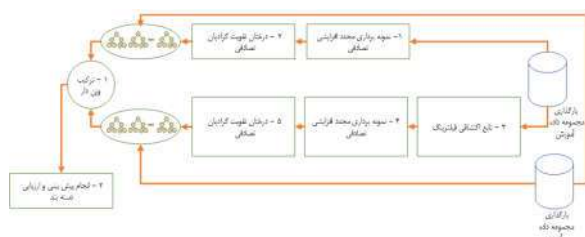
از بین تمام عملگرهای جهش موجود، نشان داده شده است که محدود کردن عملگرهای جهش به فقط عملگر حذف جمله، روشی موثر است که تعداد کمتری برنامه جهش‌یافته هم‌ارز در مقایسه با سایر عملگرها تولید می‌کند [14].

در پژوهش پاداکیس و همکاران [1]، تکنیک‌های دیگری نیز برای انتخاب زیر مجموعه‌ای از برنامه‌های جهش‌یافته بررسی شده است. به عنوان نمونه، روشی موسوم به خوشه بندی برنامه‌های جهش‌یافته را می‌توان نام برد که در آن برنامه‌های جهش‌یافته به خوشه‌هایی تقسیم بندی می‌شوند، سپس تعداد کمی برنامه جهش‌یافته از هر خوشه را می‌توان به عنوان نماینده مورد استفاده قرار داد و از سایر برنامه‌های جهش‌یافته چشم‌پوشی کرد [15].

گوپینات و همکاران [16]، پس از بررسی‌های تجربی و نظری به این نتیجه رسیدند که بین روش انتخاب تصادفی و سایر روش‌هایی که به انتخاب زیر مجموعه‌ای از برنامه‌های جهش‌یافته می‌پردازند، تفاوت چشمگیری وجود ندارد. همچنین نشان داده شده است که اکثریت برنامه‌های جهش‌یافته ربطی به خطای واقعی ندارند [3,4] و علاوه بر این آشکار سازی خطا فقط وقتی بهبود پیدا می‌کند که امتیاز جهش عددی نزدیک به ۱ باشد [7]. با در نظر گرفتن این دو نکته چکم و همکاران [10] پیشنهاد دادند تعدادی ویژگی ایستا از برنامه‌های جهش‌یافته استخراج شود و سپس به کمک الگوریتم درختان تصمیم تقویت گرادیان تصادفی^۱ سعی شود برنامه‌های جهش‌یافته آشکار کننده خطا شناسایی شود و فقط از برنامه‌هایی که به احتمال زیاد آشکار کننده خطا هستند به عنوان زیرمجموعه انتخابی از برنامه‌های جهش‌یافته کل استفاده شود. آن‌ها روش پیشنهادی خود را روی دو مجموعه داده Codefals و CoREBench آزمایش و ارزیابی کردند و به ترتیب AUC برابر ۶۲٪ و ۶۱٪ بدست آوردند. علاوه بر این آن‌ها با مقایسه روش پیشنهادی با سایر روش‌های انتخاب زیرمجموعه‌ای از برنامه‌های جهش‌یافته نشان دادند که اگر به جای انتخاب زیرمجموعه‌ای که بخواهد کل برنامه‌های جهش‌یافته را نمایندگی کند، فقط برنامه‌های جهش‌یافته آشکار کننده خطا را شناسایی و انتخاب کند، با هزینه کمتر خطاهای بیشتری نسبت به سایر روش‌های گذشته آشکار خواهد شد.

۳- روش پیشنهادی

در شکل ۲ فرآیند آموزش و فرآیند آزمون روش پیشنهادی نشان داده شده است.



شکل ۲: فرآیند کلی آموزش و آزمون روش پیشنهادی

در شکل ۲ ابتدا مجموعه داده آموزش در حافظه اصلی بارگذاری می‌شود. سپس دو مدل به صورت موازی آموزش داده می‌شود. برای آزمون مدل اول

^۱ Stochastic Gradient Boosting Decision Trees

ابتدا به کمک نمونه برداری مجدد افزایشی تصادفی^۲ مجموعه داده آموزش متوازن می‌شود و سپس الگوریتم درختان تصمیم تقویت گرادیان تصادفی اجرا می‌شود. برای آموزش مدل دوم، ابتدا به کمک یک تابع اکتشافی، تعداد قابل توجهی از برنامه‌های جهش‌یافته حذف می‌شوند. سپس به کمک نمونه برداری مجدد افزایشی تصادفی مجموعه داده آموزش متوازن می‌شود و در نهایت الگوریتم درختان تصمیم تقویت گرادیان تصادفی اجرا می‌شود. در زمان آزمون دو مدل یادگرفته شده به صورت وزن دار ترکیب می‌شوند و برای پیش‌بینی نهایی و ارزیابی استفاده می‌شوند.

۱-۳- آموزش مدل اول

مدل اول دقیقاً معادل روش چکم و همکاران [10] است که در گام‌های ۱ و ۲ شکل ۲ نشان داده شده است. از آنجایی که فقط حدود ۲٪ از برنامه‌های جهش‌یافته، برنامه‌های جهش‌یافته آشکار کننده خطا هستند و نشان داده شده است که متوازن سازی مجموعه داده دقت نهایی دسته بند را بهبود می‌دهد [17]، در گام ۱ به کمک نمونه برداری افزایشی تصادفی مجموعه داده آموزش متوازن می‌شود. سپس در گام ۲، الگوریتم درختان تصمیم تقویت گرادیان تصادفی با تنظیم ۱۰۰۰ درخت تصمیم با عمق حداکثر ۵ آموزش داده می‌شود.

در ادامه الگوریتم درختان تصمیم تقویت گرادیان شرح داده شده است. الگوریتم درختان تصمیم تقویت گرادیان روشی است که در آن گروهی از درخت‌های تصمیم به عنوان یادگیرنده‌های ضعیف به صورت ترتیبی با هم ترکیب می‌شوند. در این روش ابتدا یک درخت تصمیم آموزش داده می‌شود. سپس برای دور بعدی وزن نمونه‌ها برورسانی می‌شود. برورسانی وزن نمونه‌ها به این صورت است که از وزن نمونه‌هایی که این درخت درست پیش‌بینی می‌کند کاسته می‌شود و وزن نمونه‌هایی که اشتباه پیش‌بینی می‌شود افزایش می‌یابد. در دور بعد درخت تصمیم دیگری با توجه به وزن‌های برورسانی شده آموزش داده می‌شود و با ضربی که خطای آموزش را کمینه کند به مجموعه درختان فعلی اضافه می‌شود. سپس برورسانی وزن‌ها انجام می‌شود و همین فرآیند برای تعداد از پیش تعیین شده‌ای انجام می‌شود. رابطه (۱)، نحوه ترکیب کردن درختان تصمیم در این الگوریتم را نشان می‌دهد که در آن X مجموعه داده آموزش را نشان می‌دهد، T تعداد از پیش تعیین شده تکرارها است، DT_i درخت تصمیم مرحله i ام را نشان می‌دهد و α_i ضریب درخت i ام را نشان می‌دهد:

$$GBT(X) = \sum_{i=1}^T \alpha_i DT_i(X) \quad (1)$$

اگر به جای اینکه در هر دور از تمام مجموعه داده آموزش استفاده شود، به کمک نمونه برداری مجدد تصادفی زیرمجموعه‌ای از مجموعه داده آموزش انتخاب شود به آن الگوریتم درختان تصمیم تقویت گرادیان تصادفی گفته می‌شود که در این پژوهش نیز از همین روش استفاده شده است.

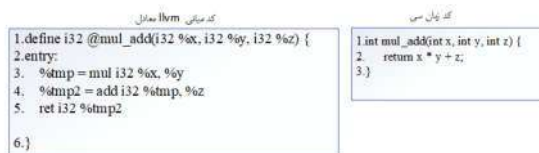
۲-۳- آموزش مدل دوم

مراحل آموزش مدل دوم در گام‌های ۳، ۴ و ۵ شکل ۲ نشان داده شده است. برای آموزش مدل دوم، ابتدا به کمک یک تابع اکتشافی تعداد زیادی از برنامه‌های جهش‌یافته از مجموعه آموزش حذف می‌شود. تابع اکتشافی

پیشنهادی به این صورت است که برنامه‌های جهش‌یافته‌ای که در خطوطی از برنامه هستند که حداقل یک برنامه‌جهش یافته آشکار ساز خطا در آن خطوط وجود دارد نگه داشته می‌شود و برنامه‌های جهش‌یافته‌ای که در خطوطی هستند که حتی یک برنامه جهش‌یافته آشکار ساز خطا در آن‌ها وجود ندارد از مجموعه داده آموزش حذف می‌شود. سپس نمونه برداری مجدد افزایشی تصادفی انجام می‌شود و در آخر درختان تصمیم تقویت گرادیان تصادفی با تنظیم ۱۰۰۰ درخت تصمیم با عمق حداکثر ۵ آموزش داده می‌شود.

یادآوری می‌گردد، مجموعه داده‌های استفاده شده در این پژوهش از برنامه‌های زبان سی که به کد میانی LLVM تبدیل شده‌اند جمع آوری شده است. هر جمله در زبان سی، معادل یک یا چند جمله در کد میانی LLVM می‌باشد. برای اعمال تابع اکتشافی پیشنهادی، برنامه‌های جهش‌یافته LLVM با توجه به این که مربوط به کدام خط برنامه زبان سی بوده‌اند گروه‌بندی می‌شوند و سپس گروه‌هایی که هیچ برنامه جهش‌یافته آشکار ساز خطایی در آن‌ها وجود ندارد از مجموعه آموزش حذف می‌شوند.

شکل ۳ مثالی را برای کد زبان سی و کد میانی LLVM معادل آن نشان می‌دهد.



شکل ۳: نمونه برنامه به زبان C و کد میانی LLVM معادل آن

مطابق شکل ۳ وقتی نمونه برنامه زبان C به LLVM تبدیل می‌شود، خط اول برنامه C، معادل خط اول کد میانی LLVM می‌شود و خط دوم برنامه C، معادل مجموعه خطوط ۳، ۴ و ۵ کد میانی LLVM می‌شود. حال فرض کنید برای خط اول کد میانی LLVM، ۵ برنامه جهش‌یافته حاصل جایگزشت متغیرهای ورودی تولید شود که از میان آن‌ها هیچ کدام برنامه جهش‌یافته آشکار کننده خطا نباشد. مطابق تابع اکتشافی در این حالت هر ۵ برنامه جهش‌یافته از مجموعه داده حذف می‌شود. علاوه بر این فرض کنید تعدادی برنامه جهش‌یافته از اعمال عملگر جهش در خطوط ۳، ۴ و ۵ کد میانی LLVM داشته باشیم که از بین آن‌ها حداقل یکی از برنامه‌های جهش‌یافته، آشکار کننده خطا باشد. در این حالت مطابق تابع اکتشافی تمام برنامه‌های جهش‌یافته خطوط ۳، ۴ و ۵ کد میانی LLVM را نگه می‌داریم چرا که این خطوط متعلق به خط یکسانی از برنامه به زبان C (یعنی خط دوم) هستند و حداقل یک برنامه جهش‌یافته آشکار کننده خطا نیز دارند.

۳-۳- آزمون

در فرآیند آزمون هر برنامه جهش یافته در مجموعه آزمون به هر دو مدل یادگیری شده داده می‌شوند. سپس هر مدل با یک عدد بین ۰ و ۱ اطمینان خود نسبت به آشکار کننده بودن نمونه داده شده را بیان می‌کند. سپس مطابق گام ۱ فرآیند آزمون نشان داده شده در شکل ۲، پیش‌بینی نهایی از جمع وزن دار این دو مدل محاسبه می‌شود به این صورت که وزن مدل اول ۰.۸۴ و وزن

^۲ Random Oversampling

مدل دوم ۰.۱۴ می‌باشد. در آخر با توجه به پیش‌بینی‌های انجام شده ارزیابی صورت می‌گیرد.

است. در نتیجه این پیش‌پردازش‌ها در مجموعه داده Codeflaws، ۱۰۷۴ ویژگی و در مجموعه داده CoREBench، ۱۰۸۲ ویژگی وجود دارد.

۴- ارزیابی روش پیشنهادی

در این بخش ابتدا مجموعه داده‌های استفاده شده معرفی شده است. سپس معیارهایی که در این پژوهش برای ارزیابی‌ها استفاده شده بررسی شده است. در آخر نیز عملکرد روش پیشنهادی و مولفه‌های آن با روش چکم و همکاران [10] مورد مقایسه قرار گرفته است.

۴-۱- معرفی مجموعه داده‌های استفاده شده

برای انجام آزمایش‌ها از دو مجموعه داده پیش‌پردازش شده Codeflaws و CoREBench استفاده شده است که توسط چکم و همکاران [10] تهیه شده و در دسترس قرار گرفته است.

Codeflaws شامل ۱۶۹۲ برنامه به زبان C برای ۷۶۸ مساله است که در مسابقات برنامه‌نویسی آنلاین نوشته شده است. در این مجموعه داده تعداد خطوط هر برنامه حداقل ۱ و حداکثر ۳۲۲ و به طور میانگین ۳۶ خط می‌باشد. همچنین تعداد ۱۷۵۶۰۳۱ برنامه جهش‌یافته موجود است که حدود ۲٪ آن‌ها برنامه‌های جهش‌یافته آشکار کننده خطا هستند.

CoREBench شامل ۴۵ نسخه^۳ از ۴ پروژه واقعی و پیچیده به زبان C می‌باشد. در این مجموعه داده تعداد خطوط هر پروژه حداقل ۹۰۰۰ و حداکثر ۸۳۰۰۰ خط می‌باشد. همچنین تعداد ۸۴۸۶۱۸ برنامه جهش‌یافته موجود است که حدود ۱٪ آن‌ها برنامه‌های جهش‌یافته آشکار کننده خطا هستند.

چکم و همکاران [10] برای تهیه این مجموعه داده‌ها، برنامه‌های به زبان C را به کد llvm تبدیل کردند و تمام مراحل بعدی پردازش را نیز روی کد llvm انجام دادند. به طور خاص آن‌ها ابزاری به نام MART توسعه دادند [18] که تمام عملگرهای جهش رایج را شامل می‌شود و به کمک آن برای هر پروژه، برنامه‌های جهش‌یافته را تولید کردند. سپس به کمک روش TCE [19] که روشی برای شناسایی برنامه‌های جهش‌یافته هم‌ارز ساده است، برنامه‌های جهش‌یافته هم‌ارز ساده را از مجموعه داده‌ها حذف کردند. سپس به کمک KLEE [20] که ابزاری برای تولید خودکار آزمون است، برای هر پروژه آزمون تولید کردند. در گام بعد برنامه‌های جهش‌یافته‌ای که خطاهای ساده را نمایان می‌کردند از مجموعه داده حذف کردند.

بعد از انجام مراحل پیش‌پردازش بالا، آن‌ها ۲۸ ویژگی از هر برنامه جهش‌یافته استخراج کردند به صورتی که این ویژگی‌ها سعی می‌کنند برنامه جهش‌یافته را از زوایای مختلفی مثل پیچیدگی جمله‌ای که عملگر جهش در آن اعمال شده، محل جهش در کد با توجه به گراف جریان کنترل، وابستگی برنامه جهش‌یافته به سایر برنامه‌های جهش‌یافته و نوع و طبیعت بلاک کد که جهش در آن رخ داده، بیان کنند. از میان این ۲۸ ویژگی ۱۵ ویژگی عددی هستند و سایر ویژگی‌ها دسته‌ای اند که به صورت one-hot کدگذاری شده‌اند. همچنین بعضی از ویژگی‌ها از جمع بردار one-hot شده برنامه‌های جهش‌یافته مختلف محاسبه شده است. در نهایت تمام مقادیر موجود در مجموعه داده به کمک روش min-max scaler به بازه بین ۰ و ۱ برده شده

۴-۲- معیارهای ارزیابی

با توجه به این که کاربردهای مشخصی برای انتخاب برنامه‌های جهش‌یافته آشکار کننده خطا متصور می‌باشد، باید معیارهایی را انتخاب کنیم که رسیدن به هدف کاربرد مورد نظر را مشخص کند. هدف اصلی انتخاب برنامه‌های جهش‌یافته آشکار کننده خطا این است که به جای تحلیل تمام برنامه‌های جهش‌یافته بالقوه، فقط برنامه‌های جهش‌یافته آشکار کننده خطا را برای تحلیل به کاربر نهایی پیشنهاد دهیم. با این دید، انتخاب برنامه‌های جهش‌یافته آشکار کننده خطا را می‌توان برای دو کاربرد در نظر گرفت: ۱- به عنوان روشی برای کاهش تعداد برنامه‌های جهش‌یافته و ۲- به صورت یک سیستم پیشنهاد دهنده که تعداد محدودی برنامه جهش‌یافته را به کاربر نهایی پیشنهاد می‌دهد.

در این مستند از معیارهایی استفاده شده که هر دو کاربرد را پوشش دهد. اگر صرفاً هدف، کاهش تعداد برنامه‌های جهش‌یافته باشد یا با نگاهی دیگر هدف، پیشنهاد تعدادی برنامه جهش‌یافته به کاربر نهایی باشد بدون اینکه محدودیت خاصی روی تعداد برنامه‌های جهش‌یافته پیشنهادی داشته باشیم معیارهای سنتی معیارهای خوبی به نظر می‌رسند. با در نظر گرفتن این نکته در این پژوهش برای کاربرد اول از AUC استفاده شده است. AUC معیاری است که در آن با حرکت دادن آستانه تصمیم‌گیری، سطح زیر نمودار True Positive Rate (TPR) و False Positive Rate (FPR) محاسبه می‌شود و می‌تواند به طور کلی توانایی دسته‌بندی را نشان دهد. رابطه (۲) و رابطه (۳) شیوه محاسبه TPR و FPR را نشان می‌دهد که TP مخفف True Positive، FN مخفف False Negative، FP مخفف False Positive و TN مخفف True Negative می‌باشد:

$$TPR = \frac{TP}{TP + FN} \quad (2)$$

$$FPR = \frac{FP}{FP + FN} \quad (3)$$

در رابطه با کاربرد دوم، با توجه به این که نشان داده شده است اکثریت برنامه‌های جهش‌یافته ربطی به خطای واقعی ندارند [3,4]، می‌دانیم در عمل درصد اندکی از برنامه‌های جهش‌یافته برای تحلیل نهایی مورد استفاده قرار می‌گیرند. بنابراین می‌توانیم خود را محدود به انتخاب درصد اندکی برنامه جهش‌یافته کنیم. به عبارت دیگر، می‌توانیم k برنامه جهش‌یافته‌ای که دسته‌بندی نسبت به آشکار کننده خطا بودنشان اطمینان بیشتری دارد را انتخاب کنیم. در این حالت برای ارزیابی می‌توانیم از معیارهایی استفاده کنیم که در سیستم‌های پیشنهاد دهنده امروزی استفاده می‌شود [21-23]. مشخصاً برای کاربرد دوم در این پژوهش از $recall@k$ استفاده شده است به صورتی که k مطابق پژوهش چکم و همکاران [10] برابر ۵٪، ۱۰٪ و ۲۰٪ در نظر گرفته شده است. رابطه (۴) نحوه محاسبه $recall@k$ را نشان می‌دهد:

$$recall@k = \frac{\#recommended\ items\ @k\ that\ are\ relevant}{\#total\ relevant\ items} \quad (4)$$

مدل ۱	۶۴,۳٪±۱٪	۱۸,۶٪±۲٪	۲۷,۸٪±۲٪	۳۹,۹٪±۲٪
مدل ۲	۶۲,۳٪±۱٪	۱۳,۳٪±۱٪	۲۰,۳٪±۲٪	۳۳,۷٪±۲٪
ترکیبی	۶۶,۱٪±۱٪	۱۹,۹٪±۲٪	۲۸,۷٪±۲٪	۴۲,۴٪±۲٪

به طور جالبی، تقریباً تمام مطالب ذکر شده در رابطه با جدول ۱، درباره جدول ۲ نیز صادق است. مطابق جدول ۲، روش پیشنهادی مقدار AUC را نسبت به روش چکم و همکاران [10]، حدود ۱,۸٪ بهبود داده است. علاوه بر این روش پیشنهادی برای recall@5% و recall@10% کمی از روش چکم و همکاران [10] بهتر عمل کرده اما وقتی درصد بیشتری یعنی ۲۰٪ برنامه جهش یافته پیشنهاد شود میزان بهبود افزایش می یابد.

۵- تهدیدهای اعتبار ارزیابی ها

دو نوع تهدید ارزیابی های این پژوهش را تهدید می کند. تهدید اول مربوط به استفاده از ابزارها و مجموعه داده های خارجی است. به این معنی که در این پژوهش از مجموعه داده های تهیه شده توسط چکم و همکاران [10] استفاده شده است و وجود اشتباهات در تولید برنامه های جهش یافته، تولید آزمون، استخراج ویژگی و تمام مراحل پیش پردازش انجام گرفته توسط آن ها می تواند در نتایج این پژوهش تاثیر بگذارد. به صورت مشابه همین مطلب را می توان برای ابزارهایی که در این پژوهش استفاده شده نیز در نظر گرفت. برای کنترل و مدیریت با این تهدید، کد در دسترس قرار گرفته توسط چکم و همکاران [10] مورد بازبینی قرار گرفته و با نویسنده مسئول مقاله مکاتبه صورت گرفته است تا ابهامات برطرف شود. در رابطه با ابزارهای استفاده شده، نیز فقط از کتابخانه های استاندارد و پذیرفته شده در میان جامعه علمی استفاده شده است. به طور خاص، برای پیاده سازی روش پیشنهادی از کتابخانه های sklearn, imbalanced-learn, pandas, numpy و xgboost زبان پایتون استفاده شده است.

تهدید دوم تهدیدهای داخلی مربوط به نتایج گزارش شده در این پژوهش است. در این رابطه، پس از انجام آزمایش ها، چندین مرتبه کدها بازبینی شد تا در صورت وجود خطا شناسایی شود. علاوه بر این از آن جایی که بهبود بدست آمده در پژوهش درصد زیادی نبود این احتمال وجود داشت که این بهبود ناشی از نویز یا شانس باشد که در این رابطه، روی دو مجموعه داده مستقل از هم آزمایش ها را با تنظیمات یکسان انجام دادیم.

علاوه بر تهدیدهای گفته شده، تهدید دیگری در نحوه ارزیابی روش پیشنهادی می باشد. چکم و همکاران [10]، برای ارزیابی روشی پیشنهادیشان همه برنامه های جهش یافته را به عنوان یک مجموعه کل مورد آزمون قرار دادند. در صورتی که اغلب در پژوهش های آزمون جهش ارزیابی پروژه به پروژه انجام می شود [24-27]، که در این پژوهش نیز از ارزیابی پروژه به پروژه استفاده شده است. با این حال برای اطمینان بیشتر، AUC به شیوه ای که چکم و همکاران [10] محاسبه کردند نیز محاسبه شد که در این حالت روش پیشنهادی AUC را روی مجموعه داده های Codeflaws و CoREBench به ترتیب از ۶۲٪ و ۶۱,۶٪ روش چکم و همکاران [10] به ۶۶٪ و ۶۵,۵٪ بهبود داد.

۶- نتیجه گیری

آزمون جهش روشی قدرتمند است که در آزمون نرم افزار استفاده می شود. با این حال، تعداد زیاد برنامه های جهش یافته کم ارزش، هزینه و اعتبار آن را

مطابق رابطه (۴)، در مسئله انتخاب برنامه های جهش یافته آشکار کننده خطا، آیت های مرتبط، برنامه های جهش یافته آشکار کننده خطا می باشند. در این حالت برای محاسبه $recall@k$ ابتدا دسته بند یادگیری شده در بخش ۳ برای هر کدام از برنامه های جهش یافته ای که به آن داده می شود یک عدد بین ۰ و ۱ را به عنوان پیش بینی برمی گرداند که در آن هر چه عدد خروجی به ۱ نزدیک تر باشد یعنی دسته بند با اطمینان بیشتری نمونه داده شده را متعلق به کلاس برنامه های جهش یافته آشکار کننده خطا می داند. سپس k آیت می انتخاب می شوند که دسته بند با اطمینان بیشتری آن ها را متعلق به کلاس جهش یافته آشکار کننده خطا می داند. در نهایت برای محاسبه صورت کسر، تعداد برنامه های جهش یافته آشکار کننده خطا از میان k آیت انتخابی شمارش می شود و برای محاسبه مخرج کسر، تعداد برنامه های جهش یافته آشکار کننده خطا در میان تعداد کل نمونه هایی که برای پیش بینی به دسته بند داده شده است، شمارش می شود.

۳-۴- نتایج ارزیابی روش پیشنهادی

پروژه های هر دو مجموعه داده به ده قسمت تقریباً برابر تقسیم شد و برای ارزیابی ها از روش 10-fold cross-validation استفاده شد. پس از آموزش مدل مربوط، یک به یک پروژه های مجموعه آزمون فولد مربوطه برای پیش بینی به دسته بند نهایی داده شد. در نهایت برای هر پروژه، مقدار معیارهای ارزیابی محاسبه گردید. یادآوری می گردد در این بخش برای اشاره کردن به هر یک از ۴۵ نسخه موجود در CoREBench و همچنین هر یک از ۱۶۹۲ برنامه موجود در Codeflaws از لفظ پروژه استفاده شده است.

جدول ۱، میانگین هر یک از معیارهای ارزیابی را به همراه انحراف از معیار آن ها برای تمامی پروژه های Codeflaws گزارش می کند.

جدول ۱: نتایج ارزیابی ها روی Codeflaws به تفکیک پروژه

روش	AUC	Recall@5%	Recall@10%	Recall@20%
مدل ۱	۷۳,۱٪±۲٪	۲۷,۷٪±۳٪	۳۸,۷٪±۳٪	۵۲,۳٪±۳٪
مدل ۲	۷۳,۰٪±۱٪	۲۲,۹٪±۲٪	۳۵,۳٪±۳٪	۵۱,۷٪±۳٪
ترکیبی	۷۴,۷٪±۱٪	۲۸,۲٪±۳٪	۳۹,۷٪±۳٪	۵۴,۱٪±۳٪

مطابق جدول ۱، مقدار AUC مدل اول و دوم بسیار به هم نزدیک هستند و این در حالی است که مدل دوم در مقایسه با مدل اول ۹۵٪ کمتر داده برای آموزش استفاده کرده است. همچنین مطابق این جدول، روش پیشنهادی (یعنی مدل ترکیبی) توانسته علاوه بر کاهش ۱٪ انحراف از معیار مقدار AUC را نسبت به روش چکم و همکاران [10]، حدود ۱,۶٪ بهبود دهد.

مطابق جدول ۱، برای recall@5% و recall@10% مدل اول بهتر از مدل دوم عمل کرده، اما وقتی در صد بیشتری یعنی ۲۰٪ برنامه جهش یافته پیشنهاد شود این اختلاف بسیار کم می شود. علاوه بر این مشاهده می شود روش پیشنهادی برای recall@5% و recall@10% کمی از روش چکم و همکاران [10] بهتر عمل کرده اما وقتی درصد بیشتری یعنی ۲۰٪ برنامه جهش یافته پیشنهاد شود میزان بهبود افزایش می یابد.

جدول ۲، میانگین هر یک از معیارهای ارزیابی را به همراه انحراف از معیار آن ها برای تمامی پروژه های CoREBench گزارش می کند.

جدول ۲: نتایج ارزیابی ها روی CoREBench به تفکیک پروژه

روش	AUC	Recall@5%	Recall@10%	Recall@20%
-----	-----	-----------	------------	------------

- 2020.
- [11] L. Zhang, S.-S. Hou, J.-J. Hu, T. Xie, and H. Mei, "Is operator-based mutant selection superior to random mutant selection?," *Proceedings of the 32nd ACM/IEEE International Conference on Software Engineering*, 2010.
 - [12] M. Papadakis and N. Malevris, "An empirical evaluation of the first and second order mutation testing strategies," *Third International Conference on Software Testing, Verification, and Validation Workshops*, 2010.
 - [13] A. J. Offutt, G. Rothermel, and C. Zapf, "An experimental evaluation of selective mutation," *Proceedings of the 15th international conference on Software Engineering*, 1993.
 - [14] L. Deng, J. Offutt, and N. Li, "Empirical evaluation of the statement deletion mutation operator," *IEEE Sixth International Conference on Software Testing, Verification and Validation*, 2013.
 - [15] S. Hussain, "Mutation Clustering," M.S. thesis, King's College London, 2008.
 - [16] R. Gopinath, M. A. Alipour, I. Ahmed, C. Jensen, and A. Groce, "On the limits of mutation reduction strategies," *Proceedings of the 38th International Conference on Software Engineering*, 2016.
 - [17] C. Tantithamthavorn, A. E. Hassan, and K. Matsumoto, "The impact of class rebalancing techniques on the performance and interpretation of defect prediction models," *IEEE trans. softw. eng.*, 2020.
 - [18] T. T. Chekam, M. Papadakis, and Y. Le Traon, "Mart: a mutant generation tool for LLVM," *Proceedings of the 27th ACM Joint Meeting on European Software Engineering Conference and Symposium on the Foundations of Software Engineering*, 2019.
 - [19] M. Papadakis, Y. Jia, M. Harman, and Y. Le Traon, "Trivial compiler equivalence: A large scale empirical study of a simple, fast and effective equivalent mutant detection technique," *IEEE/ACM 37th IEEE International Conference on Software Engineering*, 2015.
 - [20] C. Cadar, D. Dunbar, D. Engler, "KLEE: unassisted and automatic generation of high-coverage tests for complex systems programs," *8th USENIX Symposium on Operating Systems Design and Implementation*, 2008.
 - [21] W. Chen, T. Y. Liu, Y. Lan, Z. M. Ma, and H. Li, "Ranking measures and loss functions in learning to rank," *adva. in neural info. proc. syst.*, 2009.
 - [22] D. Tasche, "A plug-in approach to maximising precision at the top and recall at the top," *arXiv*, 2018.
 - [23] P. Kar, H. Narasimhan, and P. Jain, "Surrogate functions for maximizing precision at the top," *arXiv*, 2015.
 - [24] D. Mao, L. Chen, and L. Zhang, "An extensive study on cross-project predictive mutation testing," *12th IEEE Conference on Software Testing, Validation and Verification*, 2019.
 - [25] "An ensemble- based predictive mutation testing approach that considers impact of unreached mutants," *Softw. Test. Verif. Reliab.*, 2021.
 - [26] M. R. Naeem, T. Lin, H. Naeem, F. Ullah, and S. Saeed, "Scalable mutation testing using predictive analysis of deep learning model," *IEEE Access*, 2019.
 - [27] P. Zhang et al., "CBUA: A probabilistic, predictive, and practical approach for evaluating test suite effectiveness," *IEEE trans. softw. eng.*, 2020.

تهدید می‌کند. برای مواجهه شدن با این تهدیدها روش‌های مختلفی برای انتخاب زیرمجموعه‌ای از برنامه‌های جهش‌یافته پیشنهاد شده است که با وجود موثر بودن نسبی آن‌ها در کاهش هزینه آزمون جهش، در رابطه با اعتبار امتیاز جهش کمکی نمی‌کنند.

اخیراً چکم و همکاران [10] نشان دادند که اگر به جای انتخاب زیرمجموعه‌ای که بخواد کل برنامه‌های جهش‌یافته را نمایندگی کند، فقط برنامه‌های جهش‌یافته آشکار کننده خطا را شناسایی و انتخاب کنیم، با هزینه کمتر خطاهای بیشتری نسبت به سایر روش‌های گذشته آشکار خواهد شد. در نتیجه علاوه بر کاهش هزینه آزمون جهش، امتیاز جهش نیز از اعتبار بیشتری برخوردار خواهد بود.

در روش چکم و همکاران [10]، ابتدا مجموعه داده آموزش به کمک نمونه برداری تصادفی افزایشی متوازن می‌شود و سپس به کمک الگوریتم درختان تصمیم تقویت گرا دیان مدلی برای پیش‌بینی کردن برنامه‌های جهش‌یافته آشکار کننده خطا آموزش داده می‌شود.

در این پژوهش روشی پیشنهاد شد که AUC را نسبت به پژوهش چکم و همکاران [10]، حدود ۲٪ بهبود می‌دهد. ایده روش پیشنهادی این است که می‌توان با ترکیب کردن دسته‌بندی‌هایی با دیدگاه‌های مختلف نسبت به مجموعه آموزش به دقت بیشتری دست یافت. برای این منظور تابعی اکتشافی برای فیلترینگ نمونه‌های مجموعه داده آموزش، توسعه داده شد و سپس دسته‌بندی روی مجموعه فیلتر شده آموزش داده شد. در آخر دسته‌بند آموزش داده شده با دسته‌بند پیشنهادی چکم و همکاران [10] به صورت وزن‌دار ترکیب شد و برای پیش‌بینی نهایی مورد استفاده قرار گرفت. برای پژوهش‌های آینده، می‌توان توابع اکتشافی دیگری توسعه داد و در نهایت همه آن‌ها را با هم ترکیب کرد تا به دقت بیشتری دست یافت.

مراجع

- [1] M. Papadakis, M. Kintis, J. Zhang, Y. Jia, Y. L. Traon, and M. Harman, "Mutation testing advances: An analysis and survey," *Elsevier*, 2019, pp. 275–378.
- [2] J. Zhang, L. Zhang, M. Harman, D. Hao, Y. Jia, and L. Zhang, "Predictive Mutation Testing," *IEEE trans. softw. eng.*, 2019.
- [3] M. Papadakis, D. Shin, S. Yoo, and D. H. Bae, "Are mutation scores correlated with real fault detection? a large scale empirical study on the relationship between mutants and real faults," *IEEE/ACM 40th International Conference on Software Engineering*, 2018.
- [4] M. Papadakis, T. T. Chekam, and Y. Le Traon, "Mutant Quality Indicators," *IEEE International Conference on Software Testing, Verification and Validation Workshops*, 2018.
- [5] L. D. Dallilo, A. V. Pizzoleto, and F. C. Ferrari, "An evaluation of internal program metrics as predictors of mutation operator score," *Proceedings of the IV Brazilian Symposium on Systematic and Automated Software Testing*, 2019.
- [6] M. Yu and Y.-S. Ma, "Possibility of cost reduction by mutant clustering according to the clustering scope: POSSIBILITY OF COST REDUCTION BY MUTANT CLUSTERING," *Softw. Test. Verif. Reliab.*, 2019.
- [7] T. T. Chekam, M. Papadakis, Y. Le Traon, and M. Harman, "An empirical study on mutation, statement and branch coverage fault revelation that avoids the unreliable clean program assumption," *IEEE/ACM 39th International Conference on Software Engineering*, 2017.
- [8] A. J. Offutt, A. Lee, G. Rothermel, R. H. Untch, and C. Zapf, "An experimental determination of sufficient mutant operators," *ACM Trans. Softw. Eng. Methodol.*, 1996.
- [9] W. E. Wong and A. P. Mathur, "Reducing the cost of mutation testing: An empirical study," *J. Syst. Softw.*, 1995.
- [10] T. Titchou Chekam, M. Papadakis, T. F. Bissyandé, Y. Le Traon, and K. Sen, "Selecting fault revealing mutants," *Empir. Softw.*,



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

تحلیل الگوی مصرفی مشترکین برق با استفاده از خوشه بندی

منصوره خالقی^۱، شیما کاشف^۲، حسین نظام آبادی پور^۳

^۱ کارشناسی ارشد، بخش مهندسی برق، دانشکده فنی و مهندسی، دانشگاه شهید باهنر کرمان، mnsr.khaleghi@yahoo.com

^۲ دکتری، دانشکده علوم و فناوری های نوین، دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته کرمان، sh.kashef@kgut.ac.ir

^۳ استاد، بخش مهندسی برق، دانشکده فنی و مهندسی، دانشگاه شهید باهنر کرمان، nezam_h@yahoo.com

چکیده: داده کاوی امروز یکی از دانش‌های برتر در حوزه علوم بشر است که می‌تواند در کشف، جمع آوری و بهره مندی از نتایج حاصل از داده‌های فراوانی که روزانه توسط شرکت‌ها و سازمان‌ها، تولید می‌شوند، نقش برجسته‌ای داشته باشد. این علم می‌تواند با ارائه نتایج جالب توجه و کشف الگوهای پنهان موجود در داده‌ها، بستر مناسبی را برای اخذ تصمیمات بهتر، جهت کمک به مدیران فراهم آورد. این مقاله، مشترکین شرکت توزیع برق جنوب استان کرمان را با استفاده از داده‌های مربوط به مصرف آنها که به صورت سیکلی در پایگاه داده ذخیره می‌شوند، خوشه بندی کرده‌است. بدین منظور، از دو مرحله خوشه‌بندی استفاده شده‌است. ابتدا، آستانه‌های مجاز مصرف برق مشترکین در مناطق آب و هوایی مختلف تحت نظارت این شرکت، با استفاده از خوشه‌بندی تعیین شده‌است. نتایج حاصل از روش پیشنهادی نشان می‌دهد که آستانه‌های مجاز مصرف برق که در حال حاضر به صورت سراسری از سوی توانیر تعریف شده‌اند، نیاز به تجدید نظر و بازنگری به صورت منطقه‌ای دارند. در مرحله بعد، جابه جایی مشترکین از یک خوشه در یک دوره، به خوشه بعد در دوره دیگر بدست آورده شده و مورد تحلیل قرار گرفته‌اند. بدین ترتیب، الگوی مصرفی مشترکین مورد بررسی قرار می‌گیرد و موارد مشکوک شناسایی می‌شوند.

کلمات کلیدی: تحلیل الگوی مصرف مشترکین، خوشه بندی، داده کاوی، شرکت توزیع برق جنوب استان کرمان

سیستم مورد مطالعه در این پژوهش، شرکت توزیع برق جنوب استان کرمان است. این شرکت در محدوده‌ای به وسعت ۹۵۸۸۷ کیلومتر مربع واقع و ۱۶ شهرستان با جمعیت معادل ۱۶۱۱۶۷۳ نفر را تحت پوشش دارد. شرکت توزیع برق جنوب استان کرمان مشابه سایر شرکت‌های توزیع برق، یک سری آستانه‌های مجاز برای مصرف مشترکین دارد که با توجه به شرایط جغرافیایی مناطق مختلف و فصل عادی یا گرم سال، تعیین می‌شوند. حدود آستانه مجاز به صورت سراسری و کلی از سوی توانیر برای شرکت‌ها ابلاغ می‌شود. در صورتی که مشترکین از این حدود مجاز، مصرف بیشتری داشته باشند، هزینه برق مصرفی آن‌ها به صورت پلکانی محاسبه می‌شود. بنابراین، انتخاب نادرست این حدود آستانه، می‌تواند منجر به آن شود که هزینه قبوض مشترکین خیلی بالا رود و این امر می‌تواند منجر به عدم پرداخت به موقع قبوض یا اقدام برخی مشترکین به مصارف غیر مجاز برق شود. در این مقاله، سعی شده است مقادیر آستانه مجاز هر منطقه به صورت محلی و با استفاده از خوشه بندی داده‌های مربوط به مصرف مشترکین، بدست آورده شوند. تعیین حدود آستانه مجاز مصرف برای فصل‌های مختلف سال و به صورت منطقه ای است. این کار با خوشه بندی مصارف مشترکین هر منطقه انجام شده است.

۱- مقدمه

امروزه رشد داده‌های موجود در صنعت برق سبب ایجاد انگیزه در به کارگیری روش‌های داده کاوی شده است. این داده‌ها به صورت بالقوه توانایی هدایت مدیران در برنامه ریزی‌های کلان و اساسی را دارند و می‌توانند در مواردی مانند بهره برداری بهینه، اخذ راهکارهای مختلف نگهداری و تعمیرات و تنظیم بازار به کمک برنامه ریزان و مدیران بیایند. علیرغم وجود اطلاعات مفید، حجم بالا و بیش از حد این داده‌ها، یافتن قوانین و روابط پیچیده میان آن‌ها را سخت می‌کند و بدون کمک روش‌های داده کاوی نمی‌توان این اطلاعات مفید را استخراج کرد. ذکر این نکته ضروری است که در حال حاضر علی‌رغم تولید رکورد های روزانه و میلیونی از مشترکین برق به دلایل گوناگون از جمله کمبود زمان و تخصص کافی، این داده‌ها غالباً در پایگاه‌های داده‌ای ذخیره می‌شوند و استفاده از این داده‌ها در سطح ارائه گزارش‌های روزانه، ماهانه و سالانه به مدیران و سازمان‌های بالادستی می‌باشد.

پس از تعیین حدود آستانه بالا و پایین مصرف، مشترکین هر منطقه بر اساس میزان مصرفشان در سیکل‌های مختلف یک سال و همچنین، سیکل‌های مشابه سال‌های متوالی خوشه بندی شده و جابه جایی آنها در بین خوشه ها برر سی می شود. این کار، کمک بزرگی به تحلیل رفتار مصرفی مشترکین و یافتن مشترکین مشکوک به مصارف غیر مجاز برق می کند.

ساختار مقاله به این صورت است که بخش دوم به مرور تحقیقات انجام شده در این زمینه می پردازد و در بخش سوم روش پیشنهادی مقاله تشریح شده است. در ادامه در بخش چهارم به تحلیل نتایج و یافته های به دست آمده از این روش پرداخته خواهد شد. در انتها در بخش پنجم مقاله جمع بندی می شود.

۲- مروری بر تحقیقات انجام شده

خوشه بندی یک روش داده کاوی شناخته شده است که تعیین الگوهای اساسی را در مجموعه داده ها امکان پذیر می کند. در شبکه برق، خوشه بندی برای اهداف مختلف مانند شناخت پروفیل بار مشترکین، طراحی تعرفه ها و بهبود پیش بینی بار مورد استفاده قرار می گیرد. داده های مربوط به مصرف مشترکین برق که دارای کنتورهای عادی (مکانیکی یا دیجیتال) هستند، به صورت سیکلی اندازه گیری می شود (هر سال، ۶ سیکل دارد). از این رو، با توجه به کم بودن تعداد داده های مربوط به مصرف مشترکین عادی، تحقیقات کمتری برای خوشه بندی روی این داده ها انجام شده است. محققین در مقاله [۱] با استفاده از این داده ها، مشترکین شهر ایلام را با توجه به مصرف و همچنین با توجه به میزان وصول آنها، با استفاده از روش k-میانگین خوشه بندی کرده اند. سپس، رفتار هر خوشه مورد تجزیه و تحلیل قرار گرفته است. در مرجع [۲]، یک روش دو مرحله ای ارائه شده است که با استفاده از روش های تشخیص الگوی مبتنی بر خوشه بندی k-میانگین و نسخه فازی شده آن و همچنین خوشه بندی سلسله مراتبی، مشتریان دسته بندی شده اند. بدین ترتیب که در ابتدا، نمودار مصرف مشتریان تخمین زده شده و در مرحله دوم به دسته بندی آنها پرداخته اند. برای این منظور از اطلاعات بدست آمده برای پیش بینی بار و تدوین تعرفه ها استفاده شده است. این اطلاعات در واقع در عرصه رقابت شرکت ها بسیار کارگشا می باشند. در نهایت، این مدل، بر روی داده های مصرف مشتریان ولتاژ متوسط شبکه برق یونان مورد آزمون و ارزیابی قرار گرفته است. در کار دیگری [۳] ناهنجاری مصرفی مشترکین برق با استفاده از خوشه بندی مبتنی بر تراکم شناسایی شده است.

برای مشترکین دارای کنتورهای هوشمند، که میزان مصرف را در هر ۱۵ دقیقه ثبت می کنند، روش های متفاوتی برای خوشه بندی انجام شده است. با داشتن این داده ها، می توان از طریق خوشه بندی به اهدافی نظیر طراحی تعرفه، پیش بینی بار، پاسخگویی بار و طبقه بندی مشترک رسید [4]. توجه به داده های کنتور هوشمند برای شناسایی رفتار مشترکین در مطالعاتی از جمله مرجع [۵] مورد توجه قرار گرفته است. در این مقاله، با استفاده از داده های کنتورهای هوشمند، مشترکینی شناسایی شده اند که می توانند بیشترین مشارکت را در برنامه های پاسخگویی بار داشته باشند. همچنین، شناسایی الگوی مصرف مشترکین به منظور پیش بینی کوتاه مدت بار مصرفی در مرجع [6] و برای شناسایی تغییرات در رفتار مصرفی

مشترکین در مرجع [۷] مورد برر سی قرار گرفته است. در مراجع [8,9] نیز از ظرفیت داده های کنتورهای هوشمند به منظور شناسایی دقیق رفتار مشترکین استفاده شده است. در این مطالعات، هدف، استفاده از داده های سری زمانی مصرفی مشترکین برای مشارکت در برنامه های پاسخگویی بار و بهره وری انرژی است. به طور مشابه، مرجع [۱۰] نیز از داده های کنتورهای هوشمند برای تعیین تعرفه های بهینه مصرفی در شبکه استفاده نموده است.

۳- روش پیشنهادی

امور مختلف تحت مدیریت شرکت توزیع برق جنوب استان کرمان، از نظر قرارگیری در موقعیت های مختلف جغرافیایی به سه بخش مناطق عادی، گرم ۱ و گرم ۲ تقسیم می شوند.

با توجه به منطقه جغرافیایی که امور در آن واقع است، یک حد مجاز برای مصرف در سیکل های مختلف در نظر گرفته شده است، که مشترکین مقیم در این مناطق نباید مصرف بالاتر از حد مجاز داشته باشند. در غیر این صورت، قبوض این مشترکین مشمول محاسبه پلکانی صورت حساب می شود. جدول ۱، حدود آستانه مجازی که در حال حاضر در امور مختلف شرکت توزیع برق جنوب استان کرمان مورد استفاده قرار می گیرند به تفکیک منطقه جغرافیایی نشان می دهد. طبق این جدول، به طور مثال، اموری که در مناطق جغرافیایی عادی قرار گرفته اند در ماه های عادی سال (شامل سیکل های ۱، ۴، ۵ و ۶) نباید مصرف بالاتر از ۲۰۰ kWh داشته باشند. همچنین، حد آستانه مصرف ماهیانه برای مناطق گرم ۱ در ماه های گرم (شامل سیکل های ۱، ۲، ۳ و ۴) ۳۰۰ kWh تعیین شده است.

جدول (۱): الگوی مصرف مجاز برای مناطق مختلف

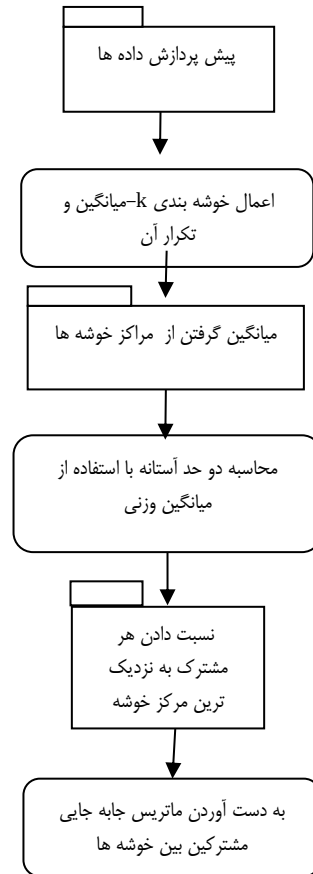
مناطق	شامل سیکل های	آستانه مصرف (kwh)
مناطق عادی	ماه های عادی ماه های گرم	۲۰۰ ۳۰۰
مناطق گرم ۱	ماه های عادی ماه های گرم	۲۰۰ ۳۰۰
مناطق گرم ۲	ماه های عادی ماه های گرم	۲۰۰ ۳۰۰

در این مقاله، حدود آستانه مجازی که برای مشترکین هر منطقه جغرافیایی تعیین شده است، مورد برر سی قرار می گیرند و با استفاده از روش های خوشه بندی، مقادیر آستانه مناسبی برای هر منطقه به دست می آید. شکل ۱، نمودار بلوکی روش پیشنهادی را برای تعیین حدود آستانه مناسب برای هر منطقه جغرافیایی را نشان می دهد.

بدین منظور از خوشه بندی k-میانگین استفاده می شود. در این روش، ابتدا k داده بصورت تصادفی به عنوان مراکز خوشه انتخاب می شوند. سپس، هر نمونه به خوشه ای که به مرکز آن خوشه کمترین فاصله را دارد، نسبت داده می شود. پس از تعلق تک تک داده ها به یکی از خوشه ها، مرکز آن خوشه به روز رسانی می شود (میانگین نقاط متعلق به هر خوشه). دو مرحله قبل تکرار می شوند تا زمانی که تغییر محسوسی در مراکز خوشه ها حاصل نشود.

ایراد مهم این روش این است که نتیجه نهایی وابسته به انتخاب اولیه مراکز است. برای حل این مشکل، ابتدا خوشه بندی را برای ده بار تکرار می کنیم. سپس، از مراکز خوشه های به دست آمده در ده تکرار، میانگین گیری کرده و

k مرکز خوشه نهایی محاسبه می شوند. در گام بعد، هر مشترک به خوشه ای تعلق می گیرد که کمترین فاصله را تا مراکز تعیین شده داشته باشد.



شکل (۱): روش پیشنهادی دو مرحله ای برای یافتن ماتریس جابه جایی مشترکین بین خوشه های مختلف مصرف

در اینجا، تعداد خوشه ها $k=3$ در نظر گرفته شده است که معرف خوشه مشترکین با مصارف پایین، مصارف متوسط و مصارف بالا است. در جدول ۱، تنها حدود آستانه بالا برای مصرف مشترکین هر منطقه، تعیین شده است. با این حال، حدود آستانه پایین نیز، می تواند برای بخش بازرسی بسیار مهم باشند؛ زیرا، مشترکین با مصارف پایین (بین ۰ تا حد آستانه پایین مصرف)، از کاندیدهای خرابی کنتور یا مصارف غیر مجاز برق هستند. از این رو، در این مقاله، آستانه های بالا و پایین برای هر منطقه بدست آورده می شوند. روش پیشنهادی برای محاسبه آستانه های بالا و پایین مصرف مشترکین هر منطقه به صورت زیر است. ابتدا مصارف پرت، نظیر مصرف های صفر و مصارف خیلی زیاد و نامتعارف، کنار گذاشته می شوند. سپس، با استفاده از

خوشه بند k-میانگین و روشی که در بالا به آن اشاره شد، مشترکین به ۳ خوشه تقسیم می شوند که معرف خوشه مشترکین با مصارف پایین، مصارف متوسط و مصارف بالا است. پس از اینکه همه مشترکین به یکی از سه خوشه تخصیص داده شدند، حدود آستانه مصرف بالا و پایین با استفاده از میانگین گیری وزن دار بین سه خوشه به دست می آید. این مطلب، در شکل ۲ نمایش داده شده است.

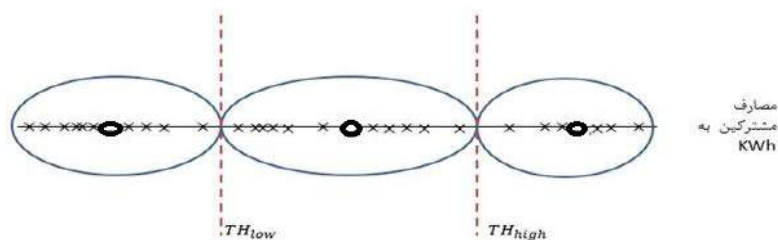
پس از تعیین حدود آستانه پایین و بالا، مشترکین، با توجه به میزان مصرفشان در یکی از سه ناحیه کم مصرف، میان مصرف یا پر مصرف جای می گیرند. این خوشه بندی، برای مشترکین در سال های مختلف با توجه به میزان مصرفشان انجام می شود. اکنون با استفاده از این اطلاعات، ماتریسی به نام ماتریس گذار مشترکین بین خوشه های مختلف تعریف می شود. سطرهای این ماتریس، خوشه بندی مشترکین را در یک دوره و ستون های آن، خوشه بندی همان جایی مشترکین یک منطقه را در بین سه خوشه کم مصرف، میان مصرف و پر مصرف در دوره های مختلف، در نظر بگیریم.

نکته ای که در ماتریس گذار مشترکین حائز اهمیت است، جابه جایی مشترکین از خوشه پر مصرف در یک دوره به خوشه کم مصرف در دوره بعد است. این گونه موارد باید از لحاظ اطمینان از صحت کنتور یا عدم مصرف غیر مجاز برق، مورد بررسی قرار گیرند. بررسی جابه جایی های بین خوشه ای مشترکین، به خصوص بعد از تاریخ اعمال تعرفه های جدید، می تواند دید مناسبی به مدیران، راجع به میزان تاثیر تعرفه ها بر رفتار مصرفی مشترکین بدهد.

۴- یافته های تحقیق

داده های استفاده شده، داده های مربوط به مصارف سیکلی مشترکین خانگی ۱۶ امور تحت نظارت شرکت توزیع برق جنوب استان کرمان در طول ۵ سال است. برای استفاده از داده ها برای خوشه بندی، مراحل پیش پردازش مورد نیاز است. بدین ترتیب که پس از مرتب کردن میزان مصرف هر مشترک از سال ۹۳ تا ۹۷، میزان مصرف انرژی هر دوره ی آن ها بر تعداد روزهای دوره تقسیم می شود تا میانگین مصرف روزانه هر مشترک در آن دوره به دست آید. سپس، مقدار به دست آمده در ۶۰ ضرب می شود تا کل مصرف انرژی هر مشترک برای یک سیکل که معادل دو ماه است تخمین زده شود.

در جدول ۲ مقایسه ای بین آستانه های تعیین شده و آستانه های به دست آمده از روش پیشنهادی انجام شده است. همانطور که مشاهده می شود برای هر کدام از سیکل های عادی و گرم، دو حد بالا و پایین محاسبه شده که نشان دهنده مصرف مشترکین در این حدود است.



شکل ۲: نحوه تعیین حدود آستانه پایین و بالای مصرف انرژی. علائم × کوچک، مصرف انرژی مشترکین مختلف، و سه علامت ●، مرکز خوشه های مختلف را نشان می دهد. حدود آستانه های مصرف انرژی برابر با عمود منصف های مراکز خوشه مجاور در نظر گرفته می شود.

به خوشه پر مصرف منتقل شده‌اند از لحاظ استفاده غیر مجاز تجاری و صنعتی از تعرفه خانگی، نیازمند بازرسی هستند.

رویکرد دوم، برر سی میزان جابه جایی مشترکین در سیکل های مختلف یک سال است. بدین منظور، جابه جایی های مشترکین در امور مختلف، در دو سیکل ۲ و ۳ که هر دو از سیکل های گرم سال هستند و همچنین، دو سیکل ۲ و ۵ که در تمام امور به ترتیب از سیکل های گرم و عادی سال است، ارزیابی می شود. جدول ۴ و ۵، نتایج حاصل از این آزمایش ها را برای سال ۹۶ و ۹۷ امور شهر X نشان می دهد.

جدول (۳): جابجایی مشترکین در سیکل ۲ سال ۹۶ و سیکل ۲ سال ۹۷ در شهر X

۹۶	کم مصرف	میان مصرف	پرمصرف
۹۷			
کم مصرف	۳۴۶۵	۰	۶۵
میان مصرف	۰	۱۹۵۱	۳۶۹
پر مصرف	۱۰۱	۳۱۸	۱۳۵۱

جدول (۴): جابجایی مشترکین در سیکل ۲ سال ۹۶ و سیکل ۳ سال ۹۶ در شهر X

۹۶	کم مصرف	میان مصرف	پرمصرف
۹۷			
کم مصرف	۲۹۳۷	۰	۲۵
میان مصرف	۰	۱۸۰۳	۱۴۷
پر مصرف	۱۳۸	۴۲۰	۱۶۴۴

جدول (۵): جابجایی مشترکین در سیکل ۲ سال ۹۷ و سیکل ۵ سال ۹۷

سیکل ۲	کم مصرف	میان مصرف	پرمصرف
سیکل ۳			
کم مصرف	۳۱۲۳	۰	۱۵۸
میان مصرف	۰	۲۰۸۳	۴۸۲
پر مصرف	۷۰	۱۸۵	۱۰۸۳

۵- نتیجه

در این مقاله، با استفاده از خوشه بندی، دو هدف مورد برر سی قرار گرفت. ابتدا روشی جدید برای تعیین حدود آستانه مجاز مصرف مشترکین، که تاکنون به صورت سراسری از سوی توانیر ابلاغ شده است، با استفاده از ابزار خوشه بندی معرفی شد. پس از به دست آمدن آستانه های جدید، مشترکین، در سه خوشه کم مصرف، میان مصرف و پر مصرف جای گرفتند. دومین هدف، این است که

جدول (۲): حدود آستانه مجاز به دست آمده از روش پیشنهادی

مناطق	شامل سیکل های	آستانه مصرف (kwh)	آستانه مصرف به دست آمده از روش پیشنهادی (kwh)
مناطق عادی	ماه های عادی	۶ و ۵، ۴، ۱	۲۰۰
	ماه های گرم	۳ و ۲	۳۰۰
مناطق گرم ۱	ماه های عادی	۶ و ۵	۲۰۰
	ماه های گرم	۴، ۳، ۲، ۱	۳۰۰
مناطق گرم ۲	ماه های عادی	۶ و ۵، ۴، ۱	۲۰۰
	ماه های گرم	۴ و ۳، ۲	۳۰۰

اعداد موجود در جدول ۲، با میانگین گیری از حدود آستانه به دست آمده برای امور واقع در مناطق جغرافیایی مختلف در طول سال های ۹۳ تا ۹۷ محاسبه شده اند. همچنین، به منظور سادگی در نمایش اعداد، از درج ارقام بعد از ممیز، صرف نظر شده است.

با برر سی اعداد به دست آمده برای حدود آستانه بالا و پایین و مقایسه این اعداد، با حدود آستانه مجاز مصرف برای امور مختلف، مشخص می شود که حدود آستانه بالای مصرف در برخی از سیکل ها و برخی از امور، باید بازنگری شوند. البته این بدان معنی نیست که هر چه مصرف مشترکین بالاتر رود، حدود آستانه مجاز مصرف نیز باید بالاتر برده شود. در اینجا نیز برای جلوگیری از چنین امری و برای رعایت انصاف، میانگین گیری مصرف را روی ۵ سال انجام دادیم.

همانطور که در روش پیشنهادی توضیح داده شد، پس از تعیین حدود آستانه بالا و پایین برای امور واقع در مناطق جغرافیایی مختلف، مشترکین هر منطقه با توجه به میزان مصرفشان و همچنین با توجه به حدود آستانه های به دست آمده، در یکی از خوشه های کم مصرف، میان مصرف و پر مصرف خوشه بندی می شوند. برای ارزیابی، رفتار مصرفی مشترکین، با استفاده از ماتریس جابجایی می توان دو رویکرد را دنبال کرد. رویکرد اول، بررسی میزان جابه جایی مشترکین در سیکل های مشابه سال های متوالی است. به طور مثال، جدول ۳ جابجایی مشترکین شهر X را در سیکل های ۲ سال ۹۶ و ۹۷ نشان می دهد. در این جدول مشاهده می شود که تعداد ۳۴۶۵ مشترک در سیکل ۲ سال ۹۶ در خوشه کم مصرف بوده اند که در سیکل ۲ سال ۹۷ نیز در همان خوشه کم مصرف مانده اند. مشترکینی که با توجه به این جدول، نیاز به بازرسی دارند، ۶۵ مشترکی هستند که در سال ۹۶ در خوشه پر مصرف قرار داشته اند و در سال ۹۷ در خوشه کم مصرف جای گرفته اند. در مورد این مشترکین، احتمال خرابی کنتور یا مصرف غیر مجاز برق وجود دارد. همچنین، تعداد ۱۰۱ مشترکی که در سال ۹۶ در خوشه کم مصرف بوده اند و در سال ۹۷

جابه جایی مشترکین در دوره های مختلف بین این سه خوشه مورد تحلیل و بررسی قرار گیرد. این تحلیل، می تواند کمک شایانی در شناخت مشترکین با کنتورهای خراب یا مصارف غیر مجاز برق و همچنین، مشترکین مشکوک به استفاده تجاری از تعرفه برق خانگی را داشته باشد.

سپاسگزاری

در اینجا از همکاری مدیرعامل محترم شرکت توزیع برق جنوب استان کرمان، کارمندان محترم این شرکت، مراتب قدردانی به عمل می آید.

مراجع

- [۱] م. دزفولیان، "کاربرد تکنیک داده کاوی در تجزیه و تحلیل اطلاعات میزان مصرف و روند بدهی مشترکین برق (مطالعه موردی: شرکت توزیع نیروی برق استان ایلام)." اولین کنفرانس بین المللی تحقیقات پیشرفته در علوم، مهندسی و فناوری تهران، ۱۳۹۸.
- [۲] م. کجوری، "مدیریت مصرف انرژی در شبکه توزیع برق هوشمند با داده کاوی در زیر ساخت اندازه گیری پیشرفته." پنجمین بین المللی فناوری و مدیریت انرژی با رویکرد پیوند انرژی آب و محیط زیست تهران انجمن انرژی ایران کنفرانس ۱۳۹۷.
- [۳] ا. جعفرپور، "شناسایی ناهنجاری های مصرفی مشتریان در سیستم های توزیع برق با استفاده از خوشه بندی مبتنی بر تراکم." پایان نامه کارشناسی ارشد دانشگاه فردوسی مشهد، ۱۳۹۷.
- [4] A. Rajabi, L. Li, J. Zhang, J. Zhu, S. Ghavidel, and M. J. Ghadi, "A review on clustering of residential electricity customers and its applications," in 2017 20th International Conference on Electrical Machines and Systems (ICEMS), 2017, pp. 1-6: IEEE.
- [۵] ع. فریدونیان، "داده کاوی در انبار داده زیر ساخت اندازه گیری پیشرفته." دانشگاه تهران، ۱۳۹۵.
- [6] M. Bevilacqua, M. Braglia, and R. Montanari, "The classification and regression tree approach to pump failure rate analysis," Reliability Engineering & System Safety, vol. 79, no. 1, pp. 59-67, 2003.
- [۷] ع. فریدونیان، "شناسایی تغییرات در رفتار مصرفی مشترکین با استفاده از خوشه بندی فازی." پنجمین کنفرانس شبکه های هوشمند ایران تهران دانشگاه علم و صنعت، ۱۳۹۴.
- [8] J. Kwac, J. Flora, and R. Rajagopal, "Household energy consumption segmentation using hourly data," IEEE Transactions on Smart Grid, vol. 5, no. 1, pp. 420-430, 2014.
- [9] M. Kojury-Naftchali, A. Fereidunian, and H. Lesani, "Identifying susceptible consumers for demand response and energy efficiency policies by time-series analysis and supplementary approaches," in 2016 24th Iranian Conference on Electrical Engineering (ICEE), 2016 pp. 1130-1135: IEEE.
- [۱۰] ح. لسانی، "استفاده از الگوریتم خوشه بندی طیفی برای شناسایی الگوی مصرفی مشترکین و تعیین تعرفه های بهینه در شبکه برق." پنجمین منطقه ای سیرد کنفرانس ۱۳۹۵.

تحلیل بورس ایران با استفاده از اندیکاتورهای باندبولینگر و MACD

آذین ملایی درختنجانی^۱، مرجان کوچکی رفسنجانی^۲، ارشام برومند سعید^۳

^۱ گروه آموزشی علوم کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، azinmolaie@math.uk.ac.ir

^۲ دانشیار، گروه آموزشی علوم کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، kuchaki@uk.ac.ir

^۳ استاد، گروه آموزشی ریاضی محض، دانشگاه شهید باهنر کرمان، کرمان، arsham@uk.ac.ir

چکیده: در پیش‌بینی بازار بورس، تعداد زیادی مدل و تئوری توسط محققان به منظور بهبود عملکرد پیش‌بینی اجرا شده است. برخی از آن‌ها از یک مدل خطی منفرد و بعضی دیگر از یک مدل واحد غیرخطی استفاده می‌کنند، برخی دیگر از یک مدل ترکیبی خطی و غیرخطی استفاده می‌کنند. در سال‌های اخیر با توجه به پیشرفت‌های قابل توجه در پردازش اطلاعات توسط رایانه‌ها و نرم‌افزارهای کاربردی، انگیزه‌ی پژوهشگران در بکارگیری مدل‌های غیرخطی به طور چشم‌گیری افزایش یافته است و نظر متخصصان اقتصاد کلان و اقتصادسنجی را به خود جلب کرد و مطالعات متعددی در زمینه استفاده از آن در پیش‌بینی متغیرهای مختلف اقتصادی صورت گرفت. هدف این مقاله، پیش‌بینی روند قیمت سهام با استفاده از اندیکاتورهای میانگین متحرک همگرایی/ واگرایی و اندیکاتور باندبولینگر است. نتایج محاسباتی نشان داد میانگین متحرک همگرایی/ واگرایی از توانایی تحلیل بالاتری نسبت به اندیکاتور باندبولینگر برخوردار است و با استفاده از خط سیگنال، نمودار سیگنال‌های خرید و فروش را صادر می‌کند و سبب می‌شود سود بیشتری برای سرمایه‌گذاران فراهم گردد.

کلمات کلیدی: پیش‌بینی بازار سهام، تحلیل تکنیکال، اندیکاتور میانگین متحرک همگرایی/ واگرایی و اندیکاتور باندبولینگر.

۱- مقدمه

یادگیری عمیق هستند که قادر به شناسایی الگوهای پنهان در داده‌ها از طریق یک فرآیند خود یادگیری می‌باشند. پیش‌بینی در جایی که رابطه بین ورودی و خروجی از نظر ماهیت غیرخطی باشد بسیار سخت است. بازار سهام از نظر طبیعت بسیار بی‌ثبات است و نوسانات بازار باعث می‌شود این نوع پیش‌بینی بسیار دشوار باشد و پیش‌بینی ارزش سهام بورس یکی از کارهای چالش برانگیز سری زمانی مالی است.

روش یادگیری ماشین، که از الگوسازی و نظریه یادگیری محاسباتی الهام گرفته شده است، مطالعه و ساخت الگوریتم‌هایی را که می‌توانند براساس داده‌ها یادگیری و پیش‌بینی انجام دهند، بررسی می‌کند. الگوریتم‌های یادگیری ماشین توانمندی بالایی در مدل کردن مسائل پیچیده و سیستم‌های غیرخطی دارند و نیاز به فرضیاتی در مورد داده‌ها ندارند و غالباً به دقت بالاتری نسبت به مدل‌های اقتصادسنجی و آماری دست می‌یابند، پس برای مسئله پیش‌بینی

پیش‌بینی قیمت سهام یکی از مسائل مهم در بازارهای مالی است که توجه بسیاری از پژوهشگران دانشگاهی و کارشناسان این حوزه را در چند دهه گذشته به خود جلب نموده است. برای پیش‌بینی قیمت سهام چهار روش متعارف وجود دارد [3]: پیش‌بینی سری زمانی^۱، الگوریتم‌های یادگیری ماشین^۲، تجزیه و تحلیل بنیادین^۳ و تجزیه و تحلیل تکنیکال^۴.

تجزیه و تحلیل داده‌های سری‌های زمانی شامل دو کلاس از الگوریتم‌ها است: مدل‌های خطی و مدل‌های غیرخطی [۱]. مدل‌های مختلف خطی مدل خودهمبسته میانگین متحرک (ARMA)^۵، میانگین متحرک خود همبسته یکپارچه (ARIMA)^۶ و مدل فیلتر کالمن و تغییرات آن هستند. مدل‌های غیرخطی شامل روش‌هایی مانند مدل‌های سری زمانی مالی و الگوریتم‌های

⁴ Technical analysis

⁵ AutoRegressive Moving Average (ARMA)

⁶ AutoRegressive Integrated Moving Average (ARIMA)

¹ Time series prediction

² Machine learning algorithms

³ Fundamental analysis

قیمت سهام به عنوان روش‌هایی سودمند شناخته شده‌اند. یکی از معروف‌ترین مدل‌ها در این زمینه شبکه‌های عصبی مصنوعی هستند. موفقیت شبکه‌های عصبی در مطالعات مربوط به حوزه‌های مالی، نظر متخصصان اقتصاد کلان و اقتصادسنجی را به خود جلب کرد و مطالعات متعددی در زمینه استفاده از شبکه‌های عصبی مصنوعی در پیش‌بینی متغیرهای مختلف اقتصادی صورت گرفت. اولین بار هالبرت وایت^۷ [4] کاربرد شبکه‌های عصبی را برای پیش‌بینی قیمت سهام در پیش‌بینی‌های اقتصادی مطرح کرد. پس از مطالعه وایت، چندین تحقیق برای بررسی اثربخشی پیش‌بینی مدل‌های شبکه‌عصبی در بازارهای سهام انجام شد از جمله مطالعات قبلی می‌توان به روش [5,6] اشاره کرد.

تحلیل بنیادی، نوعی تحلیل سرمایه‌گذاری است که در آن ارزش سهام شرکت با تحلیل فروش، درآمد، سود و عوامل دیگر اقتصادی تخمین زده می‌شود. این روش برای پیش‌بینی طولانی مدت مناسب است. برخی از معیارهایی که برای محاسبه متغیرهای تحلیل بنیادی در پژوهش‌ها مورد استفاده قرار می‌گیرد عبارتند از بازده حقوق صاحبان سهام (RoE)^۸، سود هر سهم (EPS)^۹، نسبت قیمت به درآمد (P/E)^{۱۰}، نسبت قیمت به فروش (P/S)^{۱۱}، نسبت قیمت به ارزش دفتری (P/B)^{۱۲}، نسبت بدهی به حقوق صاحبان سهام (D/E)^{۱۳}، سرمایه بازار (MC)^{۱۴}، بازده دارایی (RoA)^{۱۵} [7,8].

اکثر سرمایه‌گذارها در بازارهای مالی از تحلیل تکنیکال برای خرید و فروش سهام استفاده می‌کنند. تحلیلگران تکنیکی معتقدند تغییرات آینده قیمت سهام می‌تواند با توجه به قیمت‌های پیشین تعیین شود [9]. برخی از شاخص‌های فنی مورد استفاده در تجزیه و تحلیل فنی عبارتند از میانگین متحرک ساده (SMA)^{۱۶}، میانگین متحرک نمایی (EMA)^{۱۷}، میانگین متحرک همگرای / واگرایی (MACD)^{۱۸}، شاخص مقاومت نسبی (RSI)^{۱۹} و حجم تعادل (OBV)^{۲۰} [8,10].

مطالعات بسیاری در مورد پیش‌بینی بازار سهام با استفاده از تجزیه و تحلیل فنی یا بنیادی از طریق تکنیک‌ها و الگوریتم‌های مختلف محاسبات نرم انجام شده است. برخی از دانشمندان نیز یک رویکرد ترکیبی که ترکیبی از تحلیل تکنیکال و متغیرهای تحلیل بنیادی است را برای پیش‌بینی قیمت آینده سهام به منظور بهبود رویکردهای موجود ارائه کرده‌اند به عنوان مثال آتودل^{۲۱} و همکارانش [11] اظهار داشتند که اقدامات قبلی در پیش‌بینی بورس اوراق بهادار، تنها متغیرهای تحلیل تکنیکی را درگیر کرده است و تأثیر متغیرهای تحلیل بنیادی تا حد زیادی نادیده گرفته شده است. آن‌ها یک رویکرد ترکیبی که ترکیبی از تحلیل تکنیکی و متغیرهای تحلیل بنیادی است را برای

پیش‌بینی قیمت آینده سهام به منظور بهبود رویکردهای موجود ارائه کردند و نتایج بدست آمده نشان دهنده پیشرفت چشمگیر رویکرد ترکیبی نسبت به استفاده از تنها متغیرهای تحلیل تکنیکی است و رویکرد این امکان را دارد که با ارائه پیش‌بینی دقیق‌تر سهام نسبت به رویکرد مبتنی بر تحلیل تکنیکی، کیفیت تصمیم‌گیری سرمایه‌گذاران در بورس سهام را ارتقا بخشد. بوهن^{۲۲} و لینگ^{۲۳} [12] تجزیه و تحلیل احساسات، تحلیل بنیادی و تجزیه و تحلیل فنی را با یکدیگر ترکیب کرد و مجموعه‌ای از الگوریتم‌های یادگیری ماشین را برای پیش‌بینی بلند مدت سهام مقایسه کرد. او روشی را برای استخراج ویژگی‌های متن توسعه داد و انتخاب ویژگی را با کمک یک عملکرد ارزیابی جدید اعمال کرد. نتایج نشان داد که انتخاب ویژگی در مقایسه با مدل‌های پایه می‌تواند اعتبار و عملکرد آزمون را تا حد زیادی بهبود بخشد. دای^{۲۴} و همکارانش [13] به پیش‌بینی بازده سهام با استفاده از شاخص‌های فنی جدید پرداختند. آن‌ها دریافتند که با روند بازده سهام در حالی که بازده سهام اصلی با تبدیل موجب کم صدا می‌شود می‌توان شاخص‌های جدید فنی ساخت که قدرت و عملکرد پیش‌بینی بازده سهام را بهبود بخشد و در آن شاخص‌های فنی جدید می‌توانند روند سری بازده سهام را به طور مستقیم منعکس کنند. همچنین سرمایه‌گذاران با میانگین واریانس مایل به استفاده از شاخص‌های فنی جدید برای پیش‌بینی بازده سهام و تخصیص اوراق بهادار هستند.

۲- تجزیه و تحلیل تکنیکال

به این دیدگاه مکتب تجزیه و تحلیل فنی نیز گفته می‌شود و معروف به چارتریست‌ها هستند [۲]. آن‌ها سعی می‌کنند از طریق یادگیری نمودارهایی که قیمت‌های تاریخی بازار و شاخص‌های فنی را به تصویر می‌کشند، بازار سهام را پیش‌بینی کنند. تحلیل تکنیکال به دو قسمت کلی الگوهای نموداری و اندیکاتورها^{۲۵} تقسیم می‌شود. از آنجایی که کمی‌سازی الگوهای نموداری کمی مشکل و پیچیده است استفاده کردن از اندیکاتورها برای ایجاد معاملات خودکار مرسوم‌تر است.

۲-۱- اندیکاتورها

اندیکاتورها توابع ریاضی هستند که براساس فرمول‌های خاص در جهت تحلیل بازار بوسیله ابزارهای گرافیکی ترسیم و مورد استفاده قرار می‌گیرند. یک اندیکاتور می‌تواند با محاسبات خود، نقاطی که احتمال برگشت قیمت یا تغییر روند وجود دارد را مشخص کند [14]. اندیکاتورها از لحاظ نوع حرکت و

¹⁷ Exponential Moving Average (EMA)

¹⁸ Moving Average Convergence Divergence rules (MACD)

¹⁹ Relative-Strength Index (RSI)

²⁰ On-Balance-Volume (OBV)

²¹ Ayodele

²² Bohn

²³ Ling

²⁴ Dai

²⁵ Indicator

⁷ White

⁸ Return on Equity (RoE)

⁹ Earnings Per Share (EPS)

¹⁰ Price /Earnings ratio (P/E)

¹¹ Price/Sales ratio (P/S)

¹² Price/Book ratio (P/B)

¹³ Debt/Equity ratio (D/E)

¹⁴ Market Capitalization (MC)

¹⁵ Return on Assets (RoA)

¹⁶ Simple-Moving Average (SMA)

است. هیستوگرام نمودار زمانی و گرافیکی است که براساس تفاوت دو خط سیگنال و MACD ایجاد می‌شود.

میانگین متحرک نمایی یک نوع میانگین متحرک وزنی است که وزن و اهمیت بیشتری به داده‌های قیمت اخیر می‌دهد. این وزن به صورت «فاکتور هموارسازی» بیان و به صورت زیر محاسبه می‌شود:

$$k = \frac{2}{1+\lambda} \quad (1)$$

در فرمول (۱) مقدار λ بازه زمانی مورد نظر را مشخص می‌کند. سپس این مقدار ضریب هموارسازی با مقدار EMA قبلی جمع می‌شود تا EMA فعلی به دست آید. به این صورت، EMA به داده‌های اخیر، وزن بیشتری خواهد داد؛ در حالی که میانگین متحرک ساده وزن برابری برای همه مقادیر در نظر می‌گیرد [8].

۴- اندیکاتور باند بولینگر

باندبولینگر از باند بالایی، میانی و پایینی تشکیل شده است که توسط جان بولینگر به وجود آمده است. باند وسط یا میانی نمودار بولینگر، در حقیقت همان اندیکاتور میانگین متحرک ساده است. در تحلیل تکنیکال از میانگین متحرک ساده جهت پیش‌بینی رفتار سهام و بررسی حمایت و مقاومت با ارزیابی تحرکات قیمت دارای استفاده می‌کنند. این شاخص، داده‌های یک نمودار را با محاسبه میانگینی که دائماً به‌روز می‌شود، هموار می‌سازد و پیش‌بینی رفتار داده‌ها را راحت‌تر می‌کند. با توجه به فرمول (۲) از طریق جمع‌آوری جدیدترین قیمت‌های بسته شده سهام (پایانی) و سپس تقسیم آن بر تعداد "n" دوره‌ها میانگین متحرک محاسبه می‌شود [15].

$$SMA = \frac{A_1 + A_2 + \dots + A_n}{n} \quad (2)$$

که در آن A_n میانگین n امین دوره و n تعداد دوره‌های زمانی است. اندیکاتور باندبولینگر، با بررسی قیمت و نوسانات آن در بازه زمانی گذشته، یک محدوده‌ای را بدست می‌آورد که قیمت به احتمال زیاد در آن محدوده در آینده نوسان خواهد کرد. در حقیقت، به ما می‌گوید که در حال حاضر بازار آرام است یا پرشور! این اندیکاتور از دو باند تشکیل شده است که قیمت را محصور کرده است. زمانی که بازار آرام باشد، محدوده این اندیکاتور تنگ‌تر خواهد شد. یعنی نوسانات قیمت کم است و حد بالا و حد پایین نمودار باندبولینگر به هم نزدیک خواهند بود. اما زمانی که بازار پرتلاطم و دارای نوسانات زیاد باشد، حد بالا و حد پایین این اندیکاتور فاصله زیادی از هم گرفته و به اصطلاح دهانه آن باز خواهد شد. بسیاری از معامله‌گرها بر این باورند که هر چه حرکات قیمت به باند بالایی نزدیک‌تر باشد، بازار بیشتر در معرض بیش خرید است، و هرچه حرکات قیمت به باند پایینی نزدیک‌تر باشد، بازار بیشتر در معرض بیش‌فروش است.

۵- بررسی روش پیشنهادی

در این بخش ابتدا به معرفی و پیش‌پردازش دادگان پرداخته خواهد شد. سپس اندیکاتورهای MACD و باندبولینگر را بر روی مجموعه داده شاخص کل

رفتار به دو دسته پیشرو^{۲۶} و تأخیری^{۲۷} تقسیم می‌شوند. اندیکاتورهای پیشرو به‌طور معمول همراه با نوسانات قیمت و اندیکاتورهای تأخیری معمولاً با تأخیر، پس از حرکت قیمت و جابجایی بازار، هشدار و پیش‌بینی لازم را اعلام می‌کنند. اندیکاتورهای پیشرو اغلب مواقع سیگنال‌های بیشتری نسبت به اندیکاتورهای تأخیری ایجاد می‌کنند و به همین دلیل احتمال خطا در آن‌ها بیش‌تر است. در زمینه تحلیل تکنیکال، یک اندیکاتور از طریق محاسبات ریاضی که داده‌های آن از قیمت و حجم معاملات یک دارای بدست آمده محاسبه و روی نمودار رسم می‌شود. هدف از به‌کار بردن اندیکاتورها در تحلیل تکنیکال، یافتن فرصت‌های کسب سود است.

اندیکاتورها از نظر نوع نمایش یا مدل محاسبات به ۴ دسته روند، اسیلاتور یا نوسان‌گر، حجم و اندیکاتور بیل ویلیامز^{۲۸} تقسیم می‌شوند.

۲-۲- موارد استفاده از اندیکاتورها

به‌طور کلی از اندیکاتورها در بازارهای مالی براساس تغییرات گذشته قیمت، به سه شکل هشدار، پیش‌بینی و تایید استفاده می‌شود.

- هشدار: اندیکاتورها در برخی مواقع قبل از تغییر روند یا هم‌زمان با آن، علائم بازگشت روند را نشان می‌دهند. پس یکی از کاربردهای اندیکاتور، اعلام هشدارهای مناسب تغییر روند و جهت حرکت قیمت است.
 - پیش‌بینی: اندیکاتورها برای پیش‌بینی قیمت مناسب برای ورود به سهم یا تایید تحلیل‌ها استفاده می‌شوند. بنابراین می‌توان گفت آن‌ها در افزایش سرعت و دقت تحلیل به تحلیلگر کمک بسیاری می‌کنند.
 - تایید: یکی از موارد استفاده از اندیکاتورها گرفتن تایید از روند صحیح و جهت آن است. و زمانی مورد استفاده قرار می‌گیرد که تحلیلگر براساس داده‌های تکنیکالی یا بنیادی، جهت و قیمت ورودی مناسب در بازار را پیش‌بینی کرده و از اندیکاتور برای تایید گرفتن استفاده می‌کند.
- در این مقاله هدف تحلیل شاخص کل بورس ایران با استفاده از اندیکاتورهای مهم و کاربردی میانگین متحرک همگرای/ واگرایی و باندبولینگر است که در ادامه به معرفی آن‌ها پرداخته خواهد شد.

۳- اندیکاتور میانگین متحرک همگرای/ واگرایی

اندیکاتور MACD شامل سه عامل خط MACD، سیگنال و هیستوگرام است. این اندیکاتور معمولاً بازه زمانی ۱۲ و ۲۶ روزه را بررسی می‌کند و رابطه بین این دو خط را به صورت همگرا (هنگامی که خطوط به سمت یکدیگر کشیده می‌شوند) یا واگرا (هنگامی که خطوط از هم دور می‌شوند) بیان می‌کند. اندیکاتور MACD با کم کردن دو میانگین متحرک نمایی ۱۲ و ۲۶ روزه خط اصلی MACD را ایجاد می‌کند سپس از این خط برای محاسبه EMA دیگری استفاده می‌شود که به آن خط سیگنال می‌گویند که یک میانگین متحرک نمایی ۹ روزه و صادرکننده سیگنال‌های خرید و فروش

²⁸ Billwilliams

²⁶ Leading

²⁷ Lagging

بورس اوراق بهادار پیاده‌سازی کرده و در انتها به تحلیل هر یک از آنها پرداخته خواهد شد. چارچوب کلی روش پیشنهادی شامل گام‌های اساسی زیر می‌باشد:

۱. آماده کردن داده‌ها
۲. پیش پردازش داده‌ها
۳. پیاده‌سازی مدل پیشنهادی
۴. پیش‌بینی و تحلیل عملکرد مدل

۲-۵- پیش‌پردازش داده‌ها

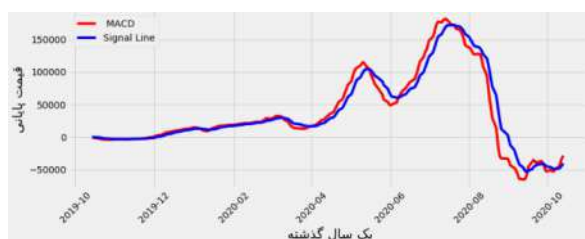
در این مرحله ستون Symbol را که فقط شامل عبارت "شاخص کل" است را از مجموعه داده حذف کرده تا مرحله پیش‌بینی به راحتی بر روی داده‌های عددی صورت گیرد.

۳-۵- پیاده‌سازی و تحلیل اندیکاتورها

با استفاده از مجموعه داده اولیه که همان شاخص کل بورس است به پیاده‌سازی و تحلیل اندیکاتورهای میانگین متحرک همگرای/ واگرایی و باندبولینگر پرداخته خواهد شد.

۱-۳-۵- پیاده‌سازی و تحلیل اندیکاتور MACD

در نمودار (۲) خط سیگنال و مکدی را برای بازه‌ی یک سال رسم کرده‌ایم.



نمودار (۲): نمودار خط سیگنال و مکدی در طی یک سال

در تحلیل این نمودار باید توجه کرد که هنگامی خط عبور MACD بالاتر از خط سیگنال باشد، سیگنال خرید صادر و هنگامی که خط عبور MACD پایین‌تر از خط سیگنال باشد، سیگنال فروش صادر می‌شود. در نمودار (۳) این دو خط با نمودار قیمت مقایسه می‌شود.

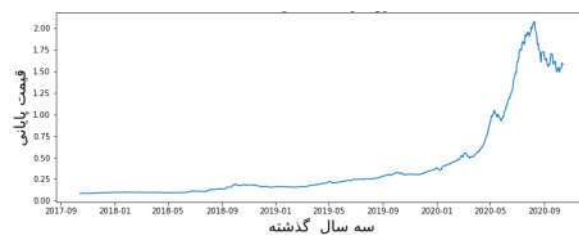


نمودار (۳): مقایسه خط سیگنال و مکدی با قیمت پایانی

همان‌طور که در نمودار (۳) مشاهده می‌شود خط MACD اغلب بین خط مرکزی (خط صفر) در نوسان است و در نقطه صفر متمرکز است. اگر خط MACD بالای خط مرکزی باشد در سیگنال مثبت قرار دارد که این اتفاق در بازه زمانی ۲۰۲۰/۶ تا ۲۰۲۰/۸ رخ داده است و با مقایسه آن با قیمت متوجه این امر نیز خواهیم شد زیرا قیمت نیز در این بازه به اوج خود رسیده است. این اتفاق زمانی می‌افتد که میانگین متحرک تصاعدی ۱۲ روزه بالاتر از میانگین متحرک تصاعدی ۲۶ روزه باشد. زمانی که خط MACD پایین‌تر از خط صفر یا خط مرکزی باشد سیگنال منفی است. زمانی سیگنال منفی است که میانگین متحرک تصاعدی ۱۲ روزه پایین‌تر از میانگین متحرک تصاعدی ۲۶ روزه باشد.

۱-۵- دادگان اولیه

داده‌های اولیه این پژوهش شامل شاخص کل بازار بورس تهران طی ۲۸۶۲ روز تجارت از سال‌های ۲۰۰۸/۱۲/۰۶ تا ۲۰۲۰/۱۰/۱۴ است. در این پژوهش از شاخص کل بازار برای انجام پیش‌بینی استفاده شده است زیرا شاخص بازار نماد کاملی از تمام شرکت‌های موجود در بازار است و کارایی بازار براساس آن تعیین می‌شود. هر داده شامل ۱۰ ویژگی عبارتند از نماد، تاریخ، اولین قیمت، بیشترین قیمت، کمترین قیمت، قیمت پایانی، حجم، ارزش، تعداد معاملات و قیمت روز گذشته سهام می‌باشد که در جدول (۱) به منظور سهولت در انجام کار عنوان انگلیسی هر یک از ویژگی‌ها قید شده است. نمودارهای قیمت سهام در بازار از دو مؤلفه قیمت و زمان تشکیل شده‌اند و نشان‌دهنده مکان هندسی نقاطی است که از دو عامل قیمت و زمان تشکیل شده‌اند. به منظور درک بهتر از داده‌ها، در نمودار (۱) روند کلی رشد قیمت شاخص کل سهام در طی سه سال اخیر در محیط پایتون نمایش داده شده است.



نمودار (۱): روند رشد قیمت شاخص کل در طی سه سال

جدول (۱): عنوان هر یک از ویژگی‌ها

فارسی	انگلیسی
نماد	Symbol
تاریخ	Date
اولین قیمت	Open Price
بیش‌ترین قیمت	High Price
کم‌ترین قیمت	Low Price
قیمت پایانی	Final Price
حجم	Volume
ارزش	Value
تعداد معاملات	Number of transactions
قیمت دیروز	Yesterday's price

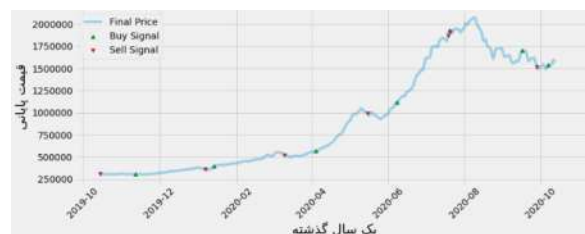
۶- نتیجه

در این پژوهش توانستیم شاخص کل بورس ایران را با استفاده از اندیکاتورهای میانگین متحرک همگرایی/ واگرایی و باندبولینگر پیاده‌سازی و تحلیل کنیم. نتایج محاسباتی نشان داد که اندیکاتور میانگین متحرک همگرایی/ واگرایی از توانایی تحلیل بالایی برخوردار است و با استفاده از خط سیگنال، نمودار سیگنال‌های خرید و فروش را صادر می‌کند و سبب می‌شود سود بیشتری برای سرمایه‌گذاران فراهم گردد.

مراجع

- [۱] موسوی چهارمی، یگانه و غلامی، الهام، مدل ترکیبی شبکه عصبی با الگوی ARIMA جهت پیش بینی مالیات بر ارزش افزوده بر مصرف بنزین در ایران، فصلنامه پژوهش‌های اقتصادی (رشد و توسعه پایدار)، شماره دوم، صفحات ۹۹-۱۱۶، تابستان ۱۳۹۵.
- [۲] مکیان، سید نظام الدین و موسوی، فاطمه السادات، پیش‌بینی قیمت سهام شرکت فرآورده‌های نفتی پارس با استفاده از شبکه عصبی و روش رگرسیونی مطالعه موردی قیمت سهام شرکت فرآورده‌های نفتی پارس، فصلنامه مدل‌سازی اقتصادی سال ششم، شماره ۲، صفحه ۱۰۵-۱۲۱، ۱۳۹۱.
- [3] Prasanna, S., Ezhilmaran, D., "An analysis on Stock Market Prediction using Data Mining Techniques", Int. J. Comput. Sci. Eng. Technol, vol. 4, no. 2, pp. 49-51, 2013.
- [4] White, H., "Economic prediction using neural networks: the case of IBM daily stock returns", in Proceedings of The IEEE International Conference on Neural Networks, pp. 451-458, 1988.
- [5] Kimoto, T., Asakawa, K., Yoda, M., and Takeoka, M., "Stock market prediction system with modular neural networks", Proc. IEEE Int. Jt. Conf. Neural Networks, pp. 1-16, 1990.
- [6] ichi Kamijo, K., Tanigawa, T., "Stock price pattern recognition--A recurrent neural network approach", Proc. IEEE Int. Jt. Conf. Neural Networks, pp. 215-221, 1990.
- [7] Renu Isidore, R., Christie, P., "Fundamental Analysis Versus Technical Analysis-a Comparative Review", Int. J. Recent Sci. Res., vol. 9, no. 1, pp. 23009-23013, 2018.
- [8] Nti, I. K., Adekoya, A.F., Weyori, B. A., "A systematic review of fundamental and technical analysis of stock market predictions", Springer, pp. 3007-3057, 2019.
- [9] Sureshkumar, K.K., Elango, N.M., "An Efficient Approach to Forecast Indian Stock Market Price and their Performance Analysis", Int. J. Comput. Appl., vol. 34, no. 5, pp. 44-49, 2011.
- [10] Anbalagan, T., Maheswari, S. U., "Classification and prediction of stock market index based on Fuzzy Metagraph", Procedia Computer Science, vol. 47, no. C, pp. 214-221, 2015.
- [11] Ayodele, A. A., Charles, A. K., Marion, A. O., Sunday, O. O., "Stock Price Prediction using Neural Network with Hybridized Market Indicators", J. Emerg. Trends Comput. Inf. Sci., vol. 3, no. 1, pp. 1-9, 2012.
- [12] Bohn, T. A., Ling, C., "Improving Long Term Stock Market Prediction with Text Analysis Recommended Citation", 2017.
- [13] Dai, Z., Zhu, H., Kang, J., "New technical indicators and stock returns predictability", Int. Rev. Econ. Financ., vol. 71, no. August 2020, pp. 127-142, 2021.

این بدان معنی است که قیمت‌های اخیر (در ۱۲ روز اخیر) نشان‌دهنده یک حرکت نزولی بیشتر در مقایسه با میانگین قیمت متوسط ۲۶ روزه است؛ بنابراین میانگین متحرک تصاعدی وزن کمتری از قیمت‌های اخیر است. اندازه حرکت قیمت‌ها در این میانگین کمتر از قبل است. پس از آن با استفاده از خط سیگنال، نمودار سیگنال‌های خرید و فروش را براساس قیمت پایانی رسم کرده و نتایج نمودار (۴) حاصل گردید.

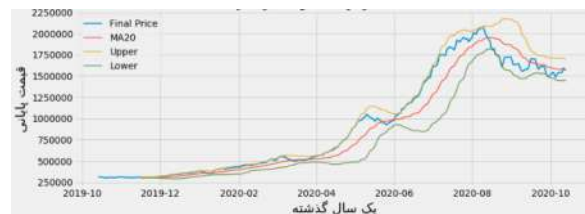


نمودار (۴): نمودار سیگنال‌های خرید و فروش براساس قیمت

همان‌طور که در نمودار بالا مشخص است در جاهایی که قیمت اوج می‌گیرد این شاخص سیگنال فروش را صادر کرده است.

۵-۳-۲- پیاده‌سازی و تحلیل اندیکاتور MACD

بسیاری از نرم افزارهای تحلیل تکنیکال و کار با نمودارهای قیمتی، به صورت پیش‌فرض دوره زمانی این اندیکاتور را بر روی ۲۰ قرار می‌دهند که حالت خوب و نرمال برای اکثر تریدرها است. اما هیچ محدودیتی در استفاده از دوره زمانی ۲۰ در این اندیکاتور وجود ندارد. در نمودار (۵) این نمودار را برای یک سال اخیر رسم کرده و نتایج به شرح زیر می‌باشد.



نمودار (۵): مقایسه نمودار قیمت و باندبولینگر در یک سال

همان‌طور که در نمودار (۵) مشاهده می‌شود تا قبل از ۲۰۲۰/۴ محدود به باند بالا و پایین به هم نزدیک بوده است و این یعنی بازار آرام و نوسانات قیمت کم است. اما پس از آن حد بالا و حد پایین این اندیکاتور فاصله زیادی از هم گرفته و به اصطلاح دهانه‌ی آن باز شده است که بیانگر این است بازار پر تلاطم و دارای نوسانات زیاد است. در بازه‌ی تقریبی ۲۰۲۰/۴ تا ۲۰۲۰/۸ حرکات قیمت به باند بالایی نزدیک‌تر است و این یعنی بازار در معرض بیش‌تر خرید قرار دارد و بهتر است در این بازه فروش صورت گیرد و پس از آن حرکات قیمت به باند پایینی نزدیک‌تر است و این یعنی بازار در معرض بیش‌تر فروش قرار دارد و بهتر است در این بازه خرید صورت گیرد.

- [14] Murphy, John J., "*Technical analysis of the financial markets: a comprehensive guide to trading methods and applications (2nd ed.)*". New York [u.a.]: New York Inst. of Finance. ISBN 0-7352-0066-1.
- [15] Anbalagan, T and Maheswari, S.U., "*Classification and prediction of stock market index based on Fuzzy Metagraph*," Procedia Computer Science, vol. 47, no. C. pp. 214–221, 2015.

تشخیص اسکیزوفرنی بر اساس سیگنال الکتروانسفالوگرام با استفاده از یادگیری عمیق

مریم الهیاری^۱، ولی درهمی^۲، فاطمه جمشیدی^۳

^۱ دانشکده مهندسی کامپیوتر، پردیس فنی و مهندسی، دانشگاه یزد، یزد، ایران، maryam.allahyari.sa@gmail.com

^۲ دانشکده مهندسی کامپیوتر، پردیس فنی و مهندسی، دانشگاه یزد، یزد، ایران، vderhami@yazd.ac.ir

^۳ گروه مهندسی برق، دانشکده مهندسی، دانشگاه فسا، فسا، ایران، jamshidi@fasau.ac.ir

چکیده: اسکیزوفرنی یک ناهنجاری در مغز است که در آن افراد واقعیت را غیر طبیعی تفسیر می‌کنند. این اختلال روانی با علائم رفتاری مانند توهم و بی‌نظمی گفتار مشخص می‌شود. سیگنال الکتروانسفالوگرام (EEG) اختلالات مغزی را نشان می‌دهد و به‌طور گسترده برای مطالعه بیماری‌های مغزی استفاده می‌شود. هدف این مقاله تشخیص خودکار اسکیزوفرنی از روی سیگنال EEG است. روش متداول در پژوهش‌ها، استخراج دستی ویژگی‌ها از سیگنال EEG است. از آنجا که الگوریتم‌های یادگیری عمیق توانایی استخراج خودکار ویژگی‌های مهم و طبقه‌بندی آنها را دارند، در این پژوهش به‌منظور استخراج ویژگی‌های مفیدتر، سیگنال EEG به یک شبکه عصبی عمیق بازگشتی کانولوشنی یازده لایه اعمال شده است. سیگنال‌های EEG جمع‌آوری شده در انیستیتو ورشو از ۱۴ فرد سالم و ۱۴ بیمار اسکیزوفرنی، در اینجا مطالعه شده است. مقدار میانگین معیارهای ارزیابی درستی مدل شامل **Accuracy**، **Sensitivity**، **Specificity** و **PPV** برای مدل پیشنهادی به ترتیب برابر ۹۸/۷۹٪، ۹۸/۷۳٪، ۹۸/۸۶٪ و ۹۹/۰۶٪ به‌دست آمد که بهبود عملکرد مدل پیشنهادی برای طبقه‌بندی بیماران اسکیزوفرنی و افراد سالم را در مقایسه با مدل‌های قبلی تایید می‌کند. مدل ارائه شده می‌تواند به‌عنوان یک ابزار تشخیصی به پزشکان برای تشخیص مراحل اولیه اسکیزوفرنی کمک کند.

کلمات کلیدی: اسکیزوفرنی، الکتروانسفالوگرام، شبکه‌های عصبی عمیق، یادگیری عمیق

۱- مقدمه

ژنتیک، انجام تحقیقات رفتاری، ضبط و بررسی نوار مغزی یا الکتروانسفالوگرام (EEG)^۱ و استفاده از تصویربرداری پیشرفته برای بررسی ساختار و عملکرد مغز، علل بیماری را کشف می‌کنند. این رویکردها نوید درمان‌های جدید و موثرتری را می‌دهند. در حال حاضر درمان قطعی برای بیماری اسکیزوفرنی وجود ندارد. اگر بیماری اسکیزوفرنی به موقع تشخیص داده نشود و درمان آن در مراحل اولیه صورت نگیرد، توانایی‌های رفتاری فرد شدیدتر آسیب می‌بیند. از این رو تشخیص زود هنگام اسکیزوفرنی نقش مهمی در درمان و محدود کردن اثرات بیماری دارد. تصویربرداری رزونانس مغناطیسی (MRI)^۲ نوعی اسکن است که از میدان‌های مغناطیسی قوی و امواج رادیویی برای تولید تصاویر دقیق از داخل بدن استفاده می‌کند. تصویربرداری رزونانس مغناطیسی تقریباً برای بررسی هر قسمتی از بدن استفاده می‌شود، از جمله: مغز و نخاع، استخوان‌ها و مفاصل، سینه‌ها، قلب

بیماری اسکیزوفرنی^۱ ناهنجاری مغزی است که ویژگی‌های رفتاری مانند تفکر، ادراک، گفتار و ... را مختل می‌کند. این اختلال روانی با علائم رفتاری مانند توهم، بی‌نظمی در گفتار، ناتوانی در درک واقعیت و انجام رفتارهای نامعقول همراه است. عارضه روانی اسکیزوفرنی به دلیل فقدان نشانه‌های بالینی در مقایسه با سایر بیماری‌های مغزی مانند صرع و تشنج کمتر شناخته شده است. در حالی که ۳ تا ۷ درصد مردم در سراسر دنیا به این بیماری مبتلا و از عوارض آن رنج می‌برند. روش متداول برای تشخیص بیماری اسکیزوفرنی این است که روان‌پزشک با صحبت کردن با بیمار، عارضه را در وی تشخیص می‌دهد. روشی است که این روش تشخیص انسانی می‌تواند خطای زیادی داشته باشد. متخصصان با مطالعه

^۲ Magnetic Resonance Imaging (MRI)

^۱ Schizophrenia (SZ)
^۲ Electroencephalogram (EEG)

و عروق خونی، اندام‌های داخلی مانند کبد، رحم یا غده پروستات. اگرچه اختلالات ساختاری متمایزی در بیماران مبتلا به اسکیزوفرنی رخ می‌دهد، تشخیص اسکیزوفرنی با تصویربرداری رزونانس مغناطیسی همچنان چالش برانگیز است. در سال ۲۰۲۰، جی‌هون و همکاران با بررسی ۵ مجموعه داده MRI شناخته شده از الگوریتم مبتنی بر یادگیری عمیق برای ترسیم ویژگی‌های ساختاری بیماری اسکیزوفرنی و ارائه اطلاعات تشخیصی در شرایط بالینی استفاده کردند [1].

سلول‌های مغز از طریق امواج الکتریکی با یکدیگر ارتباط برقرار می‌کنند و نوار مغز با سنجش این ارتباط می‌تواند به تشخیص برخی از بیماری‌ها از جمله تشنج و صرع کمک کند. سیگنال EEG الگوهای ارتباطی بین امواج الکترونیکی مغز را ثبت می‌کند. برای این منظور الکترودهایی از جنس نقره به سر متصل می‌شوند. اتصال الکترودها به سر به دو صورت تهاجمی و غیرتهاجمی انجام می‌شود. در روش غیرتهاجمی، فعالیت الکتریکی مغز از طریق نصب الکترودهای سطحی بر روی سر ثبت می‌شود. اثر الکتریکی فعالیت نرون‌های مغز از طریق الکترودها به دستگاه منتقل می‌شود و پس از تقویت و حذف نویز به صورت سیگنال زمانی ثبت و نمایش داده می‌شود. پزشک یا متخصص علوم اعصاب می‌تواند سیگنال ثبت شده را به‌طور مستقیم یا پس از پردازش کامپیوتری تجزیه و تحلیل کند. در روش‌های تهاجمی ضبط نوار مغزی که با نام نوار مغزی اینتراکرانیا (iEEG) شناخته می‌شوند، از سیستم مانیتورینگ اینتراکرانیا استفاده می‌شود تا متخصص در حین جراحی به داده‌های بیولوژیکی کامل‌تری که از طریق ضبط عادی از روی پوست سر قابل استخراج نیستند، دسترسی یابد. سیگنال EEG فعالیت‌های الکتریکی مغز را نشان می‌دهد و به دلیل داشتن اطلاعات با وضوح زمانی بالا، کاربرد گسترده‌ای در تشخیص اختلالات مغزی و پژوهش‌های مرتبط با آن دارد [2, 3]. از آنجا که امواج مغزی فرد مبتلا به بیماری اسکیزوفرنی با افراد سالم متفاوت است، با تحلیل امواج مغزی می‌توان بیماری را با دقت قابل قبولی تشخیص داد [4].

یادگیری ماشین^۲ مطالعه الگوریتم‌هایی است که می‌توانند به‌طور خودکار از طریق تجربه و با استفاده از داده‌ها بهبود یابند. این الگوریتم‌ها به عنوان بخشی از هوش مصنوعی تلقی می‌شوند. الگوریتم‌های یادگیری ماشین یک مدل را بر اساس داده‌های نمونه، معروف به «داده‌های آموزشی»، به منظور پیش‌بینی یا تصمیم‌گیری بدون برنامه‌ریزی صریح، می‌سازند. الگوریتم‌های یادگیری ماشین در طیف گسترده‌ای از برنامه‌ها مانند پزشکی، فیلترینگ ایمیل، تشخیص گفتار و بینایی رایانه‌ای استفاده می‌شوند. یادگیری عمیق^۴ بخشی از خانواده گسترده‌تری از روش‌های یادگیری ماشینی است که بر اساس شبکه‌های عصبی مصنوعی با یادگیری نمایشی ساخته شده است. یادگیری می‌تواند تحت نظارت، نیمه نظارت یا بدون نظارت باشد. معماری‌های یادگیری عمیق مانند شبکه‌های عصبی عمیق^۵، یادگیری تقویتی عمیق^۶، شبکه‌های عصبی بازگشتی^۷ و شبکه‌های عصبی

کانولوشن^۸ در زمینه‌هایی از جمله بینایی کامپیوتر^۹، تشخیص گفتار^{۱۰}، پردازش زبان طبیعی^{۱۱}، بیوانفورماتیک، طراحی دارو و پزشکی استفاده شده است و در برخی موارد عملکرد این الگوریتم‌ها از عملکرد متخصصان انسانی فراتر رفته است. از آنجا که ابزارهای هوش مصنوعی برای تحلیل امواج مغزی به منظور تشخیص بیماری اسکیزوفرنی می‌تواند کارآمد باشد، در این مقاله، تشخیص سریع و صحیح بیماری اسکیزوفرنی با استفاده از تحلیل سیگنال EEG مبتنی بر یادگیری عمیق بررسی می‌شود. ارائه‌ی یک طرح خودکار برای تشخیص زودهنگام و با خطای کم بیماری اسکیزوفرنی از روی سیگنال EEG در تشخیص و درمان به‌موقع این بیماری اثربخش است. برتری مدل پیشنهادی نسبت به کارهای گذشته در استفاده از روشی مبتنی بر یادگیری عمیق است که به استخراج ویژگی‌های بهتر از داده‌ها و بهبود شرایط یادگیری و عملکرد مدل منجر می‌شود. مقاله حاضر شامل بخش‌های زیر است: مرور اجمالی بر کارهای گذشته و روش‌های استفاده شده در این زمینه، معرفی داده مورد مطالعه در این پژوهش و نرمال‌سازی داده و آماده‌سازی آن برای ورود به مدل پیشنهادی، معرفی مدل پیشنهادی مبتنی بر یادگیری عمیق که مدل ترکیبی از شبکه‌های بازگشتی و کانولوشنی است، ارائه نتایج مدل پیشنهادی و در نهایت بخش آخر به نتیجه‌گیری می‌پردازد.

۲- مروری بر کارهای پیشین

در سال‌های اخیر متخصصین هوش مصنوعی و پزشکان با همکاری یکدیگر گام‌های بزرگی در تشخیص و پیش‌بینی بیماری‌های مختلف مغزی با استفاده از تجزیه و تحلیل سیگنال‌های EEG برداشته‌اند. الگوریتم‌های یادگیری ماشین و روش‌های یادگیری عمیق از جمله ابزارهایی هستند که استفاده شده‌اند. بستانی و همکاران در سال ۲۰۰۹ نمونه‌های سیگنال EEG ضبط شده در مرکز تحقیقات بالینی روان‌پزشکی پرت در استرالیا غربی شامل ۱۳ فرد مبتلا به اسکیزوفرنی و ۱۸ فرد سالم بررسی نمودند؛ آنها با پیاده‌سازی روش‌های باند پاور^{۱۲}، رگرسیون خودکار^{۱۳} و بعد فراکتال^{۱۴} ویژگی استخراج نموده و با به‌کارگیری الگوریتم‌های مختلف یادگیری ماشین شامل LDA^{۱۵}، BDLDA و Adaboost به ترتیب به دقت‌های ۸۷/۵۱٪ و ۸۵/۳۶٪ و ۸۵/۴۱٪ برای تشخیص فرد بیمار از فرد سالم دست یافتند [5]. سیری واتسان و همکاران در سال ۲۰۱۹ برای تشخیص بیماری اسکیزوفرنی به جای سیگنال‌های EEG، از تصاویر رزونانس مغناطیسی مغز (MRI) استفاده کردند؛ آنها ۱۴۴ نمونه که ۶۹ تای آنها مبتلا به اختلال اسکیزوفرنی بودند را مطالعه کردند و سه مدل‌های مبتنی بر یادگیری ماشین شامل ماشین بردار پشتیبان^{۱۶} (SVM)، رگرسیون منطقی^{۱۷} و جنگل تصادفی^{۱۸} و همچنین شبکه‌ی عصبی کانولوشنی (CNN) که مدلی مبتنی

^{۱۰} Speech Recognition

^{۱۱} Natural Language Processing (NLP)

^{۱۲} Band Power

^{۱۳} Autoregressive Model (AR)

^{۱۴} Fractal Dimension

^{۱۵} Latent Dirichlet Allocation (LDA)

^{۱۶} Support Vector Machine (SVM)

^{۱۷} Logistic Regression

^{۱۸} Random Forest

^۱ Intracranial Electroencephalogram (iEEG)

^۲ Accuracy

^۳ Machine Learning (ML)

^۴ Deep Learning

^۵ Deep Neural Network

^۶ Deep Reinforcement Learning

^۷ Recurrent Neural Networks (RNN)

^۸ Convolutional Neural Network (CNN)

^۹ Computer Vision

بر یادگیری عمیق است را پیاده‌سازی کردند؛ نتایج نشان داد که مدل CNN با دقت ۹۴ درصد بهترین مدل پیاده‌سازی شده در این مطالعه بوده است [6].

یان و همکاران در سال ۲۰۱۹، از تصویرهای MRI عملکردی (fMRI) مغزی به‌عنوان داده برای تشخیص اسکیزوفرنی استفاده کردند؛ آنها از شبکه‌های عصبی بازگشتی (RNN) برای تفکیک ۵۵۸ بیمار اسکیزوفرنی از ۵۴۲ فرد سالم استفاده نمودند؛ آنها برای بهبود نتایج، استراتژی leave-on-IC-out looping را به‌کار بردند تا مولفه‌های مستقلی که نقش بیشتری در تفکیک دارند را شناسایی و استفاده کنند؛ دقت‌های به‌دست آمده در این مطالعه ۸۳/۲٪ و ۸۰/۲٪ درصد به‌ترتیب برای روش‌های ارزیابی Multi-site pooling و Leave-one-site-out بود [7]. لی و همکاران در سال ۲۰۱۹ دو مدل شبکه کانولوشن عمیق ۱۱ لایه‌ای را برای تشخیص خود کار بیماری اسکیزوفرنی معرفی و پیاده‌سازی کردند؛ آنها سیگنال‌های EEG ضبط شده از ۱۴ فرد سالم و ۱۴ فرد مبتلا به اسکیزوفرنی در انستیتو وروشو در لهستان را استفاده کردند و در نهایت به دقت‌های قابل توجه ۹۸/۰۷٪ و ۸۱/۲۶٪ دست یافتند [2]. سانیل کومار و همکاران در سال ۲۰۲۰ با استخراج ویژگی از داده‌های EEG ضبط شده در وروشو لهستان [2] و انتخاب بهترین ویژگی‌ها با بهره‌گیری از روش‌های بهینه‌یابی سیاه چاله^۲ (BH)، جست‌وجوی کرم^۳ (GS)، فلورا مصنوعی^۴ (AF) و جست‌وجوی میمون^۵ (MS) و استفاده از مدل SVM-RBF به دقت ۹۲/۱۷٪ در تشخیص بیماری اسکیزوفرنی دست یافتند [8].

کریشن و همکاران در سال ۲۰۲۰ روش تحلیل چند متغیره سیگنال‌های EEG را برای تشخیص اسکیزوفرنی به‌کار بردند؛ آنها با استفاده از روش تجزیه حالت تجربی چند متغیره^۶ (MEMD) سیگنال‌های EEG را به سیگنال‌های توابع مد داخلی (IMF) تجزیه نمودند؛ سپس از ۵ معیار آنتروپی برای تفکیک سیگنال افراد بیمار و سالم به روش یادگیری ماشین SVM استفاده کردند؛ بهترین دقت به‌دست آمده ۹۳ درصد بود [9]. اسلان و همکاران در سال ۲۰۲۰، یک شبکه CNN از نوع vgg-16 را برای تشخیص خودکار اسکیزوفرنی ارائه دادند؛ آنها، به جای استفاده از سیگنال‌های EEG به‌عنوان ورودی، نخست داده EEG را با استفاده از تبدیل فوریه به فضای دو بعدی فرکانس-زمان بردند؛ سپس با استفاده از یک شبکه عصبی کانولوشنال با معماری vgg-16 به دقت‌های ۹۵ و ۹۷ درصد رسیدند؛ در این مقاله دو دسته داده استفاده شد، دسته اول را مولفان ثبت کردند و دسته دوم داده‌ی موسسه وروشو در لهستان [2] بود [10]. شالیاف و همکاران در سال ۲۰۲۰ کار روی داده‌ی موسسه وروشو در لهستان [2] و استفاده از ایده‌ی یادگیری انتقالی^۷ و پیاده‌سازی مدل SVM به دقت ۹۸/۶٪ در تشخیص بیماری اسکیزوفرنی دست یافتند [3]. شارما و همکاران در سال ۲۰۲۱ از داده‌ی موسسه وروشو در لهستان [2] استفاده کردند؛ آنها با به‌کارگیری روش مبتنی بر موجک، ابتدا سیگنال‌های EEG را طی شش تکرار که منجر به تولید هفت باند فرعی می‌شود، در معرض تجزیه موجک قرار دادند؛ هنجار l1 برای هر زیر باند محاسبه شد و ویژگی‌های نرمال استخراج شده در روش یادگیری ماشین K نزدیک‌ترین همسایه^۸ (KNN) و با روش ارزیابی 10Fold استفاده شد و به دقت قابل توجه ۹۹/۲۱٪ رسید [11].

رابط کامپیوتری مغز^۹ (BCI) یک سیستم قدرتمند برای برقراری ارتباط بین مغز و جهان خارج است. سیستم‌های سنتی BCI فقط بر اساس سیگنال‌های EEG کار می‌کنند. به تازگی، محققان از ترکیب سیگنال‌های EEG با سیگنال‌های دیگر برای بهبود عملکرد سیستم‌های BCI استفاده کرده‌اند. در میان این سیگنال‌ها، ترکیب EEG با fNIRS به نتایج مطلوبی دست یافته است. در اکثر مطالعات، فقط EEG یا fNIR ها به‌عنوان توالی‌های زنجیره‌ای در نظر گرفته شده‌اند و همبستگی پیچیده‌ای بین سیگنال‌های مجاور را در زمان و مکان کانال در نظر نمی‌گیرند. قنچی و همکاران در سال ۲۰۲۰، یک مدل شبکه عصبی عمیق برای شناسایی اهداف دقیق مغز انسان با معرفی ویژگی‌های زمانی و مکانی معرفی کردند. مدل پیشنهادی شامل ارتباط فضایی بین سیگنال‌های EEG و fNIRS است. این را می‌توان با تبدیل توالی این سیگنال‌های زنجیره‌ای به تسسورهای سه درجه‌ای سلسله مراتبی پیاده‌سازی کرد. آزمایشات نشان می‌دهد که مدل پیشنهادی آنها دارای دقت ۹۹/۶٪ است [12].

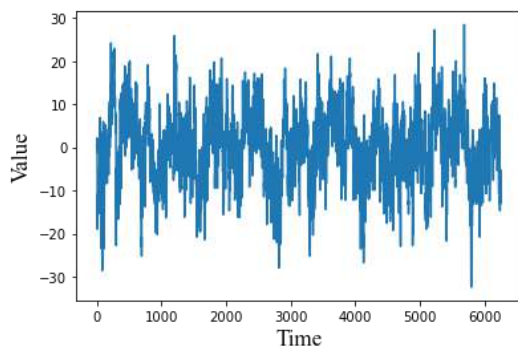
همان‌طور که دیده می‌شود چالش‌هایی هم‌چون رسیدن به دقت بالا، حجم محاسباتی بالا و زمان آموزش بالا در مساله مذکور وجود دارد. در این‌جا یک شبکه عصبی عمیق بازگشتی کانولوشنی ارائه می‌شود که سعی در فائق آمدن بر چالش‌های مذکور دارد.

۳- معرفی داده و پیش‌پردازش

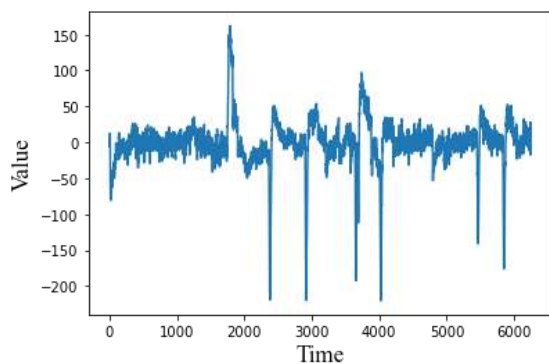
سیگنال‌های EEG استفاده شده در این مقاله، سیگنال‌های جمع‌آوری شده در انستیتو روان‌پزشکی و مغز و اعصاب واقع در شهر وروشو در کشور لهستان است. این سیگنال‌ها شامل سیگنال‌های EEG چهارده بیمار مبتلا به اختلال اسکیزوفرنی، شامل هفت مرد و هفت زن، به‌ترتیب با میانگین سنی ۲۷/۹±۳/۳ و ۲۸/۴±۳/۱ سال و سیگنال‌های EEG چهارده فرد سالم در همان گروه سنی و نسبت جمعیتی بودند [2]. پانزده دقیقه اطلاعات EEG با سرعت نمونه‌گیری ۲۵۰ هرتز جمع‌آوری شده بود. داده‌ها طبق استاندارد بین‌المللی ۱۰-۲۰ جمع‌آوری شده بودند. هر الکتروود یک حرف اختصاصی دارد: F برای لوب پیشانی، P برای لوب آهینانه‌ای، O برای لوب پس‌سری، T برای لوب گیجگاهی و Z برای خط میانی. لوب‌های مختلف مغز مسئول اعمال خاصی مانند زبان، حافظه و کلام هستند. الکتروودهای Fp1, Fp2, F7, F3, Fz, F4, F8, T3, C3, Cz, C4, T4, T5, P3, Pz, T6, O1 و O2 استفاده شده بودند و محل قرارگیری این الکتروودها روی سر طبق استاندارد ۱۰-۲۰ مطابق شکل (۱) بود. سیگنال‌های به‌دست آمده به بخش‌هایی تقسیم شدند که در آن بخش‌ها می‌توان سیگنال‌ها را ثابت در نظر گرفت. طول پنجره‌ی هر بخش ۲۵ ثانیه و شامل ۶۲۵۰ نمونه بود. قبل از تغذیه شبکه پیشنهادی برای آموزش، داده‌ها با نمره Z نرمال شدند. از ۱۱۴۲ بخش EEG استفاده شد که هر بخش ۶۲۵۰×۱۹ نقطه نمونه‌برداری داشت. از نرمال‌سازی برای مقیاس‌گذاری سیگنال‌ها در محدوده‌ی استاندارد از مقادیر استفاده شد. شکل (۲) نمونه سیگنال فرد سالم و بیمار را مقایسه می‌کند.

^۶ Multivariate Empirical Mode Decomposition
^۷ Transfer Learning
^۸ K-nearest neighbors (KNN)
^۹ Brain Computer Interface (BCI)

^۱ Functional MRI (fMRI)
^۲ Black Hole (BH)
^۳ Glowworm Search (GS)
^۴ Artificial Flora (AF)
^۵ Monkey Search (MS)



(الف)

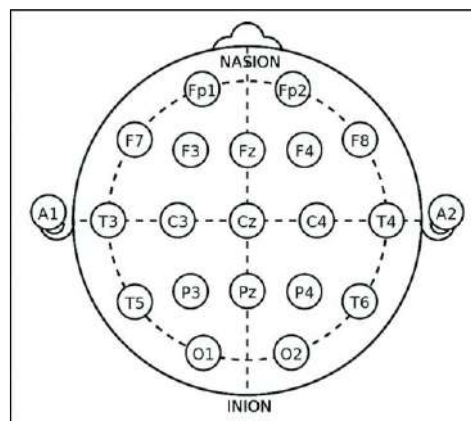


(ب)

شکل ۲: نمونه سیگنال‌های مغزی (الف) فرد سالم (ب) بیمار اسکیزوفرنی

در فرآیند یادگیری، این نوع شبکه‌ها تمایل به تکرار الگوهای مکرر را دارند که از طریق اشتراک گذاری متغیرها در تمام مراحل زمانی میسر می‌شود [15]. در این مطالعه، از مدل شبکه عصبی عمیق بازگشتی کانولوشنی ۱۱ لایه استفاده شده است. داده‌های موجود بعد از پیش‌پردازش وارد لایه‌های کانولوشن یک بعدی می‌شوند؛ سپس با گذر از لایه‌های بازگشتی، بهترین ویژگی‌ها را استخراج می‌نمایند و بر اساس ویژگی‌های استخراج شده طبقه‌بندی انجام می‌شود. توجه شود که تفاوت این ساختار با ساختار مقاله مرجع [2] از جهت آن است که در این‌جا شبکه استفاده شده یک شبکه ترکیبی با لایه‌های کانولوشن و لایه‌های بازگشتی است. برای بهبود عملکرد شبکه بازگشتی کانولوشن عمیق ۱۱ لایه، dropout به لایه‌های ۹ و ۱۰، با میزان ۰/۵ درصد (به این معنی که با احتمال ۰/۵ درصد، یک نرون در حین آموزش خارج می‌شود) اعمال شده است. پارامترهای بهینه‌سازی آدم، با میزان یادگیری ۰/۰۰۰۱ به کار رفته است و تابع فعال ساز برای لایه‌های ۱، ۳، ۵ و ۷، واحد خطی یک‌سوساز^۲ و برای لایه آخر (به منظور طبقه‌بندی)، تابع softmax استفاده شده است. سایر پارامترهای مدل پیشنهادی با آزمون و خطا به دست آمده‌اند. شبکه پیشنهادی در این مقاله، با زبان برنامه‌نویسی پایتون پیاده‌سازی و اجرا شده است. مشخصات شبکه پیاده‌سازی شده در شکل (۳) آمده است. ویژگی‌های اصلی سیستم پیشنهادی عبارتند از:

(۱) یک مدل شبکه عصبی بازگشتی کانولوشن عمیق ۱۱ لایه برای ارزیابی دقیق بیماران اسکیزوفرنی در مقابل گروه افراد سالم ارائه شده است. (۲) مدل پیشنهادی فرآیندهای استخراج و انتخاب بهترین ویژگی‌ها و طبقه‌بندی را به طور خودکار در لایه‌های کانولوشن و LSTM انجام می‌دهد. (۳) ارزیابی مدل با تکنیک



شکل ۱: محل قرارگیری الکترودها روی سر (استاندارد ۱۰-۲۰) [10]

۴- روش پیشنهادی

از آنجا که ماهیت سیگنال‌های EEG غیرخطی است، اغلب برای تشخیص سیگنال‌های EEG افراد سالم از بیماران اسکیزوفرنی تکنیک‌های استخراج ویژگی‌های غیرخطی استفاده می‌شود [13]. یادگیری ماشین عمده‌ترین روش تشخیص الگو استفاده می‌شود. با این حال، این تکنیک پیشرفته ضعیف‌هایی نیز دارد. به عنوان مثال برای کارهای ساده تشخیص، الگوریتم‌های یادگیری ماشین خوب عمل می‌کند، اما زمانی که تنوع ویژگی‌های مورد مطالعه زیاد باشد، مجموعه‌های آموزشی بزرگتری برای تشخیص صحیح نیاز دارند. در مقایسه با تکنیک‌های سنتی یادگیری ماشین، یادگیری عمیق یک مدل با ظرفیت یادگیری به طور قابل ملاحظه‌ای بالاتر ارائه می‌دهد که امکان مطالعه ویژگی‌های سطح بالاتری را از طریق یادگیری داده‌ها از مجموعه داده‌های بزرگ فراهم می‌آورد. در یادگیری عمیق برخلاف تکنیک‌های سنتی یادگیری ماشین، هر دو عمل استخراج ویژگی‌ها و طبقه‌بندی به طور خودکار انجام می‌شود [14].

حافظه کوتاه بلند مدت^۱ (LSTM) پیکربندی خاصی از شبکه‌های بازگشتی است که مسیرهایی ایجاد می‌کند که شیب گرادیان از طریق آنها قادر است در دوره‌های طولانی جریان یابد. در فرآیند یادگیری، این نوع شبکه‌ها تمایل به تکرار الگوهای مکرر را دوره‌های طولانی جریان یابد.

^۲ LeakyRelu

^۱ Long Short Term Memory (LSTM)

$$Sensitivity = \frac{TP}{TP+FN} \quad (2)$$

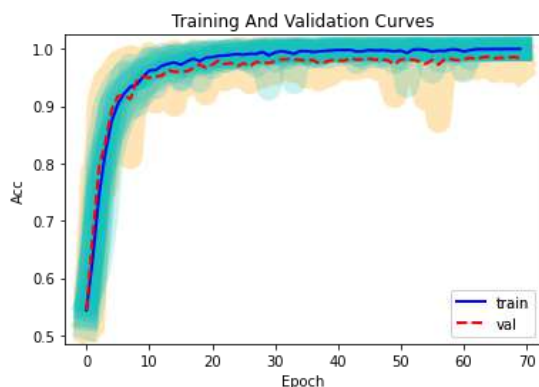
$$Specificity = \frac{TN}{TN+FP} \quad (3)$$

$$PPV = \frac{TP}{TP+FP} \quad (4)$$

جدول (۱)، نتایج طبقه‌بندی مدل پیشنهادی به ازای تعداد Fold مختلف را نشان می‌دهد. میانگین مقادیر معیارهای ارزیابی Accuracy، Sensitivity، Specificity و PPV، به دست آمده برای مدل پیشنهادی به ترتیب برابر ۹۸/۷۹٪، ۹۸/۷۳٪، ۹۸/۸۶٪ و ۹۹/۰۶٪ است. شکل (۴)، دقت حاصل از آموزش مدل ارائه شده در تکرارهای متوالی را نشان می‌دهد. جدول (۲)، روش پیشنهادی را با چند روش پیشین مقایسه می‌کند. دیده می‌شود روش‌های مبتنی بر یادگیری عمیق و به‌طور خاص روش پیشنهادی این مقاله، دقت طبقه‌بندی بالاتری دارند.

جدول ۱: نتایج طبقه‌بندی مدل پیشنهادی به ازای تعداد fold مختلف

Fold	%Accuracy	%PPV	%Sensitivity	%Specificity
1	98.77	98.57	99.19	98.25
2	99.03	99.52	98.71	99.41
3	99.12	99.36	99.04	99.22
4	98.15	98.43	98.24	98.05
5	99.03	99.21	99.03	99.03
6	98.42	98.42	98.72	98.05
7	99.12	99.68	98.72	99.61
8	98.77	99.22	98.55	99.02
9	99.03	99.36	98.87	99.22
10	98.50	99.05	98.24	98.82

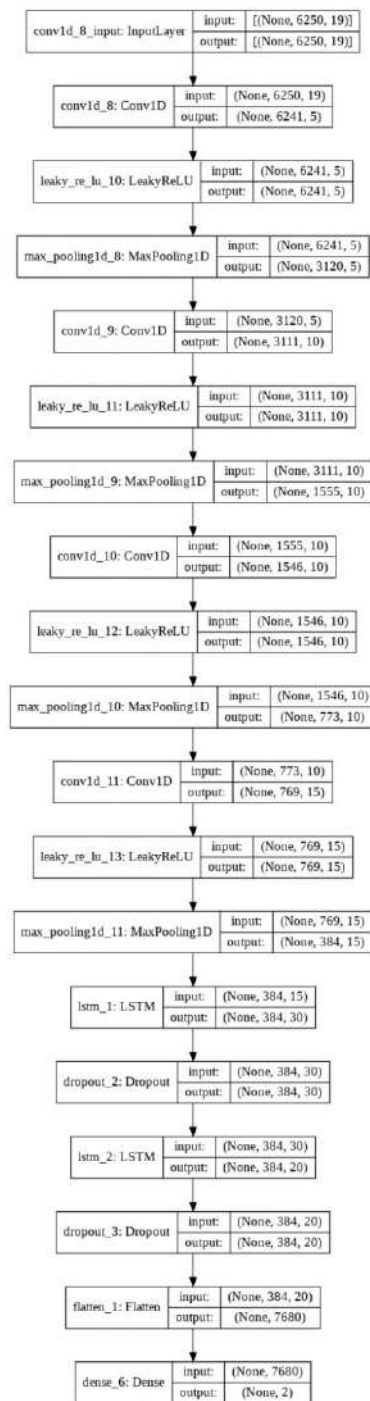


شکل ۴: دقت حاصل از آموزش مدل ارائه شده در تکرارهای متوالی

۶- نتیجه‌گیری

مقایسه روش‌های مختلف به کار رفته در تشخیص بیماری اسکیزوفرنی با روش پیشنهادی این مقاله نشان داد بالاترین دقت برای طبقه‌بندی اسکیزوفرنی با استفاده از الگوریتم‌های مبتنی بر یادگیری عمیق حاصل می‌شود. علیرغم وجود مجموعه داده کوچک مورد ارزیابی در این پژوهش، صحت طبقه‌بندی بالا نشان می‌دهد که تکنیک پیشنهادی می‌تواند به عنوان یک ابزار قابل اعتماد در تشخیص خودکار اختلال روانی اسکیزوفرنی توسط پزشکان مورد استفاده و استناد قرار بگیرد.

10Fold بسیار معتبر است. (۴) دقت بالا با حجم داده‌های کوچک صحت عملکرد سیستم را تایید می‌کند.



شکل ۳: ساختار شبکه عصبی عمیق بازگشتی کانولوشنی پیشنهادی

۵- آزمایش‌ها

در این مقاله از روش ارزیابی KFold برای ارزیابی روش پیشنهادی استفاده شده است. معیارهای ارزیابی دقت (Accuracy)، Sensitivity، Specificity و PPV^۱، که به ترتیب از روابط زیر قابل محاسبه هستند، به کار رفته‌اند:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (۱)$$

^۱ Positive Predictive Value (PPV)

جدول ۲: خلاصه مطالعات برای طبقه‌بندی اسکیزوفرنی با استفاده از EEG

نویسندگان	تعداد ویژگی‌ها	تکنیک	داده	دقت طبقه‌بندی
کیم و همکاران [16] ۲۰۱۵	-	■ قدرت طیفی EEG با تبدیل سریع فوریه با استفاده از MATLAB (متغیرهای متغیر) محاسبه می‌شود. ■ باند فرکانس دلتا، تتا، آلفا و بتا مورد تجزیه و تحلیل قرار گرفت. ■ تجزیه و تحلیل واریانس ■ سیگنال‌های مغزی P3b	نرمال: نمونه ۹۰ فرد سالم اسکیزوفرنی: نمونه ۹۰ فرد مبتلا به اختلال اسکیزوفرنی	در باند فرکانسی دلتا ٪ ۶۲/۲۰
ساتتوس و همکاران [17] ۲۰۱۷	۲۰ ویژگی از هر نمونه	■ زمان، ویژگی‌های دامنه فرکانس ■ گروه‌بندی کانال‌ها ■ طبقه‌بندی کننده‌های SVM و پرسپترون چند لایه (MLP)	نرمال: نمونه ۳۱ فرد سالم اسکیزوفرنی: نمونه ۱۶ فرد مبتلا به اختلال اسکیزوفرنی	■ در روش MLP ٪ ۹۸/۴۲ ■ در روش SVM ٪ ۹۳/۲۳
لی و همکاران [2] ۲۰۱۹	-	■ شبکه عصبی کانولوشنی عمیق ۱۱ لایه ■ ارزیابی با روش 10Fold ■ شبکه عصبی کانولوشنی عمیق ۱۱ لایه ■ ارزیابی با روش 14Fold	نرمال: نمونه ۱۴ فرد سالم اسکیزوفرنی: نمونه ۱۴ فرد مبتلا به اختلال اسکیزوفرنی	■ در شبکه اول ٪ ۹۸/۰۷ ■ در شبکه دوم ٪ ۸۱/۲۶
روش ارائه شده	-	■ شبکه عصبی بازگشتی کانولوشنی عمیق ۱۱ لایه ■ ارزیابی با روش 10Fold	نرمال: نمونه ۱۴ فرد سالم اسکیزوفرنی: نمونه ۱۴ فرد مبتلا به اختلال اسکیزوفرنی	٪ ۹۸/۷۹

مراجع

- Computing," Computational Intelligence and Neuroscience, p. 2020, 2020.
- [9] K. Palani Thanaraj, A. Noel Joseph Raj, P. Balasubramanian and Y. Chen, "Schizophrenia detection using Multivariate Empirical Mode Decomposition and entropy measures from multichannel EEG signal," Biocybernetics and Biomedical Engineering, vol. 40(3), pp. 1124-1139, 2020.
 - [10] Z. Aslan and M. Akin, "Automatic Detection of Schizophrenia by Applying Deep Learning over Spectrogram Images of EEG Signals," Traitement du Signal, vol. 37(2), pp. 235-244, 2020.
 - [11] M. & A. U. R. Sharma, "Automated detection of schizophrenia using optimal wavelet-based L1 norm features extracted from single-channel EEG," Cognitive Neurodynamics, pp. 1-14, 2021.
 - [12] H. F. M. A. V. F. S. & R. M. Ghonchi, "Deep recurrent-convolutional neural network for classification of simultaneous EEG-fNIRS signals," IET Signal Processing, vol. 14(3), pp. 142-153, 2020.
 - [13] R. A. D. J. N. S. C. I. P. J. & A. M. Hornero, "Variability, regularity, and complexity of time series generated by schizophrenic patients and control subjects," IEEE Transactions on Biomedical Engineering, vol. 53(2), pp. 210-218, 2006.
 - [14] U. R. F. H. L. O. S. A. M. T. J. H. & C. C. K. Acharya, "Automated detection of coronary artery disease using different durations of ECG segments with convolutional neural network," Knowledge-Based Systems, vol. 132, pp. 62-71, 2017.
 - [15] S. W. Z. S. J. Z. Z. W. Z. Y. T. .. & C. C. Xu, "Using a deep recurrent neural network with EEG signal to detect Parkinson's disease," Annals of Translational Medicine, vol. 8(14), 2020.
 - [16] L. S.-J.-R. L. M. & A. J. I. Santos-Mayo, "A computer-aided diagnosis system with EEG based on the P3b wave during an auditory odd-ball task in schizophrenia," IEEE transactions on biomedical engineering, vol. 64(2), pp. 395-407, 2016.
 - [1] J. O. B. L. L. K. U. C. J. H. & Y. K. Oh, "Identifying schizophrenia using structural MRI with a deep learning algorithm," Frontiers in psychiatry, vol. 11, p. 16, 2020.
 - [2] S. L. Oh, J. Vicnesh, E. J. Ciaccio, R. Yuvaraj and U. R. Acharya, "Deep convolutional neural network model for automated diagnosis of schizophrenia using EEG signals," Applied Sciences, vol. 9(14), p. 2870, 2019.
 - [3] Shalbah, S. Bagherzadeh and A. Maghsou, "Transfer learning with deep convolutional neural network for automated detection of schizophrenia from EEG signals," Physical and Engineering Sciences in Medicine, vol. 43(4), pp. 1229-1239, 2020.
 - [4] S. Siuly, S. K. Khare, V. Bajaj and H. Wang, "A computerized method for automatic detection of schizophrenia using EEG signals," IEEE Transactions on Neural Systems and Rehabilitation Engineering, vol. 28(11), pp. 2390-2400, 2020.
 - [5] R. Boostani, K. Sadatnezhad and M. Sabeti, "An efficient classifier to diagnose of schizophrenia based on the EEG signals," Expert Systems with Applications, vol. 36(3), pp. 6492-6499, 2009.
 - [6] S. Srinivasagopalana, J. Barryb, V. Gurupurc and S. Thankachan, "A deep learning approach for diagnosing schizophrenic patients," Journal of Experimental & Theoretical Artificial Intelligence, vol. 31(6), pp. 803-816, 2019.
 - [7] W. Yan, V. Calhoun, M. Song, Y. Cui, H. Yan, S. Liu, .. and J. Sui, "Discriminating schizophrenia using recurrent neural network applied on time courses of multi-site fMRI data," EBioMedicine, vol. 47, pp. 543-552, 2019.
 - [8] S. K. R. H. & K. S. H. Prabhakar, "Schizophrenia EEG Signal Classification Based on Swarm Intelligence



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بوم

تشخیص گریه نوزاد از سایر صداهای محیط با استفاده از یادگیری عمیق

پری ناز نورمحمدی^۱، مریم رحیمی هاشم آباد^۱، مهسا زمانی تراشنده^۱، محمدرضا اکبرزاده توتونچی^۱

^۱ قطب علمی رایانش نرم و پردازش هوشمند اطلاعات، گروه برق، دانشگاه فردوسی، مشهد

parinaz.nourmohammadi@mail.um.ac.ir

rahimi_maryam@elec.iust.ac.ir

mahsa.zamanitarashandeh@mail.um.ac.ir

akbazar@um.ac.ir

چکیده: مهم ترین راه ارتباطی نوزادان با دنیای اطراف گریه آنها است. علاوه بر درک نیازهای روزمره نوزادان، پیش بینی بیماری یکی دیگر از وظایف مهم در تحقیقات گریه نوزاد است. سیستم های هوشمند تشخیص گریه نوزاد زمینه ساز ساخت ربات های هوشمند مراقبتی خواهند بود. سیگنال های گریه نوزادان حاوی ویژگی های منحصر به فردی است که با استفاده از این ویژگی ها میتوان گریه نوزادان را از سایر اصوات محیط تشخیص داد. اغلب دیتا های موجود در این زمینه صداهای ضبط شده توسط افراد در NICU و یا در خانه توسط والدین است. در این پژوهش صدای گریه نوزادان از سایر اصوات محیط تشخیص داده شده است. در این مسیر از ضرایب کپسترال فرکانس مل^۱ (MFCC) بهره بردیم و عملکرد شبکه های عصبی عمیق پیچشی^۲ (CNN) و حافظه طولانی-کوتاه مدت^۳ (LSTM) را بررسی کردیم و برتری روش خود را بر اساس معیار های دقیق تر و جامع تری از جمله دقت^۴، حساسیت^۵ و ماتریس درهم ریختگی^۶ سنجیدیم. نتایج به دست آمده نشان می دهد که جهت تشخیص صدای گریه نوزادان الگوریتم شبکه عصبی حافظه طولانی-کوتاه مدت دارای دقت ۹۶/۳۰٪ و شبکه عصبی پیچشی دارای دقت ۹۷/۹۷٪ می باشد.

کلمات کلیدی: تشخیص گریه نوزادان، ضرایب کپسترال فرکانس مل، شبکه عصبی طولانی-کوتاه مدت، شبکه عصبی پیچشی

۱- مقدمه

در یکی دیگر از تحقیقات صورت گرفته در این زمینه برای تشخیص صدای گریه نوزادان از KNN^۸ استفاده کرده اند و ضرایب MFCC و LFCC^۹ استخراج شده را باهم مقایسه کردند و به این نتیجه رسیده اند که استفاده از LFCC (با دقت ۹۰٪) نسبت به MFCC (با دقت ۷۰٪) جهت تشخیص دیتاستی که شامل پنج کلاس می باشد، مناسب تر است [4]. از دیگر مطالعات صورت گرفته در این زمینه می توان به مقاله [16] اشاره کرد، که برخی از روش های استخراج ویژگی، از جمله LPC^{۱۰}، LPCC^{۱۱} و MFCC و BFCC^{۱۲} را مقایسه کرده همچنین یک روش سنجش فشرده^{۱۳} را به عنوان روش طبقه بندی با KNN و ANN^{۱۴} پیشنهاد کرده که بهترین نرخ طبقه بندی از ترکیب BFCC و ANN با دقت ۷۶/۴۷٪ به دست می آید، در حالی که روش طبقه بندی پیشنهادی آن ها وقتی که با MFCC ترکیب می شود، دقت ۷۱/۰۵٪ را نتیجه می دهد. واضح است که ترکیب های مختلفی از استخراج و طبقه بندی ویژگی ها وجود دارد و هر روش می تواند نتایج متفاوتی

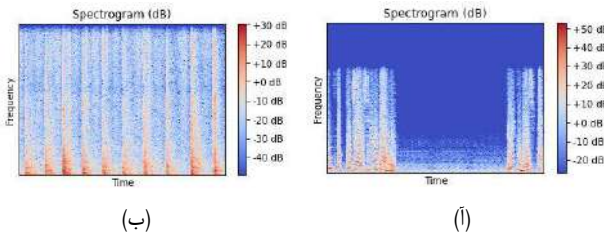
سالانه صد و سی میلیون نوزاد در جهان متولد می شوند و مراقبت از نوزادان چالش بزرگی برای پدران و مادران است. دلیل اصلی این چالش این است که مهمترین راه ارتباطی نوزادان با دنیای اطراف گریه آن هاست. این تحقیقات زمینه ساز تحقیقات گسترده تر راجع به علل گریه نوزادان و تشخیص نیازها و بیماری های احتمالی آن ها خواهد بود. تحقیقات ابتدایی تشخیص گریه نوزاد معمولاً ویژگی های سطح پایین را استخراج کرده اند و در میان برخی از تحقیقات اخیر، آزمایشی وجود دارد که توسط مرجع [15] انجام شده است و از ویژگی سطح بالای MFCC و الگوریتم یادگیری ماشین KNN در مرحله طبقه بندی استفاده می کند و دقت آن ۷۹/۹۵ درصد است که روش پیشنهادی آن از Naïve Bayes، Neural Network و SVM^۷ بهتر است. همچنین

ماشین لباس شویی و غیره می باشد و از [6,8] جمع آوری شده است. جدول شماره (۱) تعداد داده و دسته بندی پایگاه داده را نشان می دهد.

جدول ۱: مجموعه داده جمع آوری شده ۲ کلاس

تعداد داده ها	تعداد داده ها	تعداد داده ها	تعداد داده ها
۱۵	در زدن	۱۱۵	صحبت
۳۲	قدم های پا	۱۰۸	خنده نوزاد
۳۶	خنده	۵	هواپیما
۵	باران	۳۳	نفس کشیدن
۴۰	عطسه	۵	بوق ماشین
۳۳	خروپف	۵	گریه
۳۰	جاروبرقی	۵	دست زدن
۲۰	ماشین لباس شویی	۵	هشدار ساعت
۷	باد	۳۹	سرفه
۵۲۵	گریه نوزاد	۵	سگ
۵۵۸	مجموع صداهای غیر گریه	۱۵	لولای در

۲-۲ استخراج ویژگی



شکل ۲: (ا) طیف سنجی سیگنال گفتار؛ (ب) طیف سنجی سیگنال گریه نوزاد

سیگنال گفتار بزرگسالان و گریه نوزادان دارای شباهت ها و تفاوت هایی است. در مطالعات اخیر تفاوت هایی در فرکانس اساسی این صداها مشاهده شده، به طوری که سیگنال صوت گریه نوزادان دارای فرکانس اساسی بالاتری است. همچنین طیف صدای گریه نوزادان به دلیل نازک و کوتاه بودن تارهای صوتی شان بسیار مرتب است. این موضوع را می توان در شکل (۲) مشاهده کرد.

۲-۲-۱ ضرایب کپسترال فرکانس مل (MFCC)

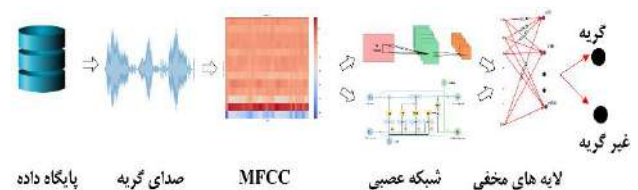
پس از پیش پردازش، قدم بعدی در تشخیص گریه نوزاد، استخراج ویژگی است. ویژگی های صدا، پس از فریم بندی سیگنال ورودی، با فریم هایی با طول ۱۰ تا ۴۰ میلی ثانیه، استخراج می شوند [4]. بسیاری از مطالعات در زمینه تشخیص گفتار از ویژگی MFCC استفاده می کنند زیرا این ویژگی با تکیه بر حساسیت گوش انسان در فرکانس های کمتر از ۱۰۰۰ هرتز ساختار یافته است [4]. برای از بین بردن اثر دگرنامی^{۲۰} که در طی فرایند فریم بندی به وجود می آید از پنجره همینگ استفاده میکنیم. معادله (۱) تابع پنجره گذاری را نشان می دهد:

$$w_n = 0.54 - 0.46 \cos \frac{2\pi n}{N-1} \quad 0 \leq n \leq N \quad (1)$$

در این معادله w_n تابع پنجره گذاری، N تعداد نمونه ها در هر فریم می باشد. با اعمال تبدیل فوریه سریع^{۲۱} سیگنال پنجره گذاری شده از حوزه زمان به

را ارائه دهد. در سال های اخیر، الگوریتم های یادگیری عمیق به دلیل توانایی تولید ویژگی های سطح بالا و طبقه بندی آن ها در مطالعات متعددی در زمینه بینایی رایانه و پردازش گفتار مورد توجه قرار گرفته اند. در میان بسیاری از رویکردهای مختلف، CNN یکی از الگوریتم های یادگیری عمیق است که به دلیل عملکرد برجسته در پردازش تصویر به طور گسترده ای مورد استفاده قرار می گیرد. برای مثال در مقاله [18] از شبکه عصبی پیچشی تک بعدی استفاده شده و نتایج بهتری در مقایسه با شبکه عصبی پیش خور^{۱۵} و ماشین بردار پشتیبان^{۱۶} به دست آمده است و یا در مقاله ای این شبکه عصبی به همراه یادگیری انتقال^{۱۷} بر طیف سنجی^{۱۸} داده های گریه نوزادان اعمال شد و نتایج قابل قبولی به دست آمد [12]. همچنین در مقاله [5] یک شبکه مبتنی بر CNN برای تشخیص علائم پیش از گریه نوزادان پیشنهاد شد و نتایج بهتری در مقایسه با ویژگی های سطح پایین به دست آوردند. علاوه بر رویکردهای مبتنی بر CNN، روش RNN^{۱۹} که عملکرد عالی آن برای داده های متوالی مانند سیگنال های گفتاری شناخته شده است نیز به طور گسترده ای توسط برخی از محققان استفاده می شود. طی مطالعات اخیر آزمایش هایی با ترکیب CNN و RNN با ساختار end-to-end انجام شده است و عملکردی بهتر از تشخیص گریه نوزادان شامل مراحل از جمله جمع آوری دیتاست، سیگنال اولیه، استخراج ویژگی و انتخاب روش طبقه بندی بهینه است.

در این پژوهش پردازش سیگنال اولیه به منظور پاکسازی نویز پس زمینه جهت ساخت پایگاه داده می باشد. سپس ویژگی MFCC را استخراج کرده و از شبکه های عصبی پیچشی و حافظه طولانی-کوتاه مدت جهت طبقه بندی استفاده می شوند. نمودار بلوکی روش تشخیص گریه نوزادان از سایر صداهای محیط با استفاده از شبکه های عصبی عمیق که در شکل (۱) نشان داده شده است، شامل سه مرحله می باشد: جمع آوری پایگاه داده مناسب، استخراج ویژگی ها و روش های طبقه بندی.

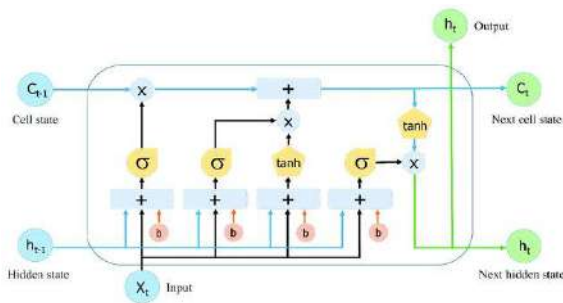


شکل ۱: نمودار بلوکی سیستم ارائه شده بر مبنای تحلیل ضرایب MFCC و شبکه های عصبی

۲-۲ روش ها

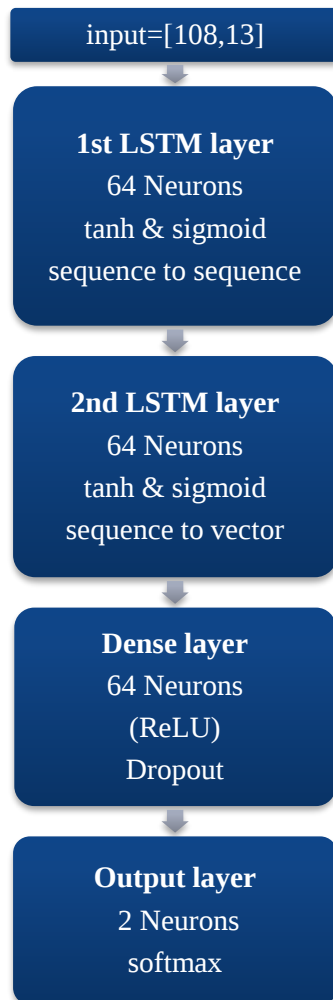
۲-۱ پایگاه داده

پایگاه داده استفاده شده در این پژوهش شامل داده های گریه و غیر گریه است که داده های گریه از [6,7] جمع آوری گردیده است. داده های غیر گریه نیز شامل صدای صحبت، صدای خنده نوزاد، صدای سرفه، صدای عطسه، صدای



شکل ۳: شبکه عصبی حافظه طولانی-کوتاه مدت [11]

پس از پردازش اولیه و استخراج ضرایب MFCC، مدل LSTM را ایجاد می کنیم که در این مدل از مدل ترتیبی^{۲۳} با یک لایه ی ورودی ۳ بعدی شامل ۱۳ ضریب MFCC، ۱۰۸ پنجره زمانی و ۲۱۶۶ بخش صدادار مدل استفاده می کنیم. ورودی شبکه عصبی LSTM دارای آرایه ای سه بعدی با ابعاد (۱۳، ۱۰۸، ۲۱۶) می باشد. شبکه عصبی پیشنهادی از دو لایه پنهان تشکیل شده و جهت جلوگیری از پدیده ی بیش برازش از روش Dropout استفاده شده است. در نهایت لایه ی خروجی را از تابع فعال ساز Softmax عبور می دهیم.



شکل ۴: نمودار پیاده سازی مدل پیشنهادی شبکه عصبی LSTM

فرکانس منتقل می شود و سپس مل فیلتر بانک ها بر روی سیگنال فرکانس لحاظ می شود. معادله (۲) مقیاس مل را نشان می دهد:

$$Mel(f) = 2595 \log_{10} \left(1 + \frac{f}{700} \right) \quad (2)$$

MFCC از فیلتربانک های لگاریتمی و مثلثی مل استفاده می کند. این فیلتر ها از نوع فیلتر میان گذر می باشند [2]. پس از اعمال تبدیل کسینوسی گسسته^{۲۴} ضرایب MFCC استخراج می شوند.

۲-۳ روش های طبقه بندی

طی این پژوهش، جهت طبقه بندی صدای گریه نوزادان از شبکه عصبی حافظه طولانی-کوتاه مدت به سبب عملکرد مناسب در یادگیری وابستگی های زمانی و شبکه عصبی پیچشی به دلیل عدم حساسیت به موقعیت الگو در تصویرهای طیف نگاری، بهره بردیم. این بخش به توصیف شبکه های عصبی و مدل پیشنهادی پرداخته است.

۲-۳-۱ شبکه عصبی حافظه طولانی-کوتاه مدت

LSTM نوع خاصی از شبکه های عصبی عمیق RNN می باشند که قادر به یادگیری الگوها و وابستگی های بلند مدت هستند. شبکه های LSTM بر خلاف شبکه های RNN که به دلیل از بین رفتن مقدار گرادیان در ساختارشان قادر به تشخیص روابط در داده های طولانی مدت نیستند، دارای بلوک های حافظه ای هستند که آنها را قادر به یادگیری روابط کوتاه و همچنین طولانی مدت می کند. هر بلوک شامل سلول حافظه و گیت هایی می باشد که مشخص کننده ی وضعیت بلوک نهایی می باشد. همچنین هر شبکه ی LSTM شامل سه نوع دروازه ورودی، فراموشی و خروجی می باشد که مقدار اطلاعات ورودی و خروجی هر سلول را کنترل میکنند [3].

لایه های LSTM دنباله ورودی $X=(x_1, x_2, \dots, x_T)$ را به دنباله خروجی $Y=(y_1, y_2, \dots, y_T)$ بر اساس معادلات زیر نگاشت می کند.

$$i_t = \text{Sig}(W_{xi}x_t + W_{yi}y_{t-1} + b_i) \quad (3)$$

$$f_t = \text{Sig}(W_{xf}x_t + W_{yf}y_{t-1} + b_f) \quad (4)$$

$$g_t = f_t \odot c_{t-1} \quad (5)$$

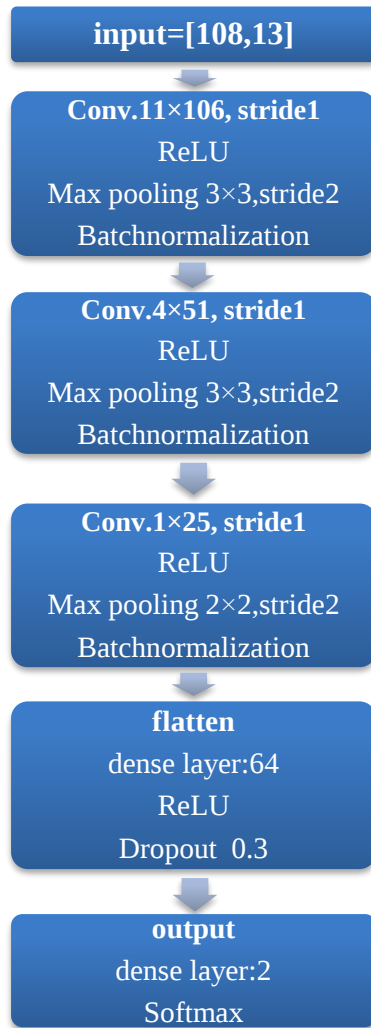
$$c_t = g_t + i_t \odot \tanh(W_{xc}x_t + W_{yc}y_{t-1} + b_c) \quad (6)$$

$$y_t = \text{Sig}(W_{xo}x_t + W_{yo}y_{t-1} + b_o) \quad (7)$$

در واقع c_t حالت سلول حافظه است و i_t ، f_t ، y_t دروازه های خروجی در زمان t هستند. W وزن های شبکه و b مقدار بایاس را نشان می دهد که در طی آموزش جهت به حداقل رساندن تابع خطا، تنظیم می شوند. در ساختار چند لایه، ورودی لایه بعدی در واقع خروجی لایه قبلی است [13].

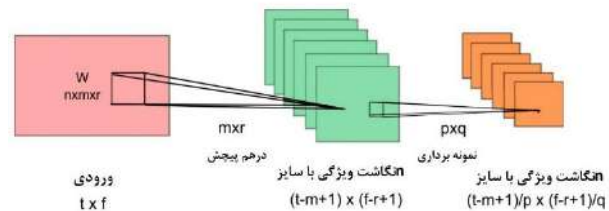
۲-۳-۲ شبکه عصبی پیچشی

و از تابع فعال ساز ReLU استفاده شده است. همچنین جهت افزایش سرعت و استانداردسازی تابع فعال ساز در هر لایه پیچشی از روش نرمال سازی دسته ای^{۲۵} استفاده شده است. سپس در لایه های مخفی متصل کامل از روش Dropout جهت جلوگیری از بیش برآزش استفاده شده است. در نهایت تابع فعالیت لایه خروجی Softmax انتخاب شده است. شبکه جهت بهینه شدن ۳۰ بار آموزش می بیند. مشخصات مدل پیشنهادی در شکل (۶) و (۷) نشان داده شده است.



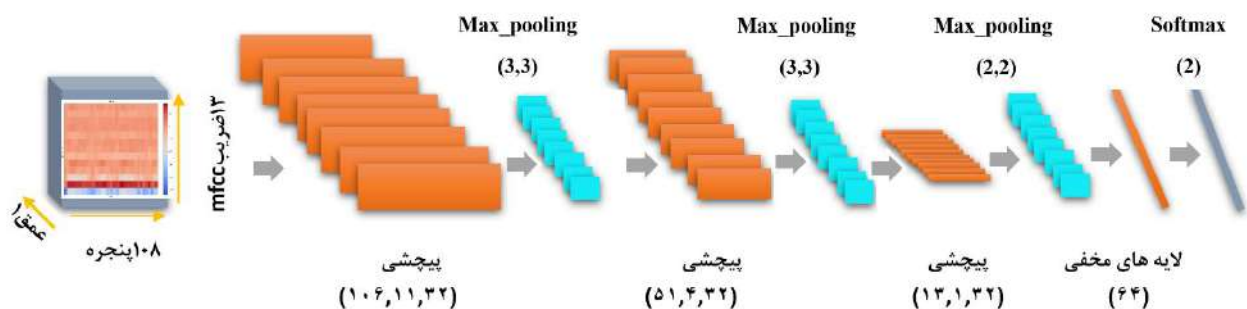
شکل ۶: نمودار بلوکی پیاده سازی شبکه عصبی پیچشی

یک شبکه عصبی پیچشی از سه بخش اصلی تشکیل می شود: لایه های در هم پیچش، لایه های ادغام^{۲۴} و لایه های مخفی متصل کامل [1]. در شکل (۵) لایه های در هم پیچش و ادغام با جزئیات بیش تر نشان داده شده اند. لایه در هم پیچش، ماتریس ورودی با ابعاد $t \times f$ را (که t نشانگر بعد زمان و f نشانگر بعد فرکانس می باشد) به عنوان ورودی دریافت می کند [20]. لایه در هم پیچش متشکل از n نورون است که هر نورون یک فیلتر است و ابعاد هر فیلتر $r \times m$ است که با کل ورودی در هم پیچ می شود [20]. هر فیلتر در بعد زمان-فرکانس روی ورودی (با فرض $r \leq f$ و $m \leq t$) حرکت می کند و در خروجی هر نورون، نگاشت های ویژگی با ابعاد $(t-m+1) \times (f-r+1)$ مشاهده خواهد شد [20]. این اشتراک گذاری وزن ها باعث کاهش تعداد وزن های یادگیری نسبت به حالت کاملاً متصل می شود، همچنین توانایی مدل سازی وابستگی ها و ارتباطات محلی در سیگنال ورودی را فراهم می آورد [1, 20]. در لایه ادغام، از خروجی های لایه در هم پیچش، نمونه برداری محلی انجام می شود و علی رغم کاهش سایز ماتریس ورودی، حاصل این نمونه برداری استخراج ویژگی های اساسی خواهد بود. در نهایت تمام ویژگی های ماتریس های ادغام در یک بردار ستونی قرار خواهد گرفت. و این بردار به عنوان ورودی یک شبکه با تعدادی لایه های مخفی متصل کامل، قرار می گیرد.



شکل ۵: نحوه عملکرد لایه در هم پیچش و ادغام

جهت تشخیص صدای گریه نوزاد ابتدا سیگنال های صدا به دو بخش تقسیم شده اند که هر بخش شامل ۱۰۸ پنجره زمانی بوده و هر پنجره حاوی ۵۱۲ نمونه از صدا می باشد. سپس ۱۳ ضریب MFCC از هر پنجره استخراج شده است. بنابراین ورودی شبکه عصبی پیچشی آرایه ای سه بعدی با ابعاد [۱۳ ۱۰۸ ۲۱۱۶] می باشد. شبکه عصبی پیشنهادی دارای ۳ لایه پیچشی است که هر لایه پیچشی شامل ۳۲ فیلتر با سایز 3×3 است.



شکل ۷: نمای کلی شبکه عصبی پیچشی تشخیص گریه

۳-۱ معیارهای ارزیابی

عملکرد روش‌ها در این پژوهش براساس معیارهای دقت و حساسیت سنجیده شده است. دقت، نرخ پیش‌بینی‌های درست به کل پیش‌بینی‌ها می‌باشد. حساسیت نیز به عنوان نسبت میزان پیش‌بینی‌های مثبتی که به درستی تشخیص داده شده اند^{۲۶} به کل پیش‌بینی‌های مثبت، تعریف شده است. فرمول حساسیت و دقت در (۸) و (۹) بیان شده است [10].

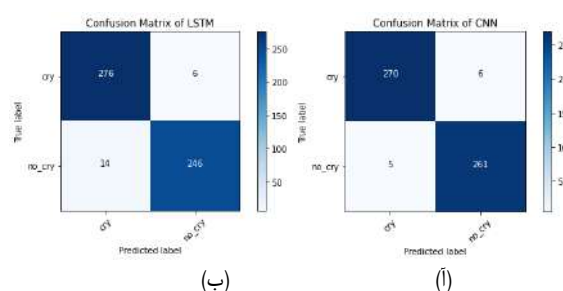
$$Accuracy = \frac{TP + TN}{TP + TN + FN + FP} \quad (8)$$

$$Sensitivity = \frac{TP}{TP + FN} \quad (9)$$

در فرمول (۸) و (۹) TP ، TN ، FN ، FP به ترتیب بیانگر تعداد پیش‌بینی‌های مثبتی که به نادرستی تشخیص داده شده اند^{۲۷}، پیش‌بینی‌های منفی ای که اشتباه تشخیص داده شده اند^{۲۸}، پیش‌بینی‌های مثبتی که درست تشخیص داده شده اند و پیش‌بینی‌های منفی ای که به درستی تشخیص داده شده اند^{۲۹} می‌باشند [9].

۳-۲ مقایسه عملکرد

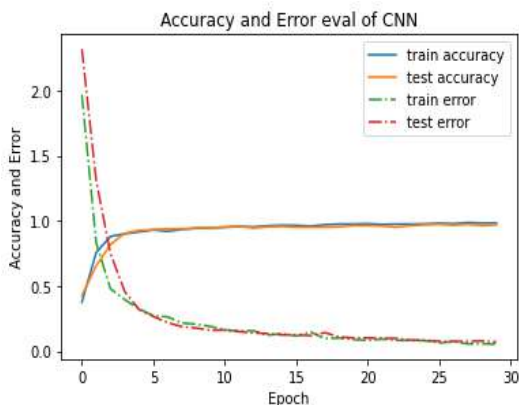
ماتریس درهم ریختگی و نتایج دقت و حساسیت شبکه‌های عصبی به ترتیب در شکل (۸) و جدول (۲) موجود می‌باشد. در شکل (۹) و (۱۰) دقت داده‌های تست و ترین شبکه‌های عصبی برحسب دفعات آموزش ارائه شده است. که با توجه به آن شبکه عصبی CNN با وجود خطای بیشتری در دفعات ابتدایی آموزش، زودتر همگرا شده و به دقت بالاتری رسیده است. در شکل (۱۱) پراکندگی دقت مدل‌ها در ۲۰ بار تکرار با استفاده از نمودار جعبه‌ای نشان داده شده است.



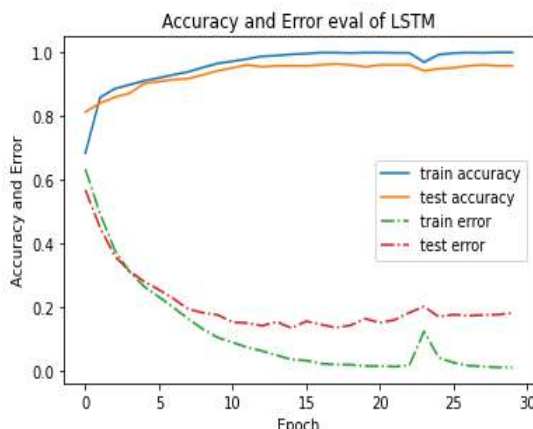
شکل ۸: (ا) ماتریس درهم ریختگی در مدل CNN
(ب) ماتریس درهم ریختگی در مدل LSTM

جدول ۲: مقایسه دقت و حساسیت

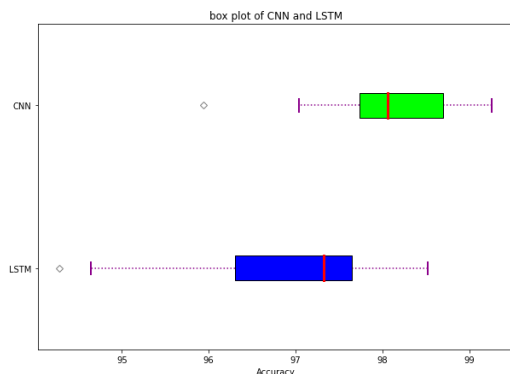
	دقت	حساسیت
CNN	۹۷/۹۷۰۴٪	۹۸/۱۸۱۸٪
LSTM	۹۶/۳۰۹٪	۹۶/۸۲٪



شکل ۹: دقت و خطای داده‌های تست و ترین مدل CNN



شکل ۱۰: دقت و خطای داده‌های تست و ترین مدل LSTM



شکل ۱۱: نمودار جعبه‌ای مدل‌های CNN و LSTM

۴- نتیجه‌گیری

در این پژوهش جهت تشخیص گریه نوزادان با بهره‌گیری از ویژگی MFCC عملکرد شبکه‌های عصبی حافظه طولانی-کوتاه مدت و پیش‌بینی را بررسی کردیم. با توجه به ارزیابی صورت گرفته شبکه عصبی پیش‌بینی با دقت ۹۷/۹۷٪ عملکرد بهتری نسبت به شبکه‌های عصبی حافظه طولانی-کوتاه مدت داشت. در مجموع می‌توان گفت شبکه‌های عصبی حافظه طولانی-کوتاه مدت و پیش‌بینی با اختلاف کمی از هم دارای دقت و حساسیت بالایی می‌باشند که حاکی از عملکرد بالقوه این دو شبکه عصبی در بررسی الگوهای گریه است. در آینده قصد داریم بر ویژگی‌ها و علایم گریه نوزادان متمرکز شویم و به طبقه بندی گریه نوزادان با هدف تشخیص علت گریه شان بپردازیم. همچنین کارایی

- [12] Le, L., Kabir, A., Ji, C., Basodi, S. and Pan, Y., 2019. *Using Transfer Learning, SVM, and Ensemble Classification to Classify Baby Cries Based on Their Spectrogram Images*. 2019 IEEE 16th International Conference on Mobile Ad Hoc and Sensor Systems Workshops (MASSW),.
- [13] Lezhenin, I., Bogach, N. and Pyshkin, E., 2019. *Urban Sound Classification using Long Short-Term Memory Neural Network*. Proceedings of the 2019 Federated Conference on Computer Science and Information Systems.
- [14] Lim, W., Jang, D. and Lee, T., 2016. *Speech emotion recognition using convolutional and Recurrent Neural Networks*. 2016 Asia-Pacific Signal and Information Processing Association Annual Summit and Conference (APSIPA),.
- [15] Limantoro, W., Fatichah, C. and Yuhana, U., 2016. *Application development for recognizing type of infant's cry sound*. 2016 International Conference on Information & Communication Technology and Systems (ICTS),.
- [16] Liu, L., Li, W., Wu, X. and Zhou, B., 2019. *Infant cry language analysis and recognition: an experimental approach*. IEEE/CAA Journal of Automatica Sinica, 6(3), pp.778-788.
- [17] Luo, D., Zou, Y. and Huang, D., 2018. *Investigation on Joint Representation Learning for Robust Feature Extraction in Speech Emotion Recognition*. Interspeech 2018,.
- [18] Manikanta, K., Soman, K. and Manikandan, M., 2019. *Deep Learning Based Effective Baby Crying Recognition Method under Indoor Background Sound Environments*. 2019 4th International Conference on Computational Systems and Information Technology for Sustainable Solution (CSITSS),.
- [19] [19] Mu, Y., Hernández Gómez, L., Montes, A., Martínez, C., Wang, X. and Gao, H., 2017. *Speech Emotion Recognition Using Convolutional- Recurrent Neural Networks with Attention Model*. DEStech Transactions on Computer Science and Engineering, (cii).
- [20] Sainath, T., Kingsbury, B., Saon, G., Soltau, H., Mohamed, A., Dahl, G. and Ramabhadran, B., 2015. *Deep Convolutional Neural Networks for Large-scale Speech Tasks*. Neural Networks, 64, pp.39-48.

زیر نویس ها

Support Vector Machine (SVM) ^{۱۱}
 Transfer Learning ^{۱۲}
 Spectrogram ^{۱۸}
 Recurrent Neural Network ^{۱۹}
 Aliasing ^{۲۰}
 Fast Fourier Transform ^{۲۱}
 Discrete Cosine Transform (DCT) ^{۲۲}
 Sequential ^{۲۳}
 Pooling layer ^{۲۴}
 Batchnormalization ^{۲۵}
 True Positive (TP) ^{۲۶}
 False Positive (FP) ^{۲۷}
 False Negative (FN) ^{۲۸}
 True Negative (TN) ^{۲۹}

شبکه عصبی CNN- LSTM و ترکیبی از شبکه های عصبی عمیق و روش های یادگیری ماشین را بررسی کنیم.

مراجع

- [1] Abdel-Hamid, O., Mohamed, A., Jiang, H., Deng, L., Penn, G. and Yu, D., 2014. *Convolutional Neural Networks for Speech Recognition*. IEEE/ACM Transactions on Audio, Speech, and Language Processing, [online] 22(10), pp.1533–1545.
- [2] Bano, S. and RaviKumar, K., 2015. *Decoding baby talk: A novel approach for normal infant cry signal classification*. 2015 International Conference on Soft-Computing and Networks Security (ICSNS),.
- [3] Choi, J.Y. and Lee, B., 2018. *Combining LSTM Network Ensemble via Adaptive Weighting for Improved Time Series Forecasting*. Mathematical Problems in Engineering, 2018, pp.1–8.
- [4] Dewi, S., Prasasti, A. and Irawan, B., 2019. *Analysis of LFCC Feature Extraction in Baby Crying Classification using KNN*. 2019 IEEE International Conference on Internet of Things and Intelligence System (IoTals),.
- [5] Franti, E., Ispas, I. and Dascalu, M., 2018. *Testing the Universal Baby Language Hypothesis - Automatic Infant Speech Recognition with CNNs*. 2018 41st International Conference on Telecommunications and Signal Processing (TSP),.
- [6] https://github.com/giulbia/baby_cry_detection (accessed on 20 November 2021)
- [7] <https://github.com/gveres/donateacry-corpus> (accessed on 20 November 2021)
- [8] <https://github.com/karoldvl/ESC-5050> (accessed on 20 November 2021)
- [9] K, A., Vincent, P., Srinivasan, K. and Chang, C., 2021. *Deep Learning Assisted Neonatal Cry Classification via Support Vector Machine Models*. Frontiers in Public Health, 9.
- [10] Lahmiri, Salim, et al. "Biomedical Diagnosis of Infant Cry Signal Based on Analysis of Cepstrum by Deep Feedforward Artificial Neural Networks." IEEE Instrumentation & Measurement Magazine, vol. 24, no. 2, Apr. 2021, pp. 24–29,.
- [11] Le, Ho, Lee and Jung, 2019. *Application of Long Short-Term Memory (LSTM) Neural Network for Flood Forecasting*. Water, 11(7), p.1387.

Mel frequency cepstral coefficients (MFCC) ^۱
 Convolutional Neural Network (CNN) ^۲
 Long Short-Term Memory (LSTM) ^۳
 Accuracy ^۴
 Sensitivity ^۵
 Confusion matrix ^۶
 Support-vector machine ^۷
 K Nearest Neighbor ^۸
 Linear Frequency Cepstral coefficients (LFCC) ^۹
 Linear Prediction Cepstral ^{۱۰}
 Linear Prediction Cepstral Coefficients ^{۱۱}
 Bark Frequency Cepstral Coefficients ^{۱۲}
 Compressed sensing method ^{۱۳}
 Artificial neural network ^{۱۴}
 Feed forward neural network ^{۱۵}

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

تعیین کارایی در تحلیل پوششی داده فازی با تقریب نزدیکترین بازه

امیر رحیمی^۱، حسن میش‌مست نهی^۲، فرانک حسین‌زاده سلجوقی^۳

^۱ دانشجوی دکتری، گروه آموزشی ریاضی، دانشگاه سیستان و بلوچستان، زاهدان، amirrahimi525@gmail.com

^۲ استاد، گروه آموزشی ریاضی، دانشگاه سیستان و بلوچستان، زاهدان، hmnehi@hamoon.usb.ac.ir

^۳ دانشیار، گروه آموزشی ریاضی، دانشگاه سیستان و بلوچستان، زاهدان، saljooghi@math.usb.ac.ir

چکیده: در مدل‌های رایج DEA، داده‌های ورودی و خروجی تنها به شرط قطعیت بررسی می‌شوند. با این حال، در محیط‌های دنیای واقعی، عدم قطعیت اغلب به شکل محیط‌های فازی و تصادفی رخ می‌دهد. هنگامی که داده‌ها به صورت نامشخص و مبهم توصیف می‌شوند، لزوم استفاده از نظریه فازی ظاهر می‌شود. از آنجایی که روش‌های حل‌کننده DEA فازی به چهار گروه اصلی تقسیم می‌شوند، در این مطالعه، رویکرد متفاوتی برای حل مدل DEA فازی ارائه شده است. مدل DEA فازی ابتدا با نزدیکترین تقریب بازه یک عدد فازی به یک DEA بازه‌ای تبدیل می‌شود. روش‌های مرسوم برای حل DEA بازه‌ای از رویکردهای مرز ثابت و مرز متغیر استفاده می‌کنند. اما در این پژوهش از رویکردهای متفاوتی برای سنجش کارایی واحدهای تصمیم‌گیری استفاده شده است. مدل‌های فازی DEA عملکرد واحدهای تصمیم‌گیری را بر اساس فاصله زمانی اندازه‌گیری می‌کنند. بنابراین، یک رویکرد پیشنهادی بر اساس فاصله مشخص شده برای مقایسه و رتبه‌بندی کارایی DMUها معرفی می‌شود. یک مثال عددی برای نشان دادن کاربردهای مدل‌های DEA فازی پیشنهادی ارائه شده است.

کلمات کلیدی: DEA فازی، DEA بازه‌ای، تقریب نزدیکترین بازه

کارایی مجموعه DMUها در محیط‌های فازی و یا بازه‌ای، می‌تواند یک مطالعه ارزشمند در بسیاری از مسائل تصمیم‌گیری باشد. چنین مطالعه‌ای نیازمند توسعه نظریه و روش‌های DEA و کاربردهای واقعی آن در دنیای واقعی است.

نخستین بار کوپر و همکاران تکنیک تحلیل پوششی داده‌ها را در کنار داده‌های غیردقیق قرار دادند [4,5]. مدل DEA در نتیجه مدل DEA نادقیق نامیده می‌شود که یک مسئله برنامه‌ریزی غیر خطی (NLP) است و از طریق یک سری تغییرات مقادیر و تغییر متغیر به یک مسئله برنامه‌ریزی خطی (LP) تبدیل می‌شود. بنابراین، مدل‌های DEA فازی شکل مدل‌های برنامه‌ریزی خطی فازی را به خود می‌گیرند. روش‌های حل DEA فازی در ۴ گروه اصلی تقسیم‌بندی می‌شود: رویکرد تلورانس، رویکرد α -برش، رویکرد رتبه‌بندی و رویکرد امکان که در این بین رویکرد α -برش مورد توجه بیشتری قرار گرفته است.

در اغلب روش‌های حل DEA فازی با رویکرد α -برش مسئله برنامه‌ریزی خطی به یک مسئله برنامه‌ریزی خطی بازه‌ای (ILP) تبدیل می‌شود. سپس با توجه به تکنیک‌های موجود به حل مسائل برنامه‌ریزی خطی بازه‌ای می‌پردازند. نتایج کارایی بصورت بازه‌ای تعریف می‌شود. ساعتی و همکاران [6] در تبدیل

۱- مقدمه

یکی از ابزارهای مناسب و کارآمد در زمینه ارزیابی کارایی واحدهای تصمیم‌گیری، تحلیل پوششی داده‌ها (DEA) می‌باشد که بعنوان یک روش غیرپارامتری به منظور محاسبه کارایی واحدهای تصمیم‌گیرنده استفاده می‌شود. اولین بار فارال مدلی برای ارزیابی و محاسبه کارایی با ورودی‌های چندگانه و یک خروجی ارائه داد [1]. تقریباً پس از دو دهه چارنز، کوپر و رودرز این تکنیک را برای چند خروجی تعمیم دادند و آن را تحلیل پوششی داده‌ها نامیدند [2]. در سال‌های اخیر در اغلب کشورهای جهان برای ارزیابی عملکرد نهادها و دیگر فعالیت‌های رایج در ارزیابی سازمان‌ها و صنایع مختلف مانند صنعت بانکداری، پست، بیمارستان‌ها، مراکز آموزشی، نیروگاه‌ها، پالایشگاه‌ها و ... کاربردهای متفاوتی از تحلیل پوششی داده‌ها دیده می‌شود [3]. در مدل‌های متداول DEA داده‌های ورودی و خروجی تنها تحت شرط قطعیت مورد بررسی قرار می‌گیرند. با این وجود، در محیط‌های واقعی اغلب عدم قطعیت در قالب محیط‌های فازی و تصادفی ظهور می‌یابد. زمانی که داده‌ها بصورت نادقیق و یا بطور مبهم توصیف شوند، ضرورت بکارگیری از نظریه فازی در نمایش این نوع از داده‌ها ایجاد می‌شود. بنابراین چگونگی مدیریت اطلاعات یا

صورت فازی مدل CCR به صورت قطعی، رویکرد α -برش را در قالب یک برنامه‌ریزی بازه‌ای بکار گرفتند و برای حل این مساله به هر بازه متغییری اختصاص دادند که در این بازه می‌توانست مقدار گیرد. بدین ترتیب پارامترهای فازی مدل که بصورت بازه‌ای درآمده بودن با متغیرهای حقیقی جایگزین می‌شوند.

Smirlis and Despotis داده‌های بازه‌ای یا کراندار در DEA را بررسی کردند، و یک طبقه‌بندی سه دسته‌ای کارایی را نشان دادند [7]. علی‌رغم انواع متفاوت داده‌ها، رویکرد فازی از داده‌های فازی داده شده، داده‌های بازه‌ای را ایجاد می‌کند. Wang و همکاران زوج جدیدی از مدل‌های DEA بازه‌ای را برای کار با داده‌های نادقیق، از قبیل داده‌های بازه‌ای، اطلاعات ترجیح ترتیبی، داده‌های فازی و مخلوط آنها تشکیل دادند [8]. مدل‌های DEA بازه‌ای آنها، در مقایسه با مدل DEA نادقیق ابداع شده توسط Cooper و همکاران، آسان‌تر و قابل فهم‌تر است [9-10]. همچنین، در مقایسه با مدل‌های DEA بازه‌ای Despotis and Smirlis، مدل‌های DEA بازه‌ای آنها مرز تولید ثابت و یکنواختی بعنوان محک برای اندازه‌گیری کارایی همه‌ی DMUها استفاده می‌کنند، که مدل‌های آنها را عقلانی و مطمئن‌تر می‌کند. به علاوه، روش کار آنها با اطلاعات ترجیح ترتیبی نسبت به روش Zhu معقول‌تر به نظر می‌رسد [11]. Entani و همکاران کارایی DEA را هم از دیدگاه خوشبینانه و هم از دیدگاه بدبینانه در نظر گرفتند [12]. در مدل‌های DEA آنها، با استفاده از کارایی‌های خوشبینانه و بدبینانه، یک بازه تشکیل می‌شود، ولی مدل آنها برای محاسبه‌ی کارایی بدبینانه دارای یک عیب اساسی است و آن این است که برخی اطلاعات ورودی و خروجی را در نظر نمی‌گیرد، زیرا عملاً فقط داده‌های یک ورودی و یک خروجی از DMU مورد ارزیابی استفاده می‌شوند و بقیه داده‌های ورودی و خروجی مورد استفاده قرار نمی‌گیرند.

هدف اصلی مقاله، ایجاد مدل جدید از مدل‌های DEA فازی است که می‌تواند کارایی را بصورت بازه‌ای بیان کند. متفاوت از سایر روش‌های حل DEA فازی، رویکرد پیشنهادی می‌تواند برخلاف رویکرد α -برش، DEA فازی را به یک مسئله برنامه‌ریزی خطی بازه‌ای تبدیل کند. برای مقایسه و رتبه‌بندی کارایی DMUها یک رویکرد مبتنی بر فاصله علامت‌دار ارائه می‌دهیم.

۲- مقدمات

۲-۱- بررسی مجموعه فازی معمولی

در اینجا برخی از مفاهیم و اصطلاحات اساسی را که در مقاله استفاده می‌شود ارائه می‌دهیم.

تعریف ۱: یک مجموعه فازی $\tilde{A}$ بصورت $\tilde{A} = \{(x, \mu_{\tilde{A}}(x)) : x \in A, \mu_{\tilde{A}}(x) \in [0, 1]\}$ در جفت مرتب $(x, \mu_{\tilde{A}}(x))$ ، عنصر نخست x متعلق به مجموعه کلاسیک A و عنصر دوم $\mu_{\tilde{A}}(x)$ متعلق به بازه $[0, 1]$ است که آن را تابع عضویت عنصر $x \in X$ می‌نامند، و یک تابع دوبعدی می‌باشد [14].

تعریف ۲: یک مجموعه فازی $\tilde{A}$ روی $\mathbb{R}$ را عدد فازی می‌نامیم هرگاه:

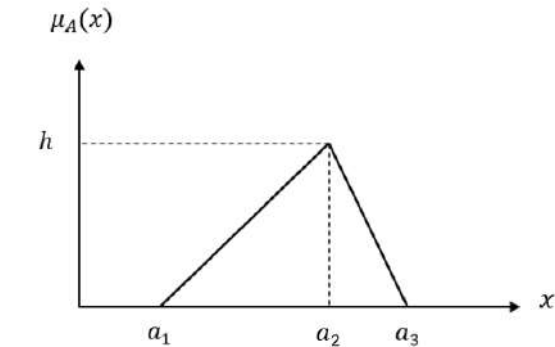
*مجموعه فازی $\tilde{A}$ نرمال باشد، یعنی $\sup_{x \in \mathbb{R}} \mu_{\tilde{A}}(x) = 1$

*فرض کنید $\tilde{A}$ یک مجموعه فازی باشد. $\tilde{A}$ محدب است اگر و تنها اگر به ازای هر $\lambda \in [0, 1]$ نامساوی زیر برقرار باشد:

$$\mu_{\tilde{A}}\{\lambda x_1 + (1 - \lambda)x_2\} \geq \min\{\mu_{\tilde{A}}(x_1), \mu_{\tilde{A}}(x_2)\}$$

*یک مجموعه فازی $\tilde{A}$ روی $\mathbb{R}$ تابع عضویت آن قطعه به قطعه پیوسته باشد. تعریف ۳: عدد فازی مثلثی عمومی که یک حالت خاصی از عدد فازی دوزنقه‌ای

عمومی ($a_2 = a_3$) بوده و روی بازه $[a_1, a_3]$ ، بصورت $A = (a_1, a_2, a_3; h)$ است که در آن $a_1 \leq a_2 \leq a_3$ و $h \in [0, 1]$ می‌باشد. تابع عضویت عدد فازی مثلثی عمومی بصورت زیر می‌باشد:



شکل ۱: عدد فازی مثلثی تخت

$$\mu_{\tilde{A}}(x) = \begin{cases} h \left(\frac{x - a_1}{a_2 - a_1} \right) & a_1 \leq x \leq a_2 \\ h \left(\frac{a_3 - x}{a_3 - a_2} \right) & a_2 \leq x \leq a_3 \\ 0 & x \leq a_1, x \geq a_3 \end{cases} \quad (۱)$$

تعریف ۴: η -برش یک عدد فازی $\tilde{A}$ ، بصورت زیر تعریف می‌شود:

$$A_\eta = \{x \in \mathbb{R} | \mu_{\tilde{A}}(x) \geq \eta\} = [\tilde{A}(\eta), \hat{A}(\eta)],$$

که در آن

$$\tilde{A}(\eta) = \inf\{x \in \mathbb{R} | \mu_{\tilde{A}}(x) \geq \eta\} \text{ و } \hat{A}(\eta) = \sup\{x \in \mathbb{R} | \mu_{\tilde{A}}(x) \geq \eta\}$$

فرض کنید $\tilde{A}$ یک عدد فازی مثلثی و $[\tilde{A}(\eta), \hat{A}(\eta)]$ یک η -برش آن باشد. تقریب نزدیکترین بازه برای $\tilde{A}$ ، بازه بسته $I(\tilde{A}) = [\tilde{I}, \hat{I}]$ است. که

$$\tilde{I} = \int_0^1 \tilde{A}(\eta) d\eta \text{ و } \hat{I} = \int_0^1 \hat{A}(\eta) d\eta, \quad [15]$$

تعریف ۵: تقریب نزدیکترین بازه برای عدد فازی مثلثی $\tilde{A} = (a_1, a_2, a_3)$ $\left[\frac{a_1 + a_2}{2}, \frac{a_3 + a_2}{2}\right]$ است.

۳- مدل‌های DEA و DEA فازی

۳-۱- مدل DEA

فرض کنید n واحد تصمیم‌گیری (DMUs) برای ارزیابی وجود دارد. هر DMU مقدار متفاوتی از m ورودی برای تولید s خروجی بکار می‌برد. به طور خاص، DMU_j مقادیر $x_{ij} (i = 1, \dots, m)$ ورودی را مصرف و مقادیر $y_{rj} (r = 1, \dots, s)$ خروجی را تولید می‌کند. در فرمول‌بندی مدل، $x_{i0} (i = 1, \dots, m)$ و $y_{r0} (r = 1, \dots, s)$ به ترتیب بیانگر بردارهای قطعی نامنفی مقادیر ورودی و خروجی DMU_0 هستند. بنابراین، مدل برنامه‌ریزی خطی DEA یعنی مدل مضربی CCR (ورودی محور) بصورت زیر می‌باشد:

فرض کنید $\tilde{y}_{rj} = (y_{rj}^1, y_{rj}^2, y_{rj}^3)$ و $\tilde{x}_{ij} = (x_{ij}^1, x_{ij}^2, x_{ij}^3)$ که برای $i = 1, \dots, m$ و $j = 1, \dots, n$ و $r = 1, \dots, s$ و $x_{ij}^1 \geq 0$ و $y_{rj}^1 \geq 0$ هستند، بنابراین مدل (۵) می‌تواند بصورت زیر بازنویسی شود:

$$\begin{aligned} \max \quad & \tilde{Z}_o = \sum_{r=1}^s u_r (y_{ro}^1, y_{ro}^2, y_{ro}^3) \\ \text{s.t.} \quad & \sum_{i=1}^m v_i (x_{io}^1, x_{io}^2, x_{io}^3) = (1^1, 1^2, 1^3), \\ & \sum_{r=1}^s u_r (y_{rj}^1, y_{rj}^2, y_{rj}^3) - \sum_{i=1}^m v_i (x_{ij}^1, x_{ij}^2, x_{ij}^3) \leq 0, \\ & v_i \geq 0 \quad \forall i, r, u_r, \end{aligned} \quad (۴)$$

که $1^1 \leq 1^2$ و $1^3 \geq 1^2$ اعداد حقیقی هستند. چندین روش برای حل مدل فوق وجود دارد، روش Π -برش بیشتر از سایر روش‌ها مورد توجه قرار گرفته است، که در آن مدل به یک مسئله برنامه‌ریزی خطی بازه‌ای تبدیل می‌شود. سپس بازه‌ها در دو طرف محدودیت‌ها با یکدیگر مقایسه می‌شوند. روش‌های زیادی برای مقایسه بازه‌ها وجود دارد.

در روش پیشنهادی خودمان، داده ورودی‌های فازی و خروجی‌های فازی، نخست با رویکرد تقریب نزدیکترین بازه، از اعداد فازی به اعداد بازه‌ای تبدیل می‌شوند. پس بنابراین، ورودی‌های فازی $\tilde{x}_{ij} = (x_{ij}^1, x_{ij}^2, x_{ij}^3)$ و خروجی‌های فازی $\tilde{y}_{rj} = (y_{rj}^1, y_{rj}^2, y_{rj}^3)$ با توجه به تقریب نزدیکترین بازه بصورت زیر می‌باشند:

$$\begin{aligned} \left[\frac{x_{ij}^1 + x_{ij}^2}{2}, \frac{x_{ij}^2 + x_{ij}^3}{2} \right] &= [\tilde{x}_{ij}, \hat{x}_{ij}] \\ (i = 1, \dots, m; j = 1, \dots, n;) \\ \left[\frac{y_{rj}^1 + y_{rj}^2}{2}, \frac{y_{rj}^2 + y_{rj}^3}{2} \right] &= [\tilde{y}_{rj}, \hat{y}_{rj}] \\ (r = 1, \dots, s; j = 1, \dots, n;) \end{aligned}$$

حال برای جلوگیری استفاده از مرزهای تولید مختلف در اندازه‌گیری کارایی DMUهای مختلف، مدل‌های جدیدی را ارائه می‌کنیم.

فرض کنید $\tilde{Z}_j = \frac{\sum_{r=1}^s u_r \tilde{y}_{rj}}{\sum_{i=1}^m v_i \tilde{x}_{ij}}$ ($j = 1, \dots, n$)، کارایی DMU_j ($j = 1, \dots, n$) باشد. پس با توجه به تقریب نزدیکترین بازه و محاسبات بازه‌ای داریم:

$$\begin{aligned} Z_j &= \frac{\sum_{r=1}^s u_r [\tilde{y}_{rj}, \hat{y}_{rj}]}{\sum_{i=1}^m v_i [\tilde{x}_{ij}, \hat{x}_{ij}]} \\ &= \frac{[\sum_{r=1}^s u_r \tilde{y}_{rj}, \sum_{r=1}^s u_r \hat{y}_{rj}]}{[\sum_{i=1}^m v_i \tilde{x}_{ij}, \sum_{i=1}^m v_i \hat{x}_{ij}]} \\ &= \left[\frac{\sum_{r=1}^s u_r \tilde{y}_{rj}}{\sum_{i=1}^m v_i \tilde{x}_{ij}}, \frac{\sum_{r=1}^s u_r \hat{y}_{rj}}{\sum_{i=1}^m v_i \hat{x}_{ij}} \right] \quad j = 1, \dots, n; \end{aligned}$$

پس Z_j نیز به یک عدد بازه‌ای بصورت $[\tilde{Z}_j, \hat{Z}_j]$ ($j = 1, \dots, n$) می‌باشد. همچنین

$$0 < Z_j = [\tilde{Z}_j, \hat{Z}_j] = \left[\frac{\sum_{r=1}^s u_r \tilde{y}_{rj}}{\sum_{i=1}^m v_i \tilde{x}_{ij}}, \frac{\sum_{r=1}^s u_r \hat{y}_{rj}}{\sum_{i=1}^m v_i \hat{x}_{ij}} \right] \leq 1, \quad j = 1, \dots, n,$$

حال برای اندازه‌گیری کران پایین و بالای DMU_o ، مدل‌های برنامه‌ریزی خطی زیر را در نظر می‌گیریم:

$$\begin{aligned} \max \quad & Z_o = \sum_{r=1}^s u_r y_{ro} \\ \text{s.t.} \quad & \sum_{i=1}^m v_i x_{io} = 1, \\ & \sum_{r=1}^s u_r y_{rj} - \sum_{i=1}^m v_i x_{ij} \leq 0, \quad \forall j \\ & u_r, v_i \geq 0 \quad \forall i, r \end{aligned} \quad (۲)$$

تعریف (کارایی): DMU_o یک مجموعه کاراست اگر $Z_o^* = 1$ باشد.

در این مقاله تمرکز بر مدل CCR خواهد بود. تمام مدل‌های دیگر، گسترشی از مدل CCR هستند که از طریق تغییر مجموعه امکان تولید مدل CCR و یا مجموع متغیرهای کمکی بدست می‌آیند.

۳-۲- مدل DEA فازی

مدل فوق تنها می‌تواند در حالتی که داده‌ها دقیق بوده مورد استفاده قرار گیرد. به دلیل کمبود دانش و اطلاعات کامل، ریاضیات دقیق برای مدل‌سازی یک سیستم پیچیده کافی نیست. در دنیای واقعی، تصمیمات براساس داده‌های کیفی و کمی است؛ از طرفی DEA فازی یک ابزار قدرتمند برای ارزیابی عملکرد DMUها با داده‌های نادقیق است. مدل‌های DEA فازی به شکل مدل‌های برنامه‌ریزی خطی فازی می‌باشند. بنابراین، مدل CCR با ضرایب فازی بصورت زیر می‌باشد.

$$\begin{aligned} \max \quad & \tilde{Z}_o = \sum_{r=1}^s u_r \tilde{y}_{ro} \\ \text{s.t.} \quad & \sum_{i=1}^m v_i \tilde{x}_{io} = \tilde{1}, \\ & \sum_{r=1}^s u_r \tilde{y}_{rj} - \sum_{i=1}^m v_i \tilde{x}_{ij} \leq 0, \quad \forall j \\ & v_i \geq 0 \quad \forall i, r, u_r, \end{aligned} \quad (۳)$$

جایی که “~” فازی را نشان می‌دهد. $\tilde{x}_{ij}$ و $\tilde{y}_{rj}$ بترتیب ورودی‌های فازی و خروجی‌های فازی هستند. چون ضرایب در مدل CCR فازی اعداد فازی هستند، مدل‌های CCR فازی نمی‌توانند توسط یک برنامه‌ریزی خطی استاندارد مانند مدل CCR قطعی حل شوند. پس بنابراین، با ورودی‌ها و خروجی‌های فازی، مدل DEA قطعی بایستی تعمیم داده شود. در بخش‌های بعدی، روشی برای حل مدل (۵) ارائه می‌دهیم.

۴- رویکرد پیشنهادی

در میان انواع مختلف اعداد فازی، اعداد فازی مثلثی از اهمیت بسیار ویژه‌ای برخوردار هستند. بنابراین، ورودی‌های فازی و خروجی‌های فازی DMUها را بصورت اعداد فازی مثلثی در نظر می‌گیریم.

۵- مثال عددی

مثال عددی استفاده شده از Guo و Tanaka گرفته شده است [13]. اطلاعات برای مثال عددی که شامل پنج DMU با دو ورودی فازی و دو خروجی فازی است در جدول ۱ بیان شده است. ورودی‌ها و خروجی‌های فازی، اعداد فازی مثلثی متقارن که یک حالت خاص از اعداد فازی مثلثی هستند در نظر گرفته شده است.

جدول ۱. داده تاناکا و گو

	DMU1	DMU2	DMU3	DMU4	DMU5
Dat					
a					
Inpu	(3.5,	(2.9,	(4.4,	(3.4,	(5.9,
t1	4.0,4.5)	2.9,	4.9,5.4	4.1,	6.5,
		2.9)	)	4.8)	7.1)
Inpu	(1.9, 2.1,	(1.4,	(2.2,	(2.2,	(3.6,
t2	2.3)	1.5,	2.6,	2.3,	4.1,
		1.6)	3.0)	2.4)	4.6)
Outp	(2.4, 2.6,	(2.2,	(2.7,	(2.5,	(4.4,
ut1	2.8)	2.2,	3.2,	2.9,	5.1,
		2.2)	3.7)	3.3)	5.8)
Outp	(3.8, 4.1,	(3.3,	(4.3,	(5.5,	(6.5,
ut2	4.4)	3.5,	5.1,	5.7,	7.4,
		3.7)	5.9)	5.9)	8.3)

نتایج رویکرد پیشنهادی خودمان در جدول ۲ بیان شده است. نتایج نشان می‌دهد که DMU#2 #4 کارایی قوی بازه‌ای هستند در حالی که DMU5 کارایی ضعیف می‌باشد. DMU3 ناکارایی بازه‌ای است. همچنین نتایج حاصل نشان می‌دهد که $\hat{\theta}_0^* = \hat{Z}_0^* \leq \hat{H}_0^* \leq \check{\theta}_0^* \leq \check{H}_0^* \leq \check{Z}_0^*$ است.

جدول ۲. نتایج کارایی بازه‌ای با مدل پیشنهادی	
DMU	$[\check{Z}_0^*, \hat{Z}_0^*]$
DMU1	[0.8129, 0.9093]
DMU2	[1, 1]
DMU3	[0.7722, 0.9577]
DMU4	[1, 1]
DMU5	[0.9208, 1]

۶- نتیجه

به دلیل کمبود دانش اطلاعات کامل، ریاضیات دقیق برای مدل‌سازی یک سیستم پیچیده کافی نیست. از این رو، مدل‌های DEA فازی به شکل مدل‌های برنامه‌ریزی خطی فازی می‌باشند. در این مقاله، مدل CCR با ضرایب فازی در نظر گرفته شده است. ما یک روش برای تبدیل این مسئله به مدل برنامه‌ریزی خطی بازه‌ای براساس تقریب نزدیکترین بازه پیشنهاد می‌کنیم. در این روش، اعداد فازی مثلثی براساس تکنیک تقریب نزدیکترین بازه به اعداد بازه‌ای تبدیل می‌شوند. لذا مدل DEA فازی با توجه به تکنیک بکار گرفته شده به دو مدل DEA معمولی تبدیل می‌شود، که کارایی را بصورت بازه‌ای نشان می‌دهد. کارایی کران بالا، کارایی DMU تحت ارزیابی را در حالتی که خود DMU و سایر DMUها در بهترین حالت ممکن قرار دارند نشان می‌دهد.

$$\begin{aligned}
 \max \quad & \hat{Z}_0 = \sum_{r=1}^s u_r \hat{y}_{r0} \\
 \text{s.t.} \quad & \sum_{i=1}^m v_i \hat{x}_{i0} = 1 \\
 \text{s.t.} \quad & \sum_{r=1}^s u_r \hat{y}_{rj} - \sum_{i=1}^m v_i \hat{x}_{ij} \leq 0, \quad j = 1, \dots, n, \\
 & v_i \geq 0 \quad \forall i, r, u_r,
 \end{aligned} \quad (5)$$

$$\begin{aligned}
 \max \quad & \check{Z}_0 = \sum_{r=1}^s u_r \check{y}_{r0} \\
 \text{s.t.} \quad & \sum_{i=1}^m v_i \check{x}_{i0} = 1, \\
 \text{s.t.} \quad & \sum_{r=1}^s u_r \check{y}_{rj} - \sum_{i=1}^m v_i \check{x}_{ij} \leq 0, \quad j = 1, \dots, n, \\
 & v_i \geq 0 \quad \forall i, r, u_r,
 \end{aligned} \quad (6)$$

که $\hat{Z}_0$ بهترین کارایی نسبی DMU₀ در حالتی که تمام DMUها در بهترین حالت فعالیت تولید هستند را نشان می‌دهد، در حالی که $\check{Z}_0$ کارایی نسبی DMU₀ در حالتی که تمام DMUها در بدترین حالت فعالیت تولید خود هستند را نشان می‌دهد. بنابراین، معانی $\hat{H}_0$ ، $\check{H}_0$ ، $\hat{\theta}_0$ و $\check{\theta}_0$ با معانی $\hat{Z}_0$ ، $\check{Z}_0$ متفاوت است.

در ادامه، کارایی DMU تحت ارزیابی را می‌توانیم بصورت عدد بازه‌ای بیان کنیم.

تعریف ۷: DMU₀، کارایی قوی بازه‌ای DEA می‌نامیم، اگر در بازه‌ی کارایی $[\check{Z}_0^*, \hat{Z}_0^*]$ ، کارایی کران پایین و بالا برابر با یک باشد؛ یعنی $\check{Z}_0^* = \hat{Z}_0^* = 1$.

تعریف ۸: DMU₀، کارایی ضعیف بازه‌ای DEA می‌نامیم، اگر در بازه‌ی کارایی $[\check{Z}_0^*, \hat{Z}_0^*]$ ، کارایی کران بالا برابر با یک باشد؛ یعنی $\hat{Z}_0^* = 1$ ؛ در غیر اینصورت، ناکارایی بازه‌ای DEA می‌نامیم، اگر $\hat{Z}_0^* < 1$ باشد.

گاهی اوقات در ارزیابی DMUها، کارایی کران پایین برخی DMUها از کارایی کران بالا DMU مورد نظر بیشتر می‌شود. در چنین شرایطی نیازمند به یک هدف ثانویه برای برقراری بازه کارایی هستیم. این هدف ثانویه به منظور برقراری بازه‌ی کارایی به مدل (۱۴) با اضافه شدن قید اضافی $\sum_{r=1}^s u_r \check{y}_{r0} \leq \sum_{r=1}^s u_r^* \hat{y}_{r0}$ محقق می‌شود.

تعریف ۹: برای هر DMU که کارایی کران پایین بیشتر از کارایی کران بالا باشد، مدل کارایی کران پایین (مدل (۱۴)) نیاز به قید اضافه $\sum_{r=1}^s u_r \check{y}_{r0} \leq \sum_{r=1}^s u_r^* \hat{y}_{r0}$ دارد. که در جواب بهینه این قید بصورت معادله برقرار است. تعریف ۱۰: اگر $\hat{Z}_0^*$ و $\check{Z}_0^*$ مقادیر بهینه تابع هدف مدل‌های (۱۴) و (۱۳)، $\hat{H}_0^*$ و $\check{H}_0^*$ مقادیر بهینه تابع هدف مدل‌های (۸) و (۷)، $\hat{\theta}_0^*$ و $\check{\theta}_0^*$ مقادیر بهینه تابع هدف مدل‌های (۱۰) و (۹)، پس $\hat{\theta}_0^* = \hat{Z}_0^* \leq \hat{H}_0^* \leq \check{\theta}_0^* \leq \check{H}_0^* \leq \check{Z}_0^*$.

European Journal of Operational Research, 136(1), 32-45, 2002.

- [13] Guo, P. and H. Tanaka, "Fuzzy DEA: A Perceptual Evaluation Method," Fuzzy Sets and Systems 119, 149–160, 2001

در حالی که کارایی کران پایین، کارایی آن را در حالتی که خود آن و سایر DMU ها در بدترین حالت ممکن قرار دارند را نشان می‌دهد. این روش می‌تواند برای مسائل DEA فازی با اعداد فازی غیر مثلثی و همچنین مسائل DEA بازه‌ای بکار گرفته شود.

در این مقاله، رویکرد پیشنهادی مورد بررسی قرار گرفت که چنین رابطه‌ای در این مقادیر $\hat{\theta}_o^* = \hat{Z}_o^* \leq \hat{H}_o^*$ و $\check{\theta}_o^* \leq \check{H}_o^* \leq \check{Z}_o^*$ در مثال برقرار بود.

سیاسگزاری

هر چند بندگی فقط شایسته ذات بی همتای خالق هستی آفرین است و لیکن به مصادیق حدیث شریفه (هر کس حرفی به من بیاموزد مرا بنده خود کرده است) معلم هم، چنان شأن و مقام والایی دارد که دریای علم و معرفت و پرچم دار توحید هم خود را بنده کسی می‌داند که حتی یک حرف به او بیاموزد. لذا بر خود لازم می‌دانم که از زحمات بی دریغ استادان ارجمند جناب آقای دکتر حسن میش مست نهی و سرکار خانم دکتر فرانک حسین زاده سلجوقی از صمیم قلب تقدیر و تشکر نمایم.

مراجع

- [1] Farrell M.J. *The Measurement of Productive Efficiency*. Journal of Royal Statistical Society, A, 120, pp. 253-281. 1957.
- [2] Charnes, A., Cooper, W. W., and Rhodes, E., *Measuring the efficiency of decision making units*, European Journal of Operational Research, 2(6), 429–444, 1978.
- [3] Cooper, W. W., Seiford, L. M., Tone, K, *Data envelopment analysis: A comprehensive text with models, applications, reference and DEA-solver software*. London: Kluwer Academic Publishers. 2000.
- [4] Cooper, W.W., Park, K.S., Yu, G., *IDEA and AR-IDEA, Models for dealing with imprecise data in DEA*, Management Science, 45:597-607, 1999.
- [5] Zhu, J., *Imprecise data envelopment analysis (IDEA) A review and improvement with an application*, European Journal of Operational Research, 144: 513-529. 2003.
- [6] Saati, S., Memariani, A., Jahanshahloo, G.R., *Efficiency analysis and ranking of DMUs with fuzzy data*. Fuzzy Optimization and Decision Making 1, 255–267. 2002.
- [7] Despotis, Dimitris K., & Smirlis, Yiannis G. *Data envelopment analysis with imprecise data*. European Journal of Operational Research, 140(1), 24-36. 2002.
- [8] Wang, Ying-Ming, Greatbanks, Richard, & Yang, Jian-Bo. *Interval efficiency assessment using data envelopment analysis*. Fuzzy Sets and Systems, 153(3), 347-370. 2005.
- [9] Cooper, W. W., Park, K. S., & Yu, G. *IDEA (Imprecise Data Envelopment Analysis) with CMDs (Column Maximum Decision Making Units)*. Journal of the Operational Research Society, 52(2), 176-181, 2001a.
- [10] Cooper, William W., Park, Kyung Sam, & Yu, Gang, *IDEA and AR-IDEA: Models for Dealing with Imprecise Data in DEA*. Management Science, 45(4), 597-607, 1999.
- [11] Zhu, Joe, *Imprecise DEA via Standard Linear DEA Models with a Revisit to a Korean Mobile Telecommunication Company*. Operations Research, 52(2), 323-329. 2004.
- [12] Entani, Tomoe, Maeda, Yutaka, & Tanaka, Hideo. *Dual models of interval DEA and its extension to interval data*.

تغییر الگوریتم GRASP برای خوشه‌بندی داده‌ها

نجمه نظری^۱ و محمدعلی یعقوبی^۲

^۱ بخش ریاضی کاربردی، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، n.nazari@math.uk.ac.ir

^۲ دانشیار، بخش ریاضی کاربردی، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، yaghoobi@uk.ac.ir

چکیده: خوشه‌بندی نقش مهمی در آنالیز داده‌ها دارد. هدف از خوشه‌بندی، افراز مجموعه‌ی متناهی از داده‌ها در گروه‌ها یا خوشه‌های مجزاست به طوری که داده‌ها با بیشترین تشابه به گروه یا خوشه یکسان اختصاص داده شوند و داده‌های غیرمتشابه در خوشه‌های مختلف قرار گیرند. آنالیز خوشه با مجموعه داده‌هایی که به صورت یک بردار از اندازه‌ها یا یک نقطه در فضای چندبعدی نشان داده می‌شوند، سر و کار دارد. در این مقاله، از معیار حداکثر فاصله‌ی بین نقاط در خوشه‌های مختلف، برای خوشه‌بندی داده‌ها استفاده می‌شود و الگوریتم GRASP برای حل مدل بکار می‌رود. نتایج آزمایش‌های عددی انجام شده، کارایی روش را به خوبی نشان می‌دهد.

کلمات کلیدی: خوشه‌بندی، GRASP، اندازه تشابه، برنامه‌ریزی عدد صحیح غیرخطی.

۱ مقدمه

پیشرفت‌های بوجود آمده در جمع‌آوری داده‌ها و قابلیت‌های ذخیره‌سازی در طی دهه‌های اخیر باعث شده است که محققان در بسیاری از علوم، با حجم بزرگی از اطلاعات روبرو شوند. محققان در زمینه‌های مختلف مانند مهندسی، اقتصاد، ستاره‌شناسی، زیست‌شناسی و ... هر روز با مشاهدات بیشتری روبرو می‌شوند. افزایش تعداد مشاهدات و همچنین افزایش ابعاد داده‌ها، امروزه چالش‌های جدیدی در تحلیل داده‌ها و فرآیند تصمیم‌گیری بوجود آورده است.

از دهه ۱۹۹۰ آنالیز خوشه به عنوان یک علم میان‌رشته‌ای مهم ظهور کرده است که شامل جوامع مختلف علمی و پژوهشی از جمله ریاضیات، آمار نظری و کاربردی، ژنتیک، بیوشیمی، زیست‌شناسی، علوم رایانه و مهندسی است. از کاربردهای آنالیز خوشه می‌توان به علوم اجتماعی، بازاریابی اطلاعات [۱]، پردازش زبان طبیعی [۲]، تجارت [۳]، تقسیم‌بندی تصویر [۴] و داده‌های بیولوژیکی [۵] اشاره کرد.

حل مسایل خوشه‌بندی از نظر محاسباتی بسیار مشکل است، زیرا نیاز به بررسی یک تعداد نمایی از افرازاها است. به طور کلی تعداد راه‌هایی که می‌توان یک مجموعه N عضوی را به K زیرمجموعه‌ی ناتهی

در مساله خوشه‌بندی فرض می‌شود یک مجموعه داده A از تعداد متناهی نقطه از فضای $\mathbb{R}^n$ داده شده است، یعنی:

$$A = \{a_1, \dots, a_N\}, \quad a_i \in \mathbb{R}^n, \quad i = 1, \dots, N.$$

هدف خوشه‌بندی، افراز مجموعه‌ی A به $K \leq N$ زیرمجموعه‌ی مجزا یا همپوشان C_k ، $k = 1, \dots, K$ ، است، به گونه‌ای که:

$$C_k \neq \emptyset, \quad \text{برای هر } k = 1, \dots, K$$

$$A = \bigcup_{k=1}^K C_k$$

مجموعه‌های C_k ، $k = 1, \dots, K$ ، خوشه‌ها نامیده می‌شوند. اگر هر نقطه داده به یک و تنها یک خوشه اختصاص داده شود، مساله را مساله خوشه‌بندی سخت می‌نامند. برخلاف خوشه‌بندی سخت، در خوشه‌بندی فازی هر داده می‌تواند به بیش از یک خوشه تعلق گیرد. در این مقاله، مساله خوشه‌بندی سخت مورد بررسی قرار می‌گیرد، یعنی علاوه بر فرض‌های فوق، فرض می‌شود:

$$C_k \cap C_j = \emptyset,$$

$$\text{برای } k, j = 1, \dots, K, \text{ و } k \neq j.$$

افراز کرد از طریق فرمول زیر بدست می‌آید:

$$\frac{1}{K!} \sum_{i=1}^K (-1)^{K-i} \binom{K}{i} i^N,$$

که حتی برای مسایل خوشه‌بندی با مقیاس کم، این عدد بسیار بزرگ می‌شود [۶، صفحه ۶۳]. از طرفی تمام این افرازهای بدست آمده با معنی نیستند و با توجه به مساله و بهینه کردن یک سری از معیارها، افرازهای مناسب انتخاب می‌شود. بنابراین با افزایش N ، یافتن یک جواب بهین برای مساله غیرعملی است و در نتیجه روش‌های حل دقیق مساله رها می‌شوند و روش‌های حل ابتکاری و فراابتکاری مورد توجه قرار می‌گیرند.

الگوریتم‌های ابتکاری یکی از روش‌ها برای حل مساله خوشه‌بندی هستند که یک جواب تقریبی برای مساله را جستجو می‌کنند. از مشهورترین و محبوب‌ترین روش‌های ابتکاری الگوریتم k -means است، که با کمینه کردن معیار مجموع مربعات خطا، یک افراز بهین از داده‌ها را محاسبه می‌کند. در مرجع [۷، بخش ۵-۲] الگوریتم k -means و نسخه‌های مختلف آن مورد بررسی قرار گرفته‌اند. یکی دیگر از الگوریتم‌های ابتکاری برای مساله خوشه‌بندی، الگوریتم k -medoids است که الگوریتم آن در مرجع [۷، بخش ۵-۴] آمده است.

اخیراً، تمایل استفاده از روش‌های فراابتکاری افزایش یافته است. روش‌های فراابتکاری، الگوریتم‌هایی هستند که قادر به تولید جواب‌های با ارزش می‌باشند. برخی از روش‌های فراابتکاری اصلی عبارتند از: Tabu Search [۷، بخش ۶-۲]، Simulated Annealing [۷، بخش ۶-۳]، Greedy Randomized Adaptive Search Procedure (GRASP) [۳] و الگوریتم ژنتیک [۸، بخش ۶-۴].

در این مقاله، تابع هدف مدل پیشنهادی توسط Nascimento و همکارانش تغییر داده می‌شود و الگوریتم GRASP برای حل مدل ریاضی ارائه می‌شود. GRASP یک الگوریتم جستجو است که یک جواب اولیه از طریق یک فرآیند نیمه حریصانه می‌سازد و سپس یک جستجو موضعی اطراف هر جوابی که از قبل ساخته شده است، اعمال می‌کند. این الگوریتم در مرجع [۹] پیشنهاد شده است و در مقاله [۱۰] روش GRASP و نسخه‌های مختلف آن مورد بررسی قرار گرفته است. اولین کتاب روی این روش در سال ۲۰۱۶ منتشر شد [۱۱].

در این مقاله، ابتدا در بخش دوم در مورد مدل ریاضی پیشنهادی بحث می‌شود و سپس یک الگوریتم برای حل آن ارائه می‌شود. در بخش سوم نتایج عددی ارائه می‌شود و در نهایت، در بخش چهارم، نتایج بدست آمده از این مقاله بررسی می‌شود.

۲ مدل ریاضی

برنامه‌ریزی ریاضی با موفقیت برای حل مسایل خوشه‌بندی مورد استفاده قرار گرفته است [۱۲]. در سال ۱۹۷۱، Rao انواع مختلف

مسایل برنامه‌ریزی عدد صحیح (۱-۰) خطی و غیرخطی برای مساله خوشه‌بندی ارائه کرد [۱۳]. در سال ۲۰۱۰، Nascimento و همکارانش مدل Rao را اصلاح می‌کنند و یک مدل ارائه می‌دهند که تابع هدف، فاصله بین اشیاء داخل خوشه یکسان را به حداقل می‌رساند [۱۴]:

$$\min \sum_{i=1}^{N-1} \sum_{j=i+1}^N \sum_{k=1}^K d_{ij} x_{ik} x_{jk} \quad (1)$$

s.t.

$$\sum_{k=1}^K x_{ik} = 1, \quad i = 1, \dots, N, \quad (2)$$

$$\sum_{i=1}^m x_{ik} \geq 1, \quad k = 1, \dots, K, \quad (3)$$

$$x_{ik} \in \{0, 1\}, \quad i = 1, \dots, N, k = 1, \dots, K. \quad (4)$$

که در آن، x_{ik} یک متغیر ۰-۱ است که اگر شیء i ام به کلاس k ام اختصاص داده شود، مقدار یک را و در غیر این صورت مقدار صفر را اختیار می‌کند. d_{ij} فاصله بین اشیاء i و j است. قیدهای (۲) تضمین می‌کنند که شیء i ام تنها به یک خوشه اختصاص داده شود و قیدهای (۳) تضمین می‌کنند که خوشه k ام ناتمام نباشد. در نهایت، قیود (۴) اطمینان می‌دهند که متغیرهای x_{ik} دودویی هستند. در این مقاله، تابع هدف مدل Nascimento و همکارانش تغییر داده می‌شود و حداکثر فاصله بین اشیاء در خوشه‌های متفاوت بهینه می‌شود.

۱-۲ الگوریتم GRASP

Nascimento و همکارانش برای حل مدل پیشنهادی خود، یک الگوریتم براساس GRASP را مورد بررسی قرار دادند. GRASP یک روش فراابتکاری است که از دو مرحله اساسی تشکیل شده است. در مرحله اول، روش یک جواب اولیه از طریق یک فرآیند نیمه حریصانه تولید می‌کند. در مرحله دوم، روش برای یافتن یک جواب بهتر در همسایگی جواب فعلی به جستجو می‌پردازد و اگر جواب بهتری پیدا شود با جواب فعلی جایگزین می‌شود [۱۰].

۱-۱-۲ مرحله اول الگوریتم GRASP

برای درک بهتر این مرحله، راحت‌تر است که خوشه‌بندی به صورت یک گراف $G = (V, E)$ در نظر گرفته شود که V مجموعه گره‌ها (داده‌ها) و E مجموعه‌ای از یال‌های بدون جهت وزن‌دار باشند که یال (i, j) داده‌های i و j را متصل می‌کند. فاصله بین اشیاء i و j به عنوان وزن یال (i, j) در نظر گرفته می‌شود. در ابتدا، گراف کامل است، یعنی همه‌ی گره‌ها دارای یک یال مرتبط با گره دیگر در گراف هستند. در گراف کامل، مجموعه داده تنها از یک خوشه تشکیل شده است و تعداد کل یال‌ها $\frac{N(N-1)}{2}$ است. در طول فرآیند، یال‌ها به تدریج از گراف حذف می‌شوند و زیرگراف‌های کامل ناهمبند (کلیک‌ها) ایجاد می‌شود. هر کلیک نشان‌دهنده یک خوشه است. فرآیند زمانی متوقف می‌شود که یک جواب شدنی برای

مساله پیدا شود. حذف یال‌ها توسط الگوریتم ۱ شرح داده شده است.

الگوریتم ۲ روش جستجوی موضعی کلی.

Input: Initial solution.

Output: best solution.

- 1: **while** x is not minimum **do**
- 2: Find $x \in V(x)$ such that $f(x) < f(x^*)$.
- 3: $x^* \leftarrow x$.
- 4: **end while**

در ادامه، الگوریتم جستجوی موضعی پیشنهاد شده توسط Nascimento و همکارانش براساس تابع معیار حداکثر فاصله برون خوشه‌ای تغییر داده خواهد شد. الگوریتم ۳، شامل انتقال یک داده از یک خوشه به خوشه دیگر به منظور بهبود جواب است. جواب فعلی بدست آمده تاکنون را در نظر بگیرید. در این جواب هر داده به یک خوشه تخصیص داده شده است. برای محاسبه $s(i, k)$ ، تنها خوشه داده i ام در جواب فعلی تغییر داده می‌شود، یعنی داده به خوشه k اختصاص داده می‌شود و بقیه داده‌ها در خوشه خود باقی می‌مانند. مقدار تابع هدف برای جواب بدست آمده، محاسبه می‌شود. توجه داشته باشید که تابع هدف در اینجا، محاسبه مجموع فاصله‌های بین داده‌ها در خوشه‌های متفاوت است. در این الگوریتم، در هر مرحله، جواب فعلی با بهترین جواب در همسایگی جایگزین می‌شود.

الگوریتم ۳ جستجوی موضعی.

Input: Initial solution and $it = 1$.

Output: best solution.

- 1: **repeat**
- 2: **for** $i = 1$ to N **do**
- 3: Calculate c such that $s(i, c) = \max_{1 \leq k \leq K} s(i, k)$.
- 4: If the data i does not belong to the cluster c then assign it to the cluster c .
- 5: $it \leftarrow it + 1$.
- 6: **end for**
- 7: **until** solution is not improved any more or $it > 100$.

الگوریتم ۱ مرحله اول الگوریتم GRASP

Input: Data set X and K .

Output: Initial solution.

- 1: **for** $i=1$ to $K-1$ **do**
- 2: **if** $N < 15$ **then**
- 3: Let L be a list with every edge of the graph.
- 4: **else**
- 5: Let
$$n_h = \min\{0.1 \frac{N(N-1)}{2}, \text{total number of actual edges}\}.$$
- 6: Let L be a list with the n_h highest weighted edges.
- 7: **end if**
- 8: Set $c_{max} = \text{the highest weighted edges of } L$.
- 9: Set $c_{min} = \text{the lowest weighted edges of } L$.
- 10: Randomly select an index $\alpha \in [0, 1]$.
- 11: Creating the RCL in a decreasing order with the edges whose weights belong to the interval
$$[c_{max} + \alpha(c_{min} - c_{max}), c_{max}].$$
- 12: Randomly select an edge $(i, j) \in G$ from RCL , $i \neq j$.
- 13: Let $S_i = \emptyset$ and $S_j = \emptyset$.
- 14: Consider all nodes v such that $(v, i) \in G$ and $(v, j) \in G$.
- 15: If the weight of (v, i) is lower than the weight of (v, j) , then $v \in S_i$, else $v \in S_j$.
- 16: Eliminate from the graph G the edges $(v_t, v_z) \in V$ where $v_t \in S_i$ and $v_z \in S_j$.
- 17: **end for**

۲-۱-۲ مرحله دوم الگوریتم GRASP

مرحله دوم از روش GRASP، از یک روش جستجوی موضعی استفاده می‌کند. روش‌های جستجوی موضعی، در همسایگی جواب فعلی به جستجوی جواب بهتر می‌پردازند و زمانی متوقف می‌شود که هیچ جواب بهتری وجود نداشته باشد. مراحل اصلی این روش‌ها در الگوریتم ۲ شرح داده شده است. در این الگوریتم $V(x)$ را به عنوان همسایگی نقطه x و $f(x)$ مقدار تابع در نقطه x در نظر بگیرید.

۳ نتایج عددی

در این بخش، ما نتایج عددی انجام شده برای ارزیابی کارایی و عملکرد مدل پیشنهادی با استفاده از هشت مجموعه داده ارائه می‌دهیم. در جدول ۱ مشخصات اصلی از هر مجموعه داده نشان داده شده است.

جدول ۱: مجموعه داده

مرجع	تعداد خوشه‌ها	تعداد ویژگی‌ها	تعداد داده‌ها	مجموعه داده
[۱۵]	2	9	699	Breast
[۱۶]	3	1213	98	BreastA
[۱۷]	4	1213	49	BreastB
[۱۸]	3	661	141	DLBCLA
[۱۹]	3	661	180	DLBCLB
[۲۰]	3	4	150	Iris
[۲۱]	4	5565	103	MultiA
[۲۱]	4	1000	103	Novartis

برای ارزیابی کارایی مدل، ما نتایج را برای سه اندازه عدم تشابه مختلف: فاصله اقلیدسی، نرم L_1 و همبستگی بدون مرکز بررسی می‌کنیم. فرض کنید a_{ik} ، k امین ویژگی از i امین داده و n تعداد ویژگی‌ها باشد. فاصله اقلیدسی، فاصله بین دو نقطه در فضای $\mathbb{R}^n$ را اندازه‌گیری می‌کنند. اغلب این متر برای خوشه بندی استفاده می‌شود:

$$d_{ij} = \sqrt{\sum_{k=1}^n (a_{ik} - a_{jk})^2}.$$

تابع فاصله بعدی، نرم L_1 است:

$$d_{ij} = \sum_{k=1}^n |a_{ik} - a_{jk}|.$$

همبستگی بدون مرکز، یک همبستگی هندسی است که با زاویه بین اشیا تعریف می‌شود:

$$D_{ij} = \frac{a_i^T a_j}{\|a_i\| \|a_j\|}.$$

همبستگی بدون مرکز عدم تشابه بین اشیا i و j را به صورت

$$d_{ij} = 1 - D_{ij},$$

در نظر می‌گیرد.

برای اعتبارسنجی نتایج با یک معیار اندازه‌گیری خارجی، ما مدل پیشنهادی را با مدل ارائه شده توسط Nascimento و همکارانش و الگوریتم شناخته شده خوشه‌بندی، K-means مقایسه کردیم. مقایسه روی سه روش براساس شاخص CRand پیشنهاد شده در مرجع [۲۲] و شاخص‌های Rand و ضریب Jaccard ارائه شده در مرجع [۶، صفحه ۲۶۶] انجام شده است. برای انجام آزمایش‌ها، از مجموعه داده‌هایی استفاده شده که از قبل برچسب هر داده مشخص است. شاخص‌های بالا بین افراز اصلی از مجموعه داده‌ها و افراز بدست آمده از سه روش مقایسه‌ای انجام می‌دهد. هر چه مقدار شاخص‌ها بزرگتر باشد دو افراز مشابه‌تر هستند و این یعنی روش بهتر عمل کرده است.

در جدول‌های زیر از متر اقلیدسی به عنوان اندازه عدم تشابه استفاده

شده است. جدول (۲) نتایج برای شاخص CRand، جدول (۳) نتایج برای شاخص Rand و در نهایت جدول (۴) نتایج برای شاخص Jaccard_coefficient ارائه می‌دهند:

جدول ۲: شاخص CRand و نرم اقلیدسی

مدل پیشنهادی	K_means	Nascimento	مجموعه داده
Breast	0.9155	0.9212	0.915
BreastA	0.5965	0.5719	0.5735
BreastB	0.2970	0.2490	0.4827
DLBCLA	0.2188	0.2419	0.2419
DLBCLB	0.4336	0.4637	0.4637
Iris	0.8401	0.8449	0.8539
MultiA	0.8096	0.8276	0.8276
Novartis	0.8544	0.7779	0.7779

جدول ۳: شاخص Rand و نرم اقلیدسی

مدل پیشنهادی	K_means	Nascimento	مجموعه داده
Breast	0.9204	0.9389	0.9336
BreastA	0.7286	0.7694	0.7707
BreastB	0.6071	0.7041	0.7789
DLBCLA	0.6049	0.6286	0.6286
DLBCLB	0.6899	0.7023	0.7023
Iris	0.8797	0.8923	0.8988
MultiA	0.8441	0.9084	0.9084
Novartis	0.9202	0.8820	0.8820

جدول ۴: شاخص Jaccard_coefficient و نرم اقلیدسی

مدل پیشنهادی	K_means	Nascimento	مجموعه داده
Breast	0.8651	0.8935	0.8847
BreastA	0.5141	0.5247	0.5263
BreastB	0.2881	0.2580	0.3882
DLBCLA	0.2799	0.28	0.28
DLBCLB	0.3828	0.4	0.4
Iris	0.6959	0.7190	0.7336
MultiA	0.5594	0.6833	0.6833
Novartis	0.7203	0.6098	0.6098

در جدول‌های زیر از متر L_1 به عنوان اندازه عدم تشابه استفاده شده است. جدول (۵) نتایج برای شاخص CRand، جدول (۶)

نتایج برای شاخص Rand و در نهایت جدول (۷) نتایج برای شاخص Jaccard_coefficient ارائه می‌دهند:

جدول ۵: شاخص CRand و نرم L_1

مدل پیشنهادی	K_means	Nascimento	مجموعه داده
Breast	0.8910	0.9243	Breast
BreastA	0.5450	0.5638	BreastA
BreastB	0.0956	0.2210	BreastB
DLBCLA	0.8428	0.8550	DLBCLA
DLBCLB	0.4465	0.5722	DLBCLB
Iris	0.8336	0.8830	Iris
MultiA	0.6915	0.7697	MultiA
Novartis	0.9033	0.9320	Novartis

جدول ۸: شاخص CRand و همبستگی بدون مرکز

مدل پیشنهادی	K_means	Nascimento	مجموعه داده
Breast	0.1442	0.1561	Breast
BreastA	0.5430	0.5677	BreastA
BreastB	0.2926	0.3778	BreastB
DLBCLA	0.4915	0.7448	DLBCLA
DLBCLB	0.3902	0.4649	DLBCLB
Iris	0.9360	0.9550	Iris
MultiA	0.7787	0.8373	MultiA
Novartis	0.7243	0.7582	Novartis

جدول ۶: شاخص Rand و نرم L_1

مدل پیشنهادی	K_means	Nascimento	مجموعه داده
Breast	0.8894	0.9389	Breast
BreastA	0.7551	0.7652	BreastA
BreastB	0.6650	0.6820	BreastB
DLBCLA	0.8869	0.8937	DLBCLA
DLBCLB	0.6952	0.7484	DLBCLB
Iris	0.8737	0.9195	Iris
MultiA	0.7738	0.8782	MultiA
Novartis	0.9532	0.9709	Novartis

جدول ۹: شاخص Rand و همبستگی بدون مرکز

مدل پیشنهادی	K_means	Nascimento	مجموعه داده
Breast	0.5458	0.5493	Breast
BreastA	0.7543	0.7673	BreastA
BreastB	0.6905	0.7398	BreastB
DLBCLA	0.6900	0.8330	DLBCLA
DLBCLB	0.6660	0.7028	DLBCLB
Iris	0.9575	0.9740	Iris
MultiA	0.8808	0.9048	MultiA
Novartis	0.8132	0.8709	Novartis

جدول ۱۰: شاخص Jaccard_coefficient و همبستگی بدون مرکز

مدل پیشنهادی	K_means	Nascimento	مجموعه داده
Breast	0.3981	0.4024	Breast
BreastA	0.5009	0.5212	BreastA
BreastB	0.2806	0.3215	BreastB
DLBCLA	0.3875	0.5961	DLBCLA
DLBCLB	0.3605	0.4007	DLBCLB
Iris	0.8790	0.9239	Iris
MultiA	0.6080	0.6787	MultiA
Novartis	0.4872	0.5820	Novartis

جدول ۷: شاخص Jaccard_coefficient و نرم L_1

مدل پیشنهادی	K_means	Nascimento	مجموعه داده
Breast	0.8196	0.8938	Breast
BreastA	0.5024	0.5179	BreastA
BreastB	0.1992	0.2505	BreastB
DLBCLA	0.7140	0.7224	DLBCLA
DLBCLB	0.3901	0.4708	DLBCLB
Iris	0.6840	0.7819	Iris
MultiA	0.4356	0.5995	MultiA
Novartis	0.8244	0.8873	Novartis

۴ نتیجه گیری

در این مقاله یک مدل ریاضی برای مساله خوشه‌بندی داده‌ها پیشنهاد شد و برای حل آن از GRASP استفاده شد. در

در جدول‌های زیر از همبستگی بدون مرکز به عنوان اندازه عدم تشابه استفاده شده است. جدول (۸) نتایج برای شاخص CRand، جدول (۹)

- [13] M. R. Rao, "Cluster analysis and mathematical programming," *American Statistical Association*, vol.66, no.335, pp.622–626, 1971.
- [14] M. C. V. Nascimento, F. M. B. Toledo, and A. C. P. L. F. de Carvalho, "Investigation of a new GRASP-based clustering algorithm applied to biological data," *American Statistical Association*, vol.37, no.8, pp.1381–1388, 2010.
- [15] K. P. Bennett and O. L. Mangasarian, "Robust linear programming discrimination of two linearly inseparable sets," *Optimization Methods and Software*, vol.1, no.1, pp.23–34, 1992.
- [16] L. J. van't Veer, H. Dai, M. J. van de Vijver, Y. D. He, A. A. Hart, M. Mao, and et al, "Gene expression profiling predicts clinical outcome of breast cancer," *SIAM Journal on Computing*, vol.415, no.6871, pp.530–536, 2002.
- [17] M. West, C. Blanchette, H. Dressman, E. Huang, S. Ishida, S. R, and et al, "Predicting the clinical status of human breast cancer by using gene expression profiles," *Proceedings of the National Academy of Sciences*, vol.98, no.20, pp.11462–11467, 2001.
- [18] S. Monti, K. J. Savage, J. L. Kutok, F. Feuerhake, P. Kurtin, M. Mihm, and et al, "Molecular profiling of diffuse large B-cell lymphoma identifies robust subtypes including one characterized by host inflammatory response," *Blood*, vol.105, no.5, pp.1851–1861, 2005.
- [19] A. Rosenwald, G. Wright, W. C. Chan, J. M. Connors, E. Campo, R. I. Fisher, and et al, "The use of molecular profiling to predict survival after chemotherapy for diffuse large-B-Cell lymphoma," *The New England Journal of Medicine*, vol.346, pp.1937–1947, 2002.
- [20] R. A. Fisher, "The use of multiple measurements in taxonomic problems," *Annals of Eugenics*, vol.7, no.2, pp.179–188, 1936.
- [21] A. I. Su, M. P. Cooke, K. A. Ching, Y. Hakak, J. R. Walker, T. Wiltshire, and et al, "Large-scale analysis of the human and mouse transcriptomes," *Proceedings of the National Academy of Sciences*, vol.99, no.7, pp.4465–4470, 2002.
- [22] L. Hubert and P. Arabie, "Comparing partitions," *Journal of Classification*, vol.2, pp.193–218, 1985.

آزمایش‌های عددی، مدل پیشنهادی تقریباً در تمام جدول‌ها نسبت به K_means نتایج بهتری بدست آورد. مدل ارائه شده با مدل Nascimento در برخی جدول‌ها برابر است ولی در نرم اقلیدسی افراز بهتری برای مجموعه داده‌ها بدست می‌آورد.

مراجع

- [1] M. Charikar, C. Chekuri, T. Feder, and R. Motwani, "Incremental clustering and dynamic information retrieval," *SIAM Journal on Computing*, vol.33, no.6, pp.1417–1440, 2004.
- [2] A. Ushioda and J. Kawasaki, "Hierarchical clustering of words and application to NLP tasks," *Association for Computational Linguistics*, pp.28–41, 1996.
- [3] M. Porter, "Location, competition, and economic development: local clusters in a global economy," *Economic Development Quarterly*, vol.14, no.1, pp.15–34, 2000.
- [4] Z. Wu and R. Leahy, "An optimal graph theoretic approach to data clustering: theory and its application to image segmentation," *IEEE*, vol.15, no.11, pp.1101–1113, 1993.
- [5] P. Festa, "A biased random-key genetic algorithm for data clustering," *Mathematical Biosciences*, vol.245, no.4, pp.76–85, 2013.
- [6] R. Xu and D. C. Wunsch. *Clustering*. Wiley & IEEE, 2009.
- [7] A. M. Bagirov, N. Karimtsa, and S. Taheri. *Partitional Clustering via Nonsmooth Optimization*. Springer, 2020.
- [8] J. R. Cano, O. Cordon, F. Herrera, and L. Sanchez, "A greedy randomized adaptive search procedure applied to the clustering problem as an initialization process using K-Means as a local search procedure," *Intelligent & Fuzzy Systems*, vol.12, no.3, pp.235–242, 2002.
- [9] T. Feo and M. Resende, "A probabilistic heuristic for a computationally difficult set covering problem," *Operations Research Letters*, vol.8, no.2, pp.67–71, 1989.
- [10] M. G. C. Resende and C. C. Ribeiro, "Greedy randomized adaptive search procedures: advances and extensions," *Handbook of Metaheuristics*, pp.169–220, 2019.
- [11] M. G. C. Resende and C. C. Ribeiro. *Optimization by GRASP: Greedy Randomized Adaptive Search Procedures*. Springer, 2016.
- [12] P. Festa, "Combinatorial optimization approaches for data clustering," *Unsupervised Learning Algorithms*, pp.109–134, 2016.

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

حل مسائل کنترل بهینه دارای عدم قطعیت

الناز حسینی^۱، مهدی الله‌دادی^۲، سمانه صردی‌زید^۳

^۱ دانشکده ریاضی، دانشگاه سیستان و بلوچستان، زاهدان، elnaz.hosseini@pgs.usb.ac.ir

^۲ دانشیار، دانشکده ریاضی، دانشگاه سیستان و بلوچستان، زاهدان، m_allahdadi@math.usb.ac.ir

^۳ استادیار، دانشکده صنعت و معدن، دانشگاه سیستان و بلوچستان، خاش، soradizeid@eng.usb.ac.ir

چکیده: در این مقاله، به منظور حل آنلاین مسائل کنترل بهینه و در شرایطی که امکان حل تحلیلی معادلات همیلتونی وجود ندارد، راهبرد کنترل پیش‌بین مدل ارائه می‌شود. در این الگوریتم کنترلی یک تابع هزینه به صورت آنلاین و با در نظر گرفتن قیدهای فیزیکی حاکم بر سیستم، بهینه‌سازی می‌شود تا سیگنال کنترلی در زمان جاری محاسبه شود. بعد از اعمال آن به سیستم و ثبت اندازه‌گیری‌های جدید در زمان بعدی، محاسبات بهینه‌سازی از ابتدا انجام خواهد شد تا سیگنال‌های کنترلی جدید حاصل شوند. برای جبران اثرات پارامترهای دارای عدم قطعیت که در دینامیک سیستم وجود دارند، مساله بهینه‌سازی رایج در کنترل پیش‌بین مدل به صورت یک بهینه‌سازی دو مرحله‌ای $\min \max$ بیان می‌شود. در گام اول بدترین تاثیر پارامتر دارای عدم قطعیت بر روی تابع هزینه محاسبه می‌شود و سپس با استفاده از سیگنال کنترلی مینیمم می‌شود.

کلمات کلیدی: مساله کنترل بهینه، کنترل پیش‌بین مدل، عدم قطعیت، کنترل مقاوم.

پیش‌بین را طوری فرموله کرد که سیستم کنترلی نسبت به این نوع از پارامترها مقاوم شود [1]. لذا در این مقاله، برای حل مسائل کنترل بهینه دارای عدم قطعیت از کنترل پیش‌بین مدل مقاوم استفاده می‌شود.

۱- مقدمه

مسائل کنترل بهینه کلاسیک نقش مهم و گسترده‌ای در طراحی سیستم‌های کنترلی مدرن برعهده دارند. در این گونه مسائل، سیگنال‌های کنترلی باید به نحوی محاسبه شوند که یک تابع هزینه مینیمم شود [2,3,4]. علاوه بر این، در تعیین سیگنال‌های کنترلی بهینه باید مجموعه‌ای از قیود فیزیکی حاکم بر سیستم هم در نظر گرفته شود. به عبارت دیگر، سیگنال‌های کنترلی هر مقدار دلخواهی نیستند و باید در عمل قابل تولید باشند. مسائل کنترل بهینه غالباً توسط اصل پونتریاگین حل می‌شوند. اما عیب این روش آنلاین بودن آن است که کاربرد آن را در سیستم‌های کنترلی واقعی محدود می‌کند. از سویی دیگر، در نظر گرفتن قیدهای فیزیکی، حل مسائل کنترل بهینه را با استفاده از اصل پونتریاگین با چالش‌های بیشتری مواجه می‌کند.

۲- کنترل پیش‌بین مدل

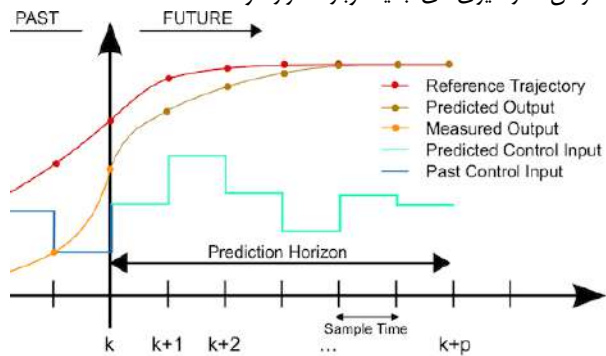
یکی از کنترل‌کننده‌های مهم که در سالیان اخیر بخصوص در صنعت، مورد توجه بسیاری قرار گرفته است، کنترل پیش‌بین مدل است. این کنترلر یک استراتژی چندمتغیره و مقید است و با بهینه‌سازی آنلاین یک تابع هزینه سیگنال‌های کنترلی بهینه را تعیین می‌کند. بنابراین با استفاده از این کنترل‌کننده می‌توان قیدهای فیزیکی حاکم بر سیستم مانند محدوده ولتاژ مجاز قابل اعمال به موتور جریان مستقیم یا محدوده مجاز موقعیت ساقه‌ی شیر برقی را در نظر گرفت. مهمترین دلایل محبوبیت این کنترل‌کننده عبارتند از: بهینه بودن، در نظر گرفتن قیدهای فیزیکی، تعمیم راحت به سیستم‌های چندمتغیره، ذخیره سازی انرژی و هزینه، جبران موثر اغتشاشات وارد بر سیستم، پیاده سازی آسان در سیستم‌های پیچیده، کاربرد در صنعت. در شکل (۱) راهبرد کنترل پیش‌بین مدل نشان داده شده است. در این کنترل‌کننده بر اساس مدل ریاضی پلانت و سیگنال‌های گذشته آن که توسط سنسورها اندازه‌گیری و ثبت شده‌اند، پیش‌بینی حالت‌های سیستم در زمان‌های آینده (در طول افق پیش‌بین) محاسبه می‌شود. در کنترل پیش‌بین مدل (با

کنترل پیش‌بین مدل یک راهبرد کنترلی نوین است که به صورت آنلاین اجرا می‌شود. این الگوریتم کنترلی در حقیقت همان کنترل بهینه است که با در نظر گرفتن یک افق زمانی کوتاه‌تر به فرم آنلاین اجرا می‌شود. همچنین در این روش قیدهای فیزیکی حاکم بر سیستم به سهولت در گام بهینه‌سازی در نظر گرفته می‌شود. از طرف دیگر، در سیستم‌های واقعی ممکن است برخی پارامترها دارای عدم قطعیت باشند. یعنی مقدار دقیق آنها معلوم نیست ولی می‌دانیم که در یک بازه مشخص قرار دارند. در این صورت می‌توان کنترل

فرض افق پیش‌بین ۰/۱ ثانیه) تابع هزینه در بازه های ۰/۱-۰۰۱ ثانیه‌ای
مینیم سازی خواهد شد. افق کوتاه‌تر باعث می‌شود که به صورت آنالین و با
دقت خوبی بتوان همان سیگنال‌های کنترلی را به دست آورد.

با این کار، نیازی به حل تئوری بهینه‌سازی تابع هزینه نداریم و توسط یک
بهینه‌ساز آنالین انجام خواهد شد. پیش‌بینی‌های انجام‌شده به صورت یک تابع
چندمتغیره از سیگنال‌های کنترلی در طول افق کنترل بدست می‌آیند. منظور از
افق کنترل، بازه تعریف‌شده برای سیگنال کنترلی است. برای مثال فرض کنید
که افق پیش‌بین برابر با ۰/۱ ثانیه و افق کنترل هم برابر با ۰/۰۵ ثانیه باشد.
این بدین معنی است که می‌خواهیم با سیگنال‌های کنترلی در فاصله زمانی t
تا $t + ۰/۰۵$ ثانیه خروجی سیستم در فاصله زمانی t تا $t + ۰/۱$ ثانیه را
به مقادیر مطلوب سوق بدهیم و یا تابع هزینه مینیمم شود. با توجه به شکل
(۱)، سیگنال کنترلی در فاصله $t + ۰/۰۵$ تا $t + ۰/۱$ ثابت فرض می‌شود.
همچنین به دلیل اینکه سیگنال کنترلی در عمل توسط یک سیستم دیجیتال
تولید خواهد شد، آن را به صورت بازه‌ای ثابت تعریف می‌کنیم.

حال یک تابع هزینه که اغلب مجموع مربعات خطای ردیابی و سیگنال
کنترلی است در نظر گرفته می‌شود که باز هم یک تابع چندمتغیره از
سیگنال‌های کنترلی در طول افق کنترل خواهد بود. سپس با استفاده از یک
الگوریتم بهینه‌ساز مانند گرادین نزولی و یا الگوریتم ژنتیک، با در نظر گرفتن
قیدهای فیزیکی حاکم بر سیستم تابع هزینه نسبت به سیگنال‌های کنترلی در
طول افق کنترل حداقل سازی می‌شود. بعد از اینکه مقادیر بهینه تعیین شدند،
اولین نمونه از سیگنال کنترلی به پلانت اعمال می‌شود و در زمان‌های بعدی
این فرایند دوباره تکرار خواهد شد. دلیل این امر جبران اثرات عدم قطعیت در
اغتاشات و سیگنال‌های تصادفی است. زیرا عدم قطعیت‌ها و اغتاشات
موجود در سیستم اغلب ماهیت تصادفی دارند و نمی‌توان آنها را در طول افق
پیش‌بینی تخمین زد. از این رو، آخرین اغتاش ثبت شده به عنوان اغتاش
در طول افق پیش‌بینی در نظر گرفته خواهد شد. سپس محاسبات بهینه‌سازی با
این فرض انجام می‌شود تا سیگنال کنترلی در زمان فعلی تعیین شود. بعد از
اعمال آن به سیستم، در نمونه زمانی بعد تمام مراحل بهینه‌سازی تابع هزینه با
در نظر گرفتن اندازه‌گیری‌های جدید دوباره تکرار خواهد شد.



شکل ۱: راهبرد کنترل پیش بین مدل

۳- الگوریتم پیشنهادی

ممساله کنترل بهینه فازی زیر را در نظر می‌گیریم:

$$\tilde{J}(\tilde{u}) = \int_{t_0}^{t_1} \tilde{g}(\tilde{x}(t), \tilde{u}(t), t) dt \quad (1)$$

$$s.t. \quad \tilde{x} = \tilde{h}(\tilde{x}(t), \tilde{u}(t), t) \quad (2)$$

$$\tilde{x}(t_0) = \tilde{x}_0, \tilde{x}(t_1) = \tilde{x}_1. \quad (3)$$

پارامتر $\tilde{p}$ به صورت پارامتر دارای عدم قطعیت در بازه $[p_L \ p_U]$ در نظر
گرفته می‌شود.

با روش‌های مختلف مانند اوایلر یا رانگ کوتای مرتبه چهارم سیستم دینامیکی
(2) به صورت زیر گسسته‌سازی می‌شود:

$$\tilde{x}_{k+1} = \tilde{h}_k(\tilde{x}(t), \tilde{u}(t), t) \quad (4)$$

که در آن منظور از $\tilde{x}_k(t)$ متغیر حالت در نمونه گسسته شده‌ی k ام است. $\tilde{h}$
مدل پیوسته سیستم ولی $\tilde{h}_k$ مدل گسسته را نشان می‌دهد که توسط روش
حل عددی معادلات دیفرانسیل تعیین می‌شود.

رابطه بین اندیس k و زمان واقعی به صورت $t = kT$ است که در آن
 T بیانگر طول گام حل عددی معادلات دیفرانسیل است. حال اگر در سیستم
(2)، T برابر با ۰/۰۰۱ ثانیه در نظر گرفته شود، در فاصله زمانی ۰ تا ۱ ثانیه
۱۰۰۰ نمونه از متغیر x و u وجود خواهد داشت. در روش‌های آفلاین که قبلاً
مطرح شدند تمامی این ۱۰۰۰ نمونه به صورت یکجا تعیین خواهد شد. اما در
این مقاله برای تعیین آنالین سیگنال‌های کنترلی $\tilde{u}$ یک آینده کوتاه‌تر مانند
۰/۰۰۳ ثانیه بعد در نظر گرفته می‌شود.

در زمان صفر، طبق معادلات گسسته سیستم داریم:

$$\begin{aligned} \tilde{x}_1(t) &= \tilde{h}_1(\tilde{x}(t), \tilde{u}(t), t, \alpha) \\ \tilde{x}_2(t) &= \tilde{h}_2(\tilde{x}(t), \tilde{u}(t), t, \alpha) \\ \tilde{x}_3(t) &= \tilde{h}_3(\tilde{x}(t), \tilde{u}(t), t, \alpha) \end{aligned} \quad (5)$$

بنابراین مشاهده می‌شود که متغیرهای $\tilde{x}$ در زمان‌های ۰/۰۰۱، ۰/۰۰۲ و
۰/۰۰۳ ثانیه بعد به سیگنال‌های کنترلی $\tilde{u}$ در زمان‌های ۰، ۰/۰۰۱ و ۰/۰۰۲
وابسته هستند. حال می‌توان با حل یک مساله بهینه‌سازی آنها را طوری تعیین
کرد که آنها تا حد امکان به صفر نزدیک شوند. علاوه بر این، تابع هزینه هم
مینیمم شود. با گسسته‌سازی تابع هزینه (1) رابطه زیر بدست می‌آید:

$$\tilde{J}(\tilde{u}) = \sum_{k=0}^3 \tilde{g}_k(\tilde{x}_k(t), \tilde{u}_k(t), t) \quad (8)$$

بنابراین مقادیر بهینه با حل مساله بهینه‌سازی زیر تعیین خواهند شد:

$$\begin{aligned} (\tilde{x}^*(t), \tilde{u}^*(t)) &= \underset{(\tilde{x}(t), \tilde{u}(t))}{\operatorname{argmin}} \sum_{k=0}^3 \tilde{g}_k(\tilde{x}_k(t), \tilde{u}_k(t), t) \\ s.t. \quad \tilde{x}(t_0) &= \tilde{x}_0, \tilde{x}(t_1) = \tilde{x}_1 \end{aligned} \quad (9)$$

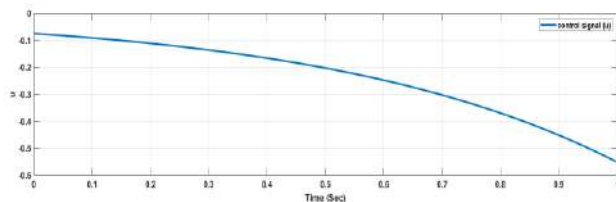
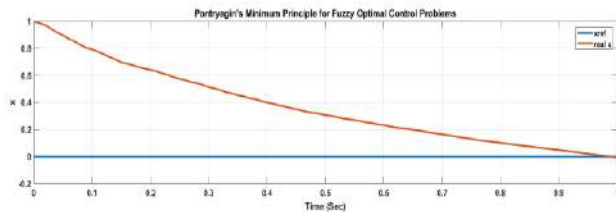
ولی باید توجه کرد که در داخل معادلات $\tilde{x}_1(t)$ تا $\tilde{x}_3(t)$ پارامتر فازی
و دارای عدم قطعیت $\tilde{p}$ وجود دارد. لذا می‌توان بدترین اثر پارامتر $\tilde{p}$ بر روی
تابع هزینه را مینیمم کرد.

$$\begin{aligned} (\tilde{x}^*(t), \tilde{u}^*(t)) &= \underset{(\tilde{x}(t), \tilde{u}(t))}{\operatorname{argmin}} \underset{\tilde{p} \in [p_L \ p_U]}{\operatorname{argmax}} \sum_{k=0}^3 \tilde{g}_k(\tilde{x}_k(t), \tilde{u}_k(t), t) \\ s.t. \quad \tilde{x}(t_0) &= \tilde{x}_0, \tilde{x}(t_1) = \tilde{x}_1 \end{aligned} \quad (10)$$

بنابراین مراحل تعیین مقادیر بهینه شامل دو مرحله بهینه‌سازی
 $\max$ است. مزیت مهم ایده مطرح شده این است که به راحتی می‌توان اثر
قیود $\tilde{x}(t)$ و $\tilde{u}(t)$ را هم در نظر گرفت. فرض کنید که $\tilde{x}(t)$ و $\tilde{u}(t)$ به
ترتیب در بازه‌های $[x_L \ x_U]$ و $[u_L \ u_U]$ تغییرات داشته باشند. پس
مساله بهینه‌سازی به صورت زیر خواهد بود:

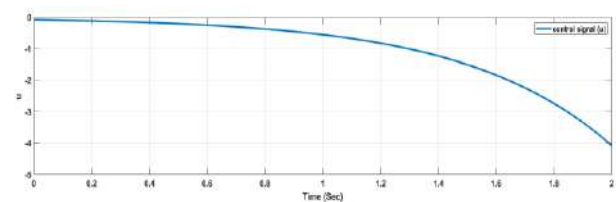
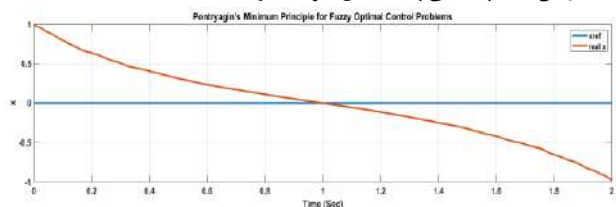
$$\begin{aligned} (\tilde{x}^*(t), \tilde{u}^*(t)) &= \underset{([x_L \ x_U], [u_L \ u_U])}{\operatorname{argmin}} \underset{\tilde{p} \in [p_L \ p_U]}{\operatorname{argmax}} \sum_{k=0}^3 \tilde{g}_k(\tilde{x}_k(t), \tilde{u}_k(t), t) \\ s.t. \quad \tilde{x}(t_0) &= \tilde{x}_0, \tilde{x}(t_1) = \tilde{x}_1 \end{aligned} \quad (11)$$

از حل مساله بهینه‌سازی بالا مقادیر بهینه $\tilde{x}(t)$ و $\tilde{u}(t)$ در زمان‌های ۰،
۰/۰۰۱ و ۰/۰۰۲ حاصل می‌شود ولی اولین نمونه در زمان ۰ به معادلات



شکل ۲: متغیر حالت و کنترل با الگوریتم [2,4]

در شکل (۲) شبیه‌سازی همان کنترل‌کننده برای مدت زمان ۲ ثانیه انجام شده است. مشخص است که متغیر حالت سیستم در حال واگرایی است و با گذشت زمان به منهای بی‌نهایت میل خواهد کرد.



شکل ۳: متغیر حالت و کنترل در [2,4] (شبیه‌سازی به مدت دو ثانیه)

حال در شکل (۳) نمودار تغییرات متغیر حالت و سیگنال کنترلی با استفاده از کنترل پیش‌بین مقاوم نشان داده شده است. $\tilde{p}$ در بازه ۰ تا ۳ تغییر می‌کند و شرایط مرزی همانند معادلات (۱۴) است. افق پیش‌بینی و کنترل به ترتیب برابر با ۳ و ۲ فرض شده است. با توجه به این شکل مشخص است که متغیر حالت از موقعیت ۱ در زمان ۰ به موقعیت تقریباً ۰ در زمان ۱ منتقل شده است و برای زمان‌های بعد از ۱ ثانیه در همان موقعیت ۰ ثابت باقی مانده است. در حالی که الگوریتم پونتریاگین واگرایی است. پس نتیجه مهم این است که کنترل پیش‌بین مقاوم علاوه بر اینکه در همان مدت زمان ۱ ثانیه متغیر حالت سیستم را به صفر هدایت می‌کند، یک کنترل‌کننده پایدار (Stable) است.

دیفرانسیل اصلی اعمال خواهد شد و مراحل بهینه‌سازی دوباره در زمان ۰/۰۱ تکرار خواهد شد.

۱-۳- اعمال الگوریتم پیشنهادی به یک سیستم

نامعین مرتبه اول

برای ارزیابی عملکرد کنترل‌کننده پیش‌بین مدل مقاوم مساله کنترل بهینه زیر در نظر گرفته شده است:

$$\tilde{J}(\tilde{u}) = \int_0^1 \tilde{u}^2(t) dt \quad (12)$$

$$s. t. \quad \tilde{\dot{x}} = \tilde{u}(t) - \tilde{p}(t)\tilde{x}(t) \quad (13)$$

$$\tilde{x}(0) = 1, \tilde{x}(1) = 0. \quad (14)$$

در $\tilde{p} = (0, 1, 3)$ ، $[1, 3]$ آن به صورت بازه زیر است:

$$\tilde{p}_\alpha = [\alpha, 3 - 2\alpha], \quad \forall \alpha \in [0, 1]. \quad (15)$$

لذا در بازه $[0, 3]$ تغییر می‌کند.

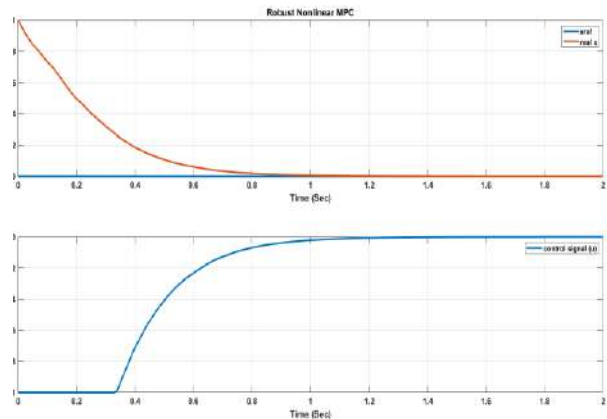
در [2,4] مقادیر بهینه این مساله به صورت یکجا برای تمام بازه زمانی تعیین می‌شود. اما پیش‌نیاز آن این است که معادلات همیلتونی حل شوند. برای این مساله ساده قادر هستیم که شکل دقیق پاسخ‌ها را تعیین کنیم اما اگر دینامیک سیستم غیر خطی و پیچیده باشد امکان حل این معادلات هم وجود ندارد. لذا باید از کنترل پیش‌بین مدل استفاده کرد. البته دلیل دیگر برای روی آوردن به کنترل پیش‌بین مدل آنالین بودن آن است. اگر فرض شود که پارامتر $\tilde{p}$ به صورت تصادفی در بازه مورد نظر تغییرات داشته باشد. معادله سیستم به فرم زیر خواهد بود:

$$\tilde{\dot{x}} = \tilde{u}(t) - \tilde{p}(t)\tilde{x}(t) \quad (16)$$

در چنین حالتی هم معادلات همیلتونی قادر به تعیین مقادیر بهینه نیست. زیرا در عمل مقدار دقیق پارامتر $\tilde{p}(t)$ مشخص نیست. بنابراین راهبرد کنترل پیش‌بین مدل مقاوم پیشنهاد می‌شود. زیرا در این راهبرد ابتدا ماکزیمم تابع هزینه نسبت به پارامتر $\tilde{p}(t)$ در بازه مورد نظر تعیین خواهد شد و سپس این مقدار ماکزیمم با انتخاب مناسب سیگنال‌های کنترلی مینیمم خواهد شد.

در شکل (۱) نمودار متغیر حالت و سیگنال کنترلی با اصل پونتریاگین بیان شده در [2,4] نشان داده شده است. واضح است که طبق صورت مساله متغیر حالت از وضعیت اولیه ۱ در زمان ۰ به وضعیت نهایی ۰ در زمان ۱ ثانیه منتقل می‌شود. در طول این شبیه‌سازی فرض شده است که پارامتر $\tilde{p}$ ثابت و برابر با ۲ و همچنین آلفا معادل با ۰/۵ باشد. یکی از ایرادات اساسی در معادلات همیلتونی این است که ممکن است سیگنال کنترلی باعث ناپایداری سیستم حلقه‌بسته در زمان‌های خارج از افق تابع هزینه شود.

- [3] Farhadinia B., *Necessary optimality conditions for fuzzy variational problems*, Information Sciences, **181** (2011), 1348-1357.
- [4] Mustafa A. M., Gong Z., Osman M., *The Solution of Fuzzy Variational Problem and Fuzzy Optimal Control Problem under Granular Differentiability Concept*, International Journal of Computer Mathematics, 98(4) (2020), 1-26.



شکل ۴: متغیر حالت و کنترل

۴- نتیجه

در این مقاله، راهبرد کنترل پیش‌بین مدل مقاوم برای حل مسائل کنترل بهینه دارای عدم قطعیت معرفی شد.

کنترل پیش‌بین مدل ویژه مسائل گسسته است. لذا ابتدا با استفاده از رانگ کوتای مرتبه چهارم مدل گسسته‌سازی شد. سپس تابع هزینه به صورت آنالین و با در نظر گرفتن قیدهای فیزیکی حاکم بر سیستم بهینه‌سازی شد تا مقادیر بهینه در زمان جاری محاسبه شوند.

برای جبران اثرات این پارامترهای دارای عدم قطعیت در سیستم حلقه‌بسته مساله بهینه‌سازی رایج در کنترل پیش‌بین مدل به صورت یک بهینه‌سازی دو مرحله‌ای $\min \max$ بیان شد. در این روش ابتدا بدترین تاثیر پارامتر دارای عدم قطعیت بر روی تابع هزینه محاسبه شد، سپس با استفاده از سیگنال کنترلی مینیمم شد.

با توجه به نتایج شبیه‌سازی می‌توان گفت با وجود اینکه دینامیک سیستم دارای عدم قطعیت شدید است و رفتار سیستم بین پایدار و ناپایدار به صورت تصادفی تغییر می‌کند اما کنترل‌کننده پیش‌بین مقاوم به خوبی قادر است متغیرهای حالت سیستم را به صفر همگرا کند. علاوه بر این، برخلاف حل تحلیلی معادلات همبستگی که فقط سیگنال کنترلی بهینه را در بازه تعریف شده برای تابع هزینه محاسبه می‌کند و ممکن است برای خارج از آن باعث واگرایی سیستم شود، کنترل پیش‌بین مقاوم اغلب منجر به کنترل‌کننده‌های پایدار ساز خواهد شد که این امر منجر به استفاده گسترده از این راهبرد کنترلی در سیستم‌های صنعتی شده است.

مراجع

- [1] Seborg, Edgar, Mellichamp, Doyle, *Process Dynamics and Control*, John Wiley & Sons, Fourth Edition, 2017, ISBN: 978-1-119-28591-5 (PBK).
- [2] Farhadinia B., *Pontryagin's Minimum Principle For Fuzzy Optimal Control Problems*, Iranian Journal Of Fuzzy Systems Vol. 11, No. 2, (2014) Pp. 27-43.

حل مسائل کنترل بهینه کسری بازه‌ای با استفاده از تبدیل لاپلاس

طاهره شکوهی^۱، مهدی الله دادی^۲ و سمانه صردی‌زید^۳

^۱ دانشجوی دکتری، دانشکده ریاضی، دانشگاه سیستان و بلوچستان، زاهدان، toktam.shokohi64@gmail.com

^۲ دانشیار، دانشکده ریاضی، دانشگاه سیستان و بلوچستان، زاهدان، m_allahdadi@math.usb.ac.ir

^۳ استادیار، دانشکده صنعت و معدن، دانشگاه سیستان و بلوچستان، زاهدان، soradizeid@eng.usb.ac.ir

چکیده: در این مقاله یک روش محاسباتی جدید برای حل مسائل کنترل بهینه غیرخطی بازه‌ای کسری بر پایه تفاضل هوکوهاری تعمیم یافته و حساب تعییرات کسری ارائه شده است. این روش به منظور تضمین شرایط لازم بهینگی جواب‌ها بیان شده است. برای بدست آوردن جواب‌های این مسائل از خواص عملگرهای مرتبه کسری در حوزه لاپلاس بر اساس توابع میتگ لفلر استفاده شده است. در انتها، روش ارائه شده را برای حل یک مسئله کنترل بهینه بازه‌ای کسری به منظور ارزیابی عملکرد و اثربخشی چارچوب پیشنهادی ارائه شده به کار می‌بریم.

کلمات کلیدی: مساله کنترل بهینه بازه‌ای کسری، حساب تعییرات کسری، تفاضل تعمیم یافته هوکوهارا، شرایط لازم بهینگی، تبدیل لاپلاس، تابع میتگ لفلر.

۱ مقدمه

مسئله کنترل بهینه کسری یک مسئله کنترل بهینه است که سیستم^(۱)
$$\min J^\pm(u^\pm) = h^\pm(x^\pm(t_f), t_f) + \int_{t_0}^{t_f} g^\pm(x^\pm(t), u^\pm(t), t) dt$$
 می‌توان

$${}_0 D_t^\alpha x^\pm(t) = a^\pm(x^\pm, u^\pm, t), \quad t \in [t_0, t_f],$$

$$x^\pm(t_0) = x_0^\pm, \quad x^\pm(t_f) = x_f^\pm,$$

که در آن $h^\pm, g^\pm \in \mathcal{K}$ توابع پیوسته بازه‌ای هستند و $\mathcal{K}$ مجموعه تمام بازه‌های فشرده ناتهی از اعداد حقیقی $\mathbb{R}$ است. $u^\pm(t)$ و $x^\pm(t)$ به ترتیب توابع بازه‌ای از متغیر حالت و متغیر کنترل برای $t \in [t_0, t_f]$ هستند و $h^\pm$ یک تابع بازه‌ای اسکالر است. ${}_0 D_t^\alpha$ مشتق کسری است که در ادامه معرفی خواهد شد. هدف یافتن یک متغیر کنترلی $u^\pm(t)$ است به گونه‌ای که تابع هزینه (۱) به حداقل برسد و محدودیت‌های مسئله برآورده شوند. همچنین، فرض می‌شود:

$$\begin{aligned} f^\pm(x^\pm(t), {}_0 D_t^\alpha x^\pm(t), u^\pm(t), t) \\ = a^\pm(x^\pm(t), u^\pm(t), t) \ominus_{gh} {}_0 D_t^\alpha x^\pm(t) = 0 \end{aligned} \quad (2)$$

دینامیکی آن با معادلات دیفرانسیل کسری توصیف می‌شود. می‌توان مسئله کنترل بهینه کسری را در تعاریف مختلف مشتق کسری مانند مشتق کسری ریمان-لیویل، مشتق کسری کاپوتو^۱ و غیره تعریف کرد. در دهه اخیر، مسائل کنترل بهینه کسری فازی توجه محققان زیادی را به خود جلب کرده است. فرد و سولاکی [۱]، شرایط لازم بهینگی از نوع پونتریاگین را برای یک رده از مسائل کنترل بهینه کسری فازی با مشتق کسری کاپوتو، اثبات کردند. فرد و صالحی در [۲] مسائل تعییرات کسری فازی محدود و نامحدود حاوی مشتقات کسری از نوع کاپوتو را با استفاده از رویکرد تمایزپذیری تعمیم یافته مطالعه کردند. محمدی و همکارانش در [۳] حل مسائل کنترل بهینه کسری فازی را با استفاده از تابع میتگ لفلر معرفی کردند. هدف کلی این مقاله تضمین شرایط لازم بهینگی برای جواب‌های عددی مسائل کنترل بهینه بازه‌ای با مشتقات کسری است. فرم کلی این مسائل به صورت زیر تعریف می‌شود:

^۱Caputo

که در آن تفاضل هوکوهاری تعمیم یافته^۲ است [۲] از رابطه (۲) نتایج زیر به دست می آید:

$$\begin{aligned} f^+(x^\pm(t), t_0 D_t^\alpha x^\pm(t), u^\pm(t), t) &\geq 0, \\ f^-(x^\pm(t), t_0 D_t^\alpha x^\pm(t), u^\pm(t), t) &\leq 0. \end{aligned} \quad (3)$$

می خواهیم روشی عددی برای حل مسئله کنترل بهینه کسری بازه ای به فرم (۱) را بررسی کنیم. برای این منظور، ابتدا شرایط لازم بهینگی را با استفاده از حساب تغییرات کسری استخراج می کنیم. سپس، برای حل این شرایط و یافتن کنترل بهینه بازه ای با استفاده از عملگرهای مرتبه کسری در حوزه لاپلاس و بر اساس تقریب توابع میتگ لفلر، جواب این مسائل را به دست می آوریم. با یافتن این جواب ها و با استفاده از شرط مرزی می توان اطلاعات خوبی در مورد کنترل بهینه بازه ای کسری به دست آورد به طوری که نتایج یک رابطه صریح برای کنترل بهینه بازه ای کسری باشد. به منظور اثبات صحت و کاربرد این روش جدید، یک مثال نیز ارائه شده است.

۲ تعاریف و مقدمات

در اینجا، تعاریف و قضیه های اساسی مورد نیاز در بخش های دیگر را بیان می کنیم. این تعاریف از مراجع [۳، ۵] استخراج شده اند.

تعریف ۱. فرض کنید $f^\pm(t)$ یک تابع بازه ای مقدار پیوسته و انتگرال پذیر لبگ در $[a, b] \subset \mathbb{R}$ باشد آنگاه انتگرال ریمان لیوویل بازه ای تابع $f^\pm$ از مرتبه کسری α ، $0 < \alpha < 1$ ، به فرم زیر تعریف می شوند:

$${}_a I_x^\alpha f^\pm(x) = \frac{1}{\Gamma(\alpha)} \int_a^x f^\pm(x)(x-t)^{\alpha-1} dt.$$

قضیه ۱. فرض کنید $f^\pm(t)$ یک تابع بازه ای مقدار پیوسته و انتگرال پذیر لبگ در $[a, b] \subset \mathbb{R}$ باشد. آنگاه انتگرال ریمان لیوویل تابع بازه ای $f^\pm$ را می توان به فرم زیر بسط داد:

$${}_a I_x^\alpha f^\pm(x) = [{}_a I_x^\alpha f^-(x), {}_a I_x^\alpha f^+(x)].$$

که در آن،

$${}_a I_x^\alpha f^-(x) = \frac{1}{\Gamma(\alpha)} \int_a^x f^-(x)(x-t)^{\alpha-1} dt,$$

$${}_a I_x^\alpha f^+(x) = \frac{1}{\Gamma(\alpha)} \int_a^x f^+(x)(x-t)^{\alpha-1} dt.$$

تعریف ۲. مشتق مرتبه کسری ریمان لیوویل تابع بازه ای مقدار $f^\pm(t)$ به صورت زیر تعریف می شود:

$${}^{RL} D_x^\alpha f^\pm(x) = \frac{d^m}{dt^m} {}_a I_x^{m-\alpha} f^\pm(x)$$

که در آن $m = [\alpha]$ و یا به عبارتی m برابر کوچکترین عدد صحیح بزرگتر مساوی α است. اگر $\alpha \in \mathbb{R}$ عدد مثبت ناصحیحی باشد و $m = [\alpha]$ ،

مشتق کسری کاپوتو برای تابع بازه ای مقدار $f^\pm(t)$ به صورت زیر تعریف می شود:

$${}_a^C D_x^\alpha f^\pm(x) = {}_a I_x^{m-\alpha} \{f^{(m)\pm}(x)\}.$$

در مشتق کسری کاپوتو بر خلاف مشتق ریمان لیوویل، ابتدا مشتق مرتبه صحیح از تابع بازه ای گرفته شده و سپس انتگرال گیری مرتبه کسری انجام می شود.

تعریف ۳. ارتباط بین مشتق مرتبه کسری کاپوتو و مشتق مرتبه کسری ریمان لیوویل تابع بازه ای مقدار $f^\pm(t)$ توسط رابطه زیر بیان می شود:

$${}_a^C D_x^\alpha f^\pm(x) = {}^{RL} D_x^\alpha f^\pm(x) - \sum_{k=0}^{m-1} \frac{f^{(k)\pm}(a)}{\Gamma(k-\alpha+1)} (t-a)^{k-\alpha}$$

قضیه ۲. مشتق کسری کاپوتو یک تابع بازه ای برابر مشتق کسری ریمان لیوویل آن تابع بازه ای است اگر و فقط اگر به ازای $k = 0, 1, \dots, [\alpha-1]$ ، $f^{(k)\pm}(a) = 0$.

۱-۲ مشتق مرتبه کسری کاپوتو در حوزه لاپلاس

در این زیربخش، تبدیل لاپلاس مشتق کسری کاپوتو برای یک تابع بازه ای را بدست می آوریم. برای این منظور ابتدا مشتق کاپوتو را به صورت زیر بازنویسی می کنیم:

$${}_0^C D_x^\alpha f^\pm(x) = {}_0 I_x^{m-\alpha} g^\pm(x)$$

که در آن $m-1 < \alpha < m$ و $f^{(m)\pm}(x) = g^\pm(x)$ است. بنابراین با توجه به تبدیل لاپلاس انتگرال مرتبه کسری $L\{{}_0 I_x^\alpha f^\pm(x)\} = L\{g^\pm(x)\}$ می توان نوشت:

$$L\{{}_0^C D_x^\alpha f^\pm(x)\} = s^{\alpha-m} G(s), \quad (4)$$

که در آن $G^\pm(s) = L\{g^\pm(x)\}$ و به صورت زیر محاسبه می شود:

$$G^\pm(s) = s^m L\{f^\pm(x)\} - \sum_{k=0}^{m-1} s^{m-k-1} f^{(k)\pm}(0) \quad (5)$$

$$L\{{}_0^C D_x^\alpha f^\pm(x)\} = s^\alpha L\{f(x)\} - \sum_{k=0}^{m-1} s^{\alpha-k-1} f^{(k)\pm}(0)$$

تابع میتگ لفلر

تابع میتگ لفلر^۳ توسط یک ریاضیدان سوئدی به نام میتگ لفلر به صورت زیر معرفی شد [۵]:

$$E_{\alpha, \beta} = \sum_{k=0}^{\infty} \frac{z^{k\pm}}{\Gamma(\alpha k + \beta)} \quad \alpha > 0, \beta > 0.$$

این تابع امروزه نقش مهمی را در حل معادلات دیفرانسیل از مرتبه کسری ایفا می کند. چند ویژه گی تابع میتگ لفلر به فرم زیر می باشد:

$$\int_0^z E_{\alpha, \beta}(\lambda t^\alpha) t^{\beta-1} dt = z^\beta E_{\alpha, \beta+1}(\lambda z^\alpha). \quad (6)$$

³Mittag Leffler function

²Generalized Hukuhara difference

۳ بررسی شرایط لازم برای مسائل کنترل بهینه کسری بازه‌ای

در این بخش یک روش کارآمد برای به دست آوردن جواب‌های بهینه یک سیستم کنترلی حلقه بسته بازه‌ای کسری، با استفاده از تفاضل هوکوهارای تعمیم یافته و حساب تغییرات بازه‌ای ارائه می‌شود.

تعریف ۵. تغییرات تابعی بازه‌ای $J^\pm$ را می‌توان به صورت زیر نوشت:

$$\Delta J^\pm(x^\pm, \delta x^\pm) = \delta J^\pm(x^\pm, \delta x^\pm) + \eta^\pm(x^\pm, \delta x^\pm) \cdot \|\delta x^\pm\|. \quad (14)$$

که در آن، $\delta J^\pm$ تابعی خطی از $\delta x^\pm$ است. اگر

$$D(\eta^\pm(x^\pm, \delta x^\pm), 0) < \varepsilon, \text{ as } \|\delta x^\pm\| \rightarrow 0, \quad (15)$$

آنگاه $J^\pm$ را بر حسب $x^\pm$ دیفرانسیل پذیر گوئیم و $\delta J^\pm$ همان تغییرات $J^\pm$ خواهد بود.

تعریف ۶. رابطه ضعیف بین دو تابعی بازه‌ای به فرم زیر تعریف می‌شود [۶]

$$J^\pm \preceq_w J^{*\pm} \Leftrightarrow J^{-*} \leq J^+. \quad (16)$$

تعریف ۷. یک تابع بازه‌ای مقدار دارای مینیمم نسبی در نقطه $x^{*\pm}$ $[x^-, x^+] \in \mathcal{K}$ است اگر تغییرات تابعی بازه‌ای مقدار $J^\pm$ نامنفی باشد، آنگاه:

$$\Delta J^\pm = J^\pm(x^\pm) \ominus_{gH} J^\pm(x^{*\pm}) \succeq_w 0, \quad (17)$$

یا به طور معادل،

$$J^+(x^\pm) \succeq_w J^-(x^{*\pm}), \quad \forall x^\pm \in \mathcal{K}, \quad (18)$$

قضیه ۴. (قضیه اساسی حساب تغییرات بازه‌ای) فرض کنید $x^{*\pm}$ یک مینیمم نسبی برای تابعی بازه‌ای $J^\pm$ باشد آنگاه تغییرات $J^\pm$ روی $x^{*\pm}$ باید صفر شود یعنی

$$\delta J^\pm(x^{*\pm}, \delta x^{*\pm}) = 0, \quad (19)$$

برای همه $\delta x^\pm$ قابل قبول به طوریکه $x^\pm + \delta x^\pm \in \mathcal{K}$.

قضیه ۵. (قضیه شرایط لازم بهینگی برای مسائل کنترل بهینه کسری بازه‌ای) فرض کنید $x^{*\pm}$ حالت بازه‌ای قابل قبول و $u^{*\pm}$ کنترل بازه‌ای قابل قبول برای مسئله کنترل بهینه کسری بازه‌ای (۱) باشد. آنگاه شرایط لازم برای اینکه $u^{*\pm}$ یک کنترل بهینه بازه‌ای برای تابعی بازه‌ای $J^{*\pm}$ باشد این است که برای همه $t \in [t_0, t_f]$ شرایط زیر برقرار باشند:

$${}_{t_0}D_t^\alpha x^-(t) \leq \frac{\partial H^\pm}{\partial \lambda^-}(x^{*\pm}(t), \lambda^{*\pm}(t), u^{*\pm}(t), t),$$

$${}_{t_0}D_t^\alpha x^+(t) \geq \frac{\partial H^\pm}{\partial \lambda^+}(x^{*\pm}(t), \lambda^{*\pm}(t), u^{*\pm}(t), t),$$

$${}_{t_0}D_t^\alpha \lambda'^{-*} \leq -\frac{\partial H^\pm}{\partial x^-}(x^{*\pm}(t), \lambda^{*\pm}(t), u^{*\pm}(t), t),$$

$$\frac{d^m}{dz^m} z^{\beta-1} E_{\alpha,\beta}(z^\alpha) = z^{\beta-m-1} E_{\alpha,\beta-m}(z^\alpha). \quad (7)$$

تعریف ۴. (تبدیل لاپلاس تابع میتگ لفلر) تبدیل لاپلاس تابع $E_{\alpha,\beta}^{(k)}(\pm at^\alpha)$ به صورت زیر محاسبه می‌شود:

$$L\{z^{\alpha k + \beta - 1} E_{\alpha,\beta}^{(k)}(\pm at^\alpha)\} = \frac{k! s^{\alpha - \beta}}{(s^\alpha \mp a)^{k+1}}. \quad (8)$$

از این رابطه می‌توان در یافتن تبدیل لاپلاس معکوس توابع تبدیلی به فرم $\frac{k! s^{\alpha - \beta}}{(s^\alpha \mp a)^{k+1}}$ استفاده کرد.

۲-۲ حل معادلات دیفرانسیل کسری

در این بخش جواب معادله دیفرانسیل کسری با توجه به خواص عمگرهای مرتبه کسری در حوزه لاپلاس بر اساس توابع میتگ لفلر را بدست می‌آوریم.

قضیه ۳. اگر معادله دیفرانسیل کسری با شرایط اولیه به صورت زیر باشد،

$$\begin{aligned} {}^{RL}D_t^\alpha x^\pm(t) &= f^\pm(t, x^\pm(t)), \quad m-1 < \alpha < m, \\ x^\pm(0) &= b, \end{aligned} \quad (9)$$

آنگاه $x^\pm(t)$ جوابی برای معادله (۹) است اگر و فقط اگر در معادله انتگرالی زیر صدق کند:

$$x^\pm(t) = \sum_{k=0}^m \frac{b_k}{\Gamma(\alpha - k + 1)} t^{\alpha - k} + \frac{1}{\Gamma(\alpha)} \int_a^x f^\pm(x)(x-t)^{\alpha-1} dt. \quad (10)$$

در حالت کلی می‌توان نشان داد که پاسخ معادله دیفرانسیل کسری بازه‌ای با شرط اولیه که به صورت زیر باشد،

$$\begin{aligned} {}^C D_t^\alpha x^\pm(t) &= \lambda x^\pm(t) + u^\pm(t), \quad m-1 < \alpha < m, \quad \lambda \in \mathbb{R} \\ x^\pm(0) &= x_0, \end{aligned} \quad (11)$$

که در آن $u^\pm(t)$ تابع بازه‌ای دلخواهی است، را می‌توان با استفاده از تبدیلات لاپلاس محاسبه کرد. برای این منظور، با تبدیل لاپلاس گرفتن از طرفین معادله دیفرانسیل کسری بازه‌ای (۱۱) و ارتباط بین $X^\pm(s)$ و $U^\pm(s)$ خواهیم داشت:

$$Y^\pm(t) = \frac{x_0 \lambda s^{\alpha-1}}{s^\alpha + 1} + \frac{U^\pm(s)}{s^\alpha + 1}, \quad (12)$$

که در آن، $X^\pm(s)$ و $U^\pm(s)$ به ترتیب تبدیل های لاپلاس $x^\pm(s)$ و $u^\pm(s)$ هستند. اکنون، با توجه به تساوی (۱۲) و در نظر گرفتن رابطه (۸) جواب معادله دیفرانسیل کسری بازه‌ای با شرط اولیه (۱۱) به صورت رابطه زیر بدست می‌آید:

$$x^\pm(t) = x_0 E_{\alpha,1}(\lambda t^\alpha) + \int_0^x u^\pm(t-T) T^{\alpha-1} E_{\alpha,\alpha}(\lambda T^\alpha) dT. \quad (13)$$

$$u^{\pm}(t) = -\frac{1}{2} \sum_{k=0}^m \frac{(r^{\pm} t^{\alpha})^k}{\Gamma(k\alpha + 1)}$$

$$x^{\pm} = \sum_{k=0}^m \frac{(-r^{\pm} t^{\alpha})^k}{\Gamma(k\alpha + 1)} - \frac{1}{2} \sum_{k=0}^m \frac{(r^{2\pm} t^{\alpha})^k t^{\alpha(k+1)} \Gamma(\alpha)}{\Gamma(\alpha(k+1) + 1)}$$

جدول ۱: جواب‌های تقریبی $x^{*\pm}$ و $u^{*\pm}$ به ازای $\alpha = 0.5$, $\beta = 0.5$

t	$(x^{\pm})$	$(u^{\pm})$
0.0	[0.7685, 1.3012]	[-0.4357, -0.2242]
0.1	[0.6749, 1.1299]	[-0.4768, -0.2503]
0.2	[0.5868, 0.9723]	[-0.5216, -0.2794]
0.3	[0.5035, 0.8265]	[-0.5708, -0.3119]
0.4	[0.4242, 0.6907]	[-0.6245, -0.3482]
0.5	[0.3484, 0.5633]	[-0.6833, -0.3887]
0.6	[0.2753, 0.4426]	[-0.7477, -0.4339]
0.7	[0.2046, 0.3274]	[-0.8181, -0.4844]
0.8	[0.1354, 0.2161]	[-0.8951, -0.5407]
0.9	[0.0674, 0.1074]	[-0.9794, -0.6036]
1.0	[0.0000, 0.0000]	[-1.0716, -0.6738]

۵ نتیجه گیری

در این مقاله، شرایط لازم برای مسئله کنترل بهینه کسری بازه‌ای با شرایط حالت نهایی ثابت با استفاده از رویکرد تغییرات بازه‌ای استخراج شده است. مسائل ما بر پایه مشتق کسری ریمن-لیوویل بر اساس تفاوت هوکوها را تعریف می‌شود. یک تکنیک عددی بر اساس تبدیل لاپلاس از مشتق کسری پیشنهاد شده است. جواب مسئله کنترل بهینه کسری بازه‌ای بر پایه خواص عملگرهای مرتبه کسری در حوزه لاپلاس بر اساس توابع میتگ لفلر بدست آمده است. در نهایت، یک مثال برای نشان دادن اثربخشی قضیه ۴ و روش عددی ارائه شده است.

۶ مراجع

- [1] O. S. Fard, J. Soolaki, and D. F. Torres, "A necessary condition of pontryagin type for fuzzy fractional optimal control problems," *arXiv preprint arXiv:1702.01093*, 2017.
- [2] O. S. Fard and M. Salehi, "A survey on fuzzy fractional variational problems," *Journal of Computational and Applied Mathematics*, vol.271, pp.71–82, 2014.

$${}^C_{t_0} D_t^{\alpha} \lambda^{*+} \geq -\frac{\partial H^{\pm}}{\partial x^+} (x^{*\pm}(t), \lambda^{*\pm}(t), u^{*\pm}(t), t),$$

$$\frac{\partial H^{\pm}}{\partial u^-} (x^{*\pm}(t), \lambda^{*\pm}(t), u^{*\pm}(t), t) \leq 0,$$

$$\frac{\partial H^{\pm}}{\partial u^+} (x^{*\pm}(t), \lambda^{*\pm}(t), u^{*\pm}(t), t) \geq 0.$$

که در آن $H^{\pm}$ تابع هامیلتونی بازه‌ای است و به صورت زیر تعریف می‌شود:

$$H^{\pm}(x^{\pm}(t), \lambda^{\pm}(t), u^{\pm}(t), t) = g^{\pm}(x^{\pm}(t), u^{\pm}(t), t) + \lambda^{\pm}(t) a^{\pm}(x^{\pm}(t), u^{\pm}(t), t),$$

(۲۰)

و $\lambda^{\pm}$ ضربگر لاگرانژ بازه‌ای است.

۴ نتایج عددی

در این بخش، برای نشان دادن کارایی روش ارائه شده، یک مسئله کنترل بهینه بازه‌ای کسری به فرم کلی (۱) را حل خواهیم کرد.

کنترل بازه‌ای کسری پیدا کنید که در سیستم دینامیک بازه‌ای ${}_0 D_t^{\alpha} x^{\pm} = u^{\pm}(t) - (\beta, 3 - 2\beta)x^{\pm}(t)$ و شرایط اولیه و نهایی $x^{\pm}(0) = 1$ و $x^{\pm}(1) = 0$ صدق کند و تابعی بازه‌ای کسری زیر را به حداقل برساند:

$$J^{*\pm}(x^{\pm}(t), t) = \frac{1}{2} x^{\pm 2}(1) + \int_0^1 u^{\pm 2}(t). \quad (۲۱)$$

می‌خواهیم روش مطرح شده در بخش قبل را برای حل این مسئله بکار ببریم. برای این منظور، ابتدا تابع هامیلتونی بازه‌ای برای این مسئله کسری به شرح زیر تعیین می‌شود:

(۲۲)

$$H^{\pm}(x^{\pm}(t), u^{\pm}(t), J^{*\pm}, t) = \lambda^{\pm}(t)(u^{\pm}(t) - r x^{\pm}(t)) + u^{\pm 2},$$

که در آن $r = [\beta, 3 - 2\beta]$. طبق قضیه ۵، شرایط لازم بهینگی برای این مسئله به صورت زیر به دست می‌آیند:

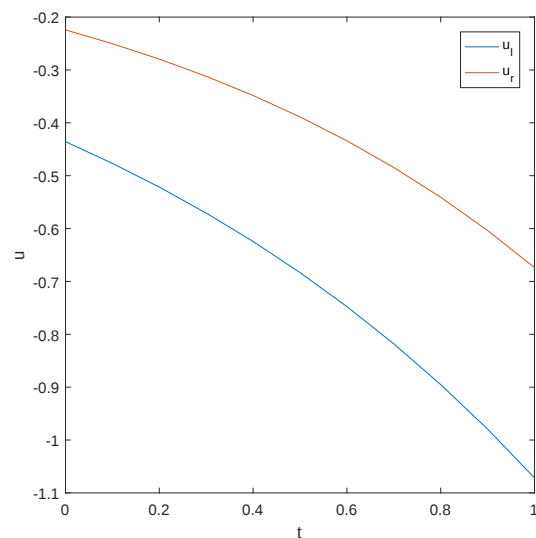
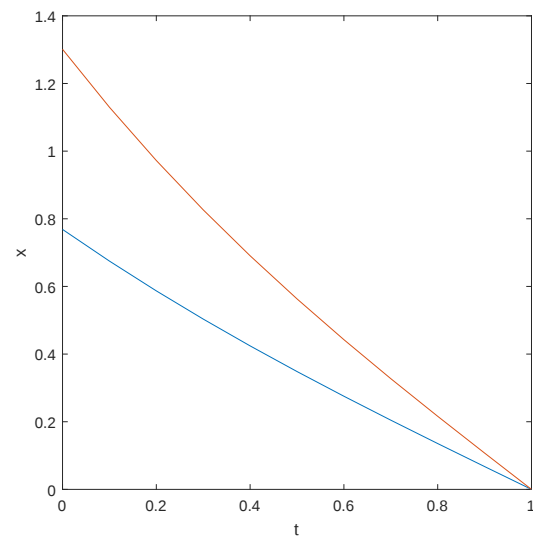
$${}_{t_0} D_t^{\alpha} x^{\pm}(t) = u^{\pm}(t) - r x^{\pm}(t), \quad (۲۳)$$

$${}^C_t D_1^{\alpha} \lambda^{\pm} = r^{\pm} \lambda^{\pm}(t),$$

$$2u^{\pm}(t) + \lambda^{\pm}(t) = 0.$$

جواب معادلات دیفرانسیل بازه‌ای کسری (۲۳) را می‌توان با توجه به خواص عملگرهای مرتبه کسری در حوزه لاپلاس و بر اساس تقریب میتگ لفلر به فرم زیر بدست آورد:

$$\lambda^{\pm}(t) = \sum_{k=0}^m \frac{(r^{\pm} t^{\alpha})^k}{\Gamma(k\alpha + 1)}$$



شکل ۱: نمودارهای تقریبی $x^{*\pm}$ و $u^{*\pm}$ به ازای $\alpha = 0.5$ و $\beta = 0.5$.

- [3] A. Mohammed, Z.-T. Gong, and M. Osman, "A numerical technique for solving fuzzy fractional optimal control problems.," *Journal of Computational Analysis & Applications*, vol.29, no.3, 2021.
- [4] L. Stefanini and B. Bede, "Generalized hukuhara differentiability of interval-valued functions and interval differential equations," *Nonlinear Analysis: Theory, Methods & Applications*, vol.71, no.3-4, pp.1311–1328, 2009.
- [5] K. Diethelm. *The analysis of fractional differential equations: An application-oriented exposition using differential operators of Caputo type*. Springer Science & Business Media, 2010.
- [6] M. Fiedler, J. Nedoma, J. Ramík, J. Rohn, and K. Zimmermann. *Linear optimization problems with inexact data*. Springer Science &

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

شبکه‌های عصبی پیچشی برای تشخیص بیماری‌های تنفسی

محمدحسن خامه‌چیان^۱، محمدرضا اکبرزاده توتونچی^۲

^۱ دانشجوی کارشناسی ارشد، رشته برق گرایش کنترل، قطب علمی رایانش نرم و پردازش هوشمند اطلاعات، دانشگاه فردوسی مشهد، مشهد، ایران،

khamehchian.mhzha@mail.um.ac.ir

^۲ استاد، گروه برق، قطب علمی رایانش نرم و پردازش هوشمند اطلاعات، دانشگاه فردوسی مشهد، مشهد، ایران، akbarzar@um.ac.ir

چکیده: بیماری‌های تنفسی حتی پیش از بروز ویروس کرونا قابل چشم‌پوشی نبودند. سالانه میزان زیادی از تلفات جانی جمعیت جهان ناشی از این نوع بیماری‌ها است. بیماری‌های تنفسی انواع بسیار وسیع و متنوعی دارند. به عنوان مثال انواع بیماری‌های مزمن انسداد مسیر عبور هوا، انواع سرطان‌های دستگاه تنفسی مانند سرطان ریه و حنجره، انواع عفونت‌های دستگاه تنفسی و حتی انواع ویروس کرونا زیر مجموعه‌ی این نوع بیماری هستند. در این پروژه با هدف تشخیص برخی از این بیماری‌ها مدل‌های متفاوتی از شبکه عصبی پیچشی و همچنین برخی از ویژگی‌های دوبعدی استخراج شده از سیگنال صوتی تنفس با یکدیگر مقایسه شده و در نهایت مناسب‌ترین و بهینه‌ترین روش پیشنهاد می‌شود. در سال‌های اخیر به خاطر دقت و مقاومت زیاد شبکه‌های عصبی پیچشی در مقابل نویز، استفاده‌ی گسترده‌ای در انواع کاربردها از جمله پردازش سیگنال و تصویر، از آن شده است. لذا از این روش برای تشخیص برخی از بیماری‌های تنفسی از روی سیگنال صوتی تنفس که توسط استتسکوپ دیجیتال دریافت می‌شود، قابل استفاده است. در این تحقیق از پایگاه داده‌ی ICBHI¹⁷ استفاده شده که بهترین و کامل‌ترین پایگاه داده‌ای از سیگنال صوتی تنفس است که تاکنون در دسترس عموم قرار گرفته است.

کلمات کلیدی: شبکه عصبی پیچشی، ICBHI، سیگنال صوتی تنفس، غربالگری پزشکی، بیماری‌های تنفسی

شبکه عصبی پیچشی و دلیل انتخاب آن برای این هدف و روش انجام این تحقیق به صورت مرحله به مرحله است. در بخش چهارم به نتیجه‌گیری نهایی و ارائه‌ی پیشنهاداتی برای بهبود نتیجه و همچنین نتایج پروژه‌های آینده پرداخته شده است.

۲- پیشینه و مقالات مرتبط

مقالات مختلف در تحلیل صدای تنفس در چهار قسمت با یکدیگر متفاوت هستند: نوع تحلیل پیشین و روش انجام آن، نوع ویژگی استخراج شده و روش استخراج آن، نوع یا روش کلاس‌بندی و در نهایت هدف کلاس‌بندی. هدف کلاس‌بندی در بعضی مقالات فقط وجود یا عدم وجود صداهای خاص تولید شده در هنگام تنفس را کلاس‌بندی می‌کنند درحالی‌که بعضی دیگر از همین صداهای خاص برای تشخیص و کلاس‌بندی بیماری‌های مختلف استفاده می‌کنند. صداهای خاص تولید شده توسط تنفس می‌تواند سکسکه، سرفه، خس خس سینه هنگام تنفس^۱، ترقی ترقی کردن سینه هنگام تنفس^۲ و صدای سوت

۱- مقدمه

هدف اصلی این پروژه ساخت و طراحی یک سیستم هوشمند برای تشخیص سیگنال مورد نظر از صدای دریافتی و تشخیص بیماری براساس آن است. سیستم هوشمند در این پروژه، مدلی از شبکه‌های عصبی پیچشی است که بهترین و بهینه‌ترین نتیجه را به ما بدهد. در این مقاله به مقایسه‌ی تعداد زیادی از مدل‌های شبکه عصبی پیچشی در راستای به دست آوردن بهترین و بهینه‌ترین مدل برای تشخیص و کلاس‌بندی برخی از بیماری‌های تنفسی، پرداخته شده است. این مقاله برای به دست آوردن این هدف از پنج نوع ویژگی استخراج شده از سیگنال صوتی تنفس استفاده کرده است.

این پروژه در ادامه این‌گونه معرفی خواهد شد: قسمت دوم پیشینه و مقالات مرتبط است که برخی از مقالاتی که روی پردازش و کلاس‌بندی سیگنال صوتی تنفس تحقیق کرده‌اند را معرفی می‌کند. قسمت سوم در مورد مواد و روش انجام تحقیق است که شامل بخش‌های معرفی پایگاه داده‌ی مورد استفاده، معرفی

² Crackle

¹ Wheeze

۳- مواد و روش انجام تحقیق:

۳-۱- پایگاه داده:

پایگاه داده‌ای که در این پروژه استفاده شده است پایگاه داده‌ی صدای تنفس کنفرانس بین‌المللی زیست‌پزشکی و خودکار آگاهانه‌ی سلامت که در سال ۲۰۱۷ مطرح شد (ICBHI, 17^{۱۶})، است. این پایگاه توسط دو گروه تحقیقاتی مستقل ساخته شده است که عبارتند از آزمایشگاه توانبخشی و تحقیقاتی تنفسی (Lab3R) که در کشور پرتغال در دانشگاه آویرو در مدرسه‌ی دانش سلامت^{۱۷} قرار دارد و گروه بیمارستان‌های ایماتیا و پاپانیکلو که در کشور یونان در دانشگاه کویمبرا و دانشگاه ارسطو^{۱۸} قرار دارند. پایگاه داده‌ی ICBHI از ۱۲۶ داوطلب، ۹۲۰ فایل صوتی ضبط کرده است. مدت هر یک از این فایل‌های صوتی با یکدیگر متفاوت است و در حدود ۱۰ تا ۹۰ ثانیه زمان می‌برد. این پایگاه داده در کل حدود ۵٫۵ ساعت فایل ضبط شده در اختیار محققان قرار می‌دهد. فرکانس نمونه برداری این فایل‌های صوتی نیز متفاوت است و سه نوع فرکانس ۴۰۰۰ هرتز و ۱۰۰۰۰ هرتز و ۴۴۱۰۰ هرتز را شامل می‌شود. این پایگاه داده به دو صورت برچسب گذاری^{۱۹} شده که در دو جدول زیر نشان داده شده است:

جدول ۱: برچسب گذاری پایگاه داده‌ی ICBHI به صورت انواع بیماری‌های تنفسی

Diseases	Number of patients
COPD	64
Bronchiectasis	7
Bronchiolitis	6
URTI	14
LRTI	2
Pneumonia	6
Asthma	1
Healthy	26
Total	126

جدول ۲: برچسب گذاری پایگاه داده‌ی ICBHI به صورت انواع صداهای موجود در سیگنال صوتی تنفس [16]

Normal or adventitious	Number of cycles
Wheeze	1864
Crackle	886
Wheeze + Crackle	506
Normal	3642
Total	6898

جدول ۱. پایگاه داده را براساس بیماری داوطلب برچسب گذاری کرده است که این برچسب گذاری می‌تواند نتیجه‌ی کلاس‌بندی طبق صداهای

مانندی هنگام تنفس^۲ باشد که البته در زمانی که از استتسکوپ استفاده می‌شود فقط خس خس و ترق ترق سینه هنگام تنفس مورد نظر است.

به عنوان نمونه مقاله‌ی [1] و [2] از ضرایب کیسترال فرکانس در مقیاس مل (MFCC^۳)، مقاله‌ی [3] از تبدیل هیلبرت^۴، مقاله‌ی [4] و [5] از نمودار طیفی^۵، مقاله‌ی [6] از ضرایب موجک و مقاله‌ی [7] از ویژگی‌های مبتنی بر انتروپی^۶ برای قسمت ویژگی استخراج شده از صوت استفاده کرده‌اند تا سه کلاس تنفس عادی، تنفس به همراه خس خس سینه و تنفس با ترق ترق سینه را تشخیص دهند. مقالاتی نیز مثل مقاله‌ی [8] وجود دارند که از تعداد زیادی ویژگی به طور همزمان استفاده کرده‌اند که یک یا چند بیماری را تشخیص و کلاس‌بندی کنند. مثلاً در همین مقاله‌ی [8] از ۱۱۶ ویژگی برای تشخیص و کلاس‌بندی بیماری COPD، بیماری التهاب ریه و فرد سالم استفاده شده است. در مقاله‌ی [9] برای تشخیص و کلاس‌بندی هفت کلاس از صداهای موجود در تنفس، از شبکه‌های عصبی مصنوعی (ANN^۷)، ماشین بردار حامی (SVM^۸) و شبکه‌های عصبی پیچشی (CNN^{۱۰}) استفاده کرده که بهترین نتیجه متعلق به کلاس‌بندی شبکه‌های عصبی پیچشی بوده است. مقالات [10] و [11] از روش کلاس‌بندی ماشین بردار حامی برای تشخیص و کلاس‌بندی صداهای موجود در تنفس که شامل صداهای اکتسابی سینه هنگام تنفس و تنفس سالم می‌شود، استفاده کرده‌اند. در این مقاله از ویژگی‌های ضرایب طیفی شبه فرکانسی و انتروپی طیفی^{۱۱} استفاده شده است. همین طور مقالات [1] و [12] از مدل ترکیبی گوسی (GMM^{۱۲}) و روش k عدد از نزدیک ترین همسایه‌ها (K-NN^{۱۳}) برای کلاس‌بندی صداهای موجود در تنفس استفاده کرده‌اند. در مقاله‌ی [13] برای تشخیص و کلاس‌بندی صداهای موجود در تنفس از شبکه‌های عصبی پیچشی و ماشین بردار حامی استفاده کرده و برتری شبکه‌های عصبی را نشان داده است. در مقاله‌ی [9] نیز برای تشخیص و کلاس‌بندی صداهای موجود در تنفس از ویژگی MFCC و طرح باینری محلی برای کلاس‌بند GMM و SVM و K-NN و CNN استفاده شده است که بهترین نتیجه را روش کلاس‌بندی CNN ارائه داده است. در مقاله‌ی [14] از افزایش داده^{۱۴} به عنوان تحلیل اولیه سیگنال صدای تنفس و از ویژگی نمودار شبه طیفی به عنوان ویژگی صدا و از یک نوع کلاس‌بندی CNN برای تشخیص و کلاس‌بندی سه کلاس قسمت بیماری تنفسی مزمن، غیر مزمن^{۱۵} و فرد سالم استفاده کرده است. همین‌طور برای تحلیل اولیه می‌توان به مقاله‌ی [15] که از شبکه‌های عصبی سازگاری برای حذف نویز در سیگنال تنفس استفاده کرده است، اشاره کرد.

¹³ K-nearest neighbor

¹⁴ Data augmentation

¹⁵ Chronic and non-chronic

¹⁶ International Conference on Biomedical and Health Informatics 2017

¹⁷ Respiratory Research and Rehabilitation Laboratory (Lab3R), School of Health Sciences, University of Aveiro, Aveiro, Portugal

¹⁸ Papanikolaou General Hospital and the General Hospital of I Mathia, Aristotle University of Thessaloniki and the University of Coimbra, Thessaloniki, Greece

¹⁹ labeling

³ Stridor

⁴ Mel-frequency cepstral coefficient

⁵ Hilbert transform

⁶ Spectrogram

⁷ Entropy based features

⁸ Artificial neural network

⁹ Support vector machine

¹⁰ Convolutional neural network

¹¹ Spectral entropy

¹² Gaussian mixture model

$$h_j = f_h \left(\sum_{i=0}^{n-1} w_{ij} x_i - \theta_j \right) \quad (3)$$

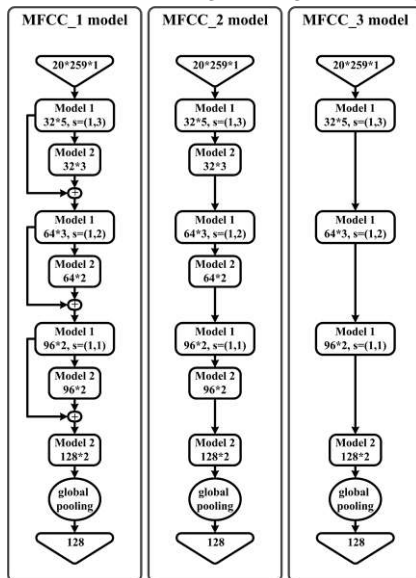
که در آن f_h تابع فعال‌سازی، w_{ij} وزن، x_i بردار یا ماتریس ورودی و θ_j بایاس است.

۳-۲-۲- روش پیشنهادی

این مقاله برای به دست آوردن بهترین و بهینه‌ترین مدل شبکه عصبی پیچشی برای تحلیل سیگنال صوتی تنفس و تشخیص ۸ نوع بیماری تنفسی با استفاده از ۳ نوع ویژگی استخراج شده از سیگنال صوتی، طراحی شده است. در این قسمت به مقایسه‌ی مدل‌های مختلفی از شبکه عصبی پیچشی و همچنین ویژگی‌های مختلف استخراج شده از صوت پرداخته شده و نتایج آن‌ها با یکدیگر مقایسه می‌شود تا بهترین و بهینه‌ترین روش به دست آید. ویژگی‌های قابل استخراج از صوت تنوع بسیار فراوانی دارند ولی متناسب با هر هدفی تنها برخی از آن‌ها قابل استفاده هستند. در این پروژه به دلیل داشتن یک ورودی صوتی نویدار، بی‌کلام و دارای ریزنکات بسیار مهم، می‌توان از برخی ویژگی‌ها مانند ضرایب کپسترال فرکانس در مقیاس مل، شبه نمودار طیفی^{۲۲}، کروما^{۲۳}، شار طیفی^{۲۴}، مرکز ثقل طیفی^{۲۵}، پخش طیفی^{۲۶}، درجه اوج^{۲۷} و بعضی از ویژگی‌های مبتنی بر انترپوی^{۲۸} برای تشخیص بیماری استفاده کرد. در این مقاله از سه نوع ویژگی مشهور و مهم ضرایب کپسترال فرکانس در مقیاس مل، شبه نمودار طیفی و کروما برای رسیدن به هدف مورد نظر استفاده شده است. تعداد لایه‌ها و مدل‌های شبکه عصبی پیچشی استفاده شده در این مقاله براساس اندازه ورودی یا به عبارتی اندازه خروجی گرفته شده از هر ویژگی صوتی انتخاب شده‌اند که به صورت زیر هستند:

A. مدل‌هایی با ابعاد ورودی ۱*۲۵۹*۲۰ که برای ویژگی ضرایب

کپسترال فرکانس در مقیاس مل استفاده شده است:



موجود در تنفس یا از روی خود سیگنال تنفس باشد. جدول ۲ پایگاه داده را براساس صداها موجود در سیگنال تنفس برچسب گذاری کرده است. چرخه^{۲۰} در این برچسب گذاری به معنی یک چرخه‌ی دم و بازدم یعنی یک نفس کامل است که می‌تواند ۰/۲ تا ۱۶/۲ ثانیه زمان ببرد.

۳-۲-۲- روش انجام تحقیق:

۳-۲-۱- شبکه‌های عصبی پیچشی (CNN):

همانطور که دیده می‌شود تقریباً در تمامی مقالاتی که مقایسه‌ای بین روش CNN و روش دیگری انجام شده است، برتری این روش را نشان می‌دهد. این برتری به خاطر محلی بودن و توجه بسیار این روش به جزئیات و مقاومت بسیار زیاد در برابر نویز است. البته در اکثر مواقع برای بهبود نتیجه در مسائل سری زمانی مثل مسائل مربوط به صدا این روش با روش شبکه‌های عصبی برگشتی (RNN) ترکیب می‌شود. شبکه‌های عصبی برگشتی بسیار در مسائل سری زمانی کاربرد دارد ولی در مقایسه با CNN از دقت کمتری برخوردار است. برای مرور CNN این گونه می‌توان توضیح داد که این الگوریتم همانند الگوریتم کلی روش DNN است؛ در هر لایه، وزنی در ورودی ضرب شده و نتایج باهم جمع شده و سپس خروجی به دست می‌آید ولی با این تفاوت که در الگوریتم CNN ورودی، خروجی و وزن یک عدد نیستند بلکه به صورت ماتریسی هستند و ضرب هم به صورت عددی نیست بلکه به صورت پیچشی است.

در ضمن لازم به ذکر است که سه نوع لایه در این الگوریتم وجود دارد: لایه پیچش که مسئولیت به دست آوردن ویژگی‌ها از ورودی خود را دارد و نحوه به دست آوردن آن برای هر نود طبق مقاله [17] به صورت معادله (۱) به دست می‌آید.

$$x_j^l = f_c \left(\sum_{i \in M_j} x_i^{l-1} * k_{ij} + \theta_j^l \right) \quad (1)$$

که در آن f_c تابع فعال‌سازی، x_i ورودی ماتریسی، k_{ij} فیلتر ماتریسی پیچش و θ_j بایاس است.

لایه pooling مسئولیت کاهش ابعاد به نحوی که اطلاعات مهم باقی بماند را دارد و نحوه به دست آوردن آن برای هر نود طبق مقاله [17] به صورت معادله (۲) به دست می‌آید.

$$x_j^l = f_p \left(\beta_j^l \text{down}(x_i^{l-1}) + \theta_j^l \right) \quad (2)$$

که در آن f_p تابع فعال‌سازی، $\text{down}(\cdot)$ روش نمونه برداری پایینی که معمولاً ماکزیمم یا میانگین است، β_j و θ_j به ترتیب بایاس ضرب‌شونده و بایاس جمع‌شونده است.

در نهایت لایه کاملاً متصل (FC^{21}) عمل طبقه‌بندی را انجام می‌دهد. نحوه محاسبه خروجی آن برای هر نود طبق مقاله [17] به شکل معادله (۳) است.

²⁵ Spectral centroid

²⁶ Spectral spread

²⁷ Kurtosis

²⁸ Entropy basis features

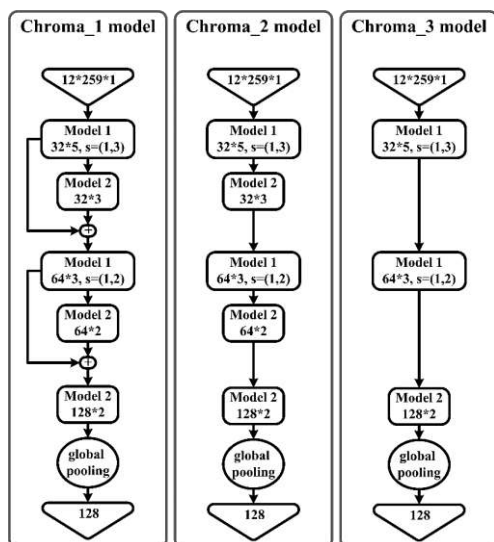
²⁰ cycle

²¹ Fully connected

²² Mel-spectrogram

²³ Chroma features

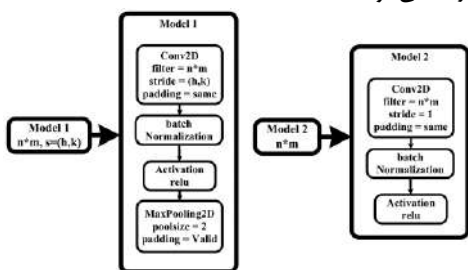
²⁴ Spectral flux



شکل ۳: مدل‌های شبکه عصبی پیچشی برای ویژگی کروما

مدل Chroma_1: مدل شبکه عصبی پیچشی رزیجوآل یا باقیمانده، مدل Chroma_2: همان مدل Chroma_1 است ولی بدون حلقه رزیجوآل یا باقیمانده، مدل Chroma_3: مدل شبکه عصبی پیچشی رزیجوآل یا باقیمانده سریع با کمترین تعداد لایه

لازم به ذکر است که ویژگی کروما دارای انواع مختلفی است که در این مقاله سه نوع cqt^{29} ، $stft^{30}$ و $cens^{31}$ به کار رفته است. هر سه مدل نشان داده شده در قسمت C برای هر سه نوع ویژگی کروما به کار رفته است زیرا خروجی این سه ویژگی هم اندازه هستند. در هر سه قسمت A، B و C بلوک‌های Model 1 و Model 2 به صورت شکل ۴ تعریف می‌شوند.



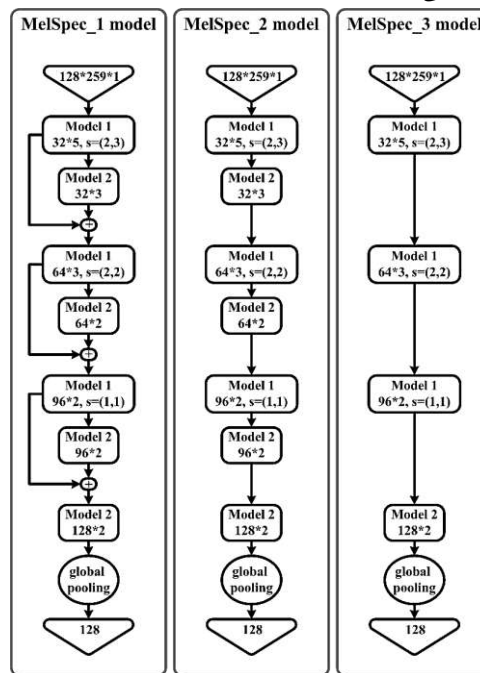
شکل ۴: مفهوم بلوک‌های مدل ۱ و مدل ۲

لازم به ذکر است در هر سه دسته A، B و C مدل‌های اول (ستون سمت چپ) یعنی مدل‌های MFCC_1، MelSpec_1 و Chroma_1 به صورت شبکه عصبی پیچشی رزیجوآل یا باقیمانده^{۳۲} ساخته شده‌اند. همچنین مدل‌های دوم (ستون‌های وسط) یعنی مدل‌های MFCC_2، MelSpec_2 و Chroma_2 از همان تعداد لایه‌های مدل اول تشکیل شده‌اند ولی قسمت رزیجوآل یا باقیمانده را دارا نیستند. در نهایت مدل‌های سوم (ستون‌های سمت راست) یعنی مدل‌های MFCC_3، MelSpec_3 و Chroma_3 از کمترین و بهینه‌ترین تعداد لایه در ساخت شبکه عصبی پیچشی استفاده کرده‌اند که باعث می‌شود این مدل‌ها سریع‌ترین و بهینه‌ترین مدل در این پروژه باشند.

شکل ۱: مدل‌های شبکه عصبی پیچشی برای ویژگی ضرایب کپسترال فرکانس در مقیاس مل

مدل MFCC_1: مدل شبکه عصبی پیچشی رزیجوآل یا باقیمانده، مدل MFCC_2: همان مدل MFCC_1 است ولی بدون حلقه رزیجوآل یا باقیمانده، مدل MFCC_3: مدل شبکه عصبی پیچشی رزیجوآل یا باقیمانده سریع با کمترین تعداد لایه

B. مدل‌هایی با ابعاد ورودی $128 \times 259 \times 1$ که برای ویژگی نمودار طیفی استفاده شده است:



شکل ۲: مدل‌های شبکه عصبی پیچشی برای ویژگی نمودار طیفی

مدل MelSpec_1: مدل شبکه عصبی پیچشی رزیجوآل یا باقیمانده، مدل MelSpec_2: همان مدل MelSpec_1 است ولی بدون حلقه رزیجوآل یا باقیمانده، مدل MelSpec_3: مدل شبکه عصبی پیچشی رزیجوآل یا باقیمانده سریع با کمترین تعداد لایه

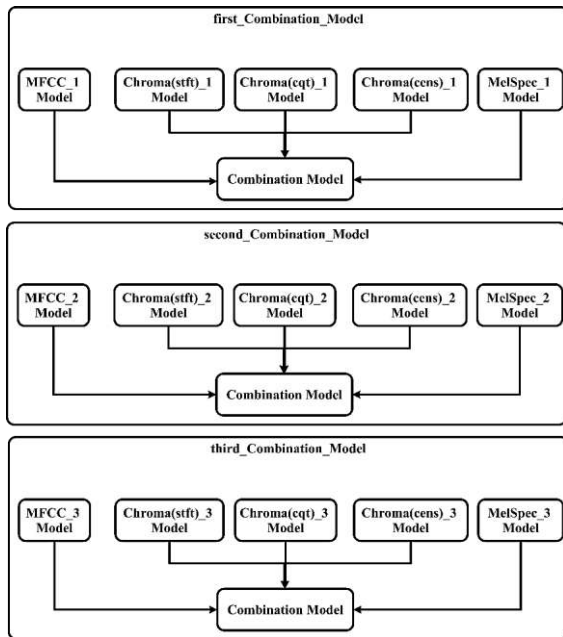
C. مدل‌هایی با ابعاد ورودی $128 \times 259 \times 1$ که برای ویژگی‌های کروما استفاده شده است:

³¹ Chroma energy normalized statics

³² Residual

²⁹ Constant-Q transform

³⁰ Short-time Fourier transform



شکل ۷: مدل‌های شبکه عصبی پیچشی ترکیبی نهایی

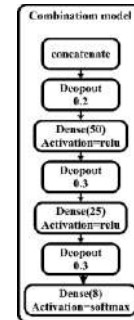
در شکل ۷ مدل ترکیبی اول یک مدل رزجوال یا باقیمانده است چون متشکل از سه مدل شبکه عصبی پیچشی رزجوال یا باقیمانده که در ستون سمت چپ در قسمت‌های A، B و C نشان داده شده برای پنج نوع ویژگی ذکر شده در قسمت‌های قبل، است. همین طور مدل ترکیبی دوم از ستون وسط در قسمت‌های A، B و C تشکیل شده است. ترکیب سوم یک مدل سریع است چون از ستون سمت راست در قسمت‌های A، B و C تشکیل شده است. لازم به ذکر است که این ترکیب سوم در مدل ذکر شده در [18] در سال ۲۰۲۰ میلادی با تعداد ورودی یا به عبارتی تعداد ویژگی‌های استخراج شده از صوت کمتری طراحی و استفاده شده است.

نتایج نهایی این مدل‌ها در جدول ۳ نشان داده شده است.

جدول ۳: نتایج مدل‌های شبکه عصبی پیچشی تکی و ترکیبی نهایی

Model name	Model detail	loss	accuracy
first_Combination_Model	{MFCC_1, Chroma(stft)_1, Chroma(cqt)_1, Chroma(cens)_1, MelSpec_1} → {Combination}	0.2565	0.902
second_Combination_Model	{MFCC_2, Chroma(stft)_2, Chroma(cqt)_2, Chroma(cens)_2, MelSpec_2} → {Combination}	0.2552	0.9043
third_Combination_Model	{MFCC_3, Chroma(stft)_3, Chroma(cqt)_3, Chroma(cens)_3, MelSpec_3} → {Combination}	0.283	0.8916
1st_single_Model	{MFCC_1} → {combination}	0.1929	0.9351
2nd_single_Model	{MFCC_2} → {combination}	0.1862	0.9316
3rd_single_Model	{MFCC_3} → {combination}	0.2	0.9443
4th_single_Model	{MelSpec_1} → {combination}	0.4096	0.8545

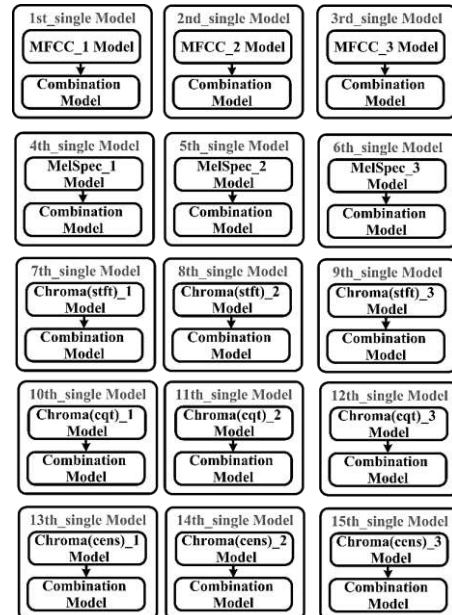
ولی باید دقت داشت که سرعت همیشه در کنار دقت بالا مفید واقع می‌شود بنابراین ممکن است این مدل‌ها بهترین نتیجه را به ما ندهند. D. مدلی که برای ترکیب نتایج مدل‌های قبلی و به دست آوردن نتیجه نهایی استفاده شده است در شکل ۵ نشان داده می‌شود.



شکل ۵: مدل شبکه عصبی پیچشی برای کلاس‌بندی و گرفتن نتیجه نهایی و در صورت نیاز ترکیب نتایج مدل‌های قبلی

این مدل برای به دست آوردن نتیجه نهایی و کلاس‌بندی و همچنین در صورت نیاز، ترکیب نتایج مدل‌های قبلی مورد استفاده قرار می‌گیرد. لازم به ذکر است که در این پروژه از بهینه‌ساز adam و کلاس‌بند softmax استفاده شده است. همان طور که مشاهده می‌شود از ترکیب این مدل‌ها می‌توان مدل‌های کاملی برای رسیدن به هدف اصلی به دست آورد. ترکیب‌هایی که در این مقاله مورد آزمایش قرار گرفته است به صورت زیر هستند:

❖ ترکیب هریک از مدل‌های قسمت A، B و C به تنهایی یا به صورت تکی با مدل قسمت D که می‌توان به صورت شکل ۶ آن‌ها را نمایش داد.



شکل ۶: مدل‌های شبکه عصبی پیچشی تکی نهایی برای هر ویژگی

همانطور که در شکل ۶ مشاهده می‌شود ۱۵ نوع ترکیب می‌توان ساخت، چون ۵ نوع ویژگی به وسیله ۳ نوع مدل سیستم شبکه عصبی پیچشی مورد آزمایش قرار گرفته است.

❖ گرفتن جواب‌های هر نوع مدل و سپس ترکیب جواب‌ها به وسیله مدل combination که به سه صورتی که در شکل ۷ نشان داده شده است به دست می‌آید:

متفاوت مثل nadam و SGD و روش‌های کلاس‌بندی متفاوت مثل LSTM و منطق فازی یا ترکیب آن‌ها با شبکه عصبی پیچشی استفاده کرد.

مراجع

- [1] M. Bahoura, "Pattern recognition methods applied to respiratory sounds classification into normal and wheeze classes," *Comput. Biol. Med.*, vol. 39, no. 9, pp. 824–843, 2009, doi: 10.1016/j.compbimed.2009.06.011.
- [2] B.-S. Lin, H.-D. Wu, and S.-J. Chen, "Automatic wheezing detection based on signal processing of spectrogram and back-propagation neural network," *J. Healthc. Eng.*, vol. 6, no. 4, pp. 649–672, 2015, doi: 10.1260/2040-2295.6.4.649.
- [3] G. Serbes, C. O. Sakar, Y. P. Kahya, and N. Aydin, "Pulmonary crackle detection using time-frequency and time-scale analysis," *Digit. Signal Process. A Rev. J.*, vol. 23, no. 3, pp. 1012–1021, 2013, doi: 10.1016/j.dsp.2012.12.009.
- [4] N. K. Gautam and S. B. Pokle, "Wavelet scalogram analysis of phonopulmonographic signals," *Int. J. Med. Eng. Inform.*, vol. 5, no. 3, pp. 245–252, 2013, doi: 10.1504/IJMEI.2013.055700.
- [5] J. Acharya, A. Basu, and W. Ser, "Feature extraction techniques for low-power ambulatory wheeze detection wearables," in *Proceedings of the Annual International Conference of the IEEE Engineering in Medicine and Biology Society, EMBS*, 2017, pp. 4574–4577, doi: 10.1109/EMBC.2017.8037874.
- [6] A. D. Orjuela-Cañón, D. F. Gómez-Cajas, and R. Jiménez-Moreno, *Artificial neural networks for acoustic lung signals classification*, vol. 8827, 2014.
- [7] J. Zhang, W. Ser, J. Yu, and T. T. Zhang, "A novel wheeze detection method for wearable monitoring systems," in *2009 International Symposium on Intelligent Ubiquitous Computing and Education, IUCE 2009*, 2009, pp. 331–334, doi: 10.1109/IUCE.2009.66.
- [8] S. Z. H. Naqvi and M. A. Choudhry, "An automated system for classification of chronic obstructive pulmonary disease and pneumonia patients using lung sound analysis," *Sensors (Switzerland)*, vol. 20, no. 22, pp. 1–23, 2020, doi: 10.3390/s20226512.
- [9] D. Bardou, K. Zhang, and S. M. Ahmad, "Lung sounds classification using convolutional neural networks," *Artif. Intell. Med.*, vol. 88, pp. 58–69, 2018, doi: 10.1016/j.artmed.2018.04.008.
- [10] D. Chamberlain, R. Kodgule, D. Ganelin, V. Miglani, and R. R. Fletcher, "Application of semi-supervised deep learning to lung sound analysis," in *Proceedings of the Annual International Conference of the IEEE Engineering in Medicine and Biology Society, EMBS*, 2016, vol. 2016-Octob, pp. 804–807, doi: 10.1109/EMBC.2016.7590823.
- [11] I. Mazić, M. Bonković, and B. Džaja, "Two-level coarse-to-fine classification algorithm for asthma wheezing recognition in children's respiratory sounds," *Biomed. Signal Process. Control*, vol. 21, pp. 105–118, 2015, doi: 10.1016/j.bspc.2015.05.002.
- [12] P. Mayorga, C. Druzgalski, R. L. Morelos, O. H. González, and J. Vidales, "Acoustics based assessment of respiratory diseases using GMM classification," in *2010 Annual International Conference of the IEEE Engineering in Medicine and Biology Society, EMBC'10*, 2010, pp. 6312–6316, doi: 10.1109/IEMBS.2010.5628092.

5th_single_Model	{MelSpec_2} → {combination}	0.4253	0.8539
6th_single_Model	{MelSpec_3} → {combination}	0.4441	0.8533
7th_single_Model	{Chroma(stft)_1} → {combination}	0.5831	0.8342
8th_single_Model	{Chroma(stft)_2} → {combination}	0.5503	0.833
9th_single_Model	{Chroma(stft)_3} → {combination}	0.5769	0.833
10th_single_Model	{Chroma(cqt)_1} → {combination}	0.671	0.8377
11th_single_Model	{Chroma(cqt)_2} → {combination}	0.6245	0.8441
12th_single_Model	{Chroma(cqt)_3} → {combination}	0.6782	0.8417
13th_single_Model	{Chroma(cens)_1} → {combination}	0.6528	0.8336
14th_single_Model	{Chroma(cens)_2} → {combination}	0.6097	0.8342
15th_single_Model	{Chroma(cens)_3} → {combination}	0.6339	0.8359

همان طور که مشاهده می‌شود بهترین نتیجه را مدل‌های ترکیبی MFCC به دست آورده‌اند و از بین آن‌ها مدل سریع و بهینه یعنی مدل MFCC_1 بهترین نتیجه را هم از لحاظ دقت و هم از لحاظ خطا به ما می‌دهد. پس از آن بهترین نتیجه را مدل‌های ترکیبی به دست آورده است.

طبق این خروجی‌ها می‌توان به این نتیجه رسید که مدل ترکیبی MFCC به تنهایی می‌تواند هم از نظر دقت و هم از نظر تعداد پارامتر و سرعت بهترین روش برای این پروژه باشد. همچنین در تمامی نتایج به دست آمده از مدل‌ها مشاهده می‌شود که اضافه کردن ویژگی ریزجوال یا باقیمانده به شبکه عصبی پیچشی کمک چندانی به پروژه نکرده و نتیجه را فقط در حد چند صدم بهبود می‌بخشد ولی محاسبات پروژه را پیچیده‌تر و زمان گرفتن نتیجه را زیادتیر می‌کند. بنابراین برای بهبود این پروژه استفاده از ویژگی ریزجوال یا باقیمانده در شبکه عصبی پیچشی پیشنهاد نمی‌شود.

۴- نتیجه

در این مقاله به مقایسه تعداد قابل توجهی از مدل‌های مهم شبکه‌های عصبی پیچشی در راستای به دست آوردن مناسب‌ترین و بهینه‌ترین روش برای تشخیص و کلاس‌بندی هشت نوع بیماری تنفسی از روی سیگنال صوتی تنفس بیمار پرداخته شد. نتایج این مقایسات نشان داد که استفاده از ویژگی ریزجوال یا باقیمانده در شبکه عصبی پیچشی برای این هدف چندان نتیجه‌نهایی را بهبود نمی‌دهد و تنها محاسبات را سنگین‌تر می‌کند در نتیجه بهتر است که از این ویژگی در پروژه‌های آتی برای این هدف استفاده نشود. همچنین به این نتیجه دست پیدا کردیم که بهترین نتیجه زمانی به ما داده می‌شود که فقط از ویژگی MFCC و مدل ساده‌ای از شبکه‌های عصبی برای تشخیص بیماری استفاده کنیم. پس از آن نیز ترکیب این ویژگی با ویژگی‌های دیگر به وسیله شبکه‌های عصبی پیچشی بهترین نتیجه را به ما می‌دهد. البته این ترکیب ویژگی‌ها باعث افزایش پیچیدگی و مقدار محاسبات می‌شود.

در پروژه‌های بعدی می‌توان به مقایسه نتایج ویژگی‌های استخراج شده‌ی دیگری از صدا مانند ویژگی‌های مبتنی بر انتروپی و ویژگی درجه اوج و غیره و ترکیب آن‌ها با MFCC استفاده کرد. همچنین می‌توان از بهینه‌سازهای

- [13] M. Aykanat, Ö. Kılıç, B. Kurt, and S. Saryal, "Classification of lung sounds using convolutional neural networks," *Eurasip J. Image Video Process.*, vol. 2017, no. 1, 2017, doi: 10.1186/s13640-017-0213-2.
- [14] M. T. García-Ordás, J. A. Benítez-Andrades, I. García-Rodríguez, C. Benavides, and H. Alaiz-Moretón, "Detecting respiratory pathologies using convolutional neural networks and variational autoencoders for unbalancing data," *Sensors (Switzerland)*, vol. 20, no. 4, 2020, doi: 10.3390/s20041214.
- [15] A. Abbas and A. Fahim, "An automated computerized auscultation and diagnostic system for pulmonary diseases," *J. Med. Syst.*, vol. 34, no. 6, pp. 1149–1155, 2010, doi: 10.1007/s10916-009-9334-1.
- [16] L. D. L. D. Pham, H. Phan, R. Palaniappan, A. Mertins, and I. Mccloughlin, "CNN-MoE based framework for classification of respiratory anomalies and lung disease detection," *IEEE J. Biomed. Heal. Informatics*, 2021, doi: 10.1109/JBHI.2021.3064237.
- [17] P. Swietojanski, A. Ghoshal, and S. Renals, "Convolutional neural networks for distant speech recognition," *IEEE Signal Process. Lett.*, vol. 21, no. 9, pp. 1120–1124, 2014, doi: 10.1109/LSP.2014.2325781.
- [18] S. Sharma, "(Part-3) Feature Extraction & Modeling (95% acc)," 2020. <https://www.kaggle.com/shivam316/part-3-feature-extraction-modeling-95-acc/notebook>.

قابلیت اعتماد بر اساس α -شک اعداد فازی

مهدیه مظفری^۱، محمد خنجری^۲ و محمد قاسم اکبری^۲

^۱ گروه آمار، دانشکده ریاضی و محاسبات، مجتمع آموزش عالی بزم، شهر بزم، mozafari@bam.ac.ir

^۲ گروه آمار، دانشکده علوم ریاضی و آمار، دانشگاه بیرجند، شهر بیرجند، mkhanjari@birjand.ac.ir، g-z-akbari@birjand.ac.ir

چکیده: در این مقاله، بر پایه α -شک اعداد فازی، یک متغیر تصادفی فازی مقیاس تعریف شده است. بر اساس این متغیر تصادفی، برخی از معیارهای قابلیت اعتماد نظیر تابع بقا و نرخ خطر مورد بررسی قرار گرفته و همچنین مثال‌هایی در این مورد، ارائه گردیده است.

کلمات کلیدی: تابع بقا، عدد فازی، قابلیت اعتماد، متغیر تصادفی فازی مقیاس، نرخ خطر.

۱ مقدمه

قابلیت اعتماد نظیر تابع بقا و تابع نرخ شکست ارائه شده است. همچنین مثال‌هایی عددی نیز برای نشان دادن محاسبه عملگرها بیان گردید.

۲ اعداد فازی

مجموعه مرجع $\mathbb{X}$ را در نظر بگیرید. مجموعه فازی $\tilde{A}$ از $\mathbb{X}$ به صورت $\{(x, \tilde{A}(x)); x \in \mathbb{X}\}$ نمایش داده می‌شود.

نگاشت $\tilde{A}: \mathbb{X} \rightarrow [0, 1]$ که به هر $x \in \mathbb{X}$ یک مقدار از بازه $[0, 1]$ نسبت می‌دهد، تابع عضویت $\tilde{A}$ گویند. برای هر $\alpha \in (0, 1]$ زیرمجموعه $\{\tilde{A}(x) \geq \alpha\}$ ، α -برش $\tilde{A}$ نامیده شده و با $\tilde{A}[\alpha]$ نشان داده می‌شود.

تعریف ۱. مجموعه فازی $\tilde{A}$ روی $\mathbb{R}$ ، یک عدد فازی می‌باشد، هرگاه (1) به ازای هر $\alpha \in [0, 1]$ ، $\tilde{A}[\alpha] = [\tilde{A}^L[\alpha], \tilde{A}^U[\alpha]]$ باشد، که در آن

$$\tilde{A}^L[\alpha] = \inf\{x \in \mathbb{R} \mid \tilde{A}(x) \geq \alpha\},$$

$$\tilde{A}^U[\alpha] = \sup\{x \in \mathbb{R} \mid \tilde{A}(x) \geq \alpha\}.$$

(2) عدد حقیقی یکتای $x^* \in \mathbb{R}$ وجود داشته باشد، بطوریکه $\tilde{A}(x^*) = 1$.

تعریف ۲. فرض کنید $\tilde{A}$ یک عدد فازی روی مجموعه مرجع $\mathbb{X}$ باشد، آن را عدد فازی LR گویند، اگر تابع عضویت آن به صورت زیر باشد:

استنباط آماری کلاسیک به مشاهدات دقیق متکی است. با این وجود، در بسیاری از رویدادهای واقعی زندگی، مشاهدات قابل استفاده نادقیق می‌باشند. پس از معرفی نظریه مجموعه‌های فازی توسط زاده [۸]، چنین عدم قطعیت‌هایی با استفاده از مجموعه‌های فازی به روشی کارآمدتر از در نظر گرفتن مقادیر دقیق مطلق مدل‌بندی شد. شاپیرو [۶] و گیل، لویز دیاز و والسکو [۲] تفسیرهای مختلفی از تعاریف مختلف متغیر تصادفی فازی را بررسی کردند. در این راستا، تلاش زیادی برای تعمیم مشخصه‌های احتمالی یک متغیر تصادفی فازی مانند: پارامترهای توزیع، توزیع تجمعی، میانگین، واریانس و کوواریانس انجام شده است. افراد دیگری، از جمله حسامیان و طاهری [۴] این مفاهیم آماری را به عنوان مجموعه‌های فازی گسترش دادند.

باید توجه داشت که، بسیاری از توزیع‌های احتمالی را می‌توان با پارامترهای مکان و مقیاس مشخص کرد. چنین پارامترهای مکان و مقیاس، اغلب در استنباط‌های آماری (لهمن و کسلا [۵]) و تئوری قابلیت اعتماد (ویلیام و اسکوبار [۷]) استفاده می‌شوند. از آنجا که، بسیاری از استنباط‌های آماری شامل عدم قطعیت می‌باشند، بنابراین در این مقاله، برخی استنباط‌های آماری یک متغیر تصادفی فازی از خانواده مقیاس نظیر توزیع تجمعی، امید ریاضی، واریانس و احتمال غیردقیق یک بازه مورد بررسی قرار گرفته است. بعلاوه، برخی از معیارهای

ریاضی اعداد فازی به صورت زیر بیان می‌شوند:

$$\begin{aligned}(\tilde{A} \oplus \tilde{B})_\alpha &= \tilde{A}_\alpha + \tilde{B}_\alpha, \\(\tilde{A} \ominus \tilde{B})_\alpha &= \tilde{A}_\alpha - \tilde{B}_{1-\alpha}, \\(\tilde{A} \oslash \tilde{B})_\alpha &= \tilde{A}_\alpha / \tilde{B}_{1-\alpha} \quad \tilde{B}_0 > 0, \\(\lambda \otimes \tilde{A})_\alpha &= \begin{cases} \lambda \tilde{A}_\alpha & \lambda > 0 \\ 0 & \lambda = 0 \\ \lambda \tilde{A}_{1-\alpha} & \lambda < 0. \end{cases}\end{aligned}$$

۳ متغیرهای تصادفی فازی

فرض کنید (Ω, A, P) یک فضای احتمال باشد.

تعریف ۳. نگاشت فازی مقدار $\tilde{X} : \Omega \rightarrow F(\mathbb{R})$ یک متغیر تصادفی فازی نامیده می‌شود، هرگاه برای هر $\alpha \in [0, 1]$ ، نگاشت حقیقی مقدار $\tilde{X}_\alpha : \Omega \rightarrow \mathbb{R}$ یک متغیر تصادفی حقیقی مقدار باشد. بعلاوه، دو متغیر تصادفی فازی $\tilde{X}$ و $\tilde{Y}$ مستقل و هم‌توزیع می‌باشند، هرگاه $\tilde{X}_\alpha$ و $\tilde{Y}_\alpha$ به ازای هر $\alpha \in [0, 1]$ مستقل و هم‌توزیع باشند. همچنین $\tilde{X}_1, \dots, \tilde{X}_n$ یک نمونه تصادفی فازی می‌باشند، اگر $\tilde{X}_i$ ها متغیرهای تصادفی فازی مستقل و هم‌توزیع باشند. مقادیر مشاهده شده یک نمونه تصادفی فازی را با $\tilde{x}_1, \dots, \tilde{x}_n$ نشان می‌دهند.

تعریف ۴. عدد فازی $\tilde{F}_{\tilde{X}}(x)$ تابع توزیع تجمعی فازی $\tilde{X}$ می‌باشد، هرگاه α -شک آن به صورت زیر بیان شود:

$$(\tilde{F}_{\tilde{X}}(x))_\alpha = \quad (۲)$$

$$P(\tilde{X}_{1-\alpha} \leq x) = \begin{cases} \inf_{\beta \in I_{2\alpha}} F_{\tilde{X}_\beta}(x) & 0.0 \leq \alpha \leq 0.5 \\ \sup_{\beta \in I_{2(1-\alpha)}} F_{\tilde{X}_\beta}(x) & 0.5 < \alpha \leq 1.0, \end{cases}$$

که $I_\alpha = [\alpha/2, 1 - \alpha/2]$

ملاحظه ۴. عدد فازی $\tilde{F}_{\tilde{X}}(x)$ را در نقطه x در نظر بگیرید. α -شک آن به صورت $(\tilde{F}_{\tilde{X}}(x))_\alpha = \tilde{F}_{\tilde{X}_\alpha}(x)$ می‌باشد.

تعریف ۵. متغیر تصادفی فازی $\tilde{X}$ را در نظر بگیرید. به ازای $(a, b) \subset \mathbb{R}$ داریم:

$$(\tilde{P}(\tilde{X} \in (a, b)))_\alpha = \begin{cases} \inf_{\beta \in I_{2\alpha}} P(a < \tilde{X}_\beta < b) & 0.0 \leq \alpha \leq 0.5 \\ \sup_{\beta \in I_{2(1-\alpha)}} P(a < \tilde{X}_\beta < b) & 0.5 < \alpha \leq 1.0. \end{cases}$$

ملاحظه ۵. اگر $(\tilde{F}_{\tilde{X}}(a))_1 \leq (\tilde{F}_{\tilde{X}}(b))_0$ ، به آسانی می‌توان نشان داد:

$$(\tilde{P}(\tilde{X} \in (a, b)))_\alpha = (\tilde{F}_{\tilde{X}}(b))_\alpha - (\tilde{F}_{\tilde{X}}(a))_{1-\alpha}.$$

تعریف ۶. امید ریاضی و واریانس $\tilde{X}$ به صورت زیر تعریف می‌شود:

$$\begin{aligned}E(\tilde{X})_\alpha &= E(\tilde{X}_\alpha), \\Var(\tilde{X}) &= E(d(\tilde{X}, \tilde{E}(\tilde{X}))),\end{aligned} \quad (۳)$$

$$\tilde{A}(x) = \begin{cases} L(\frac{a-x}{l_a}) & x \leq a \\ R(\frac{x-a}{r_a}) & x > a \\ 0 & x \in \mathbb{R} - [a-l_a, a+r_a], \end{cases}$$

که در آن L و R توابعی غیر صعودی از $\mathbb{R}^+$ به $[0, 1]$ می‌باشند و $L(0) = 1$. عدد حقیقی a مقدار نما (میانه)، l_a و r_a به ترتیب پهنای چپ و پهنای راست $\tilde{A}$ نامیده می‌شوند. عدد فازی LR با نماد $(a; l_a, r_a)_{LR}$ نمایش داده می‌شود.

همچنین $\tilde{A}$ به ازای هر $x \in [0, 1]$ یک عدد فازی مثلثی نامیده می‌شود، هرگاه $L(x)=R(x)=1-x$ و آن را با نماد $(a; l_a, r_a)_T$ نمایش می‌دهند. تابع عضویت آن به صورت زیر بیان می‌شود:

$$\tilde{A}(x) = \begin{cases} \frac{x - (a - l_a)}{l_a} & a - l_a \leq x \leq a \\ \frac{a + r_a - x}{r_a} & a \leq x \leq a + r_a \\ 0 & x \in \mathbb{R} - [a - l_a, a + r_a]. \end{cases}$$

ملاحظه ۱. فرض کنید $F(\mathbb{R})$ مجموعه اعداد فازی باشد. برای $\tilde{A} \in F(\mathbb{R})$ ، نگاشت $\tilde{A}_\alpha : [0, 1] \rightarrow \mathbb{R}$ α -شک $\tilde{A}$ نامیده می‌شود و به صورت زیر بیان می‌شود:

$$\tilde{A}_\alpha = \begin{cases} \tilde{A}^L[2\alpha] & \alpha \in [0, 0.5] \\ \tilde{A}^U[2(1-\alpha)] & \alpha \in (0.5, 1], \end{cases}$$

که $\tilde{A}^L[\alpha]$ و $\tilde{A}^U[\alpha]$ به ترتیب نشان‌دهنده کران‌های پایین و بالای α -برش‌های $\tilde{A}$ می‌باشند. به سادگی می‌توان نوشت:

$$\tilde{A}[\alpha] = [\tilde{A}^L[\alpha], \tilde{A}^U[\alpha]] = [\tilde{A}_{\alpha/2}, \tilde{A}_{1-\alpha/2}]. \quad (۱)$$

ملاحظه ۲. α -شک عدد فازی LR $\tilde{A} = (a; l_a, r_a)_{LR}$ به صورت زیر تعریف می‌شود:

$$\tilde{A}_\alpha = \begin{cases} a - l_a L^{-1}(2\alpha) & \alpha \in [0, 0.5] \\ a + r_a R^{-1}(2(1-\alpha)) & \alpha \in (0.5, 1]. \end{cases}$$

به طور مشابه، α -شک عدد فازی مثلثی $\tilde{A} = (a; l_a, r_a)_T$ به صورت زیر بیان می‌شود:

$$\tilde{A}_\alpha = \begin{cases} a - l_a + 2l_a\alpha & \alpha \in [0, 0.5] \\ a + r_a - 2r_a(1-\alpha) & \alpha \in (0.5, 1]. \end{cases}$$

ملاحظه ۳. به ازای هر $\tilde{A}, \tilde{B} \in F(\mathbb{R})$ و $\lambda \in \mathbb{R}$ و $\alpha \in [0, 1]$ عملگرهای

$$= \begin{cases} 1 - F_X\left(\frac{z}{1 + (1 - 2(1 - \alpha))b_2}\right) \\ - \int \frac{1 + (1 - 2(1 - \alpha))a_2}{z} \frac{z/x - 1 - (1 - 2(1 - \alpha))a_2}{(1 - 2(1 - \alpha))(b_2 - a_2)} \\ \times f_X(x) dx & 0 \leq 1 - \alpha \leq 0.5 \\ 1 - F_X\left(\frac{z}{1 + (1 - 2(1 - \alpha))a_1}\right) \\ - \int \frac{1 + (1 - 2(1 - \alpha))b_1}{z} \frac{z/x - 1 - (1 - 2(1 - \alpha))b_1}{(2(1 - \alpha) - 1)(b_1 - a_1)} \\ \times f_X(x) dx & 0.5 < 1 - \alpha \leq 1 \end{cases}$$

$$= \begin{cases} \bar{F}_X\left(\frac{z}{1 + (2\alpha - 1)a_1}\right) \\ - \int \frac{1 + (2\alpha - 1)b_1}{z} \frac{z/x - 1 - (2\alpha - 1)b_1}{(1 - 2\alpha)(b_1 - a_1)} f_X(x) dx \\ & 0.0 \leq \alpha \leq 0.5 \\ \bar{F}_X\left(\frac{z}{1 + (2\alpha - 1)b_2}\right) \\ - \int \frac{1 + (2\alpha - 1)a_2}{z} \frac{z/x - 1 - (2\alpha - 1)a_2}{(2\alpha - 1)(b_2 - a_2)} f_X(x) dx \\ & 0.5 < \alpha \leq 1. \end{cases}$$

□

ملاحظه ۷. فرض کنید متغیر تصادفی فازی $\tilde{X}$ یک **SFRV** باشد، بر اساس توجه ۶ داریم:

$$f_{\tilde{X}_\alpha}(z) = -\frac{d}{dz} \bar{F}_{\tilde{X}_\alpha}(z) = \frac{d}{dz} (F_{\tilde{X}}(z))_{1-\alpha}. \quad (۷)$$

لم ۳. (حسامیان، اکبری و زندهدل، [۴]) فرض کنید $\tilde{X}$ یک **SFRV** باشد، آنگاه

$$\begin{aligned} \tilde{\mathbf{E}}(\tilde{X}) &= \mathbf{E}(X) \otimes (1; \frac{a_1 + b_1}{2}, \frac{a_2 + b_2}{2})_T, \\ \mathbf{Var}(\tilde{X}) &= \mathbf{Var}(X) \left(1 + \frac{(a_1 + b_1)^2 + (a_2 + b_2)^2}{24} \right. \\ &\quad \left. - \frac{(a_2 + b_2) + (a_1 + b_1)}{4} \right) \\ &\quad + \mathbf{E}(X^2) \left(\frac{(a_1 - b_1)^2 + (a_2 - b_2)^2}{72} \right). \end{aligned} \quad (۸)$$

مثال ۱. فرض کنید $\tilde{X}$ یک **SFRV** باشد، که $X \sim \text{Exp}(0.002)$ ، $U_1 \sim U(0, 1)$ و $U_2 \sim U(0, 1)$ باشند. بنا بر (۴) داریم:

$$(\tilde{F}_{\tilde{X}}(z))_\alpha = \begin{cases} F_X\left(\frac{z}{2(1-\alpha)}\right) + \int \frac{z}{2(1-\alpha)b_2} \frac{z/x - 1}{(1-2\alpha)} f_X(x) dx & 0 \leq \alpha \leq 0.5 \\ F_X(z) + \int \frac{z}{2(1-\alpha)} \frac{z/x - 2(1-\alpha)}{(2\alpha-1)} f_X(x) dx & 0.5 < \alpha \leq 1. \end{cases}$$

که برای دو عدد فازی $\tilde{A}$ و $\tilde{B}$ ، $\mathbf{d}(\tilde{A}, \tilde{B}) = \int_0^1 (\tilde{A}_\alpha - \tilde{B}_\alpha)^2 d\alpha$ ، $\text{Var}(\tilde{X}) = \int_0^1 \text{Var}(\tilde{X})_\alpha d\alpha$

تعریف ۷. فرض کنید X یک متغیر تصادفی مکان-مقیاس با تابع چگالی زیر باشد:

$$\{f_X(x) = \frac{1}{\sigma} g((x - \mu)/\sigma) : x \in S_X \subseteq \mathbb{R}, \mu \in \mathbb{R}, \sigma > 0\},$$

می‌شود، اگر در شرایط زیر صدق کند:

$$(۱) \quad P(S_X \subseteq (0, \infty)) = 1 \text{ و } \mu = 0.$$

(۲) $U_1 \sim U(a_1, b_1)$ و $U_2 \sim U(a_2, b_2)$ ، X مستقل از U_1 و U_2 باشد. $0 \leq a_2 < b_2$ ، $0 \leq a_1 < b_1 \leq 1$

لم ۱. (حسامیان، اکبری و زندهدل [۴]) فرض کنید $\tilde{X}$ متغیر تصادفی فازی مقیاس باشد، آنگاه α -شک $\tilde{F}_{\tilde{X}}(z)$ به صورت زیر است:

$$(\tilde{F}_{\tilde{X}}(z))_\alpha = \begin{cases} F_X\left(\frac{z}{1 + (1 - 2\alpha)b_2}\right) \\ + \int \frac{1 + (1 - 2\alpha)a_2}{z} \frac{z/x - 1 - (1 - 2\alpha)a_2}{(1 - 2\alpha)(b_2 - a_2)} \\ \times f_X(x) dx & 0 \leq \alpha \leq 0.5 \\ F_X\left(\frac{z}{1 + (1 - 2\alpha)a_1}\right) \\ + \int \frac{1 + (1 - 2\alpha)b_1}{z} \frac{z/x - 1 - (1 - 2\alpha)b_1}{(2\alpha - 1)(b_1 - a_1)} \\ \times f_X(x) dx & 0.5 < \alpha \leq 1. \end{cases}$$

ملاحظه ۶. فرض کنید $\tilde{X}$ یک **SFRV** باشد، آنگاه داریم:

$$\begin{aligned} \bar{F}_{\tilde{X}_\alpha}(z) &= 1 - P(\tilde{X}_\alpha \leq z) = 1 - (\bar{F}_{\tilde{X}}(z))_{1-\alpha}, \\ \bar{F}_{\tilde{X}_{1-\alpha}}(z) &= 1 - P(\tilde{X}_{1-\alpha} \leq z) = 1 - (\bar{F}_{\tilde{X}}(z))_\alpha. \end{aligned} \quad (۵)$$

لم ۲. فرض کنید $\tilde{X}$ یک **SFRV** باشد، آنگاه α -شک $\tilde{F}_{\tilde{X}}(z)$ را می‌توان به صورت زیر نوشت:

$$(\tilde{F}_{\tilde{X}}(z))_\alpha = \begin{cases} \bar{F}_X\left(\frac{z}{1 + (2\alpha - 1)a_1}\right) \\ - \int \frac{1 + (2\alpha - 1)b_1}{z} \frac{z/x - 1 - (2\alpha - 1)b_1}{(1 - 2\alpha)(b_1 - a_1)} \\ \times f_X(x) dx & 0.0 \leq \alpha \leq 0.5 \\ \bar{F}_X\left(\frac{z}{1 + (2\alpha - 1)b_2}\right) \\ - \int \frac{1 + (2\alpha - 1)a_2}{z} \frac{z/x - 1 - (2\alpha - 1)a_2}{(2\alpha - 1)(b_2 - a_2)} \\ \times f_X(x) dx & 0.5 < \alpha \leq 1. \end{cases}$$

اثبات. بنا بر روابط (۴) و (۵) می‌توان نوشت:

$$(\tilde{F}_{\tilde{X}}(z))_\alpha = \bar{F}_{\tilde{X}_\alpha}(z) = 1 - (\bar{F}_{\tilde{X}}(z))_{1-\alpha}$$

۴ مفاهیم پایه قابلیت اعتماد فازی

در جهانی که در آن زندگی می‌کنیم، هیچ چیز از جمله اشیا و موجودات باقی و ثابت نیستند و در نهایت فانی خواهند شد. بنابراین از جمله خواسته‌های بشر نیاز به دانستن وضعیت طول عمر موجودات و اشیا اطراف اوست، تا بتواند به طور بهینه از آنها برای پیشبرد اهداف خود استفاده کند. این نیاز طبیعی بشر، امروزه منجر به ایجاد شاخه‌ای بزرگ و گسترده در تمامی علوم به نام قابلیت اعتماد یا بقا شده است.

لم ۴. فرض کنید $\tilde{T}$ که یک متغیر تصادفی فازی است، طول عمر سیستم باشد. آنگاه

$$(\tilde{R}_{\tilde{T}}(z))_{\alpha} = R_{\tilde{T}_{\alpha}}(z), \quad (9)$$

که $\tilde{R}_{\tilde{T}}(z)$ تابع بقا آن می‌باشد.

اثبات. به راحتی می‌توان نشان داد:

$$(\tilde{R}_{\tilde{T}}(z))_{\alpha} = (\bar{F}_{\tilde{T}}(z))_{\alpha} = \bar{F}_{\tilde{T}_{\alpha}}(z) = R_{\tilde{T}_{\alpha}}(z).$$

□

نکته ۱. عبارت‌های بیان شده در توجه ۶ در واقع، رابطه بین تابع توزیع و تابع بقای یک متغیر تصادفی فازی را بیان می‌کنند.

تعریف ۸. فرض کنید $\tilde{T}$ طول عمر سیستم باشد. نرخ خطر آن یعنی $\tilde{h}_{\tilde{T}}(z)$ به صورت زیر می‌باشد:

$$(\tilde{h}_{\tilde{T}}(z))_{\alpha} = \begin{cases} \inf_{\beta \in I_{2\alpha}} \frac{f_{\tilde{T}_{\beta}}(z)}{R_{\tilde{T}_{\beta}}(z)} & 0.0 \leq \alpha \leq 0.5 \\ \sup_{\beta \in I_{2(1-\alpha)}} \frac{f_{\tilde{T}_{\beta}}(z)}{R_{\tilde{T}_{\beta}}(z)} & 0.5 < \alpha \leq 1.0. \end{cases} \quad (10)$$

با توجه به (۲)، کران‌های پایین و بالای $\tilde{h}_{\tilde{T}}(z)$ به صورت زیر تعریف می‌شوند:

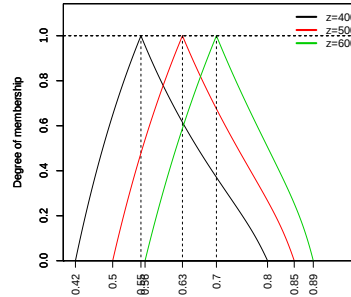
$$\begin{aligned} (\tilde{h}_{\tilde{T}}(z))^L[\alpha] &= \inf_{\beta \in [\alpha/2, 1-\alpha/2]} \frac{f_{\tilde{T}_{\beta}}(z)}{R_{\tilde{T}_{\beta}}(z)}, \\ (\tilde{h}_{\tilde{T}}(z))^U[\alpha] &= \sup_{\beta \in [\alpha/2, 1-\alpha/2]} \frac{f_{\tilde{T}_{\beta}}(z)}{R_{\tilde{T}_{\beta}}(z)}. \end{aligned} \quad (11)$$

مثال ۳. فرض کنید طول عمر سیستم $\tilde{T}$ یک SFRV باشد، که $T \sim U(0,1), \text{Exp}(0.002)$ و $U_1 \sim U(0,1)$ و $U_2 \sim U(0,1)$ باشند. برای محاسبه $\tilde{h}_{\tilde{T}}(z)$ بنا بر (۵) و (۷) می‌توان نوشت:

$$f_{\tilde{T}_{\beta}}(z) = \begin{cases} \int_{z/2\beta}^{z/2\beta} \frac{1}{t(1-2\beta)} f_T(t) dt & 0 \leq \beta \leq 0.5 \\ \int_{z/2\beta}^z \frac{1}{t(2\beta-1)} f_T(t) dt & 0.5 < \beta \leq 1, \end{cases} \quad (12)$$

$$\bar{F}_{\tilde{T}_{\beta}}(z) = \begin{cases} \bar{F}_T(z) - \int_z^{z/2\beta} \frac{z/t-2\beta}{(1-2\beta)} f_T(t) dt & 0 \leq \beta \leq 0.5 \\ \bar{F}_T(\frac{z}{2\beta}) - \int_{z/2\beta}^z \frac{z/t-1}{(2\beta-1)} f_T(t) dt & 0.5 < \beta \leq 1, \end{cases} \quad (13)$$

نمودار $\tilde{F}_{\tilde{T}}(z)$ برای نقاط ۴۰۰، ۵۰۰ و ۶۰۰ در منحنی ۱ نمایش داده شده است. بعلاوه بنا بر (۸)، $\tilde{E}(\tilde{T}) = (500; 250, 250)_T$ و $\text{Var}(\tilde{T}) = 159722.22$ می‌باشند.

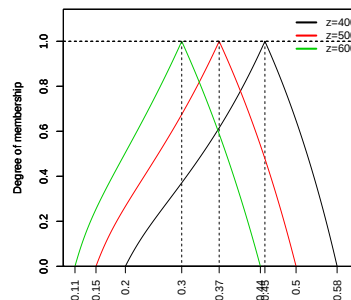


شکل ۱: منحنی $\tilde{F}_{\tilde{T}}(z)$ برای نقاط ۴۰۰، ۵۰۰ و ۶۰۰

مثال ۲. فرض کنید متغیر تصادفی فازی $\tilde{X}$ یک SFRV باشد، که $X \sim \text{Exp}(0.002)$ و $U_1 \sim U(0,1)$ و $U_2 \sim U(0,1)$ باشند. بنا بر معادله (۵) داریم:

$$(\tilde{F}_{\tilde{X}}(z))_{\alpha} = \begin{cases} \bar{F}_X(z) - \int_z^{z/2\alpha} \frac{z/x-2\alpha}{(1-2\alpha)} f_X(x) dx & 0 \leq \alpha \leq 0.5 \\ \bar{F}_X(\frac{z}{2\alpha}) - \int_{z/2\alpha}^z \frac{z/x-1}{(2\alpha-1)} f_X(x) dx & 0.5 < \alpha \leq 1, \end{cases}$$

که نمودار $\tilde{F}_{\tilde{X}}(z)$ برای نقاط ۴۰۰، ۵۰۰ و ۶۰۰ در منحنی ۲ نمایش داده شده است.



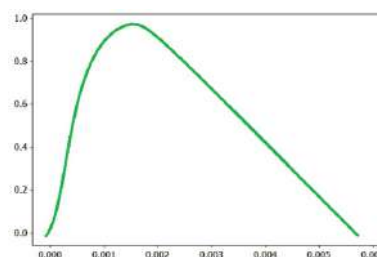
شکل ۲: منحنی $\tilde{F}_{\tilde{X}}(z)$ برای نقاط ۴۰۰، ۵۰۰ و ۶۰۰

که با جایگذاری (۱۲) و (۱۳) در (۱۱) داریم:

$$(\tilde{h}_{\bar{T}}(z))^L[\alpha] = \begin{cases} \inf_{\beta \in [\alpha/2, 0.5]} \frac{\int_z^{z/2\beta} \frac{1}{t(1-2\beta)} f_T(t) dt}{\bar{F}_T(z) - \int_z^{z/2\beta} \frac{z/t-2\beta}{(1-2\beta)} f_T(t) dt}, \\ \inf_{\beta \in [0.5, 1-\alpha/2]} \frac{\int_{z/2\beta}^z \frac{1}{t(2\beta-1)} f_T(t) dt}{\bar{F}_T(\frac{z}{2\beta}) - \int_{z/2\beta}^z \frac{z/t-1}{(2\beta-1)} f_T(t) dt}, \end{cases}$$

$$(\tilde{h}_{\bar{T}}(z))^U[\alpha] = \begin{cases} \sup_{\beta \in [\alpha/2, 0.5]} \frac{\int_z^{z/2\beta} \frac{1}{t(1-2\beta)} f_T(t) dt}{\bar{F}_T(z) - \int_z^{z/2\beta} \frac{z/t-2\beta}{(1-2\beta)} f_T(t) dt}, \\ \sup_{\beta \in [0.5, 1-\alpha/2]} \frac{\int_{z/2\beta}^z \frac{1}{t(2\beta-1)} f_T(t) dt}{\bar{F}_T(\frac{z}{2\beta}) - \int_{z/2\beta}^z \frac{z/t-1}{(2\beta-1)} f_T(t) dt}. \end{cases}$$

حال با در نظر گرفتن $z = 500$ و شبیه‌سازی الگوریتم بالا، منحنی $\tilde{h}_{\bar{T}}(z)$ در شکل ۳ نمایش داده شده است.



شکل ۳: منحنی $\tilde{h}_{\bar{T}}(z)$ برای نقطه ۵۰۰

مراجع

- [1] M. G. Akbari and G. Hesamian "Linear model with exact inputs and interval-valued fuzzy outputs," IEEE Transactions on Fuzzy Systems. 26, 518-530, 2017.
- [2] M. A. Gil, M. López-Díaz and D. A. Ralescu "Overview on the development of fuzzy random variables," Fuzzy Sets Systems, 157, 2546-2557, 2006.
- [3] G. Hesamian, M. G. Akbari and J. Zendehdel "Location and scale fuzzy random variables," International Journal of Systems Science. 229-241, 2019.
- [4] G. Hesamian and S. M. Taheri "Fuzzy empirical distribution function: Properties and application," Kybernetika, 49, 962-982, 2013.
- [5] E. L. Lehmann and G. Casella "Theory of point estimation (2nd ed.)," Springer, 1998.
- [6] F. A. Shapiro "Fuzzy random variables," Insurance: Journal of Mathematical Economics, 44, 307-314, 2009.
- [7] Q. M. William and L. A. Escobar, L. A "Statistical methods for reliability data," John Wiley and Sons, Inc, 1998.
- [8] L. A. Zadeh "Fuzzy sets," Information Control, 8, 338-356, 1965.

۵ نتیجه

در این مقاله، یک مفهوم جدید از متغیر تصادفی فازی ارائه گردید. برای این منظور، یک خانواده توزیع احتمال مقیاس مورد استفاده قرار گرفت، سپس α -شک اعداد فازی برای توصیف مبهم بودن فرایندهای قبلی ترکیب شدند و متغیرهای تصادفی فازی مقیاس را به وجود آوردند. به این ترتیب، برخی از ویژگی‌های آماری بر اساس متغیرهای تصادفی فازی مقیاس مانند توزیع تجمعی، امید ریاضی، واریانس و احتمال غیردقیق یک بازه مورد بررسی قرار گرفت. بعلاوه، برخی از معیارهای قابلیت اعتماد نظیر تابع بقا و تابع نرخ شکست ارائه شد. همچنین مثال‌هایی عددی نیز برای نشان دادن محاسبه عملگرها بیان گردید. برای مطالعه در آینده، بر اساس این مفهوم پیشنهادی متغیر تصادفی فازی،



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

کنترل شبکه‌های تنظیم‌کننده ژن P53 مبتنی بر یادگیری تقویتی عمیق و کاربرد آن در بیماری سرطان

علی صالحی^۱، محمدرضا اکبرزاده توتونچی^۲، علیرضا روحانی‌منش^۳

^۱ گروه مهندسی کامپیوتر، قطب علمی رایانش نرم و پردازش هوشمند اطلاعات، دانشگاه فردوسی مشهد، alisalehi@mail.um.ac.ir

^۲ گروه مهندسی برق، قطب علمی رایانش نرم و پردازش هوشمند اطلاعات، دانشگاه فردوسی مشهد، akbazar@um.ac.ir

^۳ گروه مهندسی برق، دانشگاه نیشابور، rowhanimanesh@neyshabur.ac.ir

چکیده: شبکه‌های تنظیم‌کننده ژن (GRN) مجموعه‌ای از تنظیم‌کننده‌های مولکولی بیان DNA هستند که با یکدیگر و سایر مواد سلول برهم‌کنش دارند. این شبکه مسئولیت کنترل فرآیند عملکرد سلولی را بر عهده دارد و با تنظیم نرخ تبدیل ژن به mRNA^۲، به‌طور غیرمستقیم روی ترجمه پروتئین اثر می‌گذارد. در پزشکی مولکولی، GRN‌ها نقشی اساسی در شناخت جاذب‌های^۳ بیماری و سلامتی دارند، و همچنین ساز و کار بیماری را از طریق روش‌های درمانی برای سوق دادن جاذب‌های سلامتی کنترل می‌کنند. از نگاه مدل‌سازی محاسباتی، GRN را می‌توان به صورت شبکه‌ای از سیستم‌های پویای غیرخطی در نظر گرفت که دارای رفتارهای تصادفی و ناپایستا هستند. شبکه‌های عصبی عمیق-Q^۴ (DQN) نسخه‌ای از یادگیری تقویتی عمیق هستند که در آن تابع Q توسط شبکه‌های عصبی عمیق تقریب زده می‌شوند. در اکثر روش‌های موجود، قابلیت پیاده‌سازی سیاست کارگزار، پیچیدگی، و عدم قطعیت در GRN‌ها کمتر مورد توجه قرار گرفته است. لذا در این مقاله قصد داریم تا سیاست درمانی هوشمندی مبتنی بر DQN برای شبکه دودویی تنظیم‌کننده ژن P53 به گونه‌ای ارائه کنیم تا با در نظر داشتن محدودیت‌های کنش در دنیای واقعی، کارگزار از مقاومت کافی در مواجهه با عدم قطعیت برخوردار باشد. نتایج حاصل نشان می‌دهد که با افزایش نرخ اغتشاش شبکه P53، درصد ماندگاری در ژن caspase که منجر به مرگ سلول سرطانی می‌شود به ترتیب برابر با ۹۹٫۸٪، ۹۷٫۸۵٪، ۹۶٫۷٪، و ۹۵٫۰۵٪ است.

کلمات کلیدی: شبکه‌های تنظیم‌کننده ژن، یادگیری تقویتی عمیق، شبکه‌های عصبی عمیق-Q، P53، سرطان.

متفاوت است [14]. نحوه یادگیری در این الگوریتم، به صورت برون-سیاست بوده و کارگزار به‌هنگام بهینه‌سازی دانش خود، این اجازه را دارد که از سیاستی متفاوت از سیاست کنونی خود پیروی کند. رابطه ۱ نحوه بهینه‌سازی این یادگیری را نشان می‌دهد [14].

$$Q(s, a, w) = Q(s, a, w) + \alpha \left[\left(R + \gamma \max_{a'} Q(s', a', w') \right) - Q(s, a, w) \right] \quad (1)$$

۱- مقدمه

در سال‌های گذشته، تحقیقات بسیاری در حوزه یادگیری تقویتی عمیق انجام پذیرفته است. این روش از یادگیری مدلی ترکیبی از یادگیری تقویتی و شبکه‌های عصبی عمیق هستند که هدف آن یافتن سیاست یک کارگزار به گونه‌ای است که بیشینه پاداش در یک محیط را به همراه داشته باشد. یادگیری Q یکی از روش‌های یادگیری مبتنی بر ارزش است که در آن کارگزار به دنبال ارزش گذاری انتخاب هر کنش در مواجهه با حالت‌های

^۱ Gene Regulatory Networks

^۲ Messenger RNA

^۳ Attractors

^۴ Deep Q-Networks

که در آن، Q مقدار ارزش حالت-کنش، α نرخ یادگیری، R پاداش دریافت شده از محیط، γ نرخ تنزیل، S حالت، a کنش‌های کارگزار، و w وزن‌های شبکه است. همان‌طور که در این رابطه بهینه‌سازی پیدا است، عبارت $\max$ نشان‌گر رفتار برون-سیاست کارگزار بوده که با نگاهی خوش‌بینانه نسبت به آینده، فرض را بر این می‌گذارد که در حالت بعدی همواره بهترین کنش را انتخاب می‌کند.

یادگیری تقویتی عمیق را می‌توان در گستره‌ای از کاربردهای گوناگون اقتصادی، رباتیک، و علوم پزشکی مشاهده نمود. همچنین تحقیق‌های بسیاری در انواع یادگیری‌های برون-سیاست انجام پذیرفته که در ادامه به بررسی چندی از آن خواهیم پرداخت. گوگل دیپ مایند^۵ توانست ۴۹ بازی آتاری را تنها با دریافت تصویر ورودی، با استفاده از **DQN** یاد بگیرد [10, 11]. پس از معرفی **DQN**، مدل‌هایی گسترش یافته‌ای از آن ارائه گردید که یادگیری دوطرفه Q یکی از آن‌ها است که به موجب بیش‌برآورد تابع Q معرفی شده است [5, 6]. علاوه بر این از آن‌جا که شبکه‌های عصبی عمیق یادگیری بهتری روی داده‌های ورودی مستقل از هم و با توزیع یکنواخت دارند، حافظه بازپخش اولویت‌بندی شده^۶ به ساختار یادگیری Q اضافه گردید [13]. همچنین، شبکه دوئل^۷ که مبتنی بر تخمین ارزش است، با استفاده از مقدار ارزش حالت، و مفهوم جدیدی به نام مزیت^۸ کنش، تابع Q را تخمین می‌زند [15]. شبکه‌های نویزی^۹ با اضافه نمودن نویز به وزن‌های شبکه حین عملیات یادگیری سعی در بهبود عملیات اکتشاف دارند تا رسیدن به سیاست بهینه را امکان‌پذیر سازند [3]. نمونه دیگری از شبکه‌های مبتنی بر **DQN** شبکه رنگین‌کمان^{۱۰} است که تمامی ویژگی‌های مؤثر یاد شده را درون یک شبکه گنجانده است [8].

یکی از کاربردهای یادگیری تقویتی عمیق در حوزه علوم پزشکی مولکولی است که کنترل شبکه‌های تنظیم‌کننده ژن نمونه‌ای از این موارد است. شبکه‌های تنظیم‌کننده ژن مسئولیت کنترل پردازش‌های سلولی را برعهده دارند که از جمله این موارد می‌توان به ترمیم **DNA**^{۱۱}، پاسخ به تحریکات شدید محیطی^{۱۲}، و یا کنترل بیماری‌های پیچیده‌ای نظیر سرطان اشاره نمود. هر ژن در این شبکه می‌تواند تأثیر ارتباطی تحریک‌کننده^{۱۳} یا بازدارنده^{۱۴} روی یک یا چند ژن دیگر را داشته باشد که موجب تغییر بیان ژن می‌شود که تأثیر غیر مستقیمی روی ترجمه پروتئین در بدن دارد. از نگاه مدل‌سازی

محاسباتی، **GRN** را می‌توان به صورت شبکه‌ای از سیستم‌های پویای غیرخطی در نظر گرفت که دارای رفتارهای تصادفی و نامانا هستند. شبکه‌های دودویی نمونه‌ای از این مدل سازی ریاضی هستند که در آن هر گره نشانگر یک ژن با مقادیر صفر یا یک است و می‌تواند رفتار این محیط را به خوبی نمایش دهد [16]. تأثیر ارتباطی ژن‌ها در این شبکه مسیره‌های زیستی گوناگونی را تشکیل می‌دهند که می‌تواند سلول را در وضعیت مطلوب سلامتی یا نامطلوب بیماری قرار دهد، بنابراین لازم است تا با مداخله‌های درمانی **GRN** را به مسیری که حالت مطلوب را تولید می‌کند سوق دهیم.

از آن‌جا که مداخلات درمانی دنباله‌ای از کنش‌های^{۱۵} مورد نظر برای رسیدن به هدفی معین است می‌توان از یادگیری تقویتی عمیق به عنوان یکی از روش‌های هوشمند یادگیری جهت یافتن سیاست‌های درمان استفاده نمود. به طور مثال، پایاگانیس و همکاران به کنترل شبکه‌های دودویی با استفاده از یادگیری تقویتی عمیق پرداخته و اشاره می‌کند که می‌توان کاربرد یادگیری مبتنی بر **DQN** را برای محیط‌های **GRN** استفاده نمود [12]. همچنین داورتی و همکاران به بررسی مبانی اولیه کنترل ساختاری **GRN** مانا و نحوه مداخله درمانی در روند کنترل پرداخته‌اند [2]. در ادامه این تحقیق می‌توان به کار ایمانی و همکاران اشاره نمود که به کنترل شبکه **WNT5A** برای بیماری ملانوما^{۱۶} با استفاده از یادگیری معکوس تقویتی می‌پردازند [9]. در این روش، ابتدا از دانش فرد خبره برای انتخاب مداخله‌های درمانی شامل کنش‌های حذف و اضافه نمودن گره بهره گرفته شده است و سپس با استفاده از این دانش کارگزار تقویتی یادگیری را انجام می‌دهد. علاوه بر این، آنالیز حلقه‌های بازخورد^{۱۷} شبکه P53 در شرایط آسیب **DNA** سلولی از جمله تحقیقات در حوزه شبکه‌های تنظیم‌کننده ژن است [1].

از آن‌جا که شبکه‌ی تنظیم‌کننده ژن **P53** محیطی پیچیده شامل ۱۶ گره و مجموع ۳۷۸ یال ارتباطی^۱ است، لذا استفاده از روش‌های کلاسیک یادگیری تقویتی و یا کنترل هوشمند برای آن کاربردی به نظر نمی‌رسد. همین امر سبب شده تا از مداخله فرد خبره برای همگرایی به جواب مطلوب استفاده شود [9]. علاوه بر این، نگاه سیستمی روش‌های پیشین به **GRN** سبب شده تا هزینه و یا قابلیت اجرای سیاست کارگزار در دنیای واقعی کمتر مورد توجه قرار گیرد. از طرفی از آن‌جا که عدم قطعیت‌های زیادی در سیستم انسان وجود دارد، لذا مهم است که تا مدلی درمانی ارائه شود که در مواجهه با تغییرات محیطی مقاوم باشد. در این مقاله قصد داریم تا به کمک شبکه‌های عمیق **Q**

^۵ Google Deep Mind

^۶ Replay Memory Buffer

^۷ Dueling Network

^۸ Advantage

^۹ Noisy Networks

^{۱۰} Rainbow Network

^{۱۱} DNA Repair

^{۱۲} Stress Response

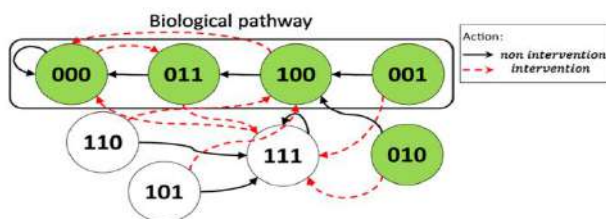
^{۱۳} Activator

^{۱۴} Suppressor

^{۱۵} Sequence of Interventions

^{۱۶} Melanoma

^{۱۷} Feedback Loop



شکل ۲: شبکه دودویی شامل سه ژن همراه با کنش‌های مداخله‌ای [7].

۲-۲- سیاست درمانی با DQN

شبکه عمیق-Q یکی از روش‌های مبتنی بر ارزش در یادگیری تقویتی عمیق است که در آن سیاست بهینه در هر حالت از طریق تخمین میزان ارزش کنش‌های تعریف شده برای کارگزار تعیین می‌شود که نمای آن در شکل ۳ نشان داده شده است. یکی از مشکلات یادگیری تقویتی عمیق، ناپایداری و بیش‌برازش آن‌ها در مواجهه با مسأله یادگیری است. در DQN برای جلوگیری از این ناپایداری و بیش‌برازش کارگزار تقویتی، از دو تکنیک حافظه بازپخش و شبکه هدف استفاده می‌شود. در تکنیک حافظه بازپخش، کارگزار همزمان با انجام کنش در محیط، مجموعه‌ای از داده‌ها را برای خود تولید می‌نماید و به کمک آن عملیات یادگیری انجام می‌شود. همچنین شبکه هدف به عنوان تجربه کارگزار در زمان‌های پیشین است، سبب می‌شود تا از بیش‌برآورد^{۱۸} مدل جلوگیری نماید. به همین منظور می‌توان رابطه ۱ را به صورت زیر بازنویسی کرد که در آن $\bar{Q}$ خروجی شبکه هدف در حالت s' خواهد بود [10]:

$$Q(s, a, w) = Q(s, a, w) + \alpha \left[(R + \gamma \max_{a'} \bar{Q}(s', a', w')) - Q(s, a, w) \right] \nabla_w Q(s, a, w) \quad (3)$$

برای یک شبکه دودویی با $V = |g_i|$ گره و E یال، تعداد حالت‌ها برابر با 2^V خواهد بود. همچنین در این محیط دو روش مداخله درمانی حذف گره^{۱۹} یا یال^{۲۰} را می‌توان در نظر گرفت که به معنای صفر نمودن گره ژن یا یال میان دو ژن توسط دارو در لحظه $t + 1$ است. تعداد کنش‌ها به ترتیب برای حذف گره و یال برابر $V + 1$ و $E + 1$ خواهد بود که کنش اضافی به معنای بدون کنش^{۲۱} خواهد بود. لازم به ذکر است که مدل‌های ترکیبی از دو روش مداخله درمانی نیز وجود دارند که در این حالت می‌توان تعداد کنش $V + E + 1$ متصور بود که عددی بسیار بزرگ است. در شبکه‌های تنظیم کننده ژن قصد بر آن است تا با انجام دنباله‌ای از مداخله‌ها، سیاستی درمانی را به نحوی بیابیم که به حالت مطلوب با فنوتایپ^{۲۲} سلامتی برسیم. یادگیری تقویتی عمیق می‌تواند سیاستی بهینه از مداخله‌های درمانی برای یک GRN را یافته و داروساز با دنبال نمودن این سیاست داروی مصرفی برای بیمار را تولید نماید. در این مقاله برای ساده‌سازی مسأله، تنها مداخله حذف گره به عنوان کنش‌های کارگزار در نظر گرفته شده است. با اضافه نمودن مداخله درمانی در

به عنوان یکی از روش‌های یادگیری هوشمند سیاست درمان، مسیری از مداخله‌های درمانی که منجر به سلامتی انسان می‌شود را بدون دخالت فرد خبره ارائه نماییم. در این روش سیاست کارگزار محدود به انجام مداخله درمانی حذف گره گردیده است تا از قابلیت اجرای آن در دنیای واقعی اطمینان حاصل گردد. سپس عملکرد مدل یادگیری تقویتی پیشنهادی در این مقاله در مواجهه با رفتار تصادفی و نامانای شبکه P53 مورد بررسی قرار گرفته است.

۲- رهیافت پیشنهادی

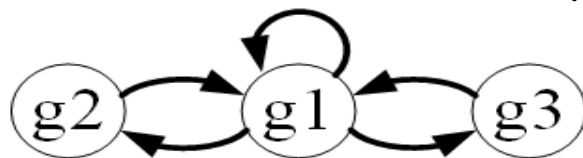
۲-۱- ساختار شبکه دودویی تنظیم کننده ژن

برای مجموعه‌ای از ژن‌های g_i در یک شبکه دودویی تنظیم کننده ژن، حالت در هر لحظه t به صورت $S_t = \{g_i: \{0,1\}; i=0,1,2,\dots\}$ تعریف می‌شود که در آن i شماره ژن است. در این شبکه بیان هر ژن می‌تواند مقدار $\{0,1\}$ که نمایانگر فعال یا غیرفعال بودن آن ژن است را داشته باشد. در هر لحظه گسسته، این شبکه با سیگنالی غیر خطی به روزرسانی می‌شود [4]:

$$S_{t+1} = f(S_t) \oplus \eta_t \quad (2)$$

که در آن $f: \{0,1\}^N \rightarrow \{0,1\}^N$ تابع گذر، و η مقدار نویز در لحظه t است.

به طور مثال، در شکل ۱ شبکه‌ای دودویی شامل سه ژن نشان داده شده است که می‌توان تابع گذر منطقی در جدول ۱ را برای آن در نظر گرفت. در این شبکه دودویی گره‌ها معادل ژن‌ها، و یال‌ها نشانگر ارتباط و تأثیرگذاری ژن‌ها روی یکدیگر هستند. به همین ترتیب می‌توان مسیرهای گوناگون زیستی را که از برهم کنش ژن‌های درون شبکه بر یکدیگر به دست می‌آیند به نمایش گذاشت. همان‌طور که در شکل ۲ مشاهده می‌شود، مسیرهای زیستی درون این شبکه دودویی تنظیم کننده ژن می‌تواند از طریق مداخله‌های درمانی و یا بدون دخالت به وجود آیند. به طور نمونه در این مثال می‌بینیم که می‌توان با مداخله‌های درمانی، از حالت نامطلوب (بیماری) "۱۱۱" مسیری به حالت مطلوب "۰۰۰" یافت [7].



شکل ۱: شبکه دودویی شامل سه ژن.

جدول ۱: تابع تنظیم شبکه دودویی شامل سه ژن [7].

ژن	تابع تنظیم
g_1	$(\bar{g}_2 \wedge \bar{g}_3) \vee (g_2 \wedge \bar{g}_3) \vee (g_1 \wedge g_3)$
g_2	g_1
g_3	g_1

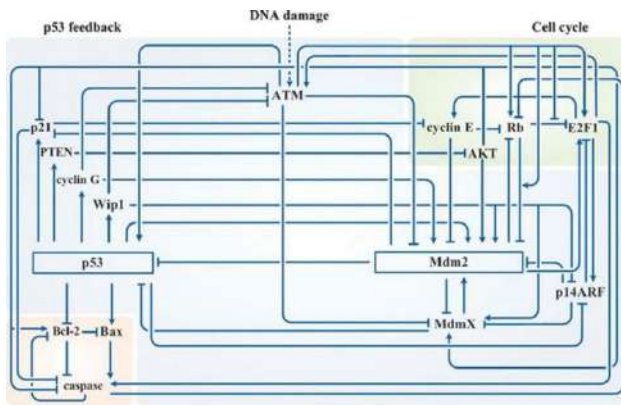
^{۱۸} Overestimation

^{۱۹} Node Removal

^{۲۰} Edge Removal

^{۲۱} No-Operation

^{۲۲} Phenotype



شکل ۴: شبکه تنظیم کننده ژن P53 [1].

۳- نتایج

از آنجا که محیط GRN استفاده شده در این تحقیق، شبکه پویای دودویی بر ساخته شده از زنجیره مارکوف است، لذا شرایط مارکوف در این مسئله همواره برقرار بوده و کارگزار یادگیر برای یافتن سیاست بهینه به تاریخچه حالت ها نیاز ندارد. به همین جهت، ساختار شبکه عصبی استفاده شده دارای ۱۶ ورودی و ۱۷ خروجی است. جزئیات بیشتری از ساختار شبکه در جدول ۲ قرار داده شده است.

جدول ۲: ساختار شبکه یادگیری تقویتی عمیق.

لایه	تعداد نورون
ورودی	۱۶
مخفی ۱	۱۲۸
مخفی ۲	۶۴
خروجی	۱۷

علاوه بر این، بهینه ساز آدام^{۲۶} با تابع خطای هوبر^{۲۷} در روند یادگیری استفاده شده است. نرخ تنزیل استفاده شده در این مسئله $\gamma = 0.99$ ، و نرخ یادگیری برابر با $\alpha = 10^{-4}$ است. به جهت بررسی مقاوم بودن مدل پیشنهادی، نرخ اغتشاش^{۲۸} با مقادیر $\rho_0 = 0$ ، $\rho_1 = 625 \times 10^{-5}$ ، $\rho_2 = 1250 \times 10^{-5}$ ، $\rho_3 = 1875 \times 10^{-5}$ در نظر گرفته شده است که احتمال اعمال نویز η_t را مشخص می نماید. منظور از نویز در یک شبکه دودویی GRN عبارت است از تغییر مقدار بیان یک ژن از صفر به یک یا بالعکس.

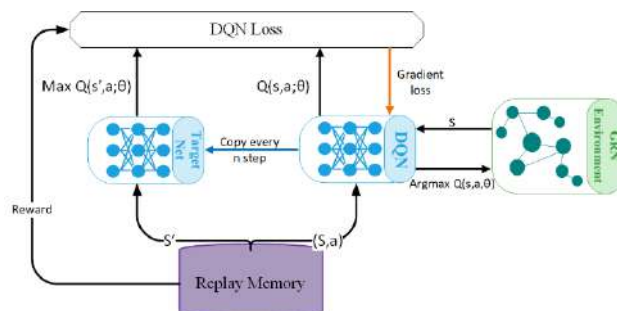
رهیافت پیشنهادی این مقاله با زبان برنامه نویسی پایتون ۳^{۲۹} توسعه داده شده که به دلیل دسترسی آن به کتابخانه های مرتبط با حوزه یادگیری ماشین و

ساختار شبکه دودویی تنظیم کننده ژن، عبارت ۲ را می توان به صورت زیر بازنویسی کرد:

$$S_{t+1} = f(S_t, i_t) \oplus \eta_t \quad (4)$$

که در آن i_t مداخله انجام گرفته در لحظه t است. در مدل های مبتنی بر یادگیری تقویتی برای GRN ها می توان تابع پاداش را به صورت زیر در نظر گرفت که در آن g_t ژن هدف مورد نظر در مسئله و r پاداشی با مقدار مثبت در لحظه t است:

$$R_t = \begin{cases} r & \text{if } r > 0 \\ 0 & \text{o.w.} \end{cases} \quad (5)$$



شکل ۳: نمای کلی DQN استفاده شده برای یادگیری سیاست درمانی.

۳-۲- شبکه تنظیم کننده ژن P53

سرطان یکی از بیماری های پیچیده و خطرناکی است که در فرآیند چرخه حیات، تولید مثل، و مرگ سلول های قدیمی یا آسیب دیده اختلال ایجاد کرده و باعث رشد غیرعادی سلول ها در بدن می شود. شبکه تنظیم کننده ژن P53 شامل ۱۶ گره با مجموع ۳۷۸ یال تحریک کننده و بازدارنده یکی از مهم ترین شبکه های سرکوب کننده سرطان است. تعداد کل حالت های این شبکه شامل $|S| = 2^{16}$ با توجه به دو مقدار صفر و یک برای هر گره آن خواهد بود. کارگزار این محیط دارای ۱۷ کنش که شامل ۱۶ دخالت حذف گره^{۲۳} و یک کنش بدون عملیات^{۲۴} است. شکل ۴ نمایی از این شبکه را نشان می دهد. این شبکه شامل جاذب های مختلفی نظیر جاذب رشد سلول، جاذب آپوپتوز، سیز، مهار سلول، و مرگ سلولی است. از آنجا که هدف درمانی در بیماری سرطان از بین بردن سلول سرطانی و جلوگیری از تولید مثل آن است، لذا ژن caspase به عنوان از بین برنده سلول سرطانی به عنوان هدف یادگیری تقویتی عمیق در نظر گرفته می شود. بنابراین، با توجه به تابع پاداش در رابطه ۵ کارگزار تقویتی عمیق به دنبال نگه داشتن ژن caspase در مقدار یک خواهد بود که به موجب آن سلول سرطانی کشته می شود. شبکه تنظیم کننده ژن P53 از آن جهت مورد اهمیت است که یکی از شبکه های محک^{۲۵} در مدل سازی و کنترل مورد استفاده قرار گرفته شده است [1].

^{۲۳} Node deletion Intervention

^{۲۴} No-operation

^{۲۵} Benchmark

^{۲۶} Adam Optimizer

^{۲۷} Huber Loss Function

^{۲۸} Perturbation Rate

^{۲۹} Python 3

جدول ۳: بیشترین پاداش جمع‌آوری شده توسط کارگزار در یک قسمت.

نرخ اغتشاش	بیشینه پاداش جمع‌آوری شده در یک قسمت	درصد ماندگاری در جاذب ژن هدف caspase
ρ_0	۱۹۹,۶	۹۹,۸٪
ρ_1	۱۹۵,۷	۹۷,۸۵٪
ρ_2	۱۹۳,۴	۹۶,۷٪
ρ_3	۱۹۰,۱	۹۵,۰۵٪

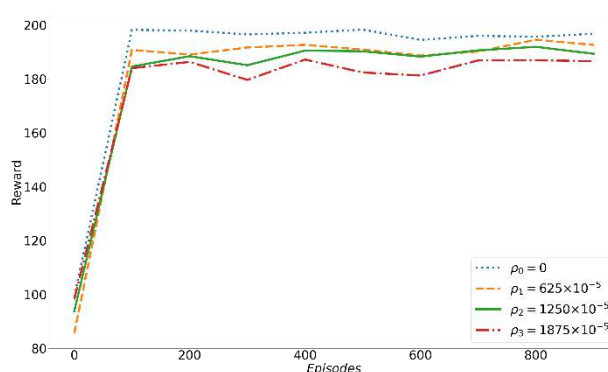
۴- نتیجه‌گیری

شبکه‌های تنظیم‌کننده ژن یکی از مدل‌های تبیین‌کننده وضعیت سلول در بدن انسان است. قرار گرفتن این شبکه درون حالت‌های مطلوبی که نشان از سلامتی است، بسیار مورد اهمیت است و می‌تواند از بیماری‌هایی نظیر سرطان جلوگیری نماید. همچنین تحلیل این شبکه برای یافتن سیاست‌های مداخله درمانی برای یافتن مسیری بهینه به حالت‌های هدف مهم خواهد بود. یادگیری تقویتی عمیق یکی از روش‌های هوشمند یافتن سیاست است که می‌تواند برای این کار مورد استفاده قرار گیرد. در این مقاله، کارایی یادگیری تقویتی **DQN** روی شبکه تنظیم‌کننده ژن **P53** سلول عادی مورد بررسی قرار گرفت. نتایج به دست آمده در این مقاله نشان می‌دهد که سیاست کارگزار سرعت همگرایی و مقاومت بالایی در مواجهه با اغتشاش برای سوق دادن شبکه به جاذب‌های مرگ سلول سرطانی دارد.

شبکه‌های عصبی مصنوعی نظیر پایتورچ^{۳۰} است. یکی از ویژگی‌های مثبت استفاده از کتابخانه پایتورچ، قابلیت اجرای آن روی واحدهای مرکزی پردازش^{۳۱} یا واحدهای گرافیکی پردازش^{۳۲} است. فضای حافظه **RAM** در این آزمایش ۱۶ گیگا بایت و از پردازنده گرافیکی **Nvidia GeForce GTX 980** برای یادگیری استفاده شده است.

نتایج به دست آمده در این بخش براساس میانگین پاداش دریافت شده توسط کارگزار در مرحله آزمون است. در این آزمایش با توجه به رابطه ۵، مقدار پاداش، زمانی که ژن هدف *caspase* فعال شود، برابر با $r = 0.1$ در نظر گرفته شده است. آزمون یک کارگزار شامل ۳۰ قسمت^{۳۳} است که به طور متناوب پس از هر ۱۰۰ قسمت یادگیری انجام پذیرفته است. همان‌طور که در شکل ۵ نشان داده شده است، کارگزار یادگیری تقویتی در محیط‌های با نرخ اغتشاش متفاوت توانسته است که در کمتر از ۱۰۰ قسمت به جواب مطلوب همگرا شود. از طرفی روند یادگیری در **DQN** بسیار ناپایدار است که برای حل این چالش از دو تکنیک حافظه بازپخش و شبکه هدف استفاده می‌شود. نتایج به دست آمده در شکل ۵ نشان می‌دهد که کارگزار توانسته است پایداری خود را حفظ نموده و به نتیجه نامطلوب و اگر نشود. از طرفی، افزایش نرخ اغتشاش در محیط پیرامونی محیط، رابطه‌ای معکوس با میانگین پاداش جمع‌آوری شده توسط کارگزار دارد؛ به این معنا که با افزایش نرخ اغتشاش، مقدار پاداش دریافت شده کمتر خواهد شد. همچنین، بیشینه پاداش دریافت شده توسط کارگزار برای نرخ اغتشاش ρ_0, ρ_1, ρ_2 و ρ_3 به ترتیب برابر با ۱۹۹,۶، ۱۹۵,۷، ۱۹۳,۴ و ۱۹۰,۱ است.

همان‌طور که در جدول ۳ نشان داده شده، درصد ماندگاری در ژن هدف *caspase* به ترتیب برابر ۹۹,۸٪، ۹۷,۸۵٪، ۹۶,۷٪ و ۹۵,۰۵٪ است که نشان از قرار گرفتن شبکه در حالت‌هایی است که جاذب‌های آپوپتوسیز و مرگ سلولی را فعال می‌کند. پروفایل زمانی بیان ژن^{۳۴} آورده شده در شکل ۶ بیانگر همین امر بوده و نشان می‌دهد که کنش‌های از نوع حذف گر توسط کارگزار، شبکه دودویی تنظیم‌کننده **P53** به سرعت به سمت حالت‌های با *caspase* فعال حرکت کرده و در آن باقی بماند.



شکل ۵: میانگین پاداش جمع‌آوری شده توسط کارگزار.

^{۳۰} Pytorch

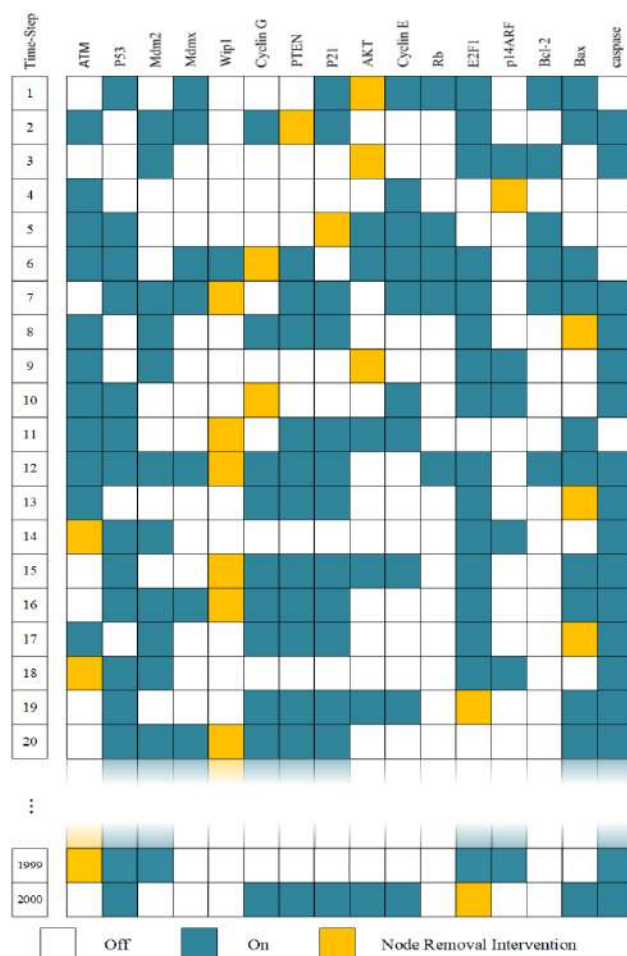
^{۳۱} Central Processing Unit (CPU)

^{۳۲} Graphics Processing Unit (GPU)

^{۳۳} Episode

^{۳۴} Temporal Gene Expression Profile

- [6] Hasselt, H. van, Guez, A. and Silver, D. 2016. Deep Reinforcement Learning with Double Q-Learning. *Proceedings of the AAAI Conference on Artificial Intelligence*. 30, 1 (Mar. 2016).
- [7] Hayama Nishida, C.E., Costa Bianchi, R.A. and Reali Costa, A.H. 2020. A framework to shift basins of attraction of gene regulatory networks through batch reinforcement learning. *Artificial Intelligence in Medicine*. 107, (Jul. 2020). DOI:https://doi.org/10.1016/J.ARTMED.2020.101853.
- [8] Hessel, M., Modayil, J., Van Hasselt, H., Schaul, T., Ostrovski, G., Dabney, W., Horgan, D., Piot, B., Azar, M. and Deepmind, D.S. 2018. Rainbow: Combining Improvements in Deep Reinforcement Learning. *Thirty-Second AAAI Conference on Artificial Intelligence*. (Apr. 2018).
- [9] Imani, M. and Braga-Neto, U.M. 2019. Control of gene regulatory networks using bayesian inverse reinforcement learning. *IEEE/ACM Transactions on Computational Biology and Bioinformatics*. 16, 4 (Jul. 2019), 1250–1261. DOI:https://doi.org/10.1109/TCBB.2018.2830357.
- [10] Mnih, V. et al. 2015. Human-level control through deep reinforcement learning. *Nature* 2015 518:7540. 518, 7540 (Feb. 2015), 529–533. DOI:https://doi.org/10.1038/nature14236.
- [11] Mnih, V., Kavukcuoglu, K., Silver, D., Graves, A., Antonoglou, I., Wierstra, D. and Riedmiller, M. 2013. Playing Atari with Deep Reinforcement Learning. (Dec. 2013).
- [12] Papagiannis, G. and Moschogiannis, S. 2019. Learning to Control Random Boolean Networks: A Deep Reinforcement Learning Approach. *Complex Networks 2019*. 881 SCI, (Jan. 2019), 721–734. DOI:https://doi.org/10.1007/978-3-030-36687-2_60.
- [13] Schaul, T., Quan, J., Antonoglou, I. and Silver, D. 2015. Prioritized Experience Replay. *4th International Conference on Learning Representations, ICLR 2016 - Conference Track Proceedings*. (Nov. 2015).
- [14] Sutton, R.S. and Barto, A.G. 2018. Reinforcement Learning, Second Edition: An Introduction - Complete Draft. *The MIT Press*. (2018), 1–3.
- [15] Wang, Z., Schaul, T., Hessel, M., Hasselt, H., Lanctot, M. and Freitas, N. 2016. Dueling Network Architectures for Deep Reinforcement Learning. PMLR.
- [16] Yang, T. and Sun, Y.F. 2011. The reconstruction of gene regulatory network based on multi-agent system by fusing multiple data sources. *Proceedings - 2011 IEEE International Conference on Computer Science and Automation Engineering, CSAE 2011*. 2, (2011), 126–130. DOI:https://doi.org/10.1109/CSAE.2011.5952438.



شکل ۶: پروفاایل زمانی بیان ژن کنترل شده توسط DQN در شبکه دودویی تنظیم کننده ژن P53.

مراجع

- [1] Choi, M., Shi, J., Jung, S.H., Chen, X. and Cho, K.H. 2012. Attractor landscape analysis reveals feedback loops in the p53 network that control the cellular response to DNA damage. *Science signaling*. 5, 251 (Nov. 2012). DOI:https://doi.org/10.1126/SCISIGNAL.2003363.
- [2] Dougherty, E.R., Pal, R., Qian, X., Bittner, M.L. and Datta, A. 2010. Stationary and structural control in gene regulatory networks: basic concepts. *https://doi.org/10.1080/00207720903144560*. 41, 1 (Jan. 2010), 5–16. DOI:https://doi.org/10.1080/00207720903144560.
- [3] Fortunato, M., Azar, M.G., Piot, B., Menick, J., Hessel, M., Osband, I., Graves, A., Mnih, V., Munos, R., Hassabis, D., Pietquin, O., Blundell, C. and Legg, S. 2017. Noisy Networks for Exploration. *6th International Conference on Learning Representations, ICLR 2018 - Conference Track Proceedings*. (Jun. 2017).
- [4] Hajiramezanali, E., Imani, M., Braga-Neto, U., Qian, X. and Dougherty, E.R. 2019. Scalable optimal Bayesian classification of single-cell trajectories under regulatory model uncertainty. *BMC Genomics*. 20, 6 (Jun. 2019), 1–11. DOI:https://doi.org/10.1186/S12864-019-5720-3/TABLES/2.
- [5] Van Hasselt, H. 2010. Double Q-learning. *Advances in*

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

مروری بر انواع اعداد فازی و دسته بندی آنها با رویکرد رتبه بندی

سیما سرگزی^۱، حسن میش مست نهی^۲

^۱دانشکده ریاضی و آمار، دانشگاه سیستان و بلوچستان، زاهدان، Sima.sargazi2020@gmail.com

^۲دانشکده ریاضی و آمار، دانشگاه سیستان و بلوچستان، زاهدان، hmnehi@hamoon.usb.ac.ir

چکیده: در این مقاله، پس از معرفی انواع اعداد فازی شناخته شده تاکنون، با توجه به ماهیت ساختاری و توابع عضویت وابسته اقدام به دسته بندی آنها نموده و سپس در راستای امکان مقایسه، رتبه بندی و استفاده کاربردی، به معرفی تعدادی از روش های رتبه بندی هر یک از انواع مذکور می پردازیم.

کلمات کلیدی: مجموعه های فازی، اعداد فازی، توابع رتبه بندی

استدلال تقریبی، سیستم های اجتماعی- اقتصادی و ... کاربردهای بسیاری دارند. از سال ۱۹۷۶ تاکنون مدل های زیادی در زمینه رتبه بندی اعداد فازی ارائه شده که بسیاری از این مدل ها توسط بورتن و دگانی [15] در سال ۱۹۸۵، چن و هانگ [16] در سال ۱۹۹۲ و وانگ و کر [17] در سال ۲۰۰۱ مورد مقایسه و بررسی قرار گرفته اند تا در نهایت استانداردهای منطقی و معقولی ارائه نمودند که توانسته مبنای تعریف مدل های رتبه بندی اعداد فازی قرار گیرند. این مقاله در بخش اول به معرفی و دسته بندی اعداد فازی با توجه به ماهیت مدل و ساختار هر نوع از اعداد فازی میپردازد و در بخش بعدی تعدادی از مدل های رتبه بندی جدید و کاربردی را ارائه میدهد. با توجه به محدودیت حجم نوشتاری مقاله از ارائه مدل ریاضی عملگرهای حسابی انواع مختلف اعداد فازی معذوریم، لذا جهت آشنایی با مدل های رتبه بندی بیشتر و عملگرهای حسابی مربوطه نگاهی اجمالی به منابع قید شده در انتهای مقاله مفید خواهد بود.

۲- مفاهیم و تعارف مقدماتی

تعریف ۱-۲: فرض کنید X نشان دهنده مجموعه مرجع باشد. در این صورت مجموعه فازی A تعریف شده بر X ، مجموعه ای از جفت مرتب های $A = \{(x, \mu_A(x)) | x \in X\}$ است که در آن $\mu_A(x)$ درجه عضویت x می باشد. $\mu_A(x)$ در بازه صفر و یک قرار دارد، صفر نشان دهنده عدم عضویت و یک نشان دهنده عضویت قطعی در A می باشد. در حالتی که مجموعه مرجع X پیوسته باشد، مجموعه فازی A بصورت

۱- مقدمه

مبنای ریاضیات کلاسیک منطق ارسطویی است. در این منطق پدیده های مختلف ارزش "درست یا نادرست" یا عبارتی صفر و یک دارند. اما در دنیای واقعی با پدیده های مختلفی روبرو می شویم که دارای ابهام ارزشی و عبارتی عدم قطعیت صفر و یک می باشند. برای گریز از جزمیت منطق دو ارزشی، نخستین بار منطق چند ارزشی در دهه ۱۹۳۰ توسط لوکا سیویچ پایه گذاری شد. در منطق چند ارزشی، ارزش هر مقدار عددی بین ۰ و ۱ است. بعدها در سال ۱۹۶۵ پروفیسور لطفی زاده مفهومی تحت عنوان فازی را مطرح نمود که جایگزین مناسبی برای مفاهیم ابهام یا چند ارزشی قلمداد گردید. منطق فازی نیز یک منطق چند ارزشی است که طیفی بین اعداد ۰ و ۱ را دربر میگیرد. لطفی زاده در اولین مقاله خود برای تشریح مفهوم فازی از مثال قد انسان استفاده کرد و در این مثال وی از مفهوم "منحنی عضویت" استفاده نمود. این منحنی برای هر اندازه قد، درجه عضویتی را مشخص میکند. لطفی زاده با این تئوری عدم قطعیت ناشی از ابهامات تفکرات و دانش بشری را بیان نمود. اصلی ترین حسن این تئوری توانایی ارائه داده هایی است که غیر قطعی هستند، همچنین این روش قادر به بکارگیری عملگرهای ریاضی و امکان مقایسه در حوزه داده های فازی می باشد. در مواردی که دلیل پیچیدگی سیستم نمی توان با دقت و صراحت در مورد پارامترها و مشخصه های آن قضاوت کرد، لذا رتبه بندی اعداد فازی در این مواقع به دلیل ماهیت غیر قطعی این اعداد اهمیت و کاربردهای فراوانی خواهد داشت. رتبه بندی اعداد فازی در مسائلی چون تجزیه و تحلیل داده ها، بهینه سازی، تصمیم گیری،

$\int_{x \in X} \frac{\mu_A(x)}{x}$ و اگر گسسته باشد مجموعه فازی A بصورت $\sum_{x \in X} \frac{\mu_A(x)}{x}$ نشان داده می شود. منظور از نمادهای $\sum$ و $\int$ همان اجتماع است.

تعریف ۲-۲: مجموعه فازی A تعریف شده بر مجموعه اعداد حقیقی $\square$ را عدد فازی نامند، هرگاه A نرمال، محدب و نیز $\mu_A(x)$ پیوسته باشد.

۳- دسته بندی اعداد فازی

۳-۱- دسته اول: اعداد فازی نوع-۱ شامل (اعداد فازی مثلثی، اعداد فازی دوزنقه ای، اعداد فازی تعمیم یافته و اعداد فازی چند ضلعی)

تعریف ۳-۱-۱: عدد فازی تعریف شده بر بازه حقیقی بسته $[a_1, a_4]$ را عدد فازی دوزنقه ای کلی نامند و با نماد $A = (a_1, a_2, a_3, a_4; h_1, h_2)$ که در آن $h_1, h_2 \in [0, 1]$ نمایش می دهند، هرگاه تابع عضویت آن به فرم زیر باشد:

$$\mu_A(x) = \begin{cases} h_2 \left(\frac{x-a_1}{a_2-a_1} \right), & a_1 \leq x \leq a_2 \\ (h_1-h_2) \frac{x-a_2}{a_3-a_2} + h_2, & a_2 \leq x \leq a_3 \\ h_1 \left(\frac{a_4-x}{a_4-a_3} \right), & a_3 \leq x \leq a_4 \\ 0, & x \leq a_1, x \geq a_4 \end{cases}$$

تعریف ۳-۱-۲: اگر در تعریف بالا $a_2 = a_3$ در این صورت عدد فازی دوزنقه ای، یک عدد فازی مثلثی کلی نامیده می شود که بر بازه حقیقی بسته $[a_1, a_3]$ تعریف می شود و آنرا با نماد اختصاری $A = (a_1, a_2, a_3; h)$ نمایش می دهیم که در آن $h \in [0, 1]$ و تابع عضویت آن نیز به صورت زیر می باشد:

$$\mu_A(x) = \begin{cases} h \left(\frac{x-a_1}{a_2-a_1} \right), & a_1 \leq x \leq a_2 \\ h \left(\frac{a_3-x}{a_3-a_2} \right), & a_2 \leq x \leq a_3 \\ 0, & x \leq a_1, x \geq a_3 \end{cases}$$

انواع دیگر اعداد فازی، نظیر اعداد فازی مثلثی، دوزنقه ای و اعداد فازی چند ضلعی هستند که به دلیل مشابهت با اعداد فازی تعریف شده و محدودیت در حجم نوشتاری مقاله از قید آن در اینجا معذوریم، جهت آشنایی با این اعداد فازی به ترتیب به منابع [3,8] مراجعه نمایید.

۲-۳- دسته دوم: اعداد فازی نوع-۲ و اعداد فازی نوع-۲ بازه ای

تعریف ۳-۲-۱: به مجموعه ای فازی که درجه عضویت هر عنصر متعلق به مجموعه مرجع معرف آن، یک مجموعه فازی نوع-۱ باشد، مجموعه فازی نوع-۲ گویند. بر خلاف مجموعه فازی نوع-۱ که درجه عضویت هر عنصر آن عددی حقیقی در بازه $[0, 1]$ است، هر عنصر در مجموعه فازی نوع-۲ توسط یک تابع عضویت فازی مشخص می شود که درجه عضویت هر عنصر از این مجموعه را نشان می دهد. مجموعه فازی نوع-۲، $\tilde{A}$ توسط تابع و تابع عضویت زیر روی مجموعه مرجع X تعریف می شوند:

$$\tilde{A}: X \rightarrow F(U)$$

$$\mu_{\tilde{A}}: X \times U \rightarrow V$$

که در آن $U = [0, 1]$ دامنه عضویت، X دامنه اولیه و $V = [0, 1]$ عضویت ثانویه می باشد. مجموعه فازی نوع-۲، $\tilde{A}$ بوسیله تابع عضویت نوع-۲ مشخص می شود و می توان آن را به صورت زیر نمایش داد:

$$\tilde{A} = \{((x, u), \mu_{\tilde{A}}(x, u)) \mid \forall x \in X, \forall u \in J_x \subseteq [0, 1]\}$$

که در آن $0 \leq \mu_{\tilde{A}}(x, u) \leq 1$ و X دامنه $\tilde{A}$ می باشد و هر مقدار u عضوی از $U = [0, 1]$ است که به آن درجه اولیه گویند. هر درجه اولیه، درجه عضویت $\mu_{\tilde{A}}(x, u)$ را تعیین می کند که به آن درجه عضویت ثانویه گویند و در بازه $V = [0, 1]$ قرار دارد. همچنین J_x مجموعه عضویت های اولیه نظیر x می باشد. مجموعه فازی نوع-۲، $\tilde{A}$ با مجموعه مرجع پیوسته به صورت زیر نمایش داده می شود:

$$\tilde{A} = \int_{x \in X} \int_{u \in J_x} \frac{\mu_{\tilde{A}}(x, u)}{(x, u)} \quad \text{و} \quad J_x \subseteq [0, 1]$$

برای مجموعه مرجع گسسته، بجای نماد $\int$ از نماد $\sum$ استفاده می شود.

تعریف ۳-۲-۲: یک مجموعه پیوسته فازی نوع-۲ تعریف شده بر مجموعه اعداد حقیقی $\square$ را عدد فازی نوع-۲ گویند هرگاه:

۱- جای پای عدم قطعیت (FOU) و مجموعه اصلی (توابع عضویت بالایی و پایینی تشکیل دهنده آنها) و تمام برش های عمودی آن محدب باشند.
۲- تابع عضویت بالایی تشکیل دهنده جای پای عدم قطعیت (FOU) و مجموعه اصلی آن، نرمال باشند.

برای آشنایی با مفاهیم ذکر شده به منبع [2] مراجعه نمایید.

با توجه به مشکلات و پیچیدگی های محاسباتی که در پردازش فرم کلی مجموعه های فازی نوع-۲ وجود دارد، امروزه نوع خاصی از آنها که تعمیمی بر مجموعه های فازی نوع-۱ نیز می باشند، بنام مجموعه های فازی نوع-۲ بازه ای معرفی گردیده اند. در این مجموعه ها به هر عنصر از مجموعه مرجع

معرف آن یک بازه نسبت داده می شود، بنابراین هر مجموعه فازی نوع-۲ بازه ای تحت عنوان مجموعه فازی بازه ای مقدار نیز شناخته می شود.

تعریف ۳-۲-۳: فرض کنید $J_x \subset [0,1]$ و

$$\tilde{A} = \int_{x \in X} \frac{\mu_{\tilde{A}}(x)}{(x)} = \int_{x \in X} \int_{u \in J_x} \frac{\mu_{\tilde{A}}(x, u)}{x}$$

نوع-۲ است که در آن برای همه درجات ثانویه، رابطه $\mu_{\tilde{A}}(x, u) = 1$ برقرار می باشد، در این صورت آن را یک مجموعه فازی نوع-۲ بازه ای پیوسته گویند اگر بر مجموعه مرجع پیوسته تعریف شود و نیز به آن مجموعه فازی نوع-۲ بازه ای گسسته گفته می شود هرگاه بر مجموعه مرجع گسسته تعریف شود و به صورت زیر نمایش داده می شود:

$$\tilde{A} = \int_{x \in X} \frac{\mu_{\tilde{A}}(x)}{(x)} = \frac{\int_{x \in X} \left[\int_{u \in J_x} \frac{1}{u} \right]}{x}$$

که $J_x \in U \subset [0,1]$ ، به طور معادل مجموعه فازی نوع-۲ بازه

ای $\tilde{A}$ به کمک رابطه زیر نیز تعریف می شود:

$$\{(x, \tilde{A}(x)) | x \in X, \tilde{A}(x) \subset U\} = \{(x, [\underline{u}_x, \bar{u}_x]) | x \in X, [\underline{u}_x, \bar{u}_x] \subset U\}$$

تعریف ۴-۲-۳: یک مجموعه پیوسته فازی نوع-۲ بازه ای تعریف

شده بر مجموعه اعداد حقیقی $\square$ را عدد فازی نوع-۲ بازه ای گویند هرگاه:

۱- جای پای عدم قطعیت (FOU) و مجموعه اصلی (توابع عضویت بالایی و پایینی تشکیل دهنده آن ها) و تمام برش های عمودی آن محدب باشند.

۲- تابع عضویت بالایی تشکیل دهنده جای پای عدم قطعیت (FOU) و مجموعه اصلی آن ، نرمال باشد.

جهت آشنایی با مفاهیم و تعاریف بیشتر در خصوص اعداد فازی دسته دوم به منبع [2] مراجعه نمایید.

۳-۳-۳ دسته سوم: اعداد فازی نوتروزوفیک شامل اعداد فازی

نوتروزوفیک مثلثی و دوزنقه ای

تعریف ۱-۳-۳: فرض کنید X یک مجموعه مرجع باشد و همچنین x

عضوی از مجموعه X باشد، یک مجموعه نوتروزوفیک روی X با تابع

عضویت $T_A(x)$ و تابع خنثی $I_A(x)$ و تابع عدم عضویت $F_A(x)$

نمایش داده می شود. جائیکه $T_A(x)$ ، $I_A(x)$ ، $F_A(x)$ ، زیر

مجموعه های حقیقی استاندارد یا غیر استاندارد از $\left[0,1\right]^+$ می باشند.

یعنی، $I_A(x): X \rightarrow \left[0,1\right]^+$ ، $T_A(x): X \rightarrow \left[0,1\right]^+$

و $F_A(x): X \rightarrow \left[0,1\right]^+$ (بطوریکه در عدد متناهی غیر استاندارد

$\varepsilon = 1 + 1^+$ ، بخش استاندارد نامیده می شود و $\varepsilon > 0$ یک اعشاری نامتناهی است که بخش غیر استاندارد نامیده می شود)

و $0 \leq \sup T_A(x) + \sup I_A(x) + \sup F_A(x) \leq 3^+$

لذا مجموعه فازی نوتروزوفیک را به صورت زیر نمایش میدهند:

$$\tilde{N} = \left\{ \langle x, T_{\tilde{N}}(x), I_{\tilde{N}}(x), F_{\tilde{N}}(x) \rangle | x \in X \right\}.$$

تعریف ۳-۳-۲: اگر داشته باشیم:

$$T_{\tilde{N}}(x) \subset [0,1], I_{\tilde{N}}(x) \subset [0,1], F_{\tilde{N}}(x) \subset [0,1]$$

$$T_{\tilde{N}}(x) = \{t_{\tilde{N}}^1(x), t_{\tilde{N}}^2(x), t_{\tilde{N}}^3(x), t_{\tilde{N}}^4(x)\}: X \rightarrow [0,1]$$

$$I_{\tilde{N}}(x) = \{i_{\tilde{N}}^1(x), i_{\tilde{N}}^2(x), i_{\tilde{N}}^3(x), i_{\tilde{N}}^4(x)\}: X \rightarrow [0,1]$$

$$F_{\tilde{N}}(x) = \{f_{\tilde{N}}^1(x), f_{\tilde{N}}^2(x), f_{\tilde{N}}^3(x), f_{\tilde{N}}^4(x)\}: X \rightarrow [0,1]$$

$$0 \leq t_{\tilde{N}}^4(x) + i_{\tilde{N}}^4(x) + f_{\tilde{N}}^4(x) \leq 3, x \in X$$

$$T_{\tilde{N}}(x) = \begin{cases} \frac{x-a_1}{a_2-a_1}, a_1 \leq x \leq a_2 \\ 1, a_2 \leq x \leq a_3 \\ \frac{a_4-x}{a_4-a_3}, a_3 \leq x \leq a_4 \\ 0, otherwise \end{cases}.$$

$$I_{\tilde{N}}(x) = \begin{cases} \frac{b_2-x}{b_2-b_1}, b_1 \leq x \leq b_2 \\ 0, b_2 \leq x \leq b_3 \\ \frac{x-b_3}{b_4-b_3}, b_3 \leq x \leq b_4 \\ 1, otherwise \end{cases}.$$

$$F_{\tilde{N}}(x) = \begin{cases} \frac{c_2-x}{c_2-c_1}, c_1 \leq x \leq c_2 \\ 0, c_2 \leq x \leq c_3 \\ \frac{x-c_3}{c_4-c_3}, c_3 \leq x \leq c_4 \\ 1, otherwise \end{cases}.$$

و اگر اعداد فازی فوق، اعداد فازی دوزنقه ای باشند آنگاه:

$$\tilde{n} = \langle (a_1, a_2, a_3, a_4), (b_1, b_2, b_3, b_4), (c_1, c_2, c_3, c_4) \rangle,$$

یک عدد فازی دوزنقه ای نوتروزوفیک است. بطوریکه:

$$a_1 \leq a_2 \leq a_3 \leq a_4, b_1 \leq b_2 \leq b_3 \leq b_4, c_1 \leq c_2 \leq c_3 \leq c_4.$$

تعریف ۳-۳-۳: با توجه به روابط بالا هرگاه:

$\pi_A(z) = 1 - \mu_A(z) - \eta_A(z) - \nu_A(z)$ درجه عضویت تردید z در A است [11].

تعریف ۳-۴-۴: یک عدد فازی تصویری $\tilde{n}$ در مجموعه اعداد حقیقی $\square$ به صورت یک (PFN) و به شکل زیر تعریف می شود:

$$\tilde{n} = \{ (z, \mu_{\tilde{n}}(z), \eta_{\tilde{n}}(z), \nu_{\tilde{n}}(z)) | z \in \square \}$$

جائیکه:

$$\mu_{\tilde{n}}(z) = \begin{cases} f_{\tilde{n}}^L(z), p_1 \leq z \leq q \\ \alpha_{\tilde{n}}, q \leq z \leq r \\ f_{\tilde{n}}^R(z), r \leq z \leq s_1 \\ 0, \text{otherwise.} \end{cases}$$

$$\nu_{\tilde{n}}(z) = \begin{cases} h_{\tilde{n}}^L(z), p_3 \leq z \leq q \\ \gamma_{\tilde{n}}, q \leq z \leq r \\ h_{\tilde{n}}^R(z), r \leq z \leq s_3 \\ 0, \text{otherwise.} \end{cases}$$

$$\eta_{\tilde{n}}(z) = \begin{cases} g_{\tilde{n}}^L(z), p_2 \leq z \leq q \\ \beta_{\tilde{n}}, q \leq z \leq r \\ g_{\tilde{n}}^R(z), r \leq z \leq s_2 \\ 0, \text{otherwise.} \end{cases}$$

تعریف ۳-۴-۵: یک $TPFN$ یک عدد فازی تصویری دوزنقه ای است در صورتیکه:

$$f_{\tilde{n}}^L(z) = \frac{z-p_1}{q-p_1} \alpha_{\tilde{n}}^q, f_{\tilde{n}}^R(z) = \frac{s_1-z}{s_1-r} \alpha_{\tilde{n}} \\ g_{\tilde{n}}^L(z) = \frac{q-z+\beta_{\tilde{n}}(z-p_2)}{q-p_2}, g_{\tilde{n}}^R(z) = \frac{z-r+\beta_{\tilde{n}}(s_2-z)}{s_2-r} \\ h_{\tilde{n}}^L(z) = \frac{q-z+\gamma_{\tilde{n}}(z-p_3)}{q-p_3}, h_{\tilde{n}}^R(z) = \frac{z-r+\gamma_{\tilde{n}}(s_3-z)}{s_3-r}$$

و برای $0 \leq \alpha_{\tilde{n}}, \beta_{\tilde{n}}, \gamma_{\tilde{n}} \leq 1$ رابطه $\alpha_{\tilde{n}} + \beta_{\tilde{n}} + \gamma_{\tilde{n}} \leq 1$ برقرار باشد، عدد فازی تصویری دوزنقه ای $\tilde{n}$ را به صورت زیر نمایش می دهند:

$$\tilde{n} = \langle ([p_1, q, r_1, s_1]; \alpha_{\tilde{n}}), ([p_2, q, r_2, s_2]; \beta_{\tilde{n}}), ([p_3, q, r_3, s_3]; \gamma_{\tilde{n}}) \rangle.$$

تعریف ۳-۴-۶: یک عدد فازی تصویری دوزنقه ای هنگامیکه $q = r$ به یک عدد فازی تصویری مثلثی به صورت زیر تبدیل می شود:

$$\tilde{n} = \langle [p_1, q, r_1]; \alpha_{\tilde{n}}, [p_2, q, r_2]; \beta_{\tilde{n}}, [p_3, q, r_3]; \gamma_{\tilde{n}} \rangle.$$

$$\tilde{T}_A(x) = \{T_A^1(x), T_A^2(x), T_A^3(x)\}$$

$$\tilde{I}_A(x) = \{I_A^1(x), I_A^2(x), I_A^3(x)\}$$

$$\tilde{F}_A(x) = \{F_A^1(x), F_A^2(x), F_A^3(x)\}$$

$$\tilde{A} = \langle (a, b, c), (e, f, g), (r, s, t) \rangle \text{ و}$$

$$(T_A^1(x), T_A^2(x), T_A^3(x)) = (a, b, c)$$

$$(I_A^1(x), I_A^2(x), I_A^3(x)) = (e, f, g)$$

$$(F_A^1(x), F_A^2(x), F_A^3(x)) = (r, s, t)$$

آنگاه $\tilde{A}$ یک مقدار فازی مثلثی نوتروزوفیک $(TFNV)$ است [13].

۳-۴-۴ دسته چهارم: شامل اعداد فازی شهودی، اعداد فازی فیثاغورثی و فیثاغورثی بازه ای و اعداد فازی تصویری شامل اعداد فازی تصویر مثلثی و تصویر دوزنقه ای

تعریف ۳-۴-۱: فرض کنید X یک مجموعه مرجع باشد. مجموعه

فازی شهودی $\tilde{A}^I (IFS)$ در X توسط سه تایی مرتب

$$\tilde{A}^I = \{ \langle x, \mu_{\tilde{A}^I}(x), \nu_{\tilde{A}^I}(x) \rangle | x \in X \}$$

جائیکه توابع $\mu_{\tilde{A}^I}(x): X \rightarrow [0, 1]$ و $\nu_{\tilde{A}^I}: X \rightarrow [0, 1]$

به ترتیب درجه عضویت و عدم عضویت x در $\tilde{A}^I$ می باشند [10]. بطوریکه

$$\text{برای هر } x \in X, 0 \leq \mu_{\tilde{A}^I}(x) + \nu_{\tilde{A}^I}(x) \leq 1 \text{ و درجه تردید}$$

$$\pi_{\tilde{A}^I}(x) = 1 - \mu_{\tilde{A}^I}(x) - \nu_{\tilde{A}^I}(x) \text{ می باشد.}$$

تعریف ۳-۴-۲: مجموعه فازی شهودی در مجموعه اعداد حقیقی $\square$ یک عدد فازی شهودی است اگر:

$$1- \tilde{A}^I, \text{ فازی شهودی نرمال باشد.}$$

$$2- \tilde{A}^I, \text{ فازی شهودی محدب باشد.}$$

$$3- \mu_{\tilde{A}^I}, \text{ نیم پیوسته بالایی و } \nu_{\tilde{A}^I} \text{ نیم پیوسته پایینی باشد.}$$

$$4- \tilde{A}^I = \{ x \in X | \nu_{\tilde{A}^I}(x) < 1 \} \text{ کراندار باشد.}$$

تعریف ۳-۴-۳: یک مجموعه فازی تصویری A روی مجموعه مرجع X به صورت زیر تعریف شده است:

$$A = \{ \langle z, \mu_A(z), \eta_A(z), \nu_A(z) | z \in X \rangle \}$$

$$\mu_A(z) \text{ و } \eta_A(z) \text{ و } \nu_A(z) \text{ عضو } [0, 1], \text{ به ترتیب توابع}$$

عضویت مثبت، خنثی و منفی عنصر z در A هستند، به طوریکه:

$$0 \leq \mu_A(z) + \eta_A(z) + \nu_A(z) \leq 1 \text{ برای هر } z \in X$$

۴ Picture Fuzzy Number (PFN)

۵ Triangular Picture Fuzzy Number (TPFN)

۲ Triangular Fuzzy Neutrosophic Value (TFNV)

۳ Intuitionistic Fuzzy Set (IFS)

استفاده نتیجه شود. در این بخش با توجه به محدودیت نوشتاری، تنها چند نمونه از توابع رتبه بندی برای اعداد فازی را معرفی می کنیم. برای آشنایی با توابع بیشتر می توانید به منابع ذکر شده در بخش مراجع مراجعه نمایید.

۱-۴-رتبه بندی اعداد فازی نوع-۱

۱-۱-۴-توابع رتبه بندی اعداد فازی ذوزنقه ای و مثلثی

در کتاب مجموعه های فازی ساکاو [8] یکسری روش های مرسوم نظیر توابع رتبه بندی روبنز، چانگ، شاخص آدامو، شاخص یاگر و شاخص کانپوس و مانوز و روش مرکز ثقل، برای رتبه بندی اعداد فازی مثلثی و ذوزنقه ای معرفی گردیده اند. جهت مشاهده توابع رتبه بندی بیشتر به [7-6-15-1] مراجعه نمایید. ما نیز در اینجا تابع رتبه بندی مقدار انتگرال کلی [5] را جهت رتبه بندی اعداد فازی مثلثی و ذوزنقه ای بصورت زیر معرفی می نماییم:

$$I_T^\alpha(\tilde{A}) = \alpha I(\tilde{A})^R + (1-\alpha)I(\tilde{A})^L = \frac{1}{2}[\alpha a_3 + a_2 + (1-\alpha)a_1]$$

$$I(\tilde{A})^L = \int_0^1 (\mu_{\tilde{A}}(x)^L)^{-1} dy = \frac{1}{2}(a_1 + a_2), I(\tilde{A})^R = \int_0^1 (\mu_{\tilde{A}}(x)^R)^{-1} dy = \frac{1}{2}(a_2 + a_3)$$

۲-۴-رتبه بندی اعداد فازی نوع-۲ و نوع-۲ بازه ای

شاخص زیر که مبتنی بر نقطه مرکزی می باشد برای رتبه بندی اعداد فازی نوع-۲ و نوع-۲ بازه ای پر کاربرد است [14].

$$C(\tilde{A}) = \frac{c_l + c_r}{2}$$

$$c_l = \min_l \text{centroid}(A_e(L))$$

$$c_r = \max_r \text{centroid}(A_e(R))$$

$$\text{centroid}(A_e(L)) = \frac{\sum_{i=1}^L x^i \bar{\mu}_{\tilde{A}}(x_i) + \sum_{i=L+1}^N x^i \underline{\mu}_{\tilde{A}}(x_i)}{\sum_{i=1}^L \bar{\mu}_{\tilde{A}}(x_i) + \sum_{i=L+1}^N \underline{\mu}_{\tilde{A}}(x_i)}$$

$$\text{centroid}(A_e(R)) = \frac{\sum_{i=1}^R x^i \underline{\mu}_{\tilde{A}}(x_i) + \sum_{i=R+1}^N x^i \bar{\mu}_{\tilde{A}}(x_i)}{\sum_{i=1}^R \underline{\mu}_{\tilde{A}}(x_i) + \sum_{i=R+1}^N \bar{\mu}_{\tilde{A}}(x_i)}$$

جهت آشنایی با شاخصهای رتبه بندی بیشتر می توانید به [2] مراجعه نمایید.

۳-۴-توابع رتبه بندی اعداد نوتروزوفیک

$$S(\tilde{n}) = \frac{1}{3} \left(2 + \frac{a_1 + a_2 + a_3 + a_4}{4} - \frac{b_1 + b_2 + b_3 + b_4}{4} - \frac{c_1 + c_2 + c_3 + c_4}{4} \right)$$

$$\text{if } S(\tilde{n}_1) > S(\tilde{n}_2), \text{ then } \tilde{n}_1 \succ \tilde{n}_2$$

$$\text{if } S(\tilde{n}_1) < S(\tilde{n}_2), \text{ then } \tilde{n}_1 \prec \tilde{n}_2$$

$$\text{if } S(\tilde{n}_1) = S(\tilde{n}_2), \text{ then}$$

$$\text{defined, } H(\tilde{n}) = \frac{a_1 + a_2 + a_3 + a_4}{4} - \frac{c_1 + c_2 + c_3 + c_4}{4}$$

تعریف ۳-۴-۷: مجموعه فازی فیثاغورثی $\tilde{A}^P$ (PFS) در مجموعه

$$\tilde{A}^P = \left\{ \langle x, \mu_{\tilde{A}^P}(x), \nu_{\tilde{A}^P}(x) \mid x \in X \rangle \right\}$$

مرجع X بوسیله

$$\mu_{\tilde{A}^P}(x): X \rightarrow [0,1], \nu_{\tilde{A}^P}(x): X \rightarrow [0,1]$$

تعریف شده است. جائیکه توابع عضویت

$$[\mu_{\tilde{A}^P}(x)]^2 + [\nu_{\tilde{A}^P}(x)]^2 \leq 1$$

عدم عضویت

$$\pi_{\tilde{A}^P}(x) = \sqrt{1 - [\mu_{\tilde{A}^P}(x)]^2 - [\nu_{\tilde{A}^P}(x)]^2}$$

صدق می کنند و درجه تردید آن نیز

تعریف ۳-۴-۸: یک عدد فازی فیثاغورثی $\tilde{A}^P$ در مجموعه اعداد

حقیقی $\square$ به صورت یک PFN و به شکل زیر تعریف میشود:

$$\tilde{A}^P = (\mu_{\tilde{A}^P}(x), \nu_{\tilde{A}^P}(x))$$

تعریف ۳-۴-۹: یک مجموعه فازی فیثاغورثی مقدار -بازه ای $\tilde{A}^P$

$$\tilde{A}^P = \left\{ \langle x, \bar{\mu}_{\tilde{A}^P}(x), \bar{\nu}_{\tilde{A}^P}(x) \mid x \in X \rangle \right\}$$

در مجموعه مرجع X ، تعریف شده است، جائیکه:

$$\bar{\mu}_{\tilde{A}^P}(x) = [\bar{\mu}_{\tilde{A}^P}(x), \underline{\mu}_{\tilde{A}^P}(x)] \subseteq [0,1], \bar{\nu}_{\tilde{A}^P}(x) = [\bar{\nu}_{\tilde{A}^P}(x), \underline{\nu}_{\tilde{A}^P}(x)] \subseteq [0,1]$$

اعداد بازه ای صادق در شرط

$$[\bar{\mu}_{\tilde{A}^P}(x)]^2 + [\bar{\nu}_{\tilde{A}^P}(x)]^2 \leq 1$$

تعریف ۳-۴-۱۰: $\tilde{A}^P$ یک عدد فازی فیثاغورثی مقدار-بازه ای

(IVPFN) نامیده می شود و بوسیله رابطه زیر نمایش داده می شود:

$$\tilde{A}^P = \left\{ \left[\bar{\mu}_{\tilde{A}^P}(x), \underline{\mu}_{\tilde{A}^P}(x) \right], \left[\bar{\nu}_{\tilde{A}^P}(x), \underline{\nu}_{\tilde{A}^P}(x) \right] \right\}.$$

۴-توابع رتبه بندی

برای اینکه بتوانیم اعداد فازی را مرتب کنیم باید بتوانیم آنها را با هم مقایسه کنیم، مقایسه یا رتبه بندی اعداد فازی از جمله کاربردهای مهم آنها در دنیای واقعی است. فرض کنید مجموعه تمام اعداد فازی $F(\square)$ باشد. یک روش مرتب سازی عناصر $F(\square)$ استفاده از توابع رتبه بندی مانند $R: F(\square) \rightarrow \square$ می باشد که به هر عدد فازی یک عدد حقیقی نسبت می دهد. در واقع هر تابع مرتب ساز مانند R برای هر دو عدد فازی بصورت زیر عمل میکند:

$$A \geq_R B \Leftrightarrow R(A) \geq R(B)$$

$$A \leq_R B \Leftrightarrow R(A) \leq R(B)$$

$$A =_R B \Leftrightarrow R(A) = R(B)$$

توابع مرتب کننده فازی، معمولاً یک مشخصه واحد را از هر عدد فازی استخراج نموده و سپس بر اساس آن اقدام به مرتب نمودن اعداد فازی می نمایند و این باعث پیچیده شدن موضوع رتبه بندی اعداد فازی شده است، زیرا برای یک مجموعه ثابت ممکن است رتبه های متفاوت وابسته به روش مورد

مسئله برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با ابهام از نوع وگنس در بردار منابع

شکوه سرگلزائی^۱ و حسن میش مست نهی^۲

^۱ گروه ریاضی، دانشگاه سیستان و بلوچستان، زاهدان، sargolzaei.sh@pgs.usb.ac.ir

^۲ دانشیار، گروه ریاضی، دانشگاه سیستان و بلوچستان، زاهدان، hmnehi@hamoon.usb.ac.ir

چکیده: اکثر داده‌ها در جهان واقعی به صورت نادقیق و مبهم هستند؛ یکی از روش‌های مدلسازی مفاهیم نادقیق، استفاده از مجموعه‌های فازی است. یک تعمیم از مجموعه‌های فازی به صورت مجموعه‌های فازی نوع-۲ هستند که درجه عضویت‌هایی از نوع مجموعه‌های فازی دارند. مجموعه‌های فازی نوع-۲ بازه‌ای حالت خاصی از مجموعه‌های فازی نوع-۲ می‌باشند که از پیچیدگی کمتر و فهم آسانتری برخوردار هستند. در مسائل برنامه‌ریزی خطی فازی برحسب موقعیت ابهامات در مسئله، حالت‌های مختلفی ایجاد می‌شود. در این پژوهش به بررسی مسئله برنامه‌ریزی خطی با بردار منبع فازی که دارای ابهام از نوع وگنس می‌باشد، پرداخته می‌شود که این نوع ابهامات در مسئله با توابع عضویت بیان می‌شوند. سه روش جدید برای حل این نوع مسائل ارائه می‌کنیم و با مثالی کارایی روش‌ها را در مقایسه با سه مطالعه پیشین، مورد بررسی قرار می‌دهیم.

کلمات کلیدی: برنامه‌ریزی خطی فازی نوع-۲، برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای، برنامه‌ریزی خطی بازه‌ای، عدم قطعیت، بردار منابع.

۱ مقدمه

اثر عدم قطعیت و مدل سازی مسائل دارند [۵]. یکی از حالت‌های خاص مجموعه‌های فازی نوع-۲، مجموعه‌های فازی نوع-۲ بازه‌ای می‌باشد که تعمیمی بر مجموعه‌های فازی نوع-۱ است. بنابراین نسبت به مجموعه‌های فازی نوع-۱، اطلاعات نادقیق بیشتری را بیان و نسبت به مجموعه‌های فازی نوع-۲ از پیچیدگی کمتر و فهم آسانتری برخوردار هستند. به همین دلیل در سال‌های اخیر، مدل برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای از اهمیت و کاربرد بیشتری برخوردار شده است و تحقیقات بسیاری روی روش‌های حل آنها صورت گرفته است. مسائل برنامه‌ریزی خطی فازی بر حسب نوع ابهام موجود در مسئله به سه دسته انعطاف‌پذیر^۱، امکانی^۲ و استوار^۳ تقسیم می‌شوند. همچنین در هر یک از این تقسیم‌بندی‌ها، بر حسب موقعیت قرارگیری ابهام‌ها در مسئله حالت‌های مختلفی به وجود می‌آید که برای هر حالت، روش‌های حل گوناگونی ارائه شده است. در این پژوهش، به بررسی مسئله برنامه‌ریزی خطی فازی تک هدفه انعطاف‌پذیر پرداخته می‌شود که دارای ابهام از

در مدل‌های جهان واقعی، داده‌ها به صورت قطعی بیان می‌شوند در حالی که بسیاری از این داده‌ها نادقیق هستند. برای تحلیل و بررسی مدل‌های واقعی با داده‌های نادقیق نمی‌توان از حساب اعداد حقیقی استفاده کرد. لذا برای بیان تجزیه و تحلیل داده‌های نادقیق، از رویکردهای بازه‌ای، تصادفی یا فازی که تعمیمی از اعداد حقیقی و حساب اعداد حقیقی هستند، استفاده می‌شود. مجموعه‌های فازی، برای اولین بار توسط پرفسور لطفی عسگرزاده [۱] معرفی گردید. سیستم‌های فازی به علت دارا بودن توابع عضویت با درجات تعلق دقیق، توانایی محدودی در کاهش اثر عدم قطعیت دارند و برای حل مسائل پیچیده با ابهامات زیاد، مناسب نیستند. زاده [۲]، [۳]، [۴] مجموعه‌های فازی نوع-۲ را به عنوان تعمیمی از مجموعه‌های فازی مطرح کرد، برای آنکه بین مجموعه‌های فازی و مجموعه‌های فازی نوع-۲، تمایزی ایجاد شود به مجموعه‌های فازی معمولی، مجموعه‌های فازی نوع-۱ می‌گویند. مجموعه‌های فازی نوع-۲ دارای درجه عضویت‌های فازی می‌باشند که توانایی بیشتری برای کاهش

¹Flexible

²Possibilistic

³Robust

نوع وگنس^۴ می باشد و این نوع ابهام را با تابع عضویت نشان می دهند. یکی از حالت های موجود برای مسئله برنامه ریزی خطی تک هدفه با ابهام از نوع وگنس، مسئله برنامه ریزی خطی فازی تک هدفه با ابهام از نوع وگنس در بردار منابع فازی نوع-۲ بازه ای است. از جمله دست آوردهای پیشین برای حل این نوع مسائل می توان به دو رویکرد از فیگورا^۵ و یک روش از سارانی اشاره کرد که الگوریتم هایی را برای حل مسائل برنامه ریزی خطی فازی نوع-۲ بازه ای با ابهام از نوع وگنس در بردار منابع پیشنهاد داده اند [۶]. نمادهای به کار رفته در این مقاله، در جدول ۱ نمایش داده می شود:

جدول ۱: نمادهای مورد استفاده در مقاله به عنوان مثال برای a

نماد	توضیحات
$\tilde{a}$	مجموعه های فازی نوع-۲
$\bar{a}$	تابع عضویت بالایی
$\underline{a}$	تابع عضویت پایینی
$a^\vee$	تابع عضویت چپ
$a^\wedge$	تابع عضویت راست
$a^\pm$	عدد بازه ای

تقسیم بندی پژوهش پیش رو به صورت زیر می باشد: ابتدا در بخش ۲، مفاهیم مقدماتی و پایه ای بیان می شود. سپس در بخش ۳ مروری بر روش های موجود برای حل مسائل برنامه ریزی خطی فازی نوع-۲ بازه ای با ابهام در بردار منابع خواهیم داشت. در بخش ۴، سه روش جدید برای حل مسائل برنامه ریزی خطی فازی نوع-۲ بازه ای تک هدفه با ابهام از نوع وگنس در بردار منابع ارائه می شود. نتایج حاصل از این روش ها با نتایج حاصل از روش های مروری مقایسه خواهد شد.

۲ تعاریف و پیش نیازها

در این بخش به ارائه تعاریف و مفاهیم پایه ای مورد نیاز پرداخته می شود.

۱-۲ مجموعه های فازی نوع-۲ بازه ای

در این زیر بخش چند تعاریف اساسی از مجموعه های فازی نوع-۲، مجموعه های فازی نوع-۲ بازه ای، مفهوم عدم قطعیت و برنامه ریزی خطی بازه ای ارائه می شود. برای مشاهده جزئیات بیشتر به [۷، ۸، ۹] مراجعه شود.

تعریف ۱ (مجموعه ای فازی نوع-۲). مجموعه ای فازی نوع-۲ $\tilde{A}$ ، بر مجموعه مرجع X به صورت (۱) تعریف می شود:

$$(1) \quad \tilde{A} = \int_{x \in X} \int_{u \in J_x} f_x(u)/(x, u) = \int_{x \in X} \left[\int_{u \in J_x} f_x(u)/u \right] / x$$

⁴Vagueness

⁵Figueroa

که در آن $f_x(u) \in [0, 1]$ ، $f_x(u) \in [0, 1]$ ، $u \in [0, 1]$ ، $J_x \subseteq [0, 1]$ عضویت اولیه و $f_x(u)$ عضویت ثانویه می باشد.

تعریف ۲ (مجموعه ای فازی نوع-۲ بازه ای). مجموعه ای فازی نوع-۲ بازه ای به صورت زیر تعریف می شود:

$$(2) \quad \tilde{A} = \int_{x \in X} \int_{u \in J_x} 1/(x, u) = \int_{x \in X} \left[\int_{u \in J_x} 1/u \right] / x$$

که $u \in [0, 1]$ و $x \in X$ ، $J_x \subseteq [0, 1]$.

تعریف ۳ (برش α). برش α روی مجموعه فازی نوع-۲ بازه ای عبارت است از:

$$(3) \quad \alpha \tilde{a}_{ij} = \left\{ (\tilde{a}_{ij}, u) \mid J_{\tilde{a}_{ij}} \geq \alpha, u \in [0, 1] \right\}$$

که توسط دو ناحیه (۴) و (۵) نشان داده می شود:

$$(4) \quad \alpha \bar{\mu}_{\tilde{a}_{ij}} = \left\{ (\tilde{a}_{ij}, u) \mid \bar{\mu}_{\tilde{a}_{ij}} \geq \alpha \right\}$$

و

$$(5) \quad \alpha \underline{\mu}_{\tilde{a}_{ij}} = \left\{ (\tilde{a}_{ij}, u) \mid \underline{\mu}_{\tilde{a}_{ij}} \geq \alpha \right\}$$

که (۴) تابع عضویت بالایی $\tilde{a}_{ij}$ متناظر با $\alpha \in [0, 1]$ و (۵) تابع عضویت پایینی متناظر با $\alpha \in [0, 1]$ است.

تعریف ۴. یک مجموعه فازی نوع-۲ بازه ای با تابع عضویت بالایی و پایینی مشخص را در نظر بگیرید. $\alpha \tilde{a}_{ij}^\wedge$ و $\alpha \tilde{a}_{ij}^\vee$ را به ترتیب، به عنوان کران های بازه ای مقدار راست و چپ (۲) تعریف می کنیم. بنابراین، کران های $\alpha \tilde{a}_{ij}$ به صورت (۶) و (۷) تعریف می شوند:

$$(6) \quad \alpha \tilde{a}_{ij}^\vee = \left[\inf_{\tilde{a}_{ij}} \alpha \bar{\mu}_{\tilde{a}_{ij}}(\tilde{a}_{ij}, u); \inf_{\tilde{a}_{ij}} \alpha \bar{\mu}_{\tilde{a}_{ij}}(\tilde{a}_{ij}, u) \right] = \left[\alpha \bar{\mu}_{\tilde{a}_{ij}}^\vee, \alpha \bar{\mu}_{\tilde{a}_{ij}}^\vee \right]$$

(۷)

$$\alpha \tilde{a}_{ij}^\wedge = \left[\sup_{\tilde{a}_{ij}} \alpha \underline{\mu}_{\tilde{a}_{ij}}(\tilde{a}_{ij}, u); \sup_{\tilde{a}_{ij}} \alpha \underline{\mu}_{\tilde{a}_{ij}}(\tilde{a}_{ij}, u) \right] = \left[\alpha \underline{\mu}_{\tilde{a}_{ij}}^\wedge, \alpha \underline{\mu}_{\tilde{a}_{ij}}^\wedge \right]$$

۲-۲ عدم قطعیت در بردار منابع

در این زیر بخش، با بیان چند تعاریف اساسی، به بررسی عدم قطعیت روی بردار منابع در یک مجموعه فازی نوع-۲ بازه ای پرداخته می شود [۱۰].

تعریف ۵. مجموعه ای از پارامترهای بردار منابع یک مسئله برنامه ریزی خطی فازی که به صورت مجموعه های فازی نوع-۲ بازه ای تعریف می شوند را در نظر بگیرید. تابع عضویت $\tilde{b}$ عبارت است از رابطه (۸):

$$(8) \quad \tilde{b}_i = \int_{b_i \in \mathbb{R}} \left[\int_{u \in J_{b_i}} 1/u \right] / b_i, \quad i \in m, \quad J_{b_i} \subseteq [0, 1]$$

۳ مسئله برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با ابهام از نوع وگنس در بردار منابع

در این بخش، مسئله برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با ابهام از نوع وگنس در بردار منابع معرفی و در زیربخش ۳-۱ توابع عضویت بردار منابع فازی نوع-۲ بازه‌ای شرح داده خواهد شد. مدل برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با تابع هدف بیشینه و بردار منابع فازی نوع-۲ بازه‌ای به صورت (۱۲) نمایش داده می‌شود:

$$\begin{aligned} \max \quad & \sum_{j=1}^n c_j x_j \\ \text{s.t.} \quad & \sum_{j=1}^n a_{ij} x_j \leq \tilde{b}_i, \quad i = 1, \dots, m \\ & x_j \geq 0, \quad j = 1, \dots, n. \end{aligned} \quad (12)$$

که در آن، $\tilde{b} \in \mathbb{R}^m$ ، $c \in \mathbb{R}^n$ ، $x \in \mathbb{R}^n$ ، $a \in \mathbb{R}^{m \times n}$ یک ماتریس از اعداد حقیقی $n \times m$ است که توسط a_{ij} نمایش داده می‌شود و $\tilde{b}$ یک بردار $m \times 1$ از مجموعه‌های فازی نوع-۲ بازه‌ای است.

۳-۱-۳ توابع عضویت بردار منابع فازی نوع-۲ بازه‌ای

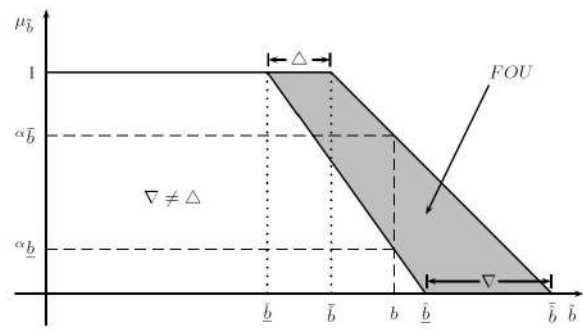
بردار $\tilde{b}$ ، یک مجموعه فازی نوع-۲ بازه‌ای است که ابهام در آن را به صورت تابع عضویت اولیه پایینی ($\mu_{\tilde{b}}^-(x)$) و تابع عضویت بالایی ($\mu_{\tilde{b}}^+(x)$) نمایش می‌دهیم. برای قیود کوچکتر یا مساوی دو تابع عضویت بالایی و پایینی و قیود بزرگتر یا مساوی دو تابع عضویت بالایی و پایینی وجود دارد. یعنی در کل چهار رابطه ممکن برای قیود کوچکتر یا مساوی و بزرگتر یا مساوی با بردار منابع فازی نوع-۲ بازه‌ای داریم که در اینجا به بررسی توابع عضویت بردار منابع فازی نوع-۲ بازه‌ای با قیود کوچکتر یا مساوی پرداخته می‌شود. برای مشاهده حالت‌های دیگر به مرجع [۶] مراجعه شود.

- اگر قیود کوچکتر مساوی، ($\leq$) باشند آنگاه تابع عضویت پایینی برابر است با (۱۳):

$$\mu_{\tilde{b}_i}^-\left((Ax)_i; \underline{b}_i^\vee; \bar{b}_i^\wedge\right) = \begin{cases} 1 & (Ax)_i \leq \underline{b}_i^\vee \\ \frac{\bar{b}_i^\wedge - (Ax)_i}{\bar{b}_i^\wedge - \underline{b}_i^\vee} & \underline{b}_i^\vee \leq (Ax)_i \leq \bar{b}_i^\wedge \\ 0 & (Ax)_i \geq \bar{b}_i^\wedge \end{cases} \quad (13)$$

- تابع عضویت بالایی قیود ($\leq$) به صورت (۱۴) بیان می‌شود:

$$\mu_{\tilde{b}_i}^+\left((Ax)_i; \bar{b}_i^\vee; \underline{b}_i^\wedge\right) = \begin{cases} 1 & (Ax)_i \leq \bar{b}_i^\vee \\ \frac{\bar{b}_i^\vee - (Ax)_i}{\bar{b}_i^\vee - \underline{b}_i^\wedge} & \bar{b}_i^\vee \leq (Ax)_i \leq \underline{b}_i^\wedge \\ 0 & (Ax)_i \geq \underline{b}_i^\wedge \end{cases} \quad (14)$$



شکل ۱: مجموعه‌های فازی نوع-۲ بازه‌ای با عدم قطعیت‌های Δ و ∇ .

توجه داشته باشید که $\tilde{b}$ به ترتیب، توسط دو تابع عضویت بالایی و پایینی اولیه به نام‌های $\mu_{\tilde{b}}^-(x)$ با مولفه‌های $\underline{b}^\vee$ و $\bar{b}^\wedge$ و $\mu_{\tilde{b}}^+(x)$ با مولفه‌های $\bar{b}^\vee$ و $\underline{b}^\wedge$ کراندار می‌شود.

تعریف ۶ (عدم قطعیت Δ و ∇). یک مسئله برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با قیود کوچکتر مساوی را در نظر بگیرید. بنابراین، Δ را به صورت $\Delta = \bar{b}^\vee - \underline{b}^\vee$ و ∇ را به صورت $\nabla = \bar{b}^\wedge - \underline{b}^\wedge$ تعریف می‌کنیم. شکل ۱ را ببینید.

۳-۲ برنامه‌ریزی خطی بازه‌ای

در این زیر بخش، مسئله برنامه‌ریزی خطی بازه‌ای با تابع هدف بیشینه‌سازی و قیود کوچکتر مساوی بیان و روش بدست آوردن بهترین و بدترین مقادیر تابع هدف برای این نوع مسائل تشریح می‌شود [۱۱].

تعریف ۷. مسئله برنامه‌ریزی خطی بازه‌ای (۹) را در نظر بگیرید:

$$\begin{aligned} \max \quad & z = \sum_{j=1}^n c_j^\pm x_j \\ \text{s.t.} \quad & \sum_{j=1}^n a_{ij}^\pm x_j \leq b_i^\pm \quad i = 1, 2, \dots, m \\ & x_j \geq 0 \quad j = 1, 2, \dots, n. \end{aligned} \quad (9)$$

که در آن $c_j^\pm$ ، $a_{ij}^\pm$ و $b_i^\pm$ اعداد بازه‌ای هستند.

قضیه ۱. بهترین و بدترین مقادیر تابع هدف بهینه مسئله (۹) را می‌توان به ترتیب با حل دو مسئله برنامه‌ریزی خطی (۱۰) و (۱۱) بدست آورد:

$$\begin{aligned} \max \quad & z^+ = \sum_{j=1}^n c_j^+ x_j \\ \text{s.t.} \quad & \sum_{j=1}^n a_{ij}^+ x_j \leq b_i^+ \quad i = 1, 2, \dots, m \\ & x_j \geq 0 \quad j = 1, 2, \dots, n, \end{aligned} \quad (10)$$

و

$$\begin{aligned} \max \quad & z^- = \sum_{j=1}^n c_j^- x_j \\ \text{s.t.} \quad & \sum_{j=1}^n a_{ij}^- x_j \leq b_i^- \quad i = 1, 2, \dots, m \\ & x_j \geq 0 \quad j = 1, 2, \dots, n. \end{aligned} \quad (11)$$

۴ ارائه سه روش جدید برای حل مسئله برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با ابهام از نوع وگنس در بردار منابع

در این بخش، سه روش جدید برای حل مسائل برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با ابهام از نوع وگنس در بردار منابع پیشنهاد می‌شود [۶].

۱-۴ روش اول

از آنجا که تابع هدف قطعی است این روش، یک روش نامتقارن است و با توجه به اینکه که در مسئله (۱۲) مقادیر سمت راست فازی هستند، تابع عضویت قیود مسئله با استفاده از عملگر بلمن و زاده به صورت (۱۵) تعریف می‌شوند:

$$\mu_D(\mathbf{x}^*) = \text{MaxMin}\{\mu((A\mathbf{x})_i, u)\} \quad (15)$$

برای خطی سازی کافی است $\alpha = \text{Min}\{\mu((A\mathbf{x})_i, u)\}$ در نظر گرفته شود. بنابراین می‌توان مسئله برنامه‌ریزی خطی (۱۶) را نوشت:

$$\begin{aligned} \max \quad & \sum_{j=1}^n c_j x_j \\ \text{s.t.} \quad & \alpha \leq \mu((A\mathbf{x})_i, u) \\ & \mathbf{x} \geq 0, \alpha \in [0, 1]. \end{aligned} \quad (16)$$

در این حالت که تابع هدف ماکزیمم سازی و قیود از نوع کوچکتری ($\leq$) هستند، توابع عضویت قیود به صورت زیر بیان می‌شوند:

$$\mu(A\mathbf{x})_i = \begin{cases} 1, & (A\mathbf{x})_i \leq \mathbf{b}_i \\ \frac{\mathbf{b}_i + \Delta_i - (A\mathbf{x})_i}{\Delta_i}, & \mathbf{b}_i \leq (A\mathbf{x})_i \leq \mathbf{b}_i + \Delta_i \\ 0, & (A\mathbf{x})_i \geq \mathbf{b}_i + \Delta_i \end{cases} \quad (17)$$

که در آن $\mathbf{b}_i \in [\underline{b}_i, \bar{b}_i]$ و $\Delta_i \in [\underline{\Delta}_i, \bar{\Delta}_i]$ برقرار هستند. می‌توان مدل (۱۶) را به صورت زیر باز نویسی کرد:

$$\begin{aligned} \max \quad & \sum_{j=1}^n c_j x_j \\ \text{s.t.} \quad & \alpha \leq \frac{\mathbf{b}_i + \Delta_i - (A\mathbf{x})_i}{\Delta_i}, \quad i = 1, \dots, m \\ & \mathbf{x} \geq 0, \alpha \in [0, 1]. \end{aligned} \quad (18)$$

در نتیجه با استفاده از مفاهیم برنامه‌ریزی خطی بازه‌ای، بهترین جواب مسئله (۱۸) با حل مسئله برنامه‌ریزی خطی (۱۹) بدست می‌آید:

$$\begin{aligned} \max \quad & \sum_{j=1}^n c_j x_j \\ \text{s.t.} \quad & (A\mathbf{x})_i + \underline{\Delta}_i \alpha \leq \bar{b}_i, \quad i = 1, \dots, m \\ & \mathbf{x} \geq 0, \alpha \in [0, 1]. \end{aligned} \quad (19)$$

دقت شود در این روش، برای هر $\alpha \in [0, 1]$ یک جواب بهینه بدست می‌آید.

مسئله برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با هدف ماکزیمم سازی و قیود کوچکتر مساوی را با داده‌های زیر در نظر بگیرید:

$$\begin{aligned} \underline{\mathbf{b}}^\vee &= \begin{bmatrix} 50 \\ 70 \\ 40 \\ 60 \\ 40 \end{bmatrix} \quad ; \mathbf{c} = \begin{bmatrix} 12 \\ 17 \\ 9 \end{bmatrix} \quad ; A = \begin{bmatrix} 5 & 3 & 7 \\ 10 & 4 & 9 \\ 4 & 6 & 3 \\ 2 & 7 & 7 \\ 5 & 6 & 11 \end{bmatrix} \\ \bar{\mathbf{b}}^\wedge &= \begin{bmatrix} 95 \\ 110 \\ 77 \\ 102 \\ 98 \end{bmatrix} \quad ; \bar{\mathbf{b}}^\vee = \begin{bmatrix} 60 \\ 80 \\ 55 \\ 75 \\ 57 \end{bmatrix} \quad ; \underline{\mathbf{b}}^\wedge = \begin{bmatrix} 72 \\ 104 \\ 65 \\ 95 \\ 80 \end{bmatrix} \end{aligned}$$

با حل مثال با روش پیشنهادی اول و با توجه به مقادیر مختلف برای $\alpha \in [0, 1]$ ، نتایج بدست آمده، در جدول ۲ نمایش داده شده‌اند:

جدول ۲: مقایسه نتایج حاصل از اولین روش پیشنهادی ما و نتایج حاصل از روش‌های مروری

روش‌ها	α^*	$\mathbf{x}_1^*$	$\mathbf{x}_2^*$	$\mathbf{x}_3^*$	z^*
روش اول	0	8/000	7/5000	0	223/5
	0/1	7/7909	7/2727	0	217/1273
	0/2	7/5815	7/0455	0	210/7545
	0/3	7/3727	6/8182	0	204/3818
	0/4	7/1634	6/5909	0	198/0091
	0/5	6/9545	6/3636	0	191/6364
	0/6	6/7455	6/1364	0	185/2636
	0/7	6/5364	5/9091	0	178/8909
	0/8	5/8000	6/0333	0	172/1667
	0/9	3/9000	6/9333	0	164/6667
	1	2/0000	7/8333	0	157/1667
روش اول فیگورا	0/51	-	-	-	190/99
روش دوم فیگورا	0/5106	-	-	-	175/76
روش سارانی	0/741	6/2495	5/5812	0	169/8743

با مقایسه نتایج بدست آمده از جدول ۲ و نتایج حاصل از روش اول فیگورا (که در $\alpha^* = 0/51$ دارای جواب بهینه $z^* = 190/99$) و روش دوم فیگورا که در آن با $\alpha^* = 0/5106$ مقدار بهینه تابع هدف برابر $z^* = 175/76$ و همچنین نتایج حاصل از روش مان، با $\alpha^* = 0/5$ مقدار بهینه تابع هدف برابر $z^* = 191/636$ بدست آمده است، مشاهده می‌شود که با مقدار بهینه تابع هدف حاصل از روش اول فیگورا در یک محدوده قرار دارند و از مقدار بدست آمده برای تابع هدف بهینه در روش دوم فیگورا بهتر می‌باشد. همچنین با مقایسه نتایج بدست آمده از این روش با نتایج حاصل از روش سارانی (مقدار بهینه تابع هدف $z^* = 169/874$ و

۰/۷۴۱ α^*)، مشاهده می‌شود که مقدار بهینه تابع هدف در روش اول پیشنهادی ما مقداری بهتر بدست آمده است و بهینه سازی بهتری حاصل شده است.

۲-۴ روش دوم

همانطور که مشاهده می‌شود روش ۴-۱، یک روش نامتقارن است، در این روش، ابتدا مسئله را به یک مسئله متقارن تبدیل می‌کنیم و سپس به حل آن می‌پردازیم. فرم کلی مسئله برنامه‌ریزی خطی فازی با بردار منابع فازی نوع-۲ بازه‌ای (مدل (۱۲)) را در نظر بگیرید. از آنجا که تابع عضویت بردار منبع فازی دارای دو تابع عضویت بالایی و پایینی است، $(\underline{b}_i^V, \bar{b}_i^V, \underline{b}_i^A, \bar{b}_i^A)$ را داریم. لذا چهار مقدار بهینه تابع هدف، حاصل از چهار مدل برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای $(\underline{z}^V, \bar{z}^V, \underline{z}^A, \bar{z}^A)$ خواهیم داشت. بنابراین، بنابر قواعد برنامه‌ریزی خطی فازی، کمترین و بیشترین مقدار تابع هدف به ترتیب عبارتند از $\underline{z}^V$ و $\bar{z}^A$ که به ترتیب به صورت (۲۰) و (۲۱) محاسبه می‌کنیم:

$$\begin{aligned} \max \quad & \sum_{j=1}^n c_j x_j \\ \text{s.t.} \quad & \sum_{j=1}^n a_{ij} x_j \leq \bar{b}_i^A, \quad i = 1, \dots, m \\ & x \geq 0. \end{aligned} \quad (20)$$

و

$$\begin{aligned} \max \quad & \sum_{j=1}^n c_j x_j \\ \text{s.t.} \quad & \sum_{j=1}^n a_{ij} x_j \leq \underline{b}_i^V, \quad i = 1, \dots, m \\ & x \geq 0. \end{aligned} \quad (21)$$

تابع عضویت $\mu_0(z)$ را تحت عنوان تابع عضویت تابع هدف به صورت (۲۲) تعریف و نشان می‌دهیم:

$$\mu_0(z) = \begin{cases} 1, & \text{cx} > \bar{z}^A, \\ 1 - \frac{\bar{z}^A - \text{cx}}{\bar{z}^A - \bar{z}^V}, & \bar{z}^V \leq \text{cx} \leq \bar{z}^A, \\ 0, & \text{cx} < \bar{z}^V. \end{cases} \quad (22)$$

حال، می‌توان با کمک عملگر Max-Min بلمن و زاده، مسئله (۱۲) را به مسئله (۲۳) که هدف آن یافتن بیشترین میزان درجه تطابق قیود و هدف است، تبدیل کرد:

$$\begin{aligned} \max \quad & \alpha \\ \text{s.t.} \quad & \alpha \leq \mu_0(z) \\ & \alpha \leq \mu((Ax)_i, u) \\ & x \geq 0, \alpha \in [0, 1]. \end{aligned} \quad (23)$$

با جایگزینی توابع عضویت قیود و هدف در مسئله (۲۳)، داریم:

$$\begin{aligned} \max \quad & \alpha \\ \text{s.t.} \quad & \alpha \leq 1 - \frac{\bar{z}^A - \text{cx}}{\bar{z}^A - \bar{z}^V} \\ & \alpha \leq \frac{\underline{b}_i + \Delta_i - (Ax)_i}{\Delta_i}, \quad i = 1, \dots, m \\ & x \geq 0, \alpha \in [0, 1]. \end{aligned} \quad (24)$$

در نتیجه با استفاده از مفاهیم برنامه‌ریزی خطی فازی، بهترین جواب مسئله (۲۳) با حل مسئله برنامه‌ریزی خطی (۲۵) بدست می‌آید:

$$\begin{aligned} \max \quad & \alpha \\ \text{s.t.} \quad & \text{cx} \geq \bar{z}^A - (1 - \alpha)(\bar{z}^A - \bar{z}^V), \\ & (Ax)_i + \Delta_i \alpha \leq \bar{b}_i^A, \quad i = 1, \dots, m \\ & x \geq 0, \alpha \in [0, 1]. \end{aligned} \quad (25)$$

داده‌های مثال ۴-۱ را در نظر بگیرید. با حل مثال ۴-۲، مقادیر $(x_1^*, x_2^*, x_3^*) = (6/675, 6/600, 0)$ و $\alpha^* = 0/633$ بدست می‌آیند. با جایگذاری مقادیر جواب بهینه در تابع هدف مسئله اصلی، مقدار بهینه تابع هدف برابر $z^* = 183/125$ می‌شود. در این روش، مقدار بهینه آلفایی که بدست می‌آید بهترین مقدار مینیمم درجه عضویت توابع عضویت می‌باشد و مقدار تابع هدفی که از حل روش دوم بدست می‌آید از مقدار تابع هدف بدست آمده از روش سازانی بهتر است.

۳-۴ روش سوم

فرم کلی مسئله برنامه‌ریزی خطی فازی با بردار منابع فازی نوع-۲ بازه‌ای (مدل (۱۲)) را در نظر بگیرید. با حل روش دوم (۴-۲)، اگر جواب بهینه منحصر به فرد نباشد و داشته باشیم (x^*, α^*) ، آنگاه داریم:

$$\begin{aligned} \max \quad & \sum_{i=0}^m \alpha \\ \text{s.t.} \quad & \alpha_0 \geq \mu_0(z = \text{cx}^*) \\ & \alpha_i \geq \mu((Ax^*)_i, u), \quad i = 0, 1, \dots, m \\ & \alpha_0 \leq \mu_0(z = \text{cx}) \\ & \alpha_i \leq \mu((Ax)_i, u), \quad i = 0, 1, \dots, m \\ & x \geq 0, \alpha_i \in [0, 1], \quad i = 0, 1, \dots, m. \end{aligned} \quad (26)$$

در نتیجه با جایگذاری توابع عضویت در مسئله (۲۶)، مسئله (۲۷) بدست می‌آید:

$$\begin{aligned} \max \quad & \sum_{i=0}^m \alpha \\ \text{s.t.} \quad & \text{cx}^* \geq \bar{z}^A - (1 - \alpha_0)(\bar{z}^A - \bar{z}^V), \\ & (Ax^*)_i \leq \bar{b}_i + (1 - \alpha_i)\Delta_i, \quad i = 1, \dots, m \\ & \text{cx} \leq \bar{z}^A - (1 - \alpha_0)(\bar{z}^A - \bar{z}^V), \\ & (Ax)_i \geq \bar{b}_i + (1 - \alpha_i)\Delta_i, \quad i = 1, \dots, m \\ & x \geq 0, \alpha_i \in [0, 1]. \end{aligned} \quad (27)$$

در نتیجه با استفاده از مفاهیم برنامه‌ریزی خطی بازه‌ای، بهترین جواب مسئله (۲۷) با حل مسئله برنامه‌ریزی خطی (۲۸) بدست می‌آید:

$$\begin{aligned} \max \quad & \sum_{i=0}^m \alpha \\ \text{s.t.} \quad & \mathbf{c}\mathbf{x}^* \geq \bar{z}^{\wedge} - (1 - \alpha_0)(\bar{z}^{\wedge} - \underline{z}^{\vee}), \\ & (\mathbf{A}\mathbf{x}^*)_i + \underline{\Delta}_i \alpha_i \leq \bar{\mathbf{b}}_i^{\wedge}, \quad i = 1, \dots, m \\ & \mathbf{c}\mathbf{x} \leq \bar{z}^{\wedge} - (1 - \alpha_0)(\bar{z}^{\wedge} - \underline{z}^{\vee}), \\ & (\mathbf{A}\mathbf{x})_i + \alpha_i \bar{\Delta}_i \geq \underline{\mathbf{b}}_i^{\vee}, \quad i = 1, \dots, m \\ & \mathbf{x} \geq \mathbf{0}, \alpha_i \in [0, 1]. \end{aligned} \quad (28)$$

یک جواب بهینه مسئله (۲۸) می‌باشد. $(\mathbf{x}^{**}, \alpha_0^{**}, \alpha_1^{**}, \dots, \alpha_m^{**})$

با حل مسئله به روش ۴-۲، مقادیر $\mathbf{x}^* = (\mathbf{x}_1^*, \mathbf{x}_2^*, \mathbf{x}_3^*)$ و $\alpha^* = 0/633$ بدست می‌آیند، همچنین جایگذاری این مقادیر در قیدهای مسئله اصلی، مقدار $z(\mathbf{x}^*) = z^* = 183/125$ و حدود مربوط به هر یک از درجات تطابق قیود به صورت $0/433 \leq \alpha_2 \leq 1, 0/58 \leq \alpha_1 \leq 1, 0/63 \leq \alpha_0 \leq 1, 0/25 \leq \alpha_5 \leq 1, 0/89 \leq \alpha_4 \leq 1, 0/88 \leq \alpha_3 \leq 1$ بدست می‌آید که این حدود را در مسئله در نظر می‌گیریم و با حل آن نتایج به صورت $\alpha_0 = \mathbf{x}^{**} = (\mathbf{x}_1^{**}, \mathbf{x}_2^{**}, \mathbf{x}_3^{**}) = (0, 10/69, 0/83)$ و $\alpha_5 = 0/815, \alpha_3 = 0/569, \alpha_1 = \alpha_2 = \alpha_4 = 1, 0/63$ بدست می‌آیند. همچنین مقدار تابع هدف مسئله اصلی برابر $z^* = 182/726$ می‌شود. این روش نه تنها مقدار بهینه تابع هدف را بدست می‌آورد بلکه در قیدهای اول، دوم و چهارم مسئله، به درجه بالاتری از تطابق دست می‌یابد. همچنین مقدار بهینه تابع هدف بدست آمده از روش سوم بهتر از مقدار بهینه بدست آمده از روش سارانی است.

۵ نتیجه‌گیری

در این مقاله، به معرفی مسئله برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با ابهام از نوع وگنس در بردار منابع پرداخته و روش‌های موجود برای حل این نوع مسائل معرفی شده است. دو رویکرد از فیگورا و یک روش از سارانی از جمله پژوهش‌هایی بودند که برای حل این نوع مسائل الگوریتم‌هایی ارائه داده‌اند. در ادامه سه روش جدید برای حل مسائل برنامه‌ریزی خطی فازی نوع-۲ بازه‌ای با ابهام از نوع وگنس در بردار منابع ارائه کردیم. روش اول، یک روش نامتقارن است. در این روش تابع هدف قطعی و بیشترین میزان درجه عضویت قیود محاسبه می‌شود. روش دوم، ابتدا مسئله را به یک مسئله متقارن تبدیل کرده و سپس با کمک عملگر بلمن وزاده بیشترین میزان درجه تطابق هدف و قیود را در یک مسئله محاسبه می‌شود. روش سوم، مبتنی بر روش دوم است؛ یعنی اگر جواب مسئله موجود در روش دوم، بهینه منحصر به فرد نباشد مسئله جدیدی تعریف می‌کنیم و به ادامه حل می‌پردازیم. برای بررسی کارایی روش‌های پیشنهادی ما مثالی ارائه کرده و این روش‌ها را با سه روش

قبلی (دوروش از فیگورا و یک روش از سارانی) مقایسه شده‌اند. نتایج حاصل از این مقایسات را در جداولی شرح داده، این نتایج بیانگر کارایی و بهتر بودن روش‌های پیشنهادی ما می‌باشد. حالت‌های دیگر مدل‌های برنامه‌ریزی خطی فازی نوع-۲ با ابهام از نوع وگنس را در مقالات بعدی بیان می‌کنیم.

مراجع

- [1] R. E. Bellman and L. A. Zadeh, "Decision-making in a fuzzy environment," *Management science*, vol.17, no.4, pp.B-141, 1970.
- [2] L. A. Zadeh, "The concept of a linguistic variable and its application to approximate reasoning—i," *Information sciences*, vol.8, no.3, pp.199-249, 1975.
- [3] L. A. Zadeh, "The concept of a linguistic variable and its application to approximate reasoning—i," *Information sciences*, vol.8, no.3, pp.301-357, 1975.
- [4] L. A. Zadeh, "The concept of a linguistic variable and its application to approximate reasoning—i," *Information sciences*, vol.9, no.3, pp.43-80, 1975.
- [5] E. Hisdal, "The if then else statement and interval-valued fuzzy sets of higher type," *International Journal of Man-Machine Studies*, vol.15, no.4, pp.385-455, 1981.
- [6] S. Sargolzaei and H. Mishmast Nehi, "Interval type-2 fuzzy linear programming problem with vagueness in the resources vector," in *Soft Computing*.
- [7] شکوه سرگلزانی و حسن میش مست نهی، "مروری بر مجموعه‌های فازی نوع-۲"، سیستم‌های فازی و کاربردها، جلد ۲، شماره ۲، صفحات ۱۲۳-۱۵۰، ۱۳۹۸.
- [8] J. M. Mendel, "Fuzzy sets for words: a new beginning," in *The 12th IEEE International Conference on Fuzzy Systems, 2003. FUZZ'03.*, vol.1, pp.37-42, IEEE, 2003.
- [9] J. M. Mendel, R. I. John, and F. Liu, "Interval type-2 fuzzy logic systems made simple," *IEEE transactions on fuzzy systems*, vol.14, no.6, pp.808-821, 2006.
- [10] J. C. Figueroa, "Solving fuzzy linear programming problems with interval type-2 rhs," pp.262-267, 2009.
- [11] G. Klir and B. Yuan. *Fuzzy sets and fuzzy logic*, vol.4. Prentice hall New Jersey, 1995.



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

مسئله زمان‌بندی پروژه منابع محدود فازی با چندین تامین‌کننده

بهنام امینی^۱، علیرضا عیدی^۲

^۱ دانشجوی کارشناسی ارشد مهندسی صنایع، دانشکده مهندسی، دانشگاه کردستان، aminibehnam9025@gmail.com

^۲ دانشیار گروه مهندسی صنایع، دانشکده مهندسی، دانشگاه کردستان، alireza.eydi@uok.ac.ir

چکیده: مسئله زمان‌بندی پروژه منابع محدود شامل مجموعه‌ای از فعالیت‌ها می‌باشد که بایستی بر اساس محدودیت‌های پیش‌نیازی و منابع برنامه‌ریزی گردند. از طرفی پایداری فاکتوری مهم و تاثیرگذار در زمان‌بندی پروژه‌ها می‌باشد که با در نظر گرفتن سیاست‌های تخفیف و اثرات زیست‌محیطی و اجتماعی در انتخاب تامین‌کنندگان می‌توان اثرات مهم این فاکتور را مشاهده کرد. در این مقاله مسئله مورد بررسی به فرم یک مدل سه هدفه فرمول‌بندی می‌شود. هدف اول حداکثرسازی ارزش فعلی خالص پروژه و هدف دوم حداکثرسازی درجه زیست‌محیطی و هدف سوم حداکثرسازی درجه اجتماعی می‌باشد. همچنین از تئوری مجموعه فازی برای بیان نامعینی و عدم قطعیت‌هایی مانند زمان و سطح دسترس‌پذیری منابع پروژه استفاده شده است. در پایان، یک مثال با ۱۰ فعالیت تولید شده و با رویکرد اپسیلون محدودیت در نرم افزار GAMS حل شده و نتایج، صحت و اعتبار مدل نشان داده شده است.

کلمات کلیدی: پایداری، زمان‌بندی پروژه منابع محدود، تامین‌کنندگان مواد، عدم قطعیت، بهینه‌سازی چندهدفه.

۱- مقدمه و مرور ادبیات

مسئله زمان‌بندی پروژه منابع محدود به طور عمده به سه عامل بستگی دارد: منابع، روابط پیش‌نیازی و فعالیت‌ها. این مسئله به دنبال یافتن توالی مناسب برای انجام فعالیت‌های یک پروژه است، به نحوی که محدودیت‌های پیش‌نیازی شبکه پروژه و انواع مختلف محدودیت‌ها در منابع موجود در پروژه به‌طور همزمان ارضا و اهداف معینی از جمله زمان انجام پروژه، هزینه‌ی کل، تعداد فعالیت‌های تاخیردار و غیره بهینه شوند. مسئله‌ی زمان‌بندی پروژه با منابع محدود، از لحاظ رده پیچیدگی یک مسئله سخت و پیچیده به شمار می‌آید.

مسائل زمان‌بندی پروژه با در نظر گرفتن محدودیت منابع، نخستین بار توسط ویست^۱ [۱۱] مطرح گردید. کادری و بوکتور^۲ [۶] مسئله زمان‌بندی پروژه منابع محدود با توالی وابسته به زمان انتقال را مورد مطالعه قرار دادند و یک الگوریتم ژنتیک با استفاده از یک عملگر تقاطع دو نقطه‌ای پیشنهاد دادند. بیرجندی و موسوی^۳ [۲] یک مدل جدید از برنامه‌ریزی غیرخطی عدد صحیح

آمیخته فازی تحت عدم قطعیت و یک رویکرد فراابتکاری ترکیبی نیز برای به حداقل رساندن هزینه‌های تکمیل پروژه ارائه دادند. عیدی و بخش‌ی [۱] یک رویکرد فازی برای زمان‌بندی پروژه تحت محدودیت‌های منابع و عدم قطعیت را توسعه دادند. یوسفلی^۴ [۱۲] یک مدل برای مسئله زمان‌بندی پروژه منابع محدود کاملاً فازی ارائه داد که در آن مدت زمان فعالیت و مصرف منابع فعالیت‌ها با اعداد فازی برآورد می‌شود. مرادی و همکاران^۵ [۹] یک مدل فازی را برای حل مسئله زمان‌بندی پروژه منابع محدود پیشنهاد کردند که در آن فعالیت‌ها را می‌توان در چندین حالت تحت عدم قطعیت، همزمان در دسترس‌ی به منابع و مدت زمان فعالیت انجام داد. همچنین به تحلیل مفهوم پایداری در مسئله زمان‌بندی پروژه منابع محدود پرداختند. حبیبی و همکاران^۶ [۴] چارچوبی یکپارچه برای مسئله زمان‌بندی پروژه و مسئله سفارش مواد با لحاظ کردن پایداری ارائه دادند. این چارچوب پیشنهادی شامل دو مرحله است که مرحله اول آن، کمی‌سازی محاسن زیست‌محیطی و اجتماعی تامین‌کنندگان بالقوه منابع پروژه است و مرحله دوم آن، شامل ساخت و حل

⁴ Yousefli, A.

⁵ Moradi, M., Hafezalkotob, A., & Ghezavati, V.

⁶ Habibi, F., Barzinpour, F., & Sadjadi, S. J.

¹ West, J. D.

² Kadri, R. L., & Boctor, F. F.

³ Birjandi, A., & Mousavi, S. M.

۲- فرمول‌بندی مسأله

۲-۱- مفروضات مدل

در این تحقیق شبکه پروژه دارای یک گراف فعالیت بر روی گره تحت عنوان $Gr(N, A)$ می‌باشد که در آن N و A به ترتیب نشانگر گره‌ها و کمان‌ها می‌باشند. در این گراف گره‌ها به صورت فعالیت‌های پروژه عمل کرده و کمان‌ها نشانگر روابط تقدم و تاخر پایان به آغاز می‌باشند. گره‌های گراف به ترتیب صعودی از نظر تقدم و تاخر چیده شده‌اند. در این گراف‌ها، گره‌ها یا همان فعالیت‌ها دارای ۱ و n متغیر مصنوعی با مدت زمان صفر و مصرف منابع بوده و بیانگر نقاط شروع و پایان پروژه می‌باشند. تمامی فعالیت‌های پروژه تنها در یک حالت اجرا شده و هر فعالیت نیازمند یک مجموعه از منابع تجدیدپذیر و تجدیدنپذیر می‌باشد.

۲-۲- اندیس‌ها، متغیرها و پارامترها

i فعالیت‌های پروژه $i \in (1, \dots, n)$ ، J مجموعه پروژه‌ها $j \in J$ ، دوره‌های زمانی، S تامین‌کنندگان، L منابع تجدیدپذیر، F منابع تجدیدنپذیر، K_{fs} دامنه تخفیف شامل میزان تخفیفی است که از حداقل تا حداکثر از سوی تامین‌کننده ارائه می‌شود. Cf_j^+ جریان‌ات نقدی ورودی به سیستم (ناشی از انجام پروژه j)، Cf_j^- جریان‌ات نقدی خارج شده از سیستم، P هزینه جریمه ناشی از زمان تکمیل دیر هنگام پروژه برای هر دوره، B پاداش به دلیل اتمام زودهنگام پروژه برای هر دوره، $(P/F, Ir\%, t)$ عامل تخفیف p/f با نرخ بهره $ir\%$ و تعداد دوره زمانی t ، $\bar{d}_j$ زمان پردازش پروژه j (این پارامتر به صورت فازی در نظر گرفته شده است)، DD مهلت برای تکمیل پروژه، A_{fs} هزینه سفارش‌دهی مرتبط با نوع ماده f از تامین‌کننده s ، ir نرخ بهره، H_f هزینه نگهداری مرتبط با نوع ماده f در هر دوره زمانی، C_{fks} هزینه خرید (هر واحد) مرتبط با نوع مواد f در دامنه k که از تامین‌کننده s سفارش داده می‌شود. C_{fkjs} هزینه خرید (هر واحد) مرتبط با نوع مواد f در دامنه k برای پروژه j که از تامین‌کننده s سفارش داده می‌شود. LT_{fs} مدت زمان خرید مواد نوع f از تامین‌کننده s ، EST_i زودترین زمان اتمام فعالیت i ، LST_i دیرترین زمان اتمام فعالیت i ، U_i میزان مواد نوع f که برای اجرای فعالیت i لازم است، η_{fs} درجه تامین‌کننده s برای نوع مواد f در معیار محیطی، p_{fs} درجه تامین‌کننده s برای نوع ماده f در معیار اجتماعی، Y_{fks} تعداد بازه‌های تخفیف در دامنه K مربوط به نوع ماده f برای تامین‌کننده s ، r_{jl} مقدار مواد تجدیدپذیر نوع l که برای پروژه j در هر دوره پردازش لازم است، R_{maxl} میزان منابع تجدیدپذیر نوع l که در هر دوره موجود است، RF_{maxf} میزان منابع تجدیدنپذیر نوع f که در هر دوره موجود است، T افق برنامه‌ریزی، cap_s ظرفیت تامین‌کننده s ، M یک عدد بزرگ ثابت مثبت، Budget بودجه تعیین شده برای تامین مواد می‌باشد. متغیرهای تصمیم عبارتست از y_{fkst} اگر نوع مواد f در دامنه تخفیف k از تامین‌کننده s در دوره زمانی t سفارش داده شود ۱ و در غیر این صورت صفر می‌باشد. x_{jt} اگر پروژه j در دوره t انتخاب شود ۱ و در غیر این صورت صفر می‌باشد. v_{jt} اگر پروژه j در دوره t اجرا شود ۱ و در غیر این صورت صفر می‌باشد. z_{fsijt} اگر نوع مواد f در دامنه تخفیف k از تامین‌کننده s برای

یک مدل ریاضی بر اساس داده‌های بدست آمده می‌باشد. حسینی و همکاران^۷ [۱۰] یک چارچوب چندهدفه برای مسأله انتخاب پروژه پایدار ارائه دادند که این چارچوب اجازه می‌دهد پروژه‌ها در هر زمانی تقسیم شده و به تعویق بیفتند تا در زمان دیگری دوباره اجرا شوند. وابستگی پروژه‌ها به دوره انتخاب آن‌ها، رابطه متناقض بین برخی پروژه‌ها و محدودیت‌های بودجه و منابع را نیز مورد توجه قرار دادند. کیم^۸ [۷] یک روش سه مرحله‌ای مسیر بحرانی منابع محدود را برای بهبود پایداری در ساختار زمان‌بندی پروژه ارائه داد. چاولا و همکاران^۹ [۳] به شناسایی و بحث در مورد احتمالات آینده استفاده از روش‌های محاسباتی به منظور کاربرد در مسائل پایداری در مدیریت پروژه‌ها پرداخته و آن را تخمین زدند و یک چارچوب یکپارچه جدید با درج عملکرد بازخورد برای ارزیابی هر تصمیم و اقدامات انجام شده در مورد پایداری پروژه‌ها را نیز شناسایی کرده و ارائه دادند. ژانگ و همکاران^{۱۰} [۱۳] به مسأله هماهنگی بین انتخاب تامین‌کنندگان و زمان‌بندی پروژه پرداختند. برای این مسأله، یک مدل برنامه‌ریزی خطی عدد صحیح مختلط ارائه دادند و با یک الگوریتم فرابتنکاری به حل این مسأله پرداختند.

با توجه به بسیاری از عدم قطعیت‌ها، متغیرها و احتمال اینکه دسترسی به منابع یا محدوده پروژه ممکن است تغییر کند، ایجاد رویکردی که دوام داشته باشد بسیار دشوار است. در مقاله حاضر، از زمان‌بندی پروژه فازی که یکی از رویکردهایی است که با عدم قطعیت در مسأله زمان‌بندی پروژه سروکار دارد، با استفاده از نظریه اعداد فازی برای برطرف کردن این مشکل استفاده شده است. همچنین، در هر دوره میزان ثابتی از منابع تجدیدپذیر برای فعالیت‌ها موجود می‌باشد اما منابع تجدیدنپذیر می‌بایست از تامین‌کنندگان خریداری شوند. فرض بر این است که انتخاب تامین‌کننده نه تنها تابعی از تخفیفات آن‌ها بلکه تابعی از تأثیرات زیست‌محیطی و اجتماعی آن‌ها نیز می‌باشد. بنابراین، اهداف زیست‌محیطی و اهداف اجتماعی به طور غیرمستقیم از طریق انتخاب تامین‌کننده دنبال می‌شود.

موضوع مهم دیگری که در خصوص مسأله موجود در پژوهش حاضر مطرح می‌باشد در نظر گرفتن معیار پایداری می‌باشد. پایداری بر اساس سه شاخصه مهم اقتصادی، اجتماعی و زیست محیطی تعریف می‌شود که بر این اساس مدل تحقیق حاضر به یک مدل سه هدفه تبدیل می‌شود که هدف اول آن حداکثرسازی ارزش فعلی خالص پروژه و دو هدف بعدی به ترتیب حداکثرسازی درجه زیست‌محیطی و درجه اجتماعی می‌باشد. به عبارت دیگر می‌توان گفت انتخاب تامین‌کننده در تحقیق حاضر بر اساس مسائل زیست‌محیطی و اجتماعی به علاوه تخفیفات کمی می‌باشد. منظور از تخفیفات کمی، تخفیفاتی می‌باشد که بر اساس مقدار تامین شده از سوی تامین‌کننده به هزینه تامین تخصیص می‌یابد. همچنین، از رویکرد حل محدودیت افسیلون تقویت‌شده برای حل مدل ریاضی استفاده شده است.

در مجموع در این مقاله، یک فرمول‌بندی جدید برای مسأله زمان‌بندی پروژه منابع محدود با در نظر گرفتن شاخص‌های پایداری، زمان و سطح دسترسی‌پذیری فازی و انتخاب تامین‌کنندگان در شرایط عدم قطعیت ارائه می‌گردد. در بخش دوم مدل ریاضی این مسأله بیان می‌شود و در بخش سوم حل مثال عددی و در بخش چهارم نتیجه‌گیری ارائه می‌گردد.

⁹ Chawla, V., Chanda, A., Angra, S., & Chawla, G.

¹⁰ Zhang, X., Hipel, K. W., & Tan, Y.

⁷ Hoseini, A., Ghannadpour, S. F., & Hemmati, M.

⁸ Kim, K.

فعالیت i از پروژه j در دوره زمانی t سفارش داده شود ۱ و در غیر این صورت صفر می‌باشد. I_{ft} نشان‌دهنده سطح موجودی مواد f در دوره t می‌باشد. $u_{f,s,j}$ میزان مواد انتقالی f از تامین کننده s برای پروژه j می‌باشد.

۳-۲- مدل ریاضی

$$\begin{aligned} \max z1 = & \sum_{j=1}^J \sum_{t=EST_j}^{LST_j} CF_j^+ x_{jt} \left(\frac{P}{F}, ir\%, t + \tilde{d}_j \right) \\ & - \sum_{j=1}^J \sum_{t=EST_j}^{LST_j} CF_j^- x_{jt} \left(\frac{P}{F}, ir\%, t \right) \\ & - \sum_{f=1}^F \sum_{s=1}^S \sum_{t=1}^{LST_j-LT_{fs}+1} A_{fs} \sum_{k=1}^{k_{fs}} y_{fks} \left(\frac{P}{F}, ir\%, t \right) \\ & - \sum_{f=1}^F \sum_{t=1}^{LST_n-d_n-1} H_f I_{ft} \left(\frac{P}{F}, ir\%, t \right) \\ & - \sum_{f=1}^F \sum_{s=1}^S \sum_{i=1}^n \sum_{t=1}^{LST_n-LT_{fs}+1} \sum_{k=1}^{k_{fs}} \sum_{j=1}^J C_{fks} U_{if} Z_{fksijt} \\ & \left(\frac{P}{F}, ir\%, t \right) \\ & - \sum_{t=DD+1}^T P(t-DD) x_{nt} \left(\frac{P}{F}, ir\%, t \right) \\ & + \sum_{t=EST_n} B(DD-t) x_{nt} \left(\frac{P}{F}, ir\%, t \right) \end{aligned} \quad (۱)$$

$$\begin{aligned} \max z2 = & \sum_{f=1}^F \sum_{s=1}^S \sum_{k=1}^K \sum_{t=1}^{LST_n-LT_{fs}-1} \eta_{fs} \cdot \sum_{i=1}^n \sum_{j=1}^J u_{if} Z_{fksijt} \end{aligned} \quad (۲)$$

$$\begin{aligned} \max z3 = & \sum_{f=1}^F \sum_{s=1}^S \sum_{k=1}^{K_f} \sum_{t=1}^{LST_n-LT_{fs}+1} \rho_{fs} \cdot \sum_{i=1}^n \sum_{j=1}^J u_{if} Z_{fksijt} \end{aligned} \quad (۳)$$

$$\begin{aligned} s.t. & \sum_{j=1}^J \sum_{t=\max(t-d_j+1, EST_j)}^{\min(t, LST_j)} r_{jl} x_{jt'} \leq \widetilde{R_{max,l}}; \\ & \forall t = 1, 2, \dots, n, \forall l \in L \end{aligned} \quad (۴)$$

$$\begin{aligned} I_{ft} = & I_{f(t-1)} \\ & + \sum_{i=1}^n \sum_{s=1}^S \sum_{k=1}^{K_{fs}} \sum_{j=1}^J u_{if} Z_{fksijt} (t-LT_{fs}) \\ & - \sum_{j=1}^J \sum_{t=\max(t-d_j+1, EST_j)}^{\min(t, LST_j)} \frac{u_{jf}}{\tilde{d}_j} x_{jt'}; \forall f \in F \end{aligned} \quad (۵)$$

$$I_{f0} = I_{fLST_n+d_n} = 0; \forall f \in F \quad (۶)$$

$$\begin{aligned} Y_{f(k-1)s} y_{fks} & \leq \sum_{i=1}^n \sum_{j=1}^J u_{if} Z_{fksijt} \leq Y_{fks} y_{fks}; \\ \forall f \in F, \forall s \in S, \forall k \in k_{fs}, \forall t = 1, 2, \dots, LST_n \end{aligned} \quad (۷)$$

$$\sum_{k=1}^{k_{fs}} y_{fks} \leq 1; \forall f \in F, \forall s \in S, \forall t = 1, 2, \dots, LST_n \quad (۸)$$

$$\begin{aligned} \sum_{s=1}^S \sum_{k=1}^{k_{fs}} \sum_{t=1}^{T-1} (t+LT_{fs}) Z_{fksijt} & \leq \sum_{t=EST_j}^{LST_j} t x_{jt}; \\ \forall f \in F, \forall j \in J, \forall i = 1, \dots, n \end{aligned} \quad (۹)$$

$$\sum_{j=1}^J \sum_{l=1}^L r_{jl} x_{jt} \leq \widetilde{R_{max,l}} \quad (۱۰)$$

$$\sum_{j=1}^J \sum_{l=1}^L u_{fl} x_{jt} \leq \widetilde{R_{max,f}}; \forall f \in F \quad (۱۱)$$

$$\sum_{t=EST_j}^{LST_j} x_{jt} = 1; \forall j \in J \quad (۱۲)$$

$$\sum_{t=EST_j}^{LST_j} t x_{jt} + \tilde{d}_j \leq T + 1; \forall j \in J \quad (۱۳)$$

$$v_{jt} \leq x_{jt}; \forall j \in J, \forall t = 1, 2, \dots, LST_n \quad (۱۴)$$

$$\sum_{f=1}^F \sum_{j=1}^J u_{f,s,j} \leq cap_s; \forall s \in S \quad (۱۵)$$

$$u_{f,s,j} \leq M \sum_{k=1}^K \sum_{t=1}^{LST_n} \sum_{i=1}^n Z_{fksijt}; \forall f \in F, \forall s \in S, \forall j \in J \quad (۱۶)$$

$$\sum_{f=1}^F \sum_{k=1}^K \sum_{s=1}^S \sum_{j=1}^J c_{fksj} u_{f,s,j} \leq Budget \quad (۱۷)$$

$$\sum_{f=1}^F \sum_{k=1}^K \sum_{s=1}^S \sum_{i=1}^n Z_{fksijt} \leq v_{jt}; \forall j \in J, \forall t = 1, 2, \dots, LST_n \quad (۱۸)$$

$$y_{fks}, x_{jt}, Z_{fksijt} \in \{0,1\} \quad (۱۹)$$

$$I_{ft}, v_{jt}, u_{f,s,j} \geq 0 \quad (۲۰)$$

رابطه (۱) تابع هدف اول مسئله است که به دنبال حداکثر ساختن ارزش خالص فعلی پروژه می‌باشد. رابطه (۲) تابع هدف دوم مدل به دنبال حداکثرسازی درجه زیست‌محیطی با توجه به خرید مواد از تامین کنندگان است. رابطه (۳) تابع هدف سوم مدل به دنبال حداکثرسازی درجه اجتماعی کلی حاصل از خرید مواد از تامین کنندگان می‌باشد. محدودیت (۴) مانع از شروع پروژه بدون منابع تجدیدپذیر کافی می‌باشد. محدودیت (۵) توازن منابع تجدیدناپذیر را حفظ می‌کند. محدودیت (۶) بیان می‌کند که سطح موجودی در شروع دوره و انتهای دوره برای تمامی مواد صفر است. محدودیت (۷) میزان مواد خریداری شده را در دامنه حدود پایینی و فوقانی اعمال تخفیف محدود

برای اجرای این پروژه نمونه و تولید مقادیر پارامترها از توزیع تصادفی یکنواخت استفاده شده است. بعنوان نمونه تولید مقدار پارامتر A_{fs} با توزیع یکنواخت بین (۵۰-۱۰) و تولید مقدار پارامتر r_{jl} با توزیع یکنواخت بین (۳۰-۲۰) صورت گرفته است و سقف میزان مواد مصرفی مورد نیاز ۳۰۰ واحد می باشد و همچنین در هر دوره میزان منابع غیرمصرفی در جدول ۲ ذکر گردیده است و سقف بودجه تعیین شده برای تامین مواد ۱۲۰۰۰۰ واحد در نظر گرفته شده و هزینه جریمه دیرکرد تکمیل پروژه ۱۰۰ واحد و پاداش ۴۰ واحد در نظر گرفته شده و نرخ بهره ۲۰ درصد در نظر گرفته شده است. همچنین نوع روابط پیش نیازی از نوع پایان به شروع در نظر گرفته شده است و از اعداد فازی مثلثی در زمان انجام فعالیت ها استفاده شده است.

جدول ۱: داده های مربوط به فعالیت های مثال

فعالیت	زمان انجام	منابع مصرفی	منابع غیرمصرفی
۱	(15,19,25)	۳۰۰	۲۰
۲	(10,19,25)	۲۰۵	۲۴
۳	(20,28,35)	۲۵۰	۳۰
۴	(10,16,20)	۲۱۰	۲۲
۵	(5,9,15)	۲۲۵	۲۰
۶	(18,24,30)	۲۳۰	۳۰
۷	(10,17,25)	۲۴۰	۲۰
۸	(8,11,15)	۲۵۰	۲۱
۹	(10,15,20)	۲۴۵	۲۴
۱۰	(10,28,35)	۲۳۵	۲۲
۱۱	.	.	.

۳-۱- روش حل

برای حل مدل ریاضی به دلیل چندهدفه بودن از روش محدودیت اِپسیلون تقویت شده ماوروتاس^{۱۲} [۸] استفاده شده است.

این روش در حالت های چندهدفه بودن مدل، جواب های بهینه کارآمد پارتو را ارائه می نماید. در این روش، یکی از توابع هدف به منظور بهینه سازی به عنوان تابع اصلی در نظر گرفته می شود و سایر توابع هدف به عنوان محدودیت در مدل قرار می گیرند. مدل محدودیت اِپسیلون تقویت شده را می توان به صورت زیر نشان داد:

$$\text{Min/Max}(f_1(x) + \vartheta * (\frac{s_2}{r_2} + \frac{s_3}{r_3} + \dots + \frac{s_i}{r_i} \dots + \frac{s_n}{r_n}))$$

St :

$$f_2(x) - s_2 = \varepsilon_2$$

$$f_3(x) - s_3 = \varepsilon_3$$

....

$$i \in [2, n]$$

$$s_i \in R^+$$

(۲۳)

می کند. محدودیت (۸) بیان می کند که کیفیت مواد خریداری شده می تواند حداکثر یک دامنه تخفیف را شامل شود. محدودیت (۹) بیان می کند که پروژه تنها پس از دستیابی به مواد ضروری آغاز می شود. محدودیت (۱۰) نشانگر محدودیت دسترسی به منابع تجدیدپذیر می باشد. محدودیت (۱۱) نشانگر محدودیت دسترسی به منابع تجدیدناپذیر می باشد. محدودیت (۱۲) تضمین می کند که پروژه در یک دوره زمانی انتخاب می شود. محدودیت (۱۳) نشان می دهد که پروژه باید در افق برنامه ریزی خاتمه یابد. محدودیت (۱۴) بیان می کند که اگر پروژه انتخاب نشود امکان اجرای آن وجود ندارد. محدودیت (۱۵) نشانگر محدودیت ظرفیت تامین کننده برای تامین مواد در پروژه می باشد. محدودیت (۱۶) نشان می دهد که اگر تامین کننده ای برای پروژه انتخاب نشود امکان تامین نیز برای آن پروژه وجود ندارد. محدودیت (۱۷) بیانگر محدودیت بودجه برای تامین مواد پروژه می باشد. محدودیت (۱۸) بیان می کند که اگر پروژه برای اجرا انتخاب نشود امکان تامین آن نیز وجود ندارد. محدودیت (۱۹) و (۲۰) نشانگر بازه متغیرهای عدد صحیح و باینری می باشد.

اما برخی روابط فوق دارای پارامتر فازی می باشند. برای تبدیل مدل غیرقطعی فازی به مدل قطعی معادل از روش برش آلفا خیمنز و همکاران^{۱۱} [۵] استفاده شده است. فرض کنید محدودیت به شکل $\tilde{a}_i x \leq \tilde{b}_i$ که دارای پارامترهای فازی $\tilde{a}$ و $\tilde{b}$ می باشد را طبق این روش می توان به شکل رابطه (۲۱) بازنویسی کرد:

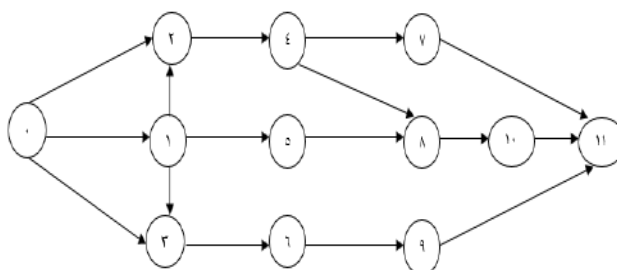
$$[(1 - \alpha)E_2^{a_i} + \alpha E_1^{a_i}]x \leq \alpha E_2^{b_i} + (1 - \alpha)E_1^{b_i} \quad (۲۱)$$

به عنوان مثال، محدودیت (۱۳) طبق روش فوق به محدودیت قطعی به صورت زیر تبدیل شده است.

$$\sum_{t=EST_j}^{LST_j} t x_{jt} + \left[(\alpha) \frac{d_j^0 + d_j^m}{2} + (1 - \alpha) \frac{d_j^p + d_j^m}{2} \right] \leq T + 1 \quad (۲۲)$$

۳-۲ مثال عددی

در این بخش برای اعتبار سنجی مدل پیشنهادی یک مثال تک پروژه ای بعنوان نمونه در نظر گرفته شده است که شامل ۱۰ فعالیت به صورت شبکه گرهی در شکل ۱ در نظر گرفته شده است. در این دیاگرام شبکه پروژه دو فعالیت مجازی ۰ و ۱۱ با میزان زمان و منابع مورد نیاز صفر لحاظ گردیده است. در جدول ۱ داده های مربوط با این مثال نشان داده شده است.



شکل ۱: دیاگرام شبکه مثال

¹² Mavrotas, G.

¹¹ Jiménez, M., Arenas, M., Bilbao, A., & Rodri, M. V.

طبق رابطه فوق راه حل های بهینه پارتو بدست می آیند. سپس با استفاده از رابطه زیر:

$$\varepsilon_i^n = PIS_{fi} + \frac{r_i}{l_i} * \eta \quad (24)$$

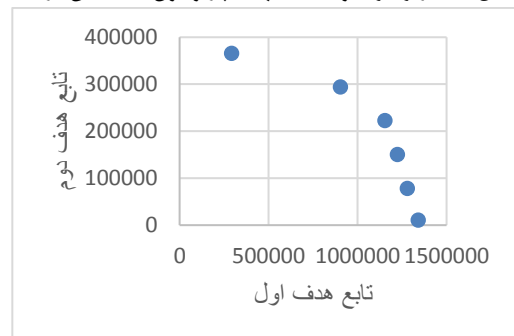
مقدار اپسیلون ها بدست می آید. پس از انجام محاسبات طبق این روش، در نهایت مدل اپسیلون تقویت شده را با استفاده از نرم افزار گمز ورژن ۲۴.۸.۲ و سالور cplex در سیستم شخصی با قدرت پردازندگی ۱۰/۲GHz و حافظه تصادفی در دسترس ۶ Gb برای هر یک از اپسیلون ها بدست آمده حل نموده و مجموعه جواب های بهینه پارتو بدست آمده مطابق جدول ۲ می باشد. مدت زمان اجرای نرم افزار برای حل این مثال نیز برابر ۷۶ ثانیه می باشد.

جدول ۲: مجموعه جواب های بهینه پارتو

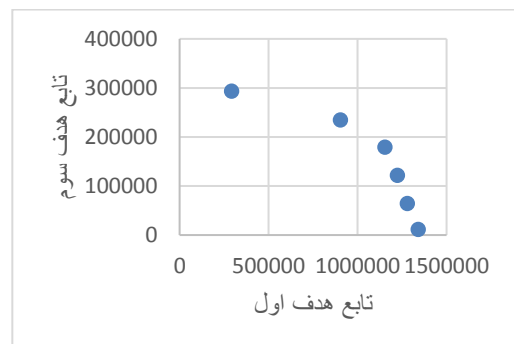
تابع هدف سوم	تابع هدف دوم	تابع هدف اول	مجموعه جواب های پارتویی
121423	۱۵۰۴۳۲	۱۲۲۴۶۱۰	۱
۱۷۹۱۱۲	۲۲۲۳۱۱	۱۱۵۵۱۶۰	۲
۲۳۴۶۹۶	۳۹۴۴۲۳	۹۰۳۸۰۴	۳
۳۹۳۱۵۳	۳۶۵۵۷۶	۲۹۳۱۲۱	۴
۶۳۹۹۷	۷۸۲۶۰	۱۲۸۱۲۱۰	۵
۱۱۰۶۹	۱۰۸۶۱	۱۳۴۲۱۷۸	۶

مقادیر فوق در سطح برش $\alpha = 0/75$ بدست آمده است.

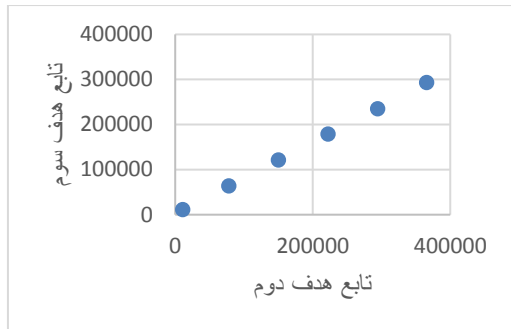
در شکل های زیر نیز، جواب های بهینه پارتو فوق به نمایش درآمده است.



شکل ۲: تضاد بین تابع هدف اول و دوم با توجه به جواب های بهینه پارتو مثال



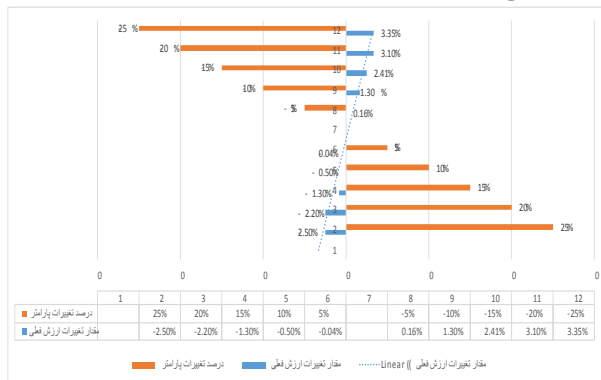
شکل ۳: تضاد بین تابع هدف اول و سوم با توجه به جواب های بهینه پارتو مثال



شکل ۴: تضاد بین تابع هدف دوم و سوم با توجه به جواب های بهینه پارتو مثال

شکل ۲ نشان می دهد که با افزایش ارزش خالص فعلی پروژه، مقدار درجه زیست محیطی کاهش می یابد. شکل ۳ نشان می دهد که با افزایش ارزش فعلی خالص پروژه، مقدار درجه اجتماعی کلی نیز کاهش می یابد. به عبارت دیگر، می توان با پذیرش افزایش هزینه پروژه به عملکرد زیست محیطی و اجتماعی بهتری دست یافت. همچنین با توجه به شکل ۴، بین مقادیر توابع هدف زیست محیطی و اجتماعی یک رابطه مستقیم وجود دارد و می توان یکی از این توابع هدف را به نفع دیگری افزایش داد.

برای بررسی رفتار مدل ارائه شده و بررسی تأثیر یکی از پارامترهای کلیدی بر اهداف مسئله، در ادامه تحلیل حساسیت بر روی پارامتر بودجه انجام شده که نتایج زیر بدست آمده است:



شکل ۵: مقدار تغییرات ارزش خالص فعلی پروژه نسبت به تغییرات در بودجه تعیین شده

با توجه به تحلیل حساسیت صورت گرفته شده، همانطور که در شکل ۵ نشان داده شده است در صورت تغییر سطح بودجه از ۵ الی ۲۵ درصد کاهش بودجه مورد نیاز، ارزش فعلی طرح به میزان ۱۶/۰ الی ۲۵/۳ افزایش می یابد و از سویی دیگر با افزایش ۵ الی ۲۵ درصدی بودجه مصوب، ارزش فعلی تا ۲/۵ درصد کاهش می یابد. در نتیجه، سطح بودجه مصوب بیش از میزان مورد نیاز بوده و ارزش خالص فعلی پروژه را کاهش داده است.

توجه به پیچیدگی مسأله تحقیق، می‌توان از الگوریتم‌های فراابتکاری برای حل مسأله با ابعاد بزرگ استفاده نمود.

مراجع

[۱] عیدی، علیرضا. بخشی، محسن. "زمان‌بندی پروژه‌های منابع محدود چندهدفه فازی با حداقل و حداکثر تأخیرهای زمانی"، چهارمین کنگره مشترک سیستم‌های فازی و هوشمند ایران، ۱۸-۲۰، شهریور، زاهدان، ایران، ۱۳۹۴.

[۲] Birjandi, A., & Mousavi, S. M. "Fuzzy resource-constrained project scheduling with multiple routes": A heuristic solution. *Automation in Construction*, ۱۰۰, ۸۴-۱۰۲, ۲۰۱۹.

[۳] Chawla, V., Chanda, A., Angra, S., & Chawla, G. "The sustainable project management: A review and future possibilities". *Journal of Project Management*, ۱۵۷-۱۷۰, ۲۰۱۸.

[۴] Habibi, F., Barzinpour, F., & Sadjadi, S. J. "A mathematical model for project scheduling and material ordering problem with sustainability considerations: A case study in Iran". *Computers & industrial engineering*, 128, 690-710, ۲۰۱۹.

[5] Jiménez, M., Arenas, M., Bilbao, A., & Rodrı, M. V. "Linear programming with fuzzy parameters: an interactive method resolution". *European journal of operational research*, 177(3), 1599-1609, ۲۰۰۷.

[6] Kadri, R. L., & Bector, F. F. "An efficient genetic algorithm to solve the resource-constrained project scheduling problem with transfer times: The single mode case". *European Journal of Operational Research*, 265(2), 454-462, ۲۰۱۸.

[7] Kim, K. "Generalized Resource-Constrained Critical Path Method to Improve Sustainability in Construction Project Scheduling". *Sustainability*, 12(21), 8918, ۲۰۲۰.

[8] Mavrotas, G. "Effective implementation of the ϵ -constraint method in multi-objective mathematical programming problems". *Applied mathematics and computation*, 213(2), 455-465, ۲۰۰۹.

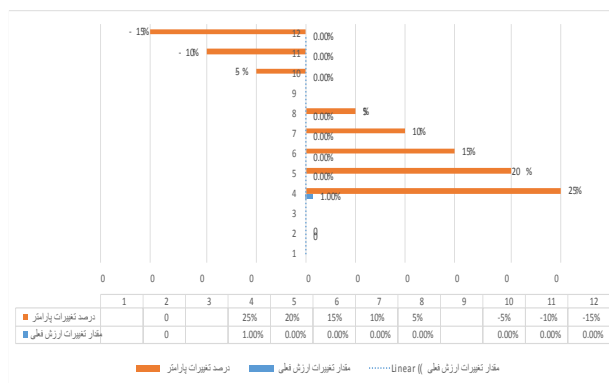
[9] Moradi, M., Hafezalkotob, A., & Ghezavati, V. "Sustainability in fuzzy resource constraint project scheduling in a cooperative environment under uncertainty: Iran's Chitgar lake case study". *Journal of Intelligent & Fuzzy Systems*, 35(6), 6255, ۲۰۱۸.

[10] Hoseini, A., Ghannadpour, S. F., & Hemmati, M. "A comprehensive mathematical model for resource-constrained multi-objective project portfolio selection and scheduling considering sustainability and projects splitting". *Journal of Cleaner Production*, 107-315-324, ۲۰۲۰.

[۱۱] Wiest, J. D., "The scheduling of large projects with limited resources". Ph.D. dissertation, Carnegie Institute of Technology, 11(4), 391-406, ۱۹۶۲.

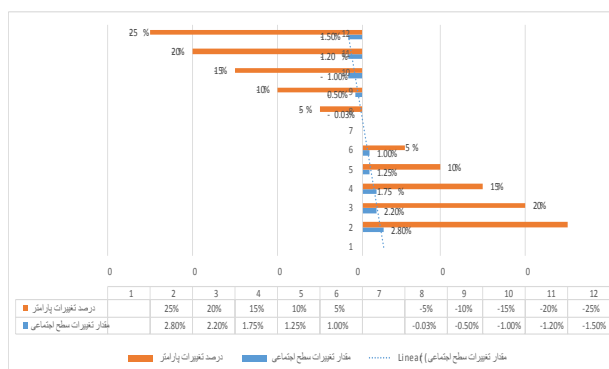
[۱۲] Yousefli, A. "A fuzzy ant colony approach to fully fuzzy resource constrained project scheduling problem". *Industrial Engineering & Management Systems*, 16(3), 307-315, 2017.

[۱۳] Zhang, X., Hipel, K. W., & Tan, Y. "Project portfolio selection and scheduling under fuzzy environment". *Memetic Computing*, 11(4), 391-406, ۲۰۱۹.



شکل ۶: مقدار تغییرات درجه زیست‌محیطی نسبت به تغییرات در بودجه تعیین‌شده

با توجه به تحلیل حساسیت صورت گرفته شده، همانطور که در شکل ۶ نشان داده شده است در اینجا درصدهای تغییراتی که بر روی سطح بودجه در نظر گرفته شده است اثر چندانی بر روی تابع هدف درجه زیست‌محیطی ندارد و هرچه به افزایش ۲۵ درصدی تغییرات در پارامتر بودجه رسیدیم به تدریج نمود تغییرات درجه زیست‌محیطی در حال نمایان شدن است و درجه زیست‌محیطی یک درصد افزایش یافته است.



شکل ۷: مقدار تغییرات درجه اجتماعی نسبت به تغییرات در بودجه تعیین‌شده

با توجه به تحلیل حساسیت صورت گرفته شده، همانطور که در شکل ۷ نشان داده شده است بودجه اثر ویژه‌ای بر روی سطح اجتماعی پروژه دارد بطوریکه با افزایش ۲۵ درصدی سطح بودجه تا ۲/۸ درصد درجه اجتماعی پروژه افزایش یافته و با کاهش سطح بودجه مورد نیاز، درجه اجتماعی نیز تا ۱/۵ درصد کاهش یافته است.

۴- نتیجه گیری

در این مقاله یک مسأله زمان‌بندی پروژه منابع محدود فازی با هدف بیشینه نمودن شاخص‌های پایداری از طریق انتخاب تأمین‌کنندگان بررسی گردید. برای نزدیک شدن مسأله به دنیای واقعی، مساله با زمان‌ها و میزان منابع فازی مدل گردید. همچنین در این مقاله، بحث تخفیفات کمی از سوی تأمین‌کنندگان نیز در مدل لحاظ گردیده است. همچنین از رویکرد برش آلفا برای ارائه مدل قطعی معادل استفاده شد. در پایان نیز با ارائه یک مثال عددی، صحت و اعتبار مدل مورد تجزیه و تحلیل قرار گرفت. برای تحقیقات آینده با



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بوم

مسیریابی مؤثر در شبکه‌های ویژه خودرویی براساس روش‌های پشتیبانی از هواپیماهای بدون سرنشین

حمیده فاطمی دخت^۱، مرجان کوچکی رفسنجانی^۲

^۱دکتری ریاضی کاربردی، بخش علوم کامپیوتر، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، h.fatemidokht@math.uk.ac.ir

^۲دانشیار، بخش علوم کامپیوتر، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان، کرمان، kuchaki@uk.ac.ir

چکیده: شبکه‌های ویژه خودرویی زیر مجموعه‌ای از شبکه‌های ویژه هستند که نیازی به زیرساخت از پیش تعیین شده‌ای ندارند. وسایل نقلیه می‌توانند با وسایل نقلیه مجاور یا تجهیزات ثابت کنار جاده‌ها ارتباط برقرار کنند. کاربردهای مهم این شبکه‌ها عبارتند از: اهداف امنیتی، کنترل ترافیک و ایجاد سرگرمی برای مسافران. به دلیل ویژگی‌های این شبکه‌ها مانند سازماندهی خودمختار، توپولوژی پویا و پهنای باند محدود، فرآیند کشف و نگهداری مسیر یکی از موضوعات تحقیقاتی مهم می‌باشد. در این مقاله با استفاده از همکاری هواپیماهای بدون سرنشین با وسایل نقلیه، پروتکلی برای مناطق شهری پیشنهاد شده است. این پروتکل از هواپیماهای بدون سرنشین برای برقراری ارتباط کارا و امن بین گره‌ها، با وسایل نقلیه استفاده می‌کند. این هواپیماها می‌توانند در برقراری ارتباط مجدد هنگام شکست اتصال همکاری کنند. در این پروتکل از الگوریتم مورچگان برای بهبود مسیریابی بین هواپیماهای بدون سرنشین استفاده شده است. همچنین از این هواپیماها برای پیدا کردن گره‌های مخرب کمک گرفته شده است. برای بررسی کارایی پروتکل پیشنهادی از شبیه‌ساز NS2 استفاده شده است. نتایج شبیه‌سازی نشان می‌دهد که این پروتکل باعث افزایش نرخ تحویل بسته و کاهش تأخیر می‌شود.

کلمات کلیدی: شبکه‌های ویژه خودرویی، شبکه‌های ویژه پرواز، مسیریابی، هوش جمعی.

ایمنی و غیرایمنی تقسیم‌بندی کرد. داده‌های ایمنی بمنظور حفظ جان مسافران منتشر می‌شوند. در حالی که داده‌های غیرایمنی راحتی مسافران را فراهم می‌کند. ارتباطات در شبکه‌های ویژه خودرویی براساس استاندارد DSRC^۳ بوده و شامل دو نوع خودرو با خودرو (V2V)^۴ و خودرو با زیرساخت (V2I)^۵ است. بررسی علت‌های وقوع تصادفات، رفع موانع ایجاد بن‌بست‌های ترافیکی و بهبود کارایی و ایمنی جاده‌ها باعث افزایش تحقیقات در زمینه‌ی شبکه‌های ویژه خودرویی شده است [۱،۲،۳،۴،۵]. شبکه‌های ویژه پروازی نوع جدیدی از شبکه‌های ویژه و زیرمجموعه‌ای از شبکه‌های ویژه خودرویی هستند که هواپیماهای بدون سرنشین یا پهپادها (UAVs)^۶ گره‌های این شبکه را تشکیل می‌دهند. این گره‌ها به

۱- مقدمه

پیشرفت سریع در حوزه تکنولوژی ارتباطات بی‌سیم، تغییرات زیادی در زمینه‌های مختلف در زندگی روزمره ما به وجود آورده است. از مهم‌ترین آن‌ها، شبکه‌های ویژه خودرویی (VANET)^۱ و شبکه‌های ویژه پروازی (FANET)^۲ هستند. شبکه‌های ویژه خودرویی، شبکه‌هایی با معماری بی‌سیم جدید و مجموعه‌ای از وسایل نقلیه می‌باشند به طوری که هر وسیله‌ی نقلیه دارای یک محدوده‌ی رادیویی است و ارتباط بین وسایل نقلیه بدون وجود زیرساخت ثابت انجام می‌شود. ارسال و دریافت پیام‌ها توسط وسایل نقلیه بیشتر برای اهداف امنیتی و کنترل ترافیک از جمله هشدار تصادف، اعلان علائم جاده‌ای، مشاهده ترافیک، تفریحات و سرگرمی و غیره، به کار می‌روند. در واقع می‌توان داده‌ها در شبکه‌های ویژه خودرویی را به داده‌های

^۳ Dedicated Short Range Communication (DSRC)

^۴ Vehicle to Vehicle (V2V)

^۵ Vehicle to Infrastructure (V2I)

^۶ Unmanned Air Vehicles (UAVs)

^۱ Vehicular Ad hoc Networks (VANETs)

^۲ Flying Ad hoc Networks (FANETs)

۲- کارهای مرتبط

اوباتی^۷ و همکارانش [۱۴] یک پروتکل مسیریابی برای شبکه‌های ویژه خودرویی طراحی کرده‌اند که از هواپیماهای بدون سرنشین در شبکه‌ی ویژه پروازی برای بهبود ارسال داده‌ها استفاده می‌کند. این پروتکل ترکیبی از دو پروتکل مسیریابی به نام‌های UVAR-G و UVAR-S می‌باشد. پروتکل UVAR-G یک پروتکل مسیریابی بین وسایل نقلیه و هواپیماهای بدون سرنشین است که از هواپیماهای بدون سرنشین برای انتخاب تدریجی بخش‌های جاده استفاده می‌شود. این هواپیماها اطلاعاتی را در مورد وضعیت اتصال بخش‌های جاده نگهداری می‌کنند. براساس اطلاعات به دست آمده از هواپیماهای بدون سرنشین، برای هر بخش جاده یک رتبه به صورت زیر محاسبه می‌شود:

$$Score_G = \delta \times \left(\frac{R_v}{(1 + \sigma) \times D_w} \right) \quad (1)$$

که R_v محدوده‌ی انتقال وسیله نقلیه‌ی V است و D نشان دهنده‌ی کوتاهترین فاصله بین وسیله نقلیه‌ی اخیر و مقصد می‌باشد. همچنین σ انحراف معیار تراکم‌های هر ناحیه و δ درجه اتصال هر بخش جاده را نشان می‌دهد. بخش جاده‌ای که بیشترین رتبه را به خود اختصاص دهد، برای انتقال بسته‌های داده به مقصد انتخاب خواهد شد. در حالی که، پروتکل UVAR-S یک پروتکل مسیریابی واکنشی است که برای پیدا کردن مسیر بین هواپیماهای بدون سرنشین استفاده می‌شود. نتایج شبیه‌سازی نشان می‌دهد پروتکل مسیریابی پیشنهاد شده نرخ تحویل بسته، کارایی و قابلیت اطمینان مسیریابی را افزایش داده و تأخیر مسیریابی را کاهش می‌دهد.

یو^۸ و همکارانش [۱۵] یک الگوریتم مسیریابی مبتنی بر الگوریتم کلونی مورچگان برای شبکه‌های ویژه پروازی ارائه دادند. با توجه به ویژگی‌های شبکه‌های ویژه پروازی مانند تحرک بالای گره‌ها، تغییر سریع توپولوژی شبکه و فرکانس بالای تبادل اطلاعات الگوریتم‌های مسیریابی سنتی نمی‌توانند ارتباطات کارآمدی را برای هواپیماهای بدون سرنشین فراهم کنند. از این رو یو و همکارانش یک الگوریتم مسیریابی پلی مورفسم براساس الگوریتم مورچگان را معرفی کردند. در این الگوریتم میزان فرومون مسیر که فرایند کشف مسیر را راهنمایی می‌کند، براساس فاصله‌ی یک مسیر، سطح ازدحام مسیر و ثبات و پایداری مسیر محاسبه می‌شود. همچنین این الگوریتم می‌تواند اختلافات شکل‌گیری هواپیماهای بدون سرنشین را تنظیم کرده و از مصالحه-ی عملکرد شبکه جلوگیری کند. نتایج شبیه‌سازی نشان می‌دهد نرخ تحویل بسته، تأخیر و سربار مسیریابی الگوریتم پیشنهادی نسبت به الگوریتم‌های سنتی بهبود یافته‌اند. همچنین این الگوریتم در یک محیط جنگی قابل اعتماد می‌باشد.

کرچه^۹ و همکارانش [۱۰] یک مدل اعتماد با استفاده از هواپیماهای بدون سرنشین برای تشخیص رفتارهای مخرب هوشمند گره‌ها در شبکه‌های ویژه خودرویی پیشنهاد کردند. این مدل شامل یک روش تشخیص آستانه‌ی تطبیقی می‌باشد که تشخیص رفتارهای نادرست هوشمند مانند تغییر و جعل

عنوان مسیریاب در شبکه عمل کرده و امکان ارسال بسته‌ها را فراهم می‌کنند. همچنین پهپادها می‌توانند برای بهبود امنیت و اعتماد شبکه‌های ویژه خودرویی مورد استفاده قرار گیرند [۱۲]. به هر حال، تفاوت‌های خاصی بین شبکه‌های ویژه پروازی و سایر شبکه‌های ویژه وجود دارد. تحرک گره‌ها در شبکه‌های ویژه پروازی خیلی بیشتر از سایر شبکه‌های ویژه می‌باشد. به همین دلیل تغییر توپولوژی در این شبکه‌های بسیار سریع می‌باشد. همچنین به دلیل فاصله‌ی زیاد بین گره‌ها در این شبکه‌ها، برای برقراری ارتباط محدوده‌ی رادیویی باید بیشتر از سایر شبکه‌های ویژه در نظر گرفته شود. به دلیل پیشرفت سریع تکنولوژی در زمینه‌ی فناوری‌های الکترونیکی و ارتباطات، امکان ساخت هواپیماهای بدون سرنشین به وجود آمده است که نوعی وسیله‌ی هدایت‌پذیر از راه دور می‌باشد و می‌تواند بطور خودمختار پرواز کند. به دلیل ویژگی‌های هواپیماهای بدون سرنشین مانند انعطاف‌پذیری، نصب آسان و هزینه‌های عملیاتی نسبتاً کم، استفاده از این وسایل کاربردهای زیادی در زمینه‌های نظامی و غیرنظامی دارد. هواپیماهای بدون سرنشین می‌توانند با یکدیگر ارتباط برقرار کنند و شبکه‌هایی ویژه که به عنوان شبکه‌های ویژه پروازی شناخته می‌شوند را ایجاد کنند [۳]. در این شبکه‌ها فرایند کشف و نگهداری مسیر یکی از موضوعات تحقیقاتی مهم می‌باشد که پروتکل‌های مسیریابی زیادی برای این شبکه‌ها معرفی شده است. همچنین هواپیماهای بدون سرنشین می‌توانند با وسایل نقلیه زمینی در یک شبکه خاص همکاری کنند که باعث بهبود مبادله‌ی اطلاعات بین آن‌ها می‌شود. این رویکرد مزایای متعددی را برای برنامه‌های کاربردی سیستم‌های هوشمند حمل و نقل از جمله عملیات نجات و سنجش از راه دور ارائه می‌دهد [۳، ۸، ۱۴، ۱۶].

به دلیل ویژگی‌ها و کاربردهای شبکه‌های ویژه خودرویی و شبکه‌های ویژه پروازی، در سال‌های اخیر سرمایه‌گذاری‌های عظیمی بر روی این شبکه‌ها ایجاد شده است. به دلیل ویژگی‌های این شبکه‌ها طراحی پروتکل‌های مسیریابی کارا در حمایت از سیستم‌های هوشمند حمل و نقل، بسیار مهم است. در سال‌های اخیر روش‌هایی برای بهبود امنیت مسیریابی در شبکه‌های ویژه ارائه شده است. از جمله این روش‌ها استفاده از الگوریتم‌های هوش جمعی می‌باشد. از الگوریتم‌های هوش جمعی می‌توان به الگوریتم کلونی مورچگان اشاره کرد. الگوریتم کلونی مورچگان براساس رفتار مورچه‌ها در طبیعت که در جستجوی پیدا کردن کوتاهترین مسیر از لانه تا منبع غذا هستند، طراحی شده است. در این مقاله با استفاده از الگوریتم‌های هوش جمعی، پروتکل مسیریابی برای مناطق شهری پیشنهاد شده است. در این پروتکل هواپیماهای بدون سرنشین می‌توانند بمنظور برقراری ارتباط کارا بین گره‌ها، با وسایل نقلیه در شبکه‌های ویژه خودرویی همکاری کرده و بسته‌های داده را ارسال کنند. همچنین این هواپیماها برقراری ارتباط مجدد هنگام شکست اتصال را فراهم می‌کنند.

بخش‌های اصلی مقاله به شرح زیر است. بخش ۲ شامل کارهای مرتبط می‌باشد. در بخش ۳ الگوریتم کلونی مورچگان توضیح داده می‌شود. در بخش ۴ پروتکل پیشنهادی ارائه می‌شود. نتایج شبیه‌سازی در بخش ۵ آورده شده است و بخش ۶ شامل نتیجه‌گیری است.

^۷ Oubbati

^۸ Yu

^۹ Kerrache

داده می‌شود. وزن اولیه‌ی اتصال‌ها می‌تواند با استفاده از فاصله‌ی فیزیکی، اعداد تصادفی یا اعداد بدست آمده از فرمول‌های ریاضی، محاسبه شود. برای یافتن کوتاهترین مسیر میان گره‌ها هر مورچه به طور تصادفی بر روی یک گره قرار گرفته و براساس قانون انتقال گره‌ی بعدی را انتخاب می‌کند. گره‌ی بعدی طبق احتمال زیر به وسیله‌ی مورچه‌ی k -ام انتخاب می‌شود:

$$P_{ij}^k(t) = \begin{cases} \frac{[\tau_{ij}(t)]^\alpha [\eta_{ij}]^\beta}{\sum_{l \in N_k} [\tau_{il}(t)]^\alpha [\eta_{il}]^\beta} & \text{if } j \in N_k \\ 0 & \text{otherwise} \end{cases} \quad (2)$$

که τ_{ij} مقدار فرومون بین گره‌های i و j و η_{ij} اطلاعات اکتشافی بدست آمده می‌باشد. α و β پارامترهایی هستند که اهمیت مربوط به مقدار فرومون را در برابر اطلاعات اکتشافی، کنترل می‌کنند. N_k نیز مجموعه گره‌های ملاقات نشده توسط مورچه‌ی k -ام را مشخص می‌کند.

بروزرسانی مقدار فرومون: همانطور که قبلاً گفته شد، مورچه‌های مصنوعی از رفتار مورچه‌های واقعی در طبیعت برای جستجوی غذا، پیروی می‌کنند. زمانی که یک گره برای ارسال داده مسیری را به یک مقصد خاص نیاز دارد، مورچه‌ها گره‌ی بعدی در مسیر را با استفاده از قانون انتقال انتخاب کرده و گره‌های ملاقات شده را در حافظه‌ی خود ذخیره می‌کنند. زمانی که یک مورچه به مقصد می‌رسد، با استفاده از مسیر معکوس بسمت گره‌ی مبدأ باز می‌گردد و مقدار فرومون اتصال‌ها را با استفاده از رابطه‌ی (۲-۳) بروزرسانی می‌کند. قانون بروزرسانی مقدار فرومون شامل تقویت و تبخیر مقدار فرومون است که نشان دهنده‌ی میزان افزایش و کاهش مقدار فرومون اتصال‌های طی شده توسط یک مورچه می‌باشند. در واقع تبخیر مقدار فرومون به مورچه‌ها کمک می‌کند تا راه‌حل‌های نامناسبی که قبلاً آموخته‌اند را فراموش کنند. بنابراین مورچه‌ها کوتاهترین مسیر از مبدأ تا مقصد را پیدا می‌کنند.

$$\tau_{ij}^{new} = (1 - \rho) \tau_{ij}^{old} + \sum_{k=1}^m \Delta \tau_{ij}^k \quad (3)$$

که m تعداد مورچه‌ها و $\rho \in (0,1)$ ضریب تبخیر فرومون است. $\Delta \tau_{ij}^k$ مقدار فرومون گذاشته شده توسط مورچه‌ی k -ام بر روی اتصال i و j می‌باشد که بصورت زیر محاسبه خواهد شد:

$$\Delta \tau_{ij}^k = \begin{cases} \frac{Q}{f_k} & \text{if the } k\text{th ant traversed link } (i,j) \\ 0 & \text{otherwise} \end{cases} \quad (4)$$

که Q یک عدد ثابت و f_k هزینه‌ی مسیر کامل شده توسط مورچه‌ی k -ام می‌باشد.

هویت را فراهم می‌سازد. در روش پیشنهاد شده به جای ارزیابی اعتماد وسایل نقلیه در زمان‌های مختلف و اندازه‌گیری تغییرات اعتماد ارزیابی شده، صداقت و درستی وسایل نقلیه در کل زمان شبیه‌سازی ارزیابی می‌شود. در نتیجه، در این روش به حداقل تعداد تعاملات و حداقل تراکم در اجرا نیازی نیست که باعث افزایش انعطاف‌پذیری روش پیشنهادی می‌شود. علاوه بر این، روش پیشنهادی از یک روش خوشه‌بندی مبتنی بر هواپیماهای بدون سرنشین استفاده می‌کند که میزان پیام‌های مبادله شده را کاهش داده و باعث حفظ و ذخیره‌ی میزان انرژی هواپیماهای بدون سرنشین می‌شود. نتایج شبیه‌سازی نشان می‌دهد روش پیشنهاد شده دارای نرخ تشخیص بالای ۸۰ درصد حتی با تعداد گره‌های مخرب بیشتر از ۲۰ درصد از گره‌های شبکه، می‌باشد. همچنین این روش با استفاده از هواپیماهای بدون سرنشین حریم خصوصی وسایل نقلیه را حفظ کرده و از استفاده‌ی غیرمجاز از نام مستعارشان جلوگیری می‌کند.

۳- الگوریتم بهینه‌ساز جمعیت مورچگان

در این بخش الگوریتم بهینه‌ساز جمعیت مورچگان (ACO) ^{۱۰} مورد بررسی قرار گرفته است. این الگوریتم جزء موفق‌ترین الگوریتم‌های مبتنی بر ازدحام می‌باشد که اولین بار توسط دوریگو ^{۱۱} در سال ۱۹۹۲ در قالب رساله‌ی دکتری برای حل مسأله‌ی فروشنده‌ی دوره گرد، ارائه شد [۵]. الگوریتم بهینه‌ساز جمعیت مورچگان، که یک روش اکتشافی می‌باشد، از رفتار کاوشگرانه مورچه‌های واقعی الهام گرفته شده است. زمانی که از لانه به سمت غذا، چندین مسیر وجود داشته باشد، مورچه‌ها ابتدا بصورت تصادفی حرکت می‌کنند. در طی مسیر حرکت، آن‌ها ماده‌ای شیمیایی به نام فرومون ^{۱۲} از خود به جا می‌گذارند. مورچه‌های بعدی مسیری را انتخاب می‌کنند که غلظت فرومون بالاتری دارد و همچنین ترشح فرومون توسط مورچه‌ها در مسیر حرکت ادامه پیدا می‌کند. با گذشت زمان، فرومون مسیرهای غیربهینه تبخیر می‌شود. به این ترتیب مورچه‌ها مسیر بهینه از لانه تا غذا را پیدا خواهند کرد [۴]. همانطور که گفته شد، تبخیر شدن فرومون امکان پیدا کردن کوتاهترین مسیر را ایجاد می‌کند. در واقع این ویژگی باعث ایجاد انعطاف در حل هر گونه مسأله‌ی بهینه‌سازی می‌شود. بعضی از کاربردهای الگوریتم بهینه‌ساز جمعیت مورچگان عبارتند از: مسیریابی داخل شهری و بین شهری، مسیریابی بین پست‌های شبکه‌های توزیع برق و لثاژ بالا، مسیریابی شبکه‌های کامپیوتری و مسیریابی تأمین مواد اولیه جهت تولید بهنگام.

برای پیاده‌سازی الگوریتم بهینه‌ساز جمعیت مورچگان از مورچه‌های مصنوعی به عنوان عناصری در بهینه‌سازی استفاده می‌شود، که برای این مورچه‌های مصنوعی یک حافظه در نظر گرفته می‌شود که مسیرهای حرکت را نگهداری می‌کند. گام‌های اصلی این الگوریتم عبارتند از [۷]:

- **مقداردهی اولیه:** الگوریتم بهینه‌ساز جمعیت مورچگان می‌تواند

برای حل مسائل گسسته استفاده شود که با استفاده از یک گراف با N گره و L اتصال نشان داده می‌شوند. هر گره شامل تعدادی از مورچه‌های مصنوعی می‌باشد و به هر اتصال یک وزن اختصاص

^{۱۰} Ant Colony Optimization (ACO)

^{۱۱} Dorigo

^{۱۲} Pheromone

۴- معرفی پروتکل مسیریابی پیشنهادی: UVRP^{۱۳}

در این بخش، پروتکل مسیریابی پیشنهادی برای مناطق شهری توضیح داده می‌شود. در این پروتکل هواپیماهای بدون سرنشین بمنظور برقراری ارتباط کارا بین گره‌ها، با وسایل نقلیه در شبکه‌های ویژه خودرویی همکاری کرده و بسته‌های داده را ارسال کنند. همچنین این هواپیماها برقراری ارتباط مجدد هنگام شکست اتصال را فراهم می‌کنند. پروتکل UVRP یک پروتکل مسیریابی واکنشی^{۱۴} است که با استفاده از الگوریتم کلونی مورچگان انجام می‌شود. این پروتکل مناسب‌ترین مسیر بین هواپیماهای بدون سرنشین را پیدا می‌کند به طوری که نرخ تحویل بسته را افزایش و تأخیر مسیریابی را کاهش می‌دهد.

۴-۱- فرضیات

در این پروتکل فرض شده است که وسایل نقلیه و هواپیماهای بدون سرنشین مجهز به یک سیستم تعیین موقعیت جهانی^{۱۵} هستند که برای به دست آوردن موقعیت جغرافیایی گره‌ها در شبکه استفاده شده است. فرض دیگر پروتکل پیشنهادی این است که هواپیماهای بدون سرنشین در ارتفاع ثابت و کم ارتفاع پرواز می‌کنند تا بتوانند با وسایل نقلیه ارتباط برقرار کنند و از IEEE 802.11p با محدوده ارتباطی بیش از ۱۰۰۰ متر استفاده می‌کنند. همچنین، هر چهار راه توسط یک هواپیمای بدون سرنشین، پوشش داده می‌شود.

۴-۲- فرآیند کشف مسیر

زمانی که یک وسیله نقلیه مسیری را برای ارسال داده به مقصد خاص نیاز داشته باشد، اما مسیری به مقصد موردنظر در جدول مسیریابی این گره وجود نداشته باشد، فرآیند کشف مسیر آغاز می‌شود. گرهی مبدأ برای کشف دنباله‌ای از هواپیماهای بدون سرنشین تا مقصد و موقعیت جغرافیایی آن‌ها، بسته‌ی Request-Ant را به در شبکه پخش فراگیر می‌کند. بسته‌ی Request-Ant شامل شناسه‌ی مبدأ، شناسه‌ی مقصد، لیستی از شناسه‌ی هواپیماهای بدون سرنشین میانی، طول، عرض و ارتفاع جغرافیایی هواپیماهای بدون سرنشین، تعداد گام، نوع و طول عمر می‌باشد. در این بسته، مقدار مربوط به فیلد نوع، ۱ در نظر گرفته می‌شود.

زمانی که یک هواپیماهای بدون سرنشین بسته‌ی کنترلی را دریافت می‌کند، فیلد نوع در این بسته را مورد بررسی قرار می‌دهد. اگر بسته‌ی کنترلی Request-Ant باشد، فیلد مربوط به لیستی از شناسه‌ی هواپیماهای بدون سرنشین میانی را بررسی می‌کند. اگر شناسه‌ی خود را در این لیست پیدا کرد، این بسته را نادیده می‌گیرد. در غیر این صورت، مقدار فیلد تعداد گام را یکی افزایش داده و شناسه‌ی خود را به لیست شناسه‌ی هواپیماهای بدون سرنشین میانی در بسته‌ی Request-Ant الحاق می‌کند. سپس این گره بسته را در شبکه پخش فراگیر می‌کند. این فرآیند توسط همه‌ی هواپیماهای بدون سرنشین میانی انجام خواهد شد.

^{۱۳} UAV-assisted VANET Routing Protocol (UVRP)

^{۱۴} Reactive

^{۱۵} Global Positioning System (GPS)

زمانی که بسته‌ی Request-Ant به مقصد می‌رسد شامل لیست هواپیماهای بدون سرنشین در طول مسیر می‌باشد. گرهی مقصد ابتدا فیلد تعداد گام را افزایش داده و مقدار فیلد نوع را از یک به صفر تغییر می‌دهد. در واقع بسته‌ی Request-Ant به بسته‌ی Reply-Ant تبدیل می‌شود. گرهی مقصد سپس معکوس لیست شناسه‌ی هواپیماهای بدون سرنشین میانی را در بسته‌ی Reply-Ant قرار داده و آن را بصورت تک بخشی به سمت گرهی مبدأ ارسال می‌کند.

زمانی که گرهی مبدأ بسته‌ی Reply-Ant را دریافت می‌کند، مسیر ذخیره شده در لیست شناسه‌ی هواپیماهای بدون سرنشین میانی و مقدار فیلد تعداد گام را در جدول مسیریابی خود ذخیره می‌کند. سپس مقدار فرومون مسیر توسط گرهی مبدأ با استفاده از رابطه‌ی زیر محاسبه می‌شود:

pheromone value of a route

$$= \prod_{i=0}^{HC-1} [*] \frac{R_U}{distance(u_i, u_{i+1})} \quad (5)$$

که HC و R_U به ترتیب مقدار فیلد تعداد گام و محدوده‌ی انتقال هواپیماهای بدون سرنشین می‌باشند. $distance(u_i, u_{i+1})$ به فاصله‌ی بین دو گرهی متوالی با شروع از گرهی مبدأ، اشاره دارد. مقدار فرومون مسیر برای انتخاب بهترین مسیر، زمانی که بیش از یک مسیر بین مبدأ و مقصد وجود داشته باشد، استفاده می‌شود. در واقع گرهی مبدأ مسیری که بیشترین مقدار فرومون را داشته باشد برای ارسال داده انتخاب می‌کند.

زمانی که شکست اتصال اتفاق بیافتد، اولین وسیله‌ی نقلیه یا هواپیمای بدون سرنشینی که این قطع اتصال را تشخیص داده است، مسیرهای موجود که قبلاً در گره‌های میانی نگه داشته شده اند را بررسی کرده و سعی می‌کند مسیر دیگری را پیدا کند. اگر مسیر دیگری پیدا شد، بسته‌های داده از طریق آن مسیر ارسال می‌شوند. در غیر این صورت، یک پیام خطا به گرهی مبدأ ارسال شده و این گره یک فرآیند کشف مسیر جدید را شروع می‌کند.

۵- شبیه‌سازی و ارزیابی نتایج

در این بخش، نتایج به دست آمده از ارزیابی پروتکل پیشنهادی را نشان داده خواهیم داد.

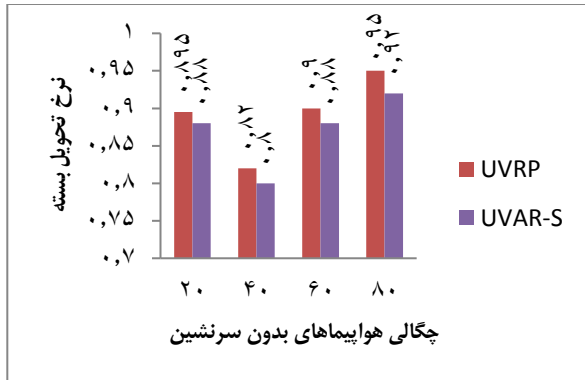
۵-۱- تنظیمات محیط شبیه‌سازی

برای انجام شبیه‌سازی پروتکل‌های پیشنهادی و ارزیابی کارایی آن‌ها از نرم‌افزار NS-2.35 تحت لینوکس Ubuntu 12.04 استفاده می‌کنیم. برای ایجاد مدل تحرک وسایل نقلیه از مولدهای مدل حرکت برای شبکه‌های ویژه خودرویی به نام MOVE و VanetMobiSim [۹] استفاده کرده‌ایم. MOVE براساس زبان برنامه‌نویسی جاوا است و مدل حرکتی ایجاد شده توسط SUMO [۱۱] را تولید خواهد کرد. MobiSim [۱۳] برای ایجاد مدل تحرک هواپیماهای بدون سرنشین مورد استفاده قرار گرفته است. پارامترهای پیش‌فرض برای انجام شبیه‌سازی در جدول ۱ ارائه شده است. در این شبیه‌سازی، یک سناریوی شهری 4000×4000 مترمربع، شامل ۴۰ جاده دوطرفه و ۲۵ تقاطع ایجاد شده است. همچنین فرض شده است که

هواپیماهای بدون سرنشین در ارتفاعی ثابت که از ۲۰۰ متر بیشتر نمی‌شود، پرواز می‌کنند.

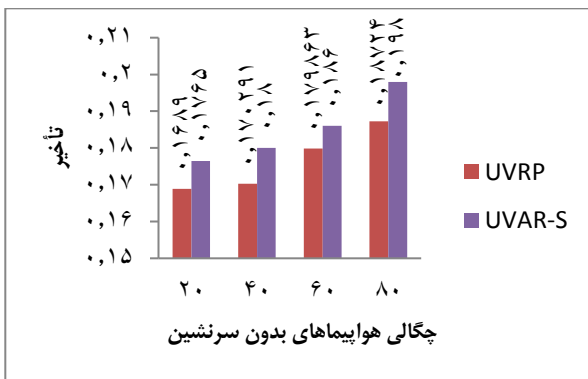
جدول ۱: پارامترهای شبیه‌سازی

پارامتر شبیه‌سازی	مقدار
زمان شبیه‌سازی	۱۸۰ ثانیه
محدوده‌ی ناحیه‌ی شبیه‌سازی	۴۰۰×۴۰۰ مترمربع
مدل ترافیکی	CBR
محدوده‌ی انتقال وسایل نقلیه	۳۰۰ متر
محدوده‌ی انتقال هواپیماهای بدون سرنشین	۱۰۰۰ متر
اندازه‌ی بسته‌ها	یک کیلو بایت
توپولوژی	منطق شهری
مولد تحرک و پویایی	[VanetMobiSim]
تعداد وسایل نقلیه	۱۰۰ تا ۳۰۰
تعداد هواپیماهای بدون سرنشین	۲۰ تا ۸۰
سرعت وسایل نقلیه	۰ تا ۵۰ کیلومتر بر ساعت
سرعت هواپیماهای بدون سرنشین	۰ تا ۶۰ کیلومتر بر ساعت
پهنای باند	۱ مگا بیت بر ثانیه
MAC/PHY	IEEE 802.11p
تعداد جاده‌ها	۴۰
تعداد تقاطع‌ها	۲۵
تعداد شبیه‌سازی اجرا شده	۱۰



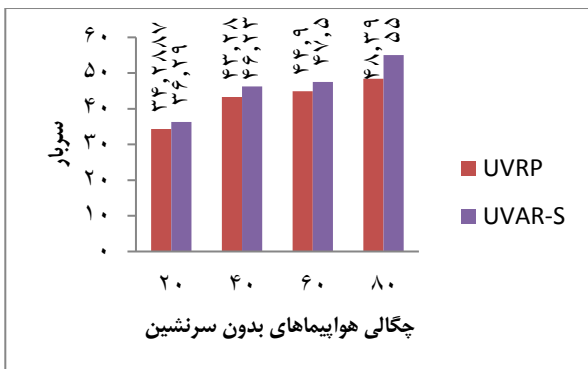
شکل ۱: نرخ تحویل بسته به عنوان تابعی از چگالی هواپیماهای بدون سرنشین

در شکل ۲ معیار تأخیر به عنوان تابعی از چگالی هواپیماهای بدون سرنشین، نشان داده شده است، که تأخیر پروتکل پیشنهادی از پروتکل مقایسه شده، کمتر است.



شکل ۲: تأخیر به عنوان تابعی از چگالی هواپیماهای بدون سرنشین

شکل ۳ نشان دهنده‌ی معیار سربار است. همانطور که در این شکل مشاهده می‌شود، زمانی که تعداد هواپیماهای بدون سرنشین افزایش می‌یابد، سربار پروتکل UVRP نیز زیاد می‌شود. این عمدتاً به دلیل بسته‌های مسیریابی اضافی (Request-Ant و Request-Ant) استفاده شده است.



شکل ۳: سربار به عنوان تابعی از چگالی هواپیماهای بدون سرنشین

۲-۵- نتایج شبیه‌سازی

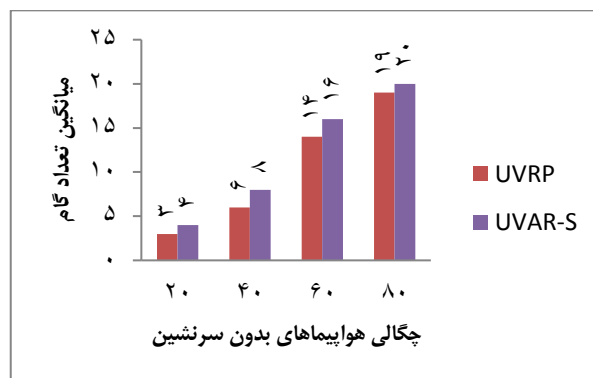
از معیارهای زیر برای مقایسه کارایی پروتکل پیشنهادی استفاده شده است.

- **نرخ تحویل بسته:** نسبت تعداد بسته‌هایی که به طور موفق توسط گرهی مقصد دریافت شده به تعداد کل بسته‌هایی که از هر گره برای آن مقصد ارسال می‌شود.
- **تأخیر:** اختلاف بین زمانی که یک بسته‌ی داده به وسیله گرهی مقصد دریافت شده و زمانی که این بسته توسط گرهی مبدأ فرستاده شده است. این اختلاف ناشی از این است که گرهی مبدأ به دلیل نداشتن مسیری به مقصد، نیاز به انجام فرآیند کشف مسیر دارد و بسته‌ی داده باید تا کشف مسیر در گرهی مبدأ منتظر بماند.
- **سربار مسیریابی:** نسبت تعداد بسته‌های منحصر به فرد ارسال شده به کل پیام‌هایی که در شبکه انتقال داده می‌شوند.
- **میانگین تعداد گام‌ها:** مجموع تعداد بسته‌های داده که به طور موفقیت آمیزی به مقصد منتقل شده‌اند تقسیم بر تعداد کل گام-هایی که توسط تمام بسته‌ها طی شده است.

نتایج شبیه‌سازی در شکل‌های ۱ تا ۴ نشان داده شده‌اند. شکل ۱ نرخ تحویل بسته را تحت چگالی‌های مختلف هواپیماهای بدون سرنشین، نشان می‌دهند. همانطور که در این شکل مشخص است، زمانی که تعداد هواپیماهای بدون سرنشین افزایش می‌یابد، نرخ تحویل بسته‌ی پروتکل UVRP نیز بالاتر می‌رود؛ که به دلیل استفاده از الگوریتم بهینه ساز جمعیت مورچگان برای انتخاب مسیر مناسب بین این هواپیماها، می‌باشد.

- [3] Bekmezci, I, Sahingoz, O.K, Temel, S., "Flying Ad-Hoc Networks (FANETs): A survey", Ad Hoc Networks, Vol. 11, pp. 1254-1270, 2013.
- [4] Dorigo, M, Stützle, T., *Ant colony optimization: overview and recent advances*, Handbook of Metaheuristics, Springer, pp. 227-263, 2010.
- [5] Dorigo, M., *Optimization, learning and natural algorithms*, Ph.D. Dissertation, Department of Electronics, Politecnico di Milano, Italy, 1992.
- [6] Fatemidokht, H, Kuchaki Rafsanjani, M., "QMM-VANET: An efficient clustering algorithm based on QoS and monitoring of malicious vehicles in vehicular ad hoc networks," The Journal of Systems and Software, Vol. 165, Art. no. 110561, 2020.
- [7] Fatemidokht, H, Kuchaki Rafsanjani, M., "F-Ant: An effective routing protocol for ant colony optimization based on fuzzy logic in vehicular ad hoc networks", Neural Computing and Applications, Vol. 29, pp. 1127-1137, 2018.
- [8] Hadiwardoyo, S.A, Hernández-Orallo, E, Calafate, C.T, Cano, J.C., "Experimental characterization of UAV-to-car communications", Computer Networks, Vol. 136, pp. 105-118, 2018.
- [9] Härrri, J, Filali, F, Bonnet, C, Fiore, M., "VanetMobiSim: generating realistic mobility patterns for VANETs", Proceedings of the 3rd International ACM Work- shop on Vehicular Ad Hoc Networks (VANET 2006), pp. 96-97, 2006.
- [10] Kerrache, C.A, Lakas, A, Lagraa, N, Barka, E., "UAV-assisted technique for the detection of malicious and selfish nodes in VANETs", Vehicular Communications, Vol. 11, pp. 1-11, 2018.
- [11] Krajzewicz, D, Erdmann, J, Behrisch, M, Bieker, L., "Recent development and applications of SUMO—Simulation of Urban Mobility", International Journal of Advances in Systems and Measurements, Vol. 5, pp. 128-138, 2012.
- [12] Mirsadeghi, F, Kuchaki Rafsanjani, M, Gupta, B. B., "A trust infrastructure based authentication method for clustered vehicular ad hoc networks", Peer Peer Netw., pp. 1-17, 2020.
- [13] Mousavi, S.M, Rabiee, H.R, Moshref, M, Dabirmoghaddam, A., "Mobisim: a framework for simulation of mobility models in mobile ad-hoc networks", Proceedings of the Third IEEE International Conference on Wireless and Mobile Computing, Networking and Communications (WiMOB), White Plains, NY, USA, 2007.
- [14] Oubbati, O.S, Lakas, A, Zhou, F, Günes, M, Lagraa, N, Yagoubi, M.B., "Intelligent UAV-assisted routing protocol for urbanVANETs", Computer Communications, Vol. 107, pp. 93-111, 2017.
- [15] Soares, V.N.G.J, Rodrigues, J.J.P.C., "Vehicular delay-tolerant networks", Advances in Delay-Tolerant Networks (DTNs), Second Edition, pp. 59-78, 2021.
- [16] Yu, Y, Ru, L, Chi, W, Liu, Y, Yu, Q, Fang, K., "Ant colony optimization based polymorphism-aware routing algorithm for ad hoc UAV network", Multimedia Tools and Applications, Vol. 75, pp. 14451-14476, 2016.
- [17] Zafar, W, Khan, B.M., "A reliable, delay bounded and less complex communication protocol for multicluster FANETs", Digital Communications and Networks, Vol. 3, pp. 30-38, 2017.

شکل ۴ میانگین تعداد گام را براساس چگالی‌های مختلف هواپیماهای بدون سرنشین، نشان می‌دهد. از شکل ۴ مشخص می‌شود که پروتکل UVRP با کمک استفاده از هواپیماهای بدون سرنشین برای ارائه بسته‌های داده، بهتر عمل می‌کند. به طوری که میانگین تعداد گام‌ها را با اجتناب از گام‌های غیر ضروری به حداقل می‌رساند، به ویژه هنگامی که شبکه بسیار متراکم است.



شکل ۴: میانگین تعداد گام به عنوان تابعی از چگالی هواپیماهای بدون سرنشین

۶- نتیجه

شبکه‌های ویژه خودرویی و شبکه‌های ویژه پرواز از مهم‌ترین کاربردهای شبکه‌های ویژه هستند. هواپیماهای بدون سرنشین که گره‌های تشکیل دهنده شبکه‌های ویژه پرواز هستند، می‌توانند با وسایل نقلیه زمینی در شبکه‌های ویژه خودرویی همکاری کرده و باعث بهبود مبادله‌ی اطلاعات بین گره‌ها شوند. در این مقاله پروتکل UVRP برای مسیریابی در مناطق شهری پیشنهاد شده است. این پروتکل برای ارتباط بین هواپیماهای بدون سرنشین از الگوریتم کلونی مورچگان استفاده می‌کند. در این پروتکل از هواپیماهای بدون سرنشین برای انتخاب مسیرهای مناسب برای انتقال بسته‌های داده هنگامی که تراکم وسایل نقلیه برای مسیریابی بسته‌ها از طریق وسایل نقلیه کافی نیست، استفاده می‌شود. نتایج شبیه‌سازی نشان می‌دهد که پروتکل UVRP برای مناطق شهری مناسب بوده و در مقایسه با سایر پروتکل مسیریابی بررسی شده، نسبت تحویل بسته را بهبود می‌دهد. همچنین این پروتکل سربار و تاخیر را کاهش می‌دهد.

با توجه به وجود گره‌های مخرب در شبکه‌های ویژه، پروتکل پیشنهادی می‌تواند برای تشخیص این گره‌ها در شبکه با استفاده از روش‌های محاسبات نرم، گسترش داده شود.

مراجع

- [۱] کاکاوند میرزایی، مریم و سیفعلی هرسینی، جلیل، "طراحی یک مکانیسم تدافعی برای بهبود امنیت در لایه فیزیکی با رویکرد نظریه بازی‌ها، کاربرد در شبکه‌های اقتضایی خودرویی"، مجله مهندسی برق دانشگاه تبریز، دوره ۴۷، شماره ۱، صفحات ۲۱۱-۲۲۰، ۱۳۹۶.
- [2] Al-Sultan, S, Al-Doori, M. M, Al-Bayatti, A. H, Zedan, H., "A comprehensive survey on vehicular Ad Hoc network", Network and Computer Applications, Vol. 37, pp. 380-392, 2014.



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

یک رویکرد ترکیبی مبتنی بر واژه‌نامه و یادگیری ماشین برای شناسایی هیجان در نظرهای کاربران فارسی زبان

وحید کیانی^۱، مهدی رسولی^۲

^۱ استادیار، گروه مهندسی کامپیوتر، دانشگاه بجنورد، بجنورد، v.kiani@ub.ac.ir
^۲ دانشجوی کارشناسی ارشد، گروه مهندسی کامپیوتر، دانشگاه بجنورد، بجنورد، mrasouli.edu@gmail.com

چکیده: شناسایی هیجان در نظرها و متون کاربران فارسی زبان یک مسئله مهم در پردازش متن است که می‌تواند برای تحلیل بازخورد و رضایت مشتریان از محصولات مورد استفاده قرار گیرد. در این مقاله، ما رویکردهای مبتنی بر واژه‌نامه و مبتنی بر یادگیری ماشین را برای شناسایی قطبیت و هیجان در نظرهای فارسی کاربران با هم ترکیب می‌کنیم. در روش پیشنهادی، بردار ویژگی نرخ هیجان که به کمک واژه‌نامه هیجانی محاسبه شده است، با بردار ویژگی رخداد کلمات ترکیب شده و یک بردار ویژگی ترکیبی با ۲۴۳۰۹ عنصر را تشکیل می‌دهند. سپس مهم‌ترین ویژگی‌ها در این بردار ویژگی ترکیبی به عنوان ورودی به یک طبقه‌بند ماشین بردار پشتیبان داده شده‌اند. همچنین، از آنجایی که در طبقات مختلف هیجانی تعداد نمونه‌های مثبت و منفی یکسان نیست، برای متوازن‌سازی نمونه‌ها حین آموزش از روش SMOTE استفاده شده است. استفاده از نرخ هیجان‌ات مبتنی بر واژه‌نامه در کنار ویژگی‌های رخداد کلمات در هر نظر، و متوازن‌سازی نمونه‌های طبقات مختلف با کمک روش SMOTE باعث شده است تا روش ترکیبی پیشنهادی نسبت به رویکردهای ساده‌تر نرخ دقت طبقه‌بندی بالاتری را فراهم آورد. روش پیشنهادی به طور میانگین به نرخ دقت ۶۶٪ برای معیار F1 در شناسایی کلاس‌های قطبیت و هیجان دست یافته است.

کلمات کلیدی: شناسایی هیجان، شناسایی قطبیت، تحلیل متن، یادگیری ماشین، تحلیل نظرهای متنی.

۱- مقدمه

ترکیبی قابل تقسیم هستند. در رویکرد تشخیص هیجان با کمک واژه‌نامه هیجانی، واژه‌های متن بعد از پیش‌پردازش در یک واژه‌نامه هیجانی جستجو می‌شوند، تا ضرایب هیجانی هر واژه در صورت وجود از واژه‌نامه استخراج گردد. سپس سطوح هیجان در متن نظر با در نظر گرفتن ضرایب هیجانی واژه‌های آن نظر محاسبه می‌شوند. به عنوان مثال، تاج و همکاران در [۴] برای تحلیل قطبیت اخبار BBC از واژه‌نامه SentiWordNet استفاده کرده‌اند و بر اساس میانگین ضرایب قطبیت کلماتی که وزن TF-IDF بالایی در هر خبر دارند، قطبیت خبر را تعیین کرده‌اند. کومار روت و همکارانش در [۵] از یک روش بدون نظارت در کنار واژه‌نامه برای تعیین قطبیت توثی‌های انگلیسی کاربران استفاده کرده‌اند. برای واژه‌های خاصی که در واژه‌نامه موجود نبوده‌اند، نرخ مثبت یا منفی بودن واژه با جستجوی آن واژه در کنار واژه‌های هم معنی با کلمه excellent و poor در موتور جستجوی گوگل تعیین شده است. سرینواسان و همکارانش در [۶] برای طبقه‌بندی هیجان و قطبیت توثی‌های کاربران ایالت‌های مختلف آمریکا در طول انتخابات سال ۲۰۱۶ این

امروزه گسترش سریع اینترنت باعث شده است متن به ر سانه مهمی برای تبادل اطلاعات بین کاربران در شبکه‌های اجتماعی، وبسایت‌های خرید و فروش، و سایت‌های خبری تبدیل شود [۱]. به خصوص در وبسایت‌های فروش کالا، مشتریان بازخوردها و دیدگاه‌های خود را از راه دور در مورد کالا و خدمات شرکت از طریق پست‌ها، نظرات، و پیام‌های متنی بیان می‌کنند [۲]. شرکت‌ها و سازمان‌ها می‌توانند با تحلیل قطبیت و تحلیل هیجان در نظرهای مشتریان خود به اطلاعات مفیدی در خصوص ویژگی‌های محصولات خود، استراتژی‌های مناسب بازاریابی، و اثرات تصمیم‌های شرکت بر بازار دست یابند [۳].

روش‌های تشخیص هیجان در متن به سه گروه کلی تشخیص هیجان با کمک واژه‌نامه هیجانی، تشخیص هیجان با یادگیری ماشین، و روش‌های

کشور از واژه‌نامه EmoLex استفاده کرده و نتایج پیش‌بینی به‌دست‌آمده را با نتایج واقعی انتخابات ۲۰۱۶ آمریکا در ایالت‌های مختلف مقایسه کرده‌اند. گر شاسی و همکاران در [۷] مسئله تشخیص هیجان را در متون فارسی در نظر گرفته، و یک رویکرد هیورستیک مبتنی بر رخداد کلمات کلیدی هیجانی، تشدیدکننده‌ها، نفی‌کننده‌ها، و ارکان جمله را برای تعیین هیجان بر اساس واژه‌نامه به کار بسته‌اند. عملکرد روش پیشنهادی روی یک مجموعه داده کوچک شامل ۱۰۰ نظر فارسی رضایت‌بخش بوده، و به نرخ طبقه‌بندی ۸۰٪ در تشخیص کلاس‌های هیجان رسیده است.

یک رویکرد دیگر برای تشخیص هیجان در متن، استفاده از مدل‌های یادگیری ماشین است. به عنوان مثال، کومار روت و همکارانش در [۵]، علاوه بر تشخیص قطبیت توئیت‌ها به کمک واژه‌نامه، طبقه‌بندی هیجان توئیت‌ها را به کمک روش‌های یادگیری ماشین با نظارت SVM و Naïve Bayes نیز انجام داده‌اند. آزمایش‌های ایشان نشان داده است که برای توئیت‌های انگلیسی مدل بیز ساده چندمتغیره می‌تواند به نرخ F1 بالای ۹۰٪ برای طبقات هیجان مختلف دست یابد. دقت این روش زمانی که از توالی‌های تک کلمه‌ای و توالی‌های دو کلمه‌ای در کنار هم برای استخراج ویژگی استفاده شده بیشتر از زمانی بوده که از توالی‌های تک کلمه‌ای، یا دو کلمه‌ای به تنهایی استفاده شود. استفاده از مدل‌های یادگیری عمیق نیز برای تشخیص هیجان در متون فارسی بررسی شده است. صادقی و همکاران در [۸] از مدل یادگیری عمیق GRU برای شناسایی هیجان در دسته‌های شادی، ناراحتی، عصبانیت، ترس، و تعجب بر روی یک مجموعه از متون فارسی شامل ۲۳ هزار سند فارسی استفاده کرده‌اند، و به نرخ F1 حدود ۸۷٪ رسیده‌اند.

بعضی محققین از ترکیب رویکرد مبتنی بر واژه‌نامه با رویکرد مبتنی بر یادگیری ماشین استفاده کرده‌اند تا روش کامل‌تری را ایجاد کرده و به دقت مناسب دست یابند. به عنوان مثال، خسروی و همکاران در [۹] از ماشین بردار پشتیبان (SVM) به عنوان طبقه‌بند استفاده کرده‌اند تا هیجان و قطبیت خبرهای فارسی را به کمک واژه‌نامه هیجانی EmoLex شناسایی کنند. شش ویژگی آماری هیجانی بر اساس تعداد واژگان هر خبر در دسته‌های هیجانی و قطبیت مختلف محاسبه شده‌اند. این ویژگی‌های هیجانی در کنار تعدادی ویژگی زمینه‌ای که بر اساس تعداد تکرار واژه‌های هر نظر، وجود تشدیدکننده، وجود علائم نگارشی، و جایگاه واژه‌ها در جمله محاسبه شده‌اند، به عنوان ویژگی‌های ورودی به طبقه‌بند ماشین بردار پشتیبان داده شده‌اند. خسروی و همکاران روش پیشنهادی خود را تنها بر روی ۱۰۰ نمونه داده خبر فارسی ارزیابی کرده‌اند، و تنها برای هیجان شادی به نرخ F1 مناسب یعنی ۸۰٪ دست یافته‌اند. برای سایر کلاس‌های هیجان نرخ F1 بین ۱۳ تا ۴۲ درصد بوده است. صادقی و همکاران در [۸] از ترکیب ویژگی‌های زبان‌شناسی، یادگیری عمیق، و یادگیری غیرعمیق برای شناسایی هیجان در مجموعه بزرگی از متون فارسی استفاده کرده‌اند، و به نتایج بسیارخوبی دست یافته‌اند. در بخش ویژگی‌های زبان‌شناسی، رخداد کلمات کلیدی کلاس‌های هیجانی مختلف، رخداد عبارات کلیدی هیجانی در نقش‌های مختلف مانند صفت و قید و ...، و رخداد عبارت‌های غیررسمی هیجانی مثل ضرب المثله‌ها در جمله بررسی شده‌اند. در بخش یادگیری عمیق از مدل GRU^۱ استفاده شده است، که نسخه پیشرفته‌تری از مدل LSTM است. در نهایت خروجی مدل عمیق،

در کنار ویژگی‌های زبان‌شناسی به عنوان ورودی به یک طبقه‌بند غیرعمیق SVM داده شده است. برای ارزیابی از یک مجموعه داده شامل ۲۳ هزار سند فارسی با طول متوسط ۲۴ کلمه استفاده شده است. روش پیشنهادی ایشان با استفاده از ویژگی‌های زبان‌شناسی به نرخ دقت F1 حدود ۷۴٪، رویکرد یادگیری عمیق به نرخ دقت F1 حدود ۸۷٪، و ترکیب این دو به نرخ دقت F1 حدود ۹۳٪ در تشخیص هیجان‌های شادی، ناراحتی، عصبانیت، ترس، و شگفتی رسیده است.

تحقیقات بسیار کمی در مورد تشخیص هیجان در متون زبان فارسی انجام شده است. در این مقاله ما دو مسئله تشخیص قطبیت و تشخیص هیجان را در متون نظرات فارسی کاربران در نظر گرفته، و با ترکیب رویکرد واژه‌نامه با رویکرد یادگیری ماشین آن‌ها را به طور همزمان حل خواهیم کرد. نسبت به تحقیق انجام شده در [۹]، روش پیشنهادی ما بر روی متون نظرات فارسی به جای اخبار فارسی کار خواهد کرد. نظرات کاربران فارسی زبان نسبت به اخبار فارسی متون چالشی‌تری هستند، زیرا کاربران معمولاً در بیان نظرات خود از واژه‌های عامیانه استفاده می‌کنند و قواعد نحوی را در سطح ضعیف‌تری رعایت می‌کنند، در حالی که در متون اخبار از شکل رسمی واژه‌ها استفاده می‌شود. از طرف دیگر، در این تحقیق ما از مجموعه داده بزرگتری شامل ۱۰۰۰ نظر برای آموزش و ارزیابی روش پیشنهادی استفاده می‌کنیم، و نتایج قابل اعتمادتری را ارائه خواهیم داد. در روش پیشنهادی ما از انتخاب ویژگی‌های مفید رخداد کلمات در هر کلاس قطبیت یا هیجان، و متوازن‌سازی مجموعه آموزش در مسئله قطبیت و هیجان به کمک روش SMOTE نیز برای آموزش بهتر طبقه‌بند و دستیابی به دقت بالاتر استفاده شده است.

۲- تعریف مسئله

یکی از اطلاعات قابل استخراج از متون نظرات کاربران، قطبیت نظرات به صورت منفی یا مثبت است [۱۰]. بعضی محققین علاوه بر دسته‌بندی به دو دسته منفی یا مثبت، دسته خنثی را نیز در قطبیت نظرات در نظر می‌گیرند [۱۱]. اکثر مشتریان دیدگاه‌های مثبتی را درباره محصول و مزایا و نقاط قوت آن بیان می‌کنند. اما بعضی نظرات نیز به نقاط ضعف و ایرادهای محصول مربوط هستند. شناسایی قطبیت نظرات یک مسئله رایج است که توسط محققین مختلفی مورد تحقیق و بررسی قرار گرفته است [۲، ۴، ۱۰، ۱۱]. این مسئله اساساً یک مسئله طبقه‌بندی نامتوازن و تک برچسبی است.

یک جنبه اطلاعاتی مهم دیگر که از نظرات مشتریان قابل استخراج است، هیجانات کاربر است که در متن نظر مستتر شده است [۱]. سطوح اولیه تشخیص هیجان شامل سه هیجان مثبت و سه هیجان منفی است. هیجان‌های مثبت پایه شامل شادی (Joy)، اعتماد (Trust)، و شگفتی (Surprise) هستند، و هیجانات منفی پایه شامل عصبانیت (Anger)، ترس (Fear)، و ناراحتی (Sadness) هستند. بعضی محققین تعداد این دسته‌ها را به هشت دسته افزایش می‌دهند، و احساسات اشتیاق (Anticipation) و نفرت (Disgust) را نیز در نظر می‌گیرند. تشخیص هیجان نسبت به تعیین قطبیت، وضعیت عاطفی مشتری را به‌طور دقیق‌تری تعیین می‌کند. مسئله تشخیص هیجان در نظرات متنی یک مسئله طبقه‌بندی نامتوازن و چند برچسبی است. اصطلاح چندبرچسبی به این معنی است که در یک نظر به طور

^۱ Gated Recurrent Unit

کلمه کلیدی	جستجو در واژه‌نامه	مدخل مناسب در واژه‌نامه
راضیم	ناموفق ❌	satisfied راضی ✅
دزدیدن	ناموفق ❌	burglary دزدی ✅
دوام نیاورد	ناموفق ❌	دوام آوردن ❌
روشن نشد	ناموفق ❌	روشن شدن ❌

شکل ۳: مشکلات جستجوی کلمات فارسی در واژه‌نامه EmoLex

یک ساله این گوشی را دارم هیچ نقطه ضعفی از او ندیدم
در خریدش شک نکنید
من گوشی را یکسال که دارم خیلی راضیم ولی الان قابش تو بازار نیست

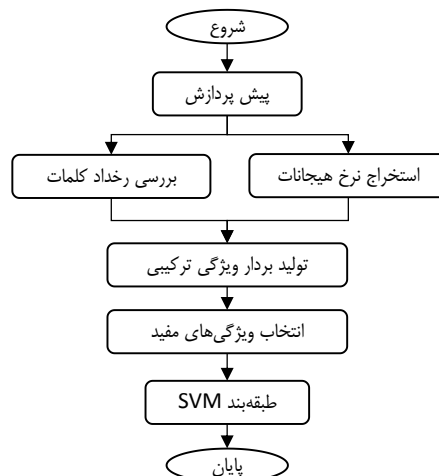
شکل ۴: چالش‌های باقی‌مانده در متون پس از پیش‌پردازش

۱-۳- پیش‌پردازش متن

نظرات کاربران فارسی زبان عموماً حاوی کلمات محاوره‌ای مثل «عالیه»، «حرف نداره»، «مشکل داره»، «کم نیاره»، و ... است. عدم استفاده از شکل صحیح کلمات می‌تواند باعث شود که یک کلمه در شکل‌های مختلفی در نظرات کاربران مختلف به کار برده شود و هر شکل توسط الگوریتم به عنوان یک کلمه مجزا لحاظ گردد. از طرف دیگر مراحل پیش‌پردازش بعدی مثل نرمال سازی و ریشه‌یابی نیز روی شکل‌های محاوره‌ای کلمات در ست عمل نمی‌کنند. به همین دلیل، برای تبدیل شکل محاوره‌ای کلمات به شکل رسمی آن‌ها در متن هر نظر از خدمات FormalConverter سامانه پردازش متن فارسی‌یار استفاده شد، که به صورت وب‌سرویس ارائه شده است [۱۲].

بعد از تبدیل کلمات محاوره‌ای به رسمی، عملیات نرمال سازی به کمک کتابخانه هضم [۱۳] در پایتون بر روی متن هر نظر اعمال شد، تا هر کلمه به شکل استانداردش در سیستم تبدیل شود و فاصله‌های اضافه نیز از جملات حذف شوند. در مرحله بعدی، عملیات توکن‌سازی (tokenization) انجام شد. در این فرآیند اعداد به عبارت NUM، و هشتگ‌ها به عبارت TAG تبدیل شدند. همچنین، کلمات افعال ترکیبی به کمک زیرخط به هم چسبانده شدند تا یک کلمه به حساب بیایند. استفاده از ریشه‌یاب (stemmer) نیز بررسی شد، اما به دلیل کاهش دقت روش پیشنهادی، از آن صرف نظر شد. نمونه‌ای از یک نظر قبل و بعد از پیش‌پردازش در شکل (۲) نمایش داده شده است.

همزمان می‌تواند چند نوع هیجان مشاهده شود. مثلاً در یک نظر مثبت ممکن است هیجانات شادی، اعتماد، و شگفتی به طور همزمان وجود داشته باشند. در این مقاله ما دو مسئله تشخیص قطبیت و تشخیص هیجان در نظرات متنی را به صورت همزمان حل خواهیم کرد. مسئله تشخیص هیجان را یک مسئله چند برچسبی و مسئله تشخیص قطبیت را یک مسئله تک برچسبی در نظر می‌گیریم.



شکل ۱: مراحل روش ترکیبی پیشنهادی

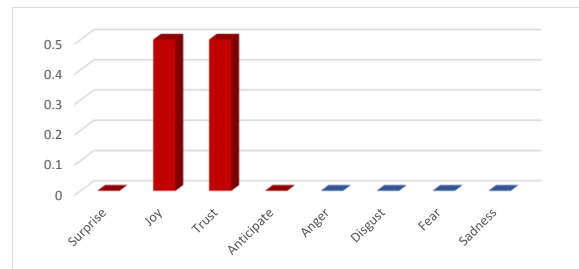
۳- روش ترکیبی پیشنهادی

در این مقاله با ترکیب روش مبتنی بر واژه‌نامه و روش مبتنی بر یادگیری ماشین، یک سیستم ترکیبی برای شناسایی انواع هیجان و قطبیت در نظرات فارسی ارائه خواهد شد. مراحل مختلف روش پیشنهادی در شکل (۱) نمایش داده شده‌اند. در اولین مرحله متن هر نظر مورد پیش‌پردازش قرار گرفته است تا کلمات از یکدیگر جدا شده، و به شکل استاندارد خود درآیند. سپس ویژگی‌های نرخ هیجانات و ویژگی‌های رخداد کلمات از متن نظر استخراج شده و در بردار ویژگی ترکیبی در کنار هم قرار می‌گیرند. در مرحله بعدی از یک روش انتخاب ویژگی برای انتخاب مرتبط‌ترین و مهم‌ترین ویژگی‌ها به هر کلاس هدف استفاده شده است. انتخاب ویژگی و ساخت مدل طبقه‌بندی برای هر کلاس هیجانی یا قطبیت به صورت مجزا از کلاس‌های دیگر انجام شده است. در نهایت برچسب‌های هیجانی و برچسب قطبیت هر نظر توسط مدل‌های طبقه‌بندی کلاس‌های مختلف پیش‌بینی شده‌اند، که این مدل‌های طبقه‌بندی همگی از نوع بردار ماشین پشتیبان (SVM) هستند.

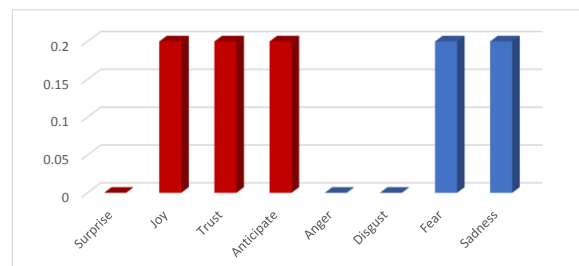
گوشیه فوق العاده ای هستش اصلاً همیشه براش عیب و ایرادی گرفت دیجی
جوون بابت هدیه‌ها ممنون
گوشی فوق العاده‌ای هستش اصلاً نمی‌شود برای آن عیب و ایرادی گرفت دیجی
جان بابت هدیه‌ها ممنون

شکل ۲: پیش‌پردازش هر نظر (بالا) متن اولیه، (پائین) متن پیش‌پردازش شده

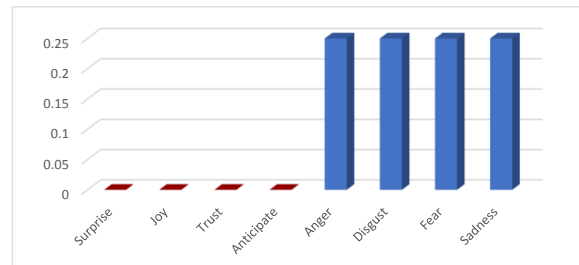
Review #15 گوشه‌ی فوق‌العاده‌ای هستش اصلاً نمی‌شود برای آن عیب و ایرادی گرفت دیجی جان بابت هدیه‌ها ممنون



Review #5 من پکیج کاملش را خریدم از هر نظر عالیه هیچ نقد و ایرادی به آن وارد نیست فقط یک گوشه لوکس باید خیلی مواظب ضربه خوردن و شکستنش باشید



Review #46 من این گوشه را از دیجی کالا تهیه کردم متأسفانه بعد از شش ماه گوشه مداوم دچار هنگ میشد و تاج گوشه از کار افتاد شوربختانه گارانتی هم به هیچ عنوان پاسخ‌گوی تماس‌ها نیست



شکل ۵: استخراج موفق نرخ هیجانات برای تعدادی از نظرات

۳-۲- استخراج ویژگی‌های هیجانی

یک رویکرد رایج برای شناسایی انواع مختلف هیجان در نظرهای کاربران استفاده از واژه‌نامه‌های هیجانی از جمله EmoLex است که توسط NRC ارائه شده است [۱۴]. اما استفاده از این واژه‌نامه برای متون فارسی مشکل است. زیرا معادل‌های فارسی هر کلمه در EmoLex به کمک یک مترجم خودکار تهیه شده‌اند [۱۵]. همین باعث شده است تا برای هر کلمه انگلیسی، یک عبارت معادل فارسی تولید شود، که شامل چند کلمه فارسی است. این ویژگی‌ها باعث می‌شوند تا استخراج برچسب‌های هیجانی هر کلمه فارسی به کمک این واژه‌نامه برای بسیاری از کلمات ممکن نباشد. نمونه‌ای از این موارد در شکل (۳) نمایش داده شده است.

از طرف دیگر، در متون نظرهای فارسی اکثر کاربران از شکل محاوره‌ای کلمات مثل «عالیه»، «همه‌چی تموم»، «حرف نداره»، و ... استفاده می‌کنند

که در واژه‌نامه یافت نمی‌شوند. اگر چه در مرحله پیش پردازش ما سعی کرده‌ایم که شکل محاوره‌ای کلمات را به شکل رسمی آن‌ها تبدیل کنیم، اما هنوز بعضی کلمات ممکن است به صورت محاوره‌ای باقی مانده باشند. تعیین ضرایب هیجانی برای توالی‌های دوتایی کلمات یک مشکل دیگر است. استفاده از واژه‌نامه EmoLex نمی‌تواند مستقیماً توالی‌های دوتایی کلمات را پوشش دهد، مگر این که یک رویکرد هیورستیک مناسب برای استخراج توالی‌های هیجانی دوتایی و تخمین هیجان آن‌ها در سیستم تعبیه گردد [۶]. نمونه‌ای از این چالش‌ها در شکل (۴) آمده است. با توجه به این مشکلات، ما در این تحقیق یک واژه‌نامه هیجانی اختصاصی برای مجموعه داده خود تهیه کردیم که شامل عبارت‌های کلیدی هیجانی معنادار در مجموعه داده است. در این مقاله ما علاوه بر تک کلمات، عبارت‌های دو کلمه‌ای را نیز از نظرها استخراج کردیم، و موارد معنادار را در واژه‌نامه هیجانی در نظر گرفتیم. این واژه‌نامه هیجانی توسط کاربر انسانی از نظر هیجان‌های مختلف برچسب زده شد. برای هر کلمه کلیدی، وجود یا عدم وجود هر کلاس هیجانی، قطبیت منفی، و قطبیت مثبت به کمک برچسب‌های باینری صفر یا یک مشخص شد. برای استخراج نرخ هیجانات مختلف، عبارت‌های تک کلمه‌ای و دو کلمه‌ای از هر نظر استخراج شدند. سپس این عبارت‌ها در واژه‌نامه هیجانی جستجو شدند، و نرخ کلمات در دسته‌های مثبت، منفی، شادی، اعتماد، اشتیاق، شگفتی، عصبانیت، ترس، ناراحتی، و نفرت از هر نظر استخراج شد. این کمیت‌ها به عنوان بردار ویژگی نرخ هیجانات هر نظر به طول ۸ عنصر لحاظ شدند. اگر چه نرخ هیجانات استخراج شده همیشه صحیح نیست، اما برای بسیاری از نظرات می‌تواند کمک کننده باشد. نرخ هیجانات مختلف استخراج شده برای یک نظر مثبت، یک نظر خنثی، و یک نظر منفی در شکل (۵) نمایش داده شده است. برای نظر شماره ۱۵ که یک نظر مثبت است، نرخ شادی و اعتماد استخراج شده بالا است. برای نظر شماره ۵ که یک نظر نسبتاً خنثی است، نرخ هیجانات مثبت شامل شادی، اعتماد، و اشتیاق بالا بوده و به طور همزمان نرخ هیجانات منفی ترس و ناراحتی نیز بالا است. در نهایت برای نظر شماره ۴۶ که یک نظر منفی است، نرخ احساسات منفی ناراحتی، عصبانیت، ترس، و نفرت بالا است. در مجموع، استفاده از واژه‌نامه برچسب‌گذاری شده به صورت دستی به استخراج قابل قبول نرخ هیجانات منجر شده است. با این وجود، در مورد بعضی نظرها، ممکن است ضرایب استخراج شده کمی خطا نیز داشته باشند.

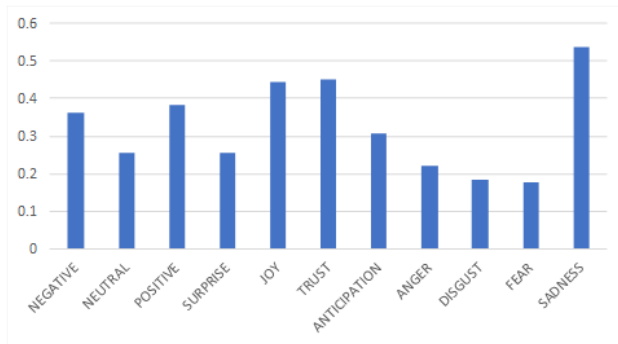
۳-۳- استخراج ویژگی‌های رخداد کلمات

رویکرد «کیف کلمات» (BOW) یک روش رایج برای استخراج ویژگی از متن است. در این رویکرد، ترتیب کلمات در متن روی ویژگی‌های استخراج شده تأثیری نمی‌گذارد، بلکه حضور یا عدم حضور کلمات و تعداد دفعات تکرار آن‌ها می‌تواند روی مقادیر ویژگی‌ها تأثیر گذار باشد. در این بخش، ما عبارت‌های تک کلمه‌ای و دو کلمه‌ای را به عنوان کلمات کلیدی در نظر گرفته و مقادیر وزن TF-IDF را برای آن‌ها محاسبه کردیم. استفاده از عبارت‌های دو کلمه‌ای باعث می‌شود تا سیستم پیشنهادی توانایی بیشتری در مدل‌سازی معنای جملات داشته باشد. در این مرحله، حذف کلمات هجو و کم اهمیت که بار معنایی خاصی ندارند نیز به طور خودکار حین استخراج ویژگی انجام شد. به این ترتیب، از هر نظر فارسی یک بردار ویژگی رخداد کلمات به طول ۲۴۳۰۱ عنصر استخراج شد.

۳-۴- انتخاب ویژگی

در این مقاله ما روش مبتنی بر واژه‌نامه را با روش یادگیری ماشینی ترکیب می‌کنیم، تا استفاده از اطلاعات واژه‌نامه به بهبود مدل یادگیری ماشینی کمک کند. برای این منظور، بردار نرخ هیجانات در هر نظر را به انتهای بردار ویژگی رخداد کلمات همان نظر متصل می‌کنیم، و بردار حاصل را بردار ویژگی ترکیبی می‌نامیم. بردار ویژگی ترکیبی در روش پیشنهادی ما شامل ۲۴۳۰۹ عنصر است. تعداد ویژگی‌ها در بردار ویژگی ترکیبی خیلی زیاد است و برای هر کلاس قطبیت یا هیجان قطعاً تنها تعداد کمی از این ویژگی‌ها مرتبط خواهند بود. به همین دلیل، برای هر کلاس هیجان و برای بخش قطبیت، ویژگی‌های مطلوب با کمک یک روش انتخاب ویژگی خود کار مبتنی بر فیلتر انتخاب شدند. ابتدا ویژگی‌هایی که واریانس آن‌ها صفر است از لیست ویژگی‌های کاندیدا حذف شدند. سپس، معیار کیفیت ویژگی ANOVA F-Value برای هر ویژگی کاندیدا محاسبه شد، و ویژگی‌هایی انتخاب شدند که در لیست ویژگی‌ها بالاترین کیفیت را داشتند. معیار کیفیت ANOVA F-Value یک معیار آماری است که نشان می‌دهد تا چه حد میانگین داده‌های طبقات مختلف در یک ویژگی خاص با هم تفاوت معنادار دارند. در نتیجه، ویژگی‌هایی امتیاز بالاتری در این معیار کسب می‌کنند که قدرت تمایز بیشتری بین کلاس‌های مختلف دارند. در این مقاله ما تعداد ویژگی‌های استخراج شده برای هر طبقه‌بند را ۳۰۰ ویژگی در نظر گرفتیم.

مثبت ۳۸۳ نمونه، خنثی ۲۵۴ نمونه، منفی ۳۶۳ نمونه، شادی ۴۶۶ نمونه، اعتماد ۴۵۰ نمونه، شگفتی ۲۵۴ نمونه، اشتیاق ۳۰۸ نمونه، ناراحتی ۵۳۸ نمونه، عصبانیت ۲۲۱ نمونه، ترس ۱۷۸ نمونه، و نفرت ۱۸۲ نمونه است. برای این مجموعه داده نرخ عدم توازن در کلاس‌های هیجانی مختلف در شکل ۶ به صورت نمودار میله‌ای نمایش داده شده است. در این نمودار واضح است که مجموعه داده در کلاس‌های مختلف نامتوازن است. به عنوان مثال در کلاس نفرت تنها ۱۸ درصد نمونه‌ها برچسب مثبت دارند، در حالی که ۸۲ درصد نمونه‌ها برچسب منفی دارند. به همین دلیل در روش پیشنهادی از الگوریتم SMOTE برای متوازن سازی کلاس‌های مختلف استفاده شد.



شکل ۶: نسبت تعداد نمونه‌های مثبت به نمونه‌های منفی در مجموعه داده برای کلاس‌های مختلف هیجان

۳-۵- طبقه‌بندی با ماشین بردار پشتیبان

برای دستیابی به نتایج مناسب از ماشین بردار پشتیبان به عنوان طبقه‌بند استفاده کردیم. ماشین بردار پشتیبان یک الگوریتم یادگیری ماشین با نظارت برای طبقه‌بندی و رگرسیون است، که بیشتر در کاربردهای طبقه‌بندی داده‌ها مورد استفاده قرار می‌گیرد. مسئله تعیین قطبیت نظر را به عنوان یک مسئله سه کلاسه شامل کلاس‌های «مثبت»، «خنثی»، و «منفی» در نظر گرفته شد. مسئله تعیین هیجان هر نظر نیز برای هر نوع هیجان شامل «شادی»، «اعتماد»، «اشتیاق»، «شگفتی»، «عصبانیت»، «ترس»، «غم»، و «نفرت» به عنوان یک مسئله دو کلاسه در نظر گرفته شد. برای هر یک از این مسائل، انتخاب ویژگی‌های مربوطه به صورت مستقل انجام شد، و یک طبقه‌بند ماشین بردار پشتیبان مستقل نیز ساخته شد. به این ترتیب، در مجموع از ۹ طبقه‌بند ماشین بردار پشتیبان استفاده شد. از آنجایی که تعداد نمونه‌های کلاس‌های مختلف در مجموعه داده مورد استفاده برای تمام این مسائل نامتوازن است، برای کاهش تأثیر عدم توازن بر نتایج و آموزش بهتر طبقه‌بندها، در مرحله آموزش از روش SMOTE برای متوازن سازی تعداد نمونه‌های مسئله در کلاس‌های مختلف استفاده شد.

۴- ارزیابی

روش پیشنهادی این مقاله در زبان برنامه‌نویسی پایتون به کمک کتابخانه‌های sklearn, imblearn, numpy و pandas پیاده‌سازی شد. برای ارزیابی روش پیشنهادی در تشخیص صحیح کلاس‌های قطبیت و هیجان در نظرهای فارسی، یک مجموعه داده شامل ۱۰۰۰ نظر از نظرهای کاربران سایت دیجی کالا در خصوص گوشی‌های موبایل و لپ‌تاپ جمع‌آوری و به صورت دستی در کلاس‌های مختلف برچسب‌گذاری شد. تعداد نمونه‌ها در کلاس‌های

جدول ۱: نتیجه ارزیابی دقت روش پیشنهادی و رویکردهای ساده‌تر بر اساس معیار F1 بر حسب درصد

روش	هیجان مثبت	هیجان خنثی	هیجان منفی	شادی	اعتماد	شگفتی	اشتیاق	عصبانیت	ترس	غم	نفرت
UNI	۷۸	۲۹	۷۳	۸۵	۸۶	۵۹	۷۷	۷۸	۰۳	۰۴	۱۴
BI	۷۷	۱۷	۷۱	۸۵	۸۹	۵۸	۷۷	۷۸	۰۹	۰۴	۰۷
UNI-FS	۷۹	۳۶	۶۹	۸۳	۸۳	۴۸	۷۲	۷۵	۳۰	۳۲	۲۱
BI-FS	۷۷	۴۳	۷۰	۸۵	۸۳	۵۶	۷۴	۷۸	۳۴	۲۵	۲۵
UNI-FS-EMO	۷۶	۳۳	۷۲	۸۰	۸۴	۵۲	۷۰	۷۷	۲۸	۲۵	۲۷
BI-FS-EMO	۷۶	۳۳	۷۵	۸۳	۸۴	۶۱	۷۳	۷۷	۲۷	۲۰	۲۷
UNI-FS-BL	۷۹	۵۰	۷۱	۸۳	۸۳	۶۱	۷۶	۷۵	۳۷	۴۵	۳۷
BI-FS-BL	۷۵	۴۳	۷۰	۸۵	۸۳	۶۱	۷۵	۷۶	۳۹	۳۳	۴۷
UNI-FS-EMO-BL	۷۷	۳۹	۷۴	۸۱	۸۵	۶۳	۷۴	۷۶	۴۴	۴۰	۵۳
BI-FS-EMO-BL	۷۷	۴۰	۷۳	۸۴	۸۶	۶۱	۷۴	۷۷	۵۳	۴۲	۵۵

نتایج ارزیابی دقت طبقه‌بندی بر اساس معیار F1 در مسائل مختلف در جدول (۱) آمده است. در این جدول نام هر سطر نشان‌دهنده جزئیات روش ارزیابی شده است. در صورت استفاده از عبارت‌های تک‌کلمه‌ای به عنوان کلمات کلیدی از عنوان UNI، و در صورت استفاده از عبارت‌های تک‌کلمه‌ای و دوکلمه‌ای به عنوان کلمات کلیدی از عنوان BI استفاده شده است. عنوان FS نشان‌دهنده استفاده از انتخاب ویژگی، عنوان EMO نشان‌دهنده ترکیب ویژگی‌های نرخ هیجان با بردار ویژگی رخداد کلمات، و عنوان BL نشان‌دهنده متوازن سازی داده‌های آموزشی با کمک روش SMOTE است. روش پیشنهادی ما با عنوان BI-FS-EMO-BL مشخص شده است که به

Text-based emotion recognition in decision support", Decision Support Systems, Vol. 115, pp. 24-35, 2018.

- [4] Taj, S., Shaikh, B.B., and Meghji, A.F., "Sentiment Analysis of News Articles: A Lexicon based Approach", 2019 2nd International Conference on Computing, Mathematics and Engineering Technologies (iCoMET), pp. 1-5, 2019.
- [5] Rout, J.K., Choo, K.-K.R., Dash, A.K., Bakshi, S., Jena, S.K., and Williams, K.L., "A model for sentiment and emotion analysis of unstructured social media text", Electronic Commerce Research, Vol. 18, No (1), pp. 181-199, 2018.
- [6] Srinivasan, S.M., Sangwan, R.S., Neill, C.J., and Zu, T., "Twitter Data for Predicting Election Results: Insights from Emotion Classification", IEEE Technology and Society Magazine, Vol. 38, No (1), pp. 58-63, 2019.
- [7] Garshasbi, M., Rais-Rohani, A., and Kabaranzadeh Ghadim, M., "Presenting a Text Mining Algorithm to Identify Emotion in Persian Corpus", Journal of Information Technology Management, Vol. 10, No. (2), pp. 375-389, 2018.
- [8] Sadeghi, S.S., Khotanlou, H., and Rasekh Mahand, M., "Automatic Persian Text Emotion Detection using Cognitive Linguistic and Deep Learning", Journal of AI and Data Mining, Vol. 9, No. (2), pp. 169-179, 2021.

[۹] خسروی، علی، کلارستانی، منوچهر، پورمحمد، مهدی، "شناسایی هیجان

در متن فارسی با استفاده از مدل یادگیری ماشینی"، روانشناسی

معاصر، دوره ۱۴، شماره (۱)، صفحه ۴۲-۴۸، ۱۳۹۸.

- [10] Al-Azani, S., and El-Alfy, E.-S.M., "Using Word Embedding and Ensemble Learning for Highly Imbalanced Data Sentiment Analysis in Short Arabic Text", Procedia Computer Science, Vol. 109, pp. 359-366, 2017.
- [11] Zhang, S., Wei, Z., Wang, Y., and Liao, T., "Sentiment analysis of Chinese micro-blog text based on extended sentiment dictionary", Future Generation Computer Systems, Vol. 81, pp. 395-403, 2018.
- [12] <https://api.text-mining.ir/index.html>
- [13] <https://www.sobhe.ir/hazm/>
- [14] Mohammad, S.M., and Turney, P.D., "Crowdsourcing a Word-Emotion Association Lexicon", Computational Intelligence, Vol. 29, No. (3), pp. 436-465, 2013.
- [15] <https://github.com/mhbashari/NRC-Persian-Lexicon>

طور میانگین بهترین نرخ F1 را در کلاس‌های مختلف فراهم آورده است، و به میانگین دقت ۶۶٪ برای کلیه کلاس‌ها دست یافته است. این موضوع به معنی نرخ Precision و Recall بالاتر است.

روش پیشنهادی در مسائل شدیداً نامتوازن مثل کلاس عصبانیت، نفرت، و ترس عملکرد بهتری را نسبت به روش‌های ساده‌تر داشته است. مقایسه روش‌ها با یکدیگر نشان‌دهنده میزان اثرگذاری هر یک از مراحل تکمیلی است. اگر کلاس نفرت را در نظر بگیریم، در صورت استفاده از عبارت‌های دوکلمه‌ای، انتخاب ویژگی دقت را در سطر BI-FS نسبت به سطر BI حدود ۲۰ درصد افزایش داده است. سپس، متوازن سازی در کنار انتخاب ویژگی در سطر BI-FS-BL دقت را حدود ۲۰ درصد دیگر نسبت به سطر BI-FS افزایش داده است. در نهایت، ترکیب ویژگی‌های نرخ هیجانات مختلف با ویژگی‌های رخداد کلمات در روش پیشنهادی در سطر BI-FS-EMO-BL دقت روش پیشنهادی را ۱۰ درصد دیگر نسبت به سطر BI-FS-BL افزایش داده، و باعث شده است روش پیشنهادی در شناسایی کلاس هیجانی نفرت به نرخ دقت ۵۵٪ برای معیار F1 برسد. نسبت به پژوهش انجام شده در (خسروی و همکاران، ۱۳۹۸)، در این تحقیق ما از مجموعه داده بزرگتری استفاده کردیم، و ارزیابی قابل اطمینان‌تری را انجام داده‌ایم. ضمن این که از انتخاب ویژگی، متوازن سازی، و بردار ویژگی ترکیبی نیز استفاده کرده‌ایم.

۵- نتیجه گیری

در این مقاله یک روش ترکیبی مبتنی بر واژه‌نامه هیجانی کلمات و ویژگی‌های رخداد کلمات کلیدی برای شناسایی کلاس‌های قطبیت و کلاس‌های هیجان در نظرهای کاربران فارسی زبان ارائه شد. این روش می‌تواند به صاحبان تجارت کمک کند تا از بازخورد و نظر مشتریان خود نسبت به کالاها، تحولات بازار، و اقدامات تجاری انجام شده توسط شرکت آگاه شوند. رویکرد ترکیبی پیشنهادی توانست نسبت به روش‌های ساده‌تر به دقت بالاتری در طبقه‌بندی هیجان و قطبیت دست یابد، و به طور میانگین به نرخ ۶۶٪ در معیار F1 برای کلیه کلاس‌ها دست یافت. در روش پیشنهادی استفاده از واژه‌نامه هیجانی کلمات باعث تزریق اطلاعات کمکی حل مسئله به سیستم شد، و دقت روش پیشنهادی را افزایش داد. از طرف دیگر استفاده از عبارت‌های یک و دوکلمه‌ای به عنوان کلمات کلیدی توانایی مدل‌سازی معنایی روش پیشنهادی را بهبود داد. در نهایت استفاده از روش SMOTE برای متوازن سازی مجموعه داده آموزش، باعث آموزش بهتر مدل‌های طبقه‌بندی و در نتیجه دستیابی به دقت بالا شد. تهیه مجموعه داده بزرگتر، برچسب‌زنی قطبیت نظرات توسط خود کاربران، استفاده از مدل‌های یادگیری عمیق، و ترجمه متون فارسی به انگلیسی از جمله راهکارهای ممکن برای توسعه این پژوهش هستند.

مراجع

- [1] Peng, S., Cao, L., Zhou, Y., Ouyang, Z., Yang, A., Li, X., Jia, W., and Yu, S., "A survey on deep learning for textual emotion analysis in social networks", Digital Communications and Networks, 2021. (In Press)
- [2] Zhang, J., Zhang, A., Liu, D., and Bian, Y., "Customer preferences extraction for air purifiers based on fine-grained sentiment analysis of online reviews", Knowledge-Based Systems, Vol. 228, pp. 107259, 2021.
- [3] Kratzwald, B., Ilić, S., Kraus, M., Feuerriegel, S., and Prendinger, H., "Deep learning for affective computing:

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بوم

یک سیستم بیومتریکی ترکیبی هایپرید مبتنی بر تلفیق ویژگی‌های چهره، هر دو عنبریه و دو اثر انگشت

محمدحسن صفوی پور^۱، محمدعلی دوستاری^۲، حامد ساجدی^۳

^۱ دانشجوی دکترا، گروه الکترونیک، دانشکده فنی و مهندسی دانشگاه شاهد، تهران، mhassan.safavipour@shahed.ac.ir

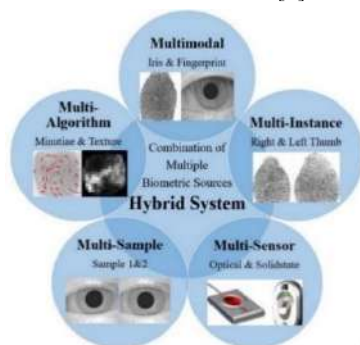
^۲ استادیار، گروه کامپیوتر، دانشکده فنی و مهندسی دانشگاه شاهد، تهران، doostari@shahed.ac.ir

^۳ دانشیار، گروه الکترونیک، دانشکده فنی و مهندسی دانشگاه شاهد، تهران، sadjedi@shahed.ac.ir

چکیده: یک سیستم چند بیومتریکی هایپریدی با ترکیب چندین صفت و منبع بیومتریکی، بازنمایی با دقت بیشتر، اطمینان از امنیت بالاتر و همچنین غلبه بر محدودیت‌هایی مانند عدم جهانشمولی، نویز داده‌های اولیه سنسور و تغییرات بزرگ درون کلاسی را فراهم می‌کند. در این مقاله دو استراتژی برای تلفیق بردارهای ویژگی با استفاده از روش‌های کاهش ابعاد و الگوریتم‌های مبتنی بر کوآترینیون در فضای هیلبرت هسته بازآفرین (RKHS) ارائه شده است تا کارایی سیستم هایپرید پیشنهادی برای تلفیق در سطح ویژگی پنج صفت بیومتریکی (چهره، هر دو عنبریه و دو اثر انگشت) از سه منبع چند نمونه، چند الگوریتم و ترکیبی نشان داده شود. برای این منظور در فضای ویژگی، با انتخاب هسته بازآفرین مناسب، هر یک از بردارهای ویژگی به صورت جداگانه به RKHS نگاشت می‌شوند. سپس در مازول تلفیق ویژگی، سه بردار ویژگی چهره، عنبریه ترکیبی و اثر انگشت ترکیبی در RKHS با متودولوژی پیشنهادی با هم تلفیق می‌شوند. نتایج آزمایش‌ها در شش پایگاه داده برای سیستم پیشنهادی به خوبی نشان می‌دهد الگوی ترکیبی هایپرید بدست آمده در عین ایمن بودن در برابر حملات تقلبی و مقاوم سازی سیستم، می‌تواند با ابعاد تنها ۱۵ ویژگی دقت سیستم را به ۱۰۰٪ برساند.

کلمات کلیدی: تلفیق در سطح ویژگی، بیومتریکی ترکیبی، فضای هیلبرت هسته بازآفرین، کوآترینیون، کاهش ابعاد.

هایپریدی گفته می‌شود [۲] (شکل ۱). علاوه بر این برای افزایش پیچیدگی تأیید اعتبار کاربر و اطمینان از امنیت بالا، بیش از یک صفت با هم ترکیب می‌شود [۳]. بنابراین سیستم بیومتریکی ترکیبی هایپریدی در عین حل مشکلات ذکر شده در بالا، به دلیل قابلیت اطمینان بالاتر، کاربرد گسترده تر و امنیت قوی تر، در زمینه بازنمایی بیومتریکی رشد یافته و محققان بیشتری را به خود جلب نموده است [۴].



شکل ۱: سیستم چندبیومتریکی هایپریدی با ترکیب منابع بیومتریکی

۱- مقدمه

سیستم‌های تک بیومتریکی که تنها به یک صفت بیومتریکی محدود هستند، دلیل وجود نویز، کیفیت پایین اطلاعات، عدم جهانشمولی و تغییرات بزرگ درون کاربر از محدودیت‌های قابل توجهی رنج می‌برند. سیستم‌های چند بیومتریکی با استفاده از تلفیق برای ترکیب چندین منبع بیومتریکی دقت شناسایی را افزایش می‌دهند [۱]. پنج حالت ممکن وجود دارد که می‌تواند منابع مختلفی از اطلاعات بیومتریکی را فراهم کند. بر اساس منابع اطلاعات، سیستم‌های چند بیومتریکی را می‌توان به سیستم چند سنسور، چند الگوریتم، چند نمونه، چند سمپل و ترکیبی (مالتی مدال) طبقه بندی کرد. در چهار سناریوی اول، اطلاعات متعدد از یک صفت بیومتریکی منشأ گرفته شده است (به عنوان مثال اثر انگشت یا عنبریه)، در حالی که در سناریوی پنجم (که سیستم مالتی مدال بیومتریکی نیز نامیده می‌شود) از چندین صفت بیومتریکی (به عنوان مثال اثر انگشت و عنبریه) استفاده می‌شود. همچنین برای یک سیستم چند بیومتریکی امکان استفاده از ترکیبی از دو یا چند مورد از پنج سناریو وجود دارد که به چنین سیستم‌هایی، سیستم‌های بیومتریکی ترکیبی

متدولوژی‌های پیشنهادی تلفیق در RKHS معرفی شده است. طراحی و پیاده سازی سیستم ترکیبی هابیرید پیشنهادی شامل ماژول استخراج کننده ویژگی، ماژول تلفیق در سطح ویژگی و ماژول طبقه بند در بخش ۳ ارائه شده است و در بخش ۴ نیز تجزیه و تحلیل و نتیجه آزمایش‌ها گزارش شده و سرانجام نتیجه گیری در بخش ۵ ارائه شده است.

۲- مبانی نظری تحقیق

۲-۱- روش‌های هسته

روش‌های هسته [۸-۱۰] یکی از قوی ترین تکنیک‌ها برای مسایل یادگیری آماری غیرخطی از جمله طبقه بندی [۱۱]، رگرسیون [۱۲]، کاهش ابعاد [۱۳] و خوشه بندی [۱۴] است. ایده اصلی در روش‌های هسته اگر $x_1, x_2 \in X \subset \mathbb{R}^d$ دو نمونه باشد آنگاه $\Phi: X \rightarrow H$ یک نگاشت از ویژگی‌های غیرخطی است که هر عنصر ویژگی را در X به یک بعد بزرگ (یا حتی بی نهایت بعدی) فضای هیلبرت هسته بازآفرین H تبدیل می‌کند. ضرب داخلی بین $\Phi(x_1)$ و $\Phi(x_2)$ در فضای جدید ویژگی H را می‌توان با استفاده از یک تابع هسته $k(x_1, x_2)$ محاسبه کرد:

$$k(x_1, x_2) = \langle \Phi(x_1), \Phi(x_2) \rangle, \quad \Phi: X \rightarrow H, x \rightarrow \Phi(x) \quad (۱)$$

در عمل این ضرب داخلی با اعمال مستقیم هسته k بدون نیاز به یافتن صریح Φ (که به عنوان ترنند هسته می‌باشد) محاسبه می‌شود. روش‌های هسته اغلب به شکل ظریف اما محدود به نظریه زیبای ریاضی (فضای هیلبرت هسته بازآفرین و غیره) و بصری (حداکثر حاشیه در فضای ویژگی، تف سیر هندسی بردارهای پشتیبانی و غیره) توصیف می‌شوند. روش‌های هسته یک رابطه خطی را در یک فضای ویژگی، با تمام غیر خطی بودن موجود در نگاشت ثابت از فضای ورودی به فضای ویژگی یاد می‌گیرند. با این حال روش‌های هسته اگرچه برای یادگیری ساختارهای غیرخطی موثر هستند، به دلیل پیچیدگی‌های فضا و زمان زیاد، اغلب از مشکلات مقیاس پذیری در مسایل با مقیاس بزرگ رنج می‌برند [۱۵].

۲-۲- فضای هیلبرت (H)

H در حقیقت، حالت تعمیم یافته‌ای از فضای اقلیدسی است. به این ترتیب جبر بردارها از فضای دو و سه بعدی به ابعاد متناهی و حتی نامتناهی سوق داده می‌شوند. توصیف و شواهد هندسی در تئوری فضای هیلبرت، بخصوص در «فضاهای تابعی بی‌نهایت بُعد» نقش اساسی ایفا می‌کند. فضای هیلبرت H یک فضای برداری مختلط است که ضرب داخلی روی آن برای هر زوج از بردارها دارای خواص زیر است:

$$\langle y, x \rangle = \overline{\langle x, y \rangle} \quad (۲)$$

$$\langle ax_1 + bx_2, y \rangle = a\langle x_1, y \rangle + b\langle x_2, y \rangle \quad (۳)$$

$$\langle x, x \rangle > 0; x \neq 0, \quad \langle x, x \rangle = 0; x = 0 \quad (۴)$$

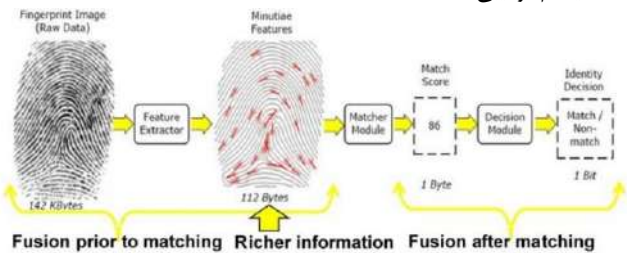
$$\langle x, ay_1 + by_2 \rangle = \bar{a}\langle x, y_1 \rangle + \bar{b}\langle x, y_2 \rangle \quad (۵)$$

همچنین H با تابع فاصله d یک فضای متریک کامل است به طوری که هر دنباله کوشی در H محدود به H است:

$$d(x, y) = \|x - y\| = \sqrt{\langle x - y, x - y \rangle} \quad (۶)$$

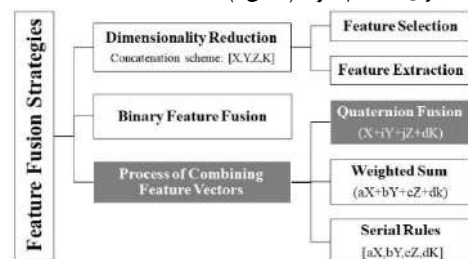
$$d(x, z) \leq d(x, y) + d(y, z) \quad (\text{نامساوی کوشی-شوارتز}) \quad (۷)$$

همانطور که در شکل ۲ نشان داده شده است تلفیق اطلاعات بیومتریک می‌تواند در چهار سطح انجام شود. اگر تلفیق در سطح سنسور [۵] اتفاق بیفتد، داده‌های خام با هم تلفیق می‌شود این نوع تلفیق برای طراحی یک سیستم ترکیبی غیرمنطقی است اما می‌تواند برای افزایش کارایی یک سیستم تک بیومتریک بسیار مفید باشد. تلفیق در سطح ویژگی، بردارهای ویژگی که از بیومتریک‌های مختلف مربوط به یک کلاس استخراج شده اند را با هم تلفیق مینماید. همچنین امتیازهای حاصل از طبقه بندی‌های مختلف، به شکلی که هر طبقه بند مربوط به یک بیومتریک باشد، می‌توانند در سطح امتیاز [۶] ترکیب شود، این روش بدلیل سادگی و کم بودن هزینه پردازش، متداول ترین نوع تلفیق برای طراحی یک سیستم چند بیومتریک است. در نهایت تلفیق در سطح تصمیم می‌تواند با تلفیق چند تصمیم که هر کدام حاصل یک سیستم تک بیومتریک است، انجام شود [۷]. تلفیق تصمیم‌ها حتی نسبت به تلفیق در سطح امتیاز هم دارای تاثیر کمتری است. چون دو سطح فوق هر دو به عملکرد تشخیص تک بیومتریک بستگی دارد لذا فضای قابل بهبود محدود است. با مقایسه چهار سطح تلفیق، سطح ویژگی می‌تواند متمایزترین اطلاعات را از مجموعه ویژگی‌های اصلی تشخیص داده و اطلاعات زائد بین مجموعه ویژگی‌های مختلف را از بین ببرد [۴]. بنابراین واضح است که تلفیق در سطح ویژگی، به دلیل غنی اطلاعات در بردارهای ویژگی، بهترین راه برای طراحی یک سیستم ترکیبی است.



شکل ۲: سطوح مختلف تلفیق در سیستم چند بیومتریک.

تلفیق بردارهای ویژگی در فضای ویژگی با هدف تبدیل دو یا چند بردار ویژگی به یک بردار واحد بشکلی که بردار نهایی نسبت به تمام بردارهای اولیه دارای قدرت تمایز بیشتری باشد می‌تواند از طریق فرآیندهای "ترکیب سریالی یا موازی"، "الگوریتم‌های کاهش ابعاد (استخراج یا انتخاب ویژگی)" و یا "تلفیق ویژگی‌های باینری" انجام شود. (شکل ۳)



شکل ۳: استراتژی‌های تلفیق دو یا چند فضای ویژگی

این مقاله دو استراتژی تلفیق فضاهای ویژگی با الگوریتم‌های کاهش ابعاد و ترکیب موازی سه بردار ویژگی چهره، عنبیه ترکیبی و اثر انگشت ترکیبی مبتنی بر کوآترنیون در فضای هیلبرت هسته بازآفرین (RKHS) را برای دستیابی به یک الگوی بیومتریک ترکیبی مقاوم و امن در یک سیستم بازشناسی بیومتریک ترکیبی هابیریدی ارائه می‌نماید.

در ادامه، در بخش ۲ مبانی نظری تحقیق شامل الگوریتم‌های کاهش ابعاد مبتنی بر هسته، فضای هیلبرت هسته بازآفرین و کوآترنیون‌ها برای توضیح

$$(8) \quad (\sum_{k=0}^{\infty} uk) \text{ دنباله های کوشی} \quad \sum_{k=0}^{\infty} \|uk\| < \infty$$

قضیه فیثاغورث و قاعده متوازی الاضلاع در فضای هیلبرت نیز صادق است. از طرفی هر عنصر از فضای هیلبرت را می توان بر اساس مختصات آن روی محورهای (پایه های متعامد نرمال) نشان داد. تصویر کردن و تغییر پایه در فضاهای با ابعاد متناهی به فضای هیلبرت از مواردی است که وسعت کاربردهای فضای هیلبرت را نشان می دهد.

۳-۲- فضای هیلبرت هسته بازآفرین (RKHS)

RKHS یک مورد مفید از فضای هیلبرت است که می تواند برای استخراج غیرخطی یا مامنت های درجه بالاتر از دیتا استفاده شود [۱۶]. برای RKHS داریم:

$$(9) \quad x \in R; \quad k_x(y) = k(x, y) \in H$$

$$(10) \quad x \in R, \quad f(x) \in H, \quad f(x) = \langle f, k_x \rangle_H$$

خاصیت بازآفرینی: برای هر $x \in R$ ، یک تابع غیر صفر k_x در H وجود دارد به طوری که $f(x) = \langle f, k_x \rangle_H$ برای هر تابع f در H ، که $H \langle \cdot, \cdot \rangle$ نشان دهنده ضرب داخلی H است. k_x فوق هسته بازآفرین H در x نامیده می شود. سپس مجموعه بازآفرین هسته ها $\{k_x\}_{x \in R}$ یک زیرمجموعه متراکم از H است.

چنانچه $(k(x, y) = \langle k_y, k_x \rangle_H)$ یک تابع دو متغیر را روی R^*R تعریف می کند.

$$(11) \quad k(x, y) = k_y(x) = \langle k_y, k_x \rangle$$

$$(12) \quad \|k_x\| = k(x, x)^{1/2} = \langle k_x, k_x \rangle^{1/2}$$

$$(13) \quad k(y, x) = \langle k_y, k_x \rangle = \langle k_x, k_y \rangle = k(x, y)$$

اگر R یک مجموعه محدود باشد، می توان $G = (k(x, y))_{x, y}$ را با یک ماتریس خودپیوست شناسایی کرد و آن را ماتریس گرام H می نامند [۱۷].

طبق خاصیت بازآفرینی برای هر $S_1, S_2, \dots, S_m$ و $S_1', S_2', \dots, S_m'$ و T و هر $a_1, \dots, a_m$ و $b_1, \dots, b_m$ در R داریم:

$$(14) \quad \left\langle \sum_{i=1}^m a_i k(\cdot, s_i), \sum_{j=1}^m b_j k(\cdot, s_j') \right\rangle = \sum_{i=1}^m \sum_{j=1}^m a_i b_j k(s_i, s_j')$$

$$(15) \quad \sum_{i,j=1}^n (a_i a_j a_{ij}) > 0, \quad A = (a_{ij}) > 0$$

خاصیت RKHS در یادگیری هسته مهم است از آنجا که آن یک تابع هسته بازآفرین $k(x, \cdot)$ را در برخی از فضای خاص هیلبرت H به طور منحصر به فرد تعیین می کند. در عمل می توان H را به عنوان تصویر $\Phi(x)$ با یک تابع نگاشت غیرخطی، احتمالاً بعد-نامحدود با ویژگی $k(x, y) = \langle \Phi(x), \Phi(y) \rangle_H$ که به ترفند هسته نیز معروف است. ترفند هسته به ما اجازه می دهد تا ضرب داخلی $\langle \Phi(x), \Phi(y) \rangle_H$ در فضای H را مستقیماً بدون ساختن Φ صریح ارزیابی کنیم. در واقع، برخی از هسته ها مانند هسته گاوسی با فضاهای ویژگی بی اندازه متناسب هستند. بیشتر الگوریتم های یادگیری مبتنی بر هسته متکی به محاسباتی هستند که فقط شامل ماتریس های گرامی هستند که روی تعداد محدودی از نقاط داده ارزیابی می شوند. یعنی ماتریس $G \in R^{n \times n}$ در یک مجموعه داده $x_1, \dots, x_n$ توسط $G_{ij} = k(x_i, x_j)$ داده می شود. بنابراین اگرچه اپراتورهای انتقال ما از نظر Φ تعریف شده اند و ممکن است در یک فضای بی نهایت قرار گرفته باشند، اما تمام عملیات مرتبط را می توان با توجه به ماتریس های گرام محدود-بعدی حاصل از

داده های آموزش انجام داد. بسته به هسته، فضای ویژگی ممکن است کم بعدی (به عنوان مثال چند جمله ای) باشد، اما همچنین می تواند بی نهایت بعدی (به عنوان مثال هسته گاوسی) باشد [۱۸].

در این مقاله بردارهای ویژگی چهره، عنبیه ها و اثر انگشتان با توابع هسته مناسب از فضای ویژگی به فضای هیلبرت (با خصوصیت های بالا) نگاشت شده و در آن فضا به منظور استخراج روابط خطی بیشتر تغییر پایه داده می شوند. برای یافتن بهترین تابع هسته در RKHS که می تواند مؤلفه اصلی یا متمایز خطی را در یک فضا با همبستگی های مرتبه بالای پیکسل های تصویر ورودی محاسبه کند، تصویر ورودی با استفاده از هسته های متعدد به فضای ویژگی مرتبه بالاتر نگاشت می شود و بر اساس کدنویسی در Matlab، هسته ای که بهتر پاسخ می دهد، انتخاب می شود. با توجه به پرکاربردترین توابع کرنل شامل هسته خطی، هسته چند جمله ای، هسته گاوسی و هسته همینگ، در این تحقیق از این انواع هسته استفاده شده است. توابع هسته در RKHS در زیر نشان داده شده است:

$$(16) \quad k(x, y) = \exp\left(-\frac{|x - y|^2}{2\sigma^2}\right)$$

$$(17) \quad k(x, y) = (xy)^d$$

$$(18) \quad k(x, y) = (xy + 1)^d$$

$$(19) \quad k(x, y) = (x'y)$$

$$(20) \quad k(x, y) = 1 - \frac{1}{mN} \sum_{i=1}^m |x_i - y_i|$$

۴-۲- کواترنیون ها

کواترنیون ها در علم ریاضیات، یک دستگاه عددنویسی برای بسط اعداد مختلط هستند. یک عدد کواترنیون به فرم:

$$(21) \quad Q = a.1 + b.i + c.j + dk$$

که در آن a, b, c, d اعداد حقیقی و i, j, k واحدهای اساسی کواترنیون هستند (نمادهایی هستند که می توانند به عنوان بردارهای واحدی که در امتداد سه محور فضایی قرار دارند تفسیر شوند). مجموعه H همه کواترنیون ها یک فضای بردار بر روی اعداد حقیقی با بعد ۴ است (R^4). (در مقایسه؛ اعداد حقیقی R دارای بعد ۱، اعداد مختلط C دارای بعد ۲ و اوتکتان ها نیز دارای بعد ۸ هستند).

H در فضای برداری ۴ بعدی با استفاده از جمع مؤلفه ای و ضرب اسکالر مؤلفه ای، بر روی اعداد حقیقی (a, b, c, d) با پایه های $1, i, j, k$ ساخته شده اند.

جدول ضرب کواترنیون، عمل ضرب کواترنیون را برای سیستم جبری کواترنیون تعریف می کند (جدول ۱) [۱۹]:

جدول ۱: ضرب عناصر پایه

*	1	i	j	k
1	1	i	j	k
i	i	-1	k	-j
j	j	-k	-1	i
k	k	j	-i	-1

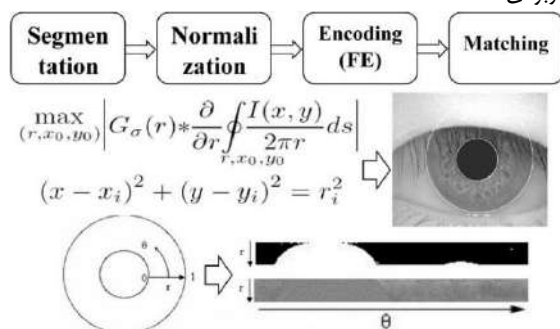
فواید بسیاری در استفاده از کواترنیون نسبت به نمایش های سنتی با مقدار حقیقی در فضای اقلیدسی وجود دارد؛ (۱) اعداد/ بردارهای کواترنیون از یک مؤلفه حقیقی و سه مؤلفه موهومی تشکیل شده اند و باعث تشریح غنی تری از بیان می شوند. (۲) به جای استفاده از ضرب نقطه ای در فضای اقلیدسی، اعداد/ بردارهای کواترنیون بر روی ضرب همپلتون کار می کنند که با اجزای

۱-۳- استخراج ویژگی‌های چهره

در کاربردهای واقعی، از یک طرف تغییرات درونی (مثل بالا رفتن سن و سال، ریش و...) و تغییرات بیرونی (مثل حالت‌های ژست گرفتن، تغییرات در تابش نور، زاویه چرخش و یا استفاده از عینک و...) در ظاهر زیاد می‌باشد و از طرف دیگر ثبت یک مجموعه از مایشی بزرگ چهره، به دلیل هزینه کاری و پیچیدگی محاسباتی می‌تواند بسیار دشوار باشد. بر این اساس بازشناسی چهره به دلیل پیچیدگی و تغییرات و تعداد و سائز نمونه کوچک تصاویر، اغلب یک مسئله غیرخطی می‌باشد. بنابراین توسعه الگوریتم‌هایی که ویژگی‌های موثر و مقیاس‌های غیرخطی دارد و بتوان با استفاده از آنها، روابط غیر خطی میان نمونه‌های چهره را توصیف نمود از اهمیت به سزایی برخوردار است. با الهام گرفتن از اینکه روش‌های هسته، توانایی ثبت تشابهات غیرخطی نمونه‌ها را به صورت موثر دارند روش‌های استخراج ویژگی چهره مبتنی بر هسته برای توسعه الگوریتم‌های خطی ارائه شده است که در آنها با استفاده از توابع هسته مناسب، نمونه‌ها به طور غیر صریح به فضای ویژگی جدیدی با ابعاد بالاتر نگاشت یافته و سپس در این فضای ویژگی جدید، روابط غیرخطی به خطی تبدیل شده و آنگاه معیار فاصله برای کاربرد مورد نظر آموزش داده می‌شود. البته برخلاف عملکرد بسیار خوب توابع هسته در الگوریتم‌های مختلف، یکی از مسائلی که در این الگوریتم‌ها وجود دارد، انتخاب هسته مناسب و یا پارامترهای مناسب برای یک هسته مشخص است [۲۱]. در این مقاله از روش‌های [۲۲] KPCA و [۲۳] KLDA برای استخراج ویژگی استفاده شده است.

۲-۱-۳- استخراج ویژگی‌های عنبیه

سیستم تشخیص عنبیه چشم یک سیستم بیومتریک قابل اعتماد و دقیق است. دو روش داگمن [۲۴] و تبدیل هاف [۲۵] برای استخراج ویژگی‌های عنبیه کاربردی هستند.



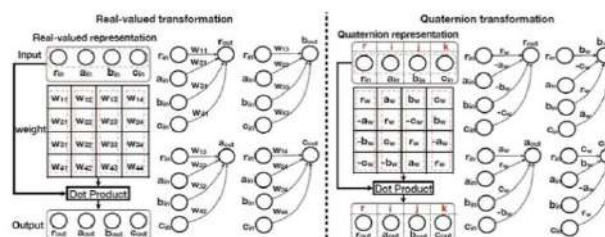
شکل ۶: مکانیابی و نرمال سازی عنبیه

مطابق شکل ۶ الگوریتم استخراج ویژگی‌های عنبیه را می‌توان در سه مرحله خلاصه نمود:

۱- اولین و مرحله حیاتی در روند تشخیص عنبیه، مکانیابی مرزهای عنبیه در تصویر چشم است. در الگوریتم داگمن $I(x, y)$ تصویر چشم و r شعاع جستجو در تصویر و $G(r)$ یک تابع نازک سازی گوسین هست که شروع به جستجو از مردمک می‌کند، تا بتواند تغییر مقادیر حداکثر پیکسل (مشتق جزئی) را تشخیص دهد. در تبدیل هاف نیز جایی که (x_i, y_i) مختصات مرکز و r شعاع است، چشم توسط دو دایره مدل می‌گردد.

۲- نرمال سازی تصویر تفکیک شده از دایره‌ها به بلوک مستطیلی، بصورت یک فرم با ابعاد ثابت می‌باشد.

کوآترنیون چندگانه (پنهان) سازگار است و فعل و انفعالات بین نهان آنها را تقویت می‌کند و منجر به یک مدل با رسایی بالاتر می‌شود. (۳) ماهیت تقسیم وزن ضرب همیلتون منجر به مدلی با تعداد پارامترهای کمتری می‌شود. برای نشان دادن این مزایای استفاده از کوآترنیون، مقایسه ای از فرآیند تبدیل با نمایش‌های کوآترنیون در مقابل نمایش‌ها با مقدار حقیقی در شکل ۴ نشان داده شده است. کوآترنیون کدگذاری تعاملی بین وابستگی بهتری را ارائه می‌دهد و ۷۵٪ تعداد پارامترها را در مقایسه با نمایش‌های با ارزش حقیقی در فضای اقلیدسی کاهش می‌دهد. [۲۰].

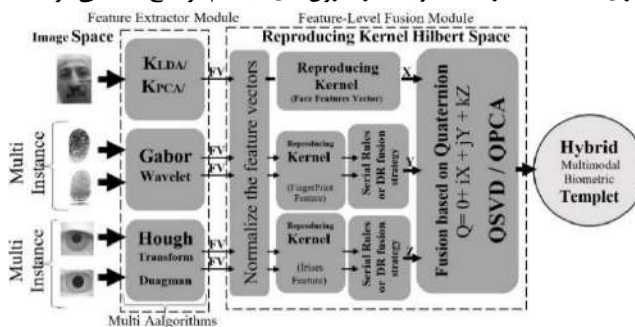


شکل ۴: مقایسه بین تبدیل با مقدار حقیقی (چپ) و تبدیل کوآترنیون (راست) [۲۰].

تلفیق موازی مبتنی بر کوآترنیون را می‌توان در RKHS توسعه داد تا مستقیماً با یک هسته سازگار با داده کار کند. این امر برای تلفیق موازی در یک سطح بالاتر در فضای ابعادی احتمالاً نامحدود H اما با الگوهای نقشه برداری $\Phi(x)$ القا شده توسط یک هسته غیرخطی مطابقت دارد. مزیت این است که با استفاده از هسته مناسب در RKHS تا حد زیادی یک تفکیک پذیری خطی از روابط غیرخطی ویژگی‌های چند بیومتریک حاصل می‌شود.

۳- پیاده سازی سیستم

همانطور که در شکل ۵ نشان داده شده است بلوک دیاگرام سیستم بیومتریک ترکیبی هابیرید پیشنهادی شامل سه ماژول اصلی استخراج ویژگی، تلفیق ویژگی و طبقه بند است در حالیکه متودولوژی پیشنهادی در لایه تلفیق پیاده سازی شده است. در ادامه هر یک از ماژول‌های سیستم توضیح داده می‌شود.



شکل ۵: بلوک دیاگرام سیستم بیومتریک ترکیبی هابیرید.

۱-۳- ماژول استخراج ویژگی

ماژول استخراج ویژگی بهترین ویژگی‌ها را برای هر یک از صفات‌های بیومتریک چهره، هر دو عنبیه و دو اثر انگشت به طور جداگانه استخراج نموده و سیستم از فضای تصویر به فضای ویژگی نگاشت می‌شود. در این بخش الگوریتم استخراج بردارهای ویژگی پنج صفت بیومتریک چهره، عنبیه‌ها و دو اثر انگشت اشاره چپ و راست مختصراً توضیح داده می‌شود.

۳- برای استخراج ویژگی عنبیه از فیلتر گابور-لگاریتم یک بعدی بر روی تصویر نرمال سازی شده برای نمایش اطلاعات بافت عنبیه استفاده می‌گردد. این فیلتر، تابع گاوری است که در مقیاس لگاریتمی عمل می‌کند و فرکانس پاسخ فیلتر گابور-لگاریتم بصورت فرمول زیر است که در آن f_0 مرکز فرکانسی و σ پهنای باند فیلتر را نشان می‌دهند:

$$G(f) = \exp\left(\frac{-(\log(\frac{f}{f_0}))^2}{2(\log(\frac{\sigma}{f_0}))^2}\right) \quad (22)$$

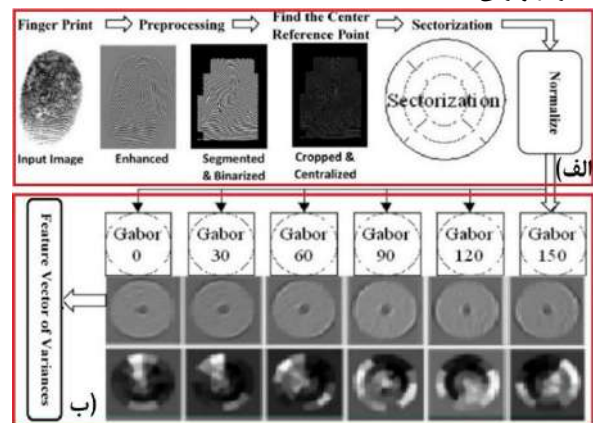
ویژگی‌های عنبیه به صورت کد ۹۶۰۰ بیتی و پلک‌های بالا و پایین به صورت ماسک ۹۶۰۰ بیتی پردازش می‌شوند.

۳-۱-۳- استخراج ویژگی‌های اثر انگشت

از ویژگی بافت اثر انگشت به عنوان فضای ویژگی اثر انگشت استفاده می‌شود. روش‌هایی مثل فیلتر بانک گابور، تطبیق مینوشیا [۲۶]، تبدیل فوریه زمان کوتاه [۲۷] (STFT) و فیلتر بانک گابور [۲۸] برای استخراج ویژگی اثر انگشت استفاده می‌گردد. یکی از روش‌های متداول استخراج ویژگی بانک فیلتر گابور می‌باشد که در شکل ۷ نمایش داده شده است:

الگوریتم استخراج ویژگی‌های اثر انگشت پس از بهبود تصویر اثر انگشت شامل چهار مرحله اصلی است:

- ۱- تعیین کردن نقطه مرجع و ناحیه موردنظر نسبت به آن
- ۲- قطعه بندی ناحیه موردنظر حول نقطه مرجع
- ۳- فیلتر کردن ناحیه موردنظر در شش جهت متفاوت با استفاده از بانک فیلتر گابور
- ۴- محاسبه قدرمطلق انحراف معیار سطوح خاکستری در هر قطعه به منظور تولید بردار ویژگی [۲۹].



شکل ۷: الف) مراحل استخراج ویژگی از اثر انگشت. ب) نتایج اعمال بانک فیلتر گابور بر روی ناحیه موردنظر

۳-۲- مازول تلفیق ویژگی

فضای ویژگی حاوی غنی‌ترین اطلاعات می‌باشد، به این معنی که بردارهای ویژگی نسبت به اطلاعات سطوح دیگر به طور همزمان از نظر کیفیت و کمیت برتری دارند. تلفیق داده در فضای ویژگی که چکیده و اجزای اصلی تشکیل دهنده و تفکیک کننده اطلاعات خام (فضای تصویر) در آن موجود می‌باشند از دو جنبه حایز اهمیت است. اول اینکه می‌تواند ترکیبی از

اطلاعات تفکیک کننده را از مجموعه ویژگی‌های اولیه مشتق کند و دوم اینکه قادر است اطلاعات اضافی و تکراری حاصل از همبستگی بین مجموعه ویژگی‌های جداگانه را حذف کند تا در کوتاهترین زمان ممکن بتواند بهترین تصمیم گیری را داشته باشد. به عبارت دیگر با تلفیق ویژگی‌ها می‌توان بهترین بردار تلفیقی از ویژگی‌ها را بدست آورد بطوریکه بیشترین تمایز را ایجاد کند و کمترین ابعاد ممکن را داشته باشد تا بهترین تصمیم گیری برای سیستم حاصل شود [۱]. در شکل ۳، استراتژی‌های تلفیق بردارها در فضای ویژگی که طی یکی از فرآیندهای "ترکیب سریالی یا موازی [۲،۵]"، "الگوریتم‌های کاهش ابعاد [۳۰-۳۳]" و یا "تلفیق ویژگی‌های باینری [۳۱،۳۳،۳۴]" انجام می‌پذیرد نمایش داده شد.

یکی از سه استراتژی تلفیق در سطح ویژگی، ترکیب بردارهای ویژگی است. به طور کلی، سه حالت اساسی برای فرآیند ترکیب بردارهای ویژگی وجود دارد: قانون سریال و قانون جمع وزنی [۳۵،۳۶] و تلفیق موازی. با این حال، مورد اول منابع محاسباتی زیادی را مصرف می‌کند. برای مورد دوم، انتخاب مقدار وزن یک موضوع بحث برانگیز است. تلفیق موازی توسط یانگ [۳۷] پیشنهاد شد که از محاسبه زیاد در قانون سریال و انتخاب مقدار وزنی در قانون جمع وزنی جلوگیری می‌کند. اما این روش دو ویژگی را به عنوان بخشی حقیقی و بخشی موهومی از یک بردار مختلط در نظر می‌گیرد و فقط می‌تواند دو مدال ویژگی واحد یا یک مدال را با دو نوع ویژگی با هم تلفیق کند. در حالی که اغلب، بازشناسی بیومتریکی ترکیبی باید بتواند از روش‌های بیشتر یا ویژگی‌های بیشتر برای عملکرد بهتر استفاده کند. به همین دلیل در این مقاله تلفیق موازی سه یا چند بردار ویژگی مبتنی بر کوتاهترین پیشنهاد شده است. البته برای عملکرد بهتر، قبل از تلفیق مبتنی بر کوتاهترین، بردارهای ویژگی با استفاده از هسته‌های بازآفرین برای استخراج روابط غیرخطی بیشتر به فضا هیلبرت نگاشت می‌شوند.

استراتژی دیگری که نتایج پیاده سازی آن در این مقاله نمایش داده شده است؛ تلفیق با الگوریتم‌های کاهش ابعاد می‌باشد. ترکیب اطلاعات ابعاد مسئله را بطور قابل ملاحظه ای بیشتر می‌کند بنابراین برخی از روش‌های کاهش ابعاد (انتخاب یا استخراج ویژگی) هم برای مدیریت محاسبات و هم به دست آوردن الگویی که کاهش یافته و به طور موثرتر قابل تفسیر باشد، کاملاً مورد نیاز است. این استراتژی باید چند منظوره باشد تا الگوی تلفیقی خلاصه تر از یک الگوی ترکیبی کامل، ضمن در هم آمیختن اطلاعات ناهمگن بتواند به بهترین تفکیک پذیری بین کلاسی دست یابد. در این استراتژی تلفیقی نیز برای استخراج روابط خطی فضاهای ویژگی، از نگاشت به RKHS بهره برده شده است.

۳-۳- مازول طبقه بندی

همانگونه که در بخش‌های قبلی توضیح داده شد در ابتدا ویژگی‌های تصاویر چهره، عنبیه و اثر انگشت استخراج می‌گردد و بعد از نرمالیزه کردن با نگاشت بردارهای ویژگی به فضای هیلبرت هسته بازآفرین با بهره گیری از الگوریتم‌های مبتنی بر کوتاهترین یا روش‌های کاهش ابعاد تلفیق ویژگی‌ها صورت می‌پذیرد و یک الگوی ترکیبی از بیومتریکی‌ها به نمایندگی از هر کلاس در پایگاه داده ذخیره می‌گردد. در نهایت مازول طبقه بند هربار در فاز بازشناسی، الگوی ترکیبی بیومتریکی جدید بدست آمده از مازول‌های قبل (استخراجگر و تلفیق ویژگی) را با الگوهای ترکیبی که قبلاً در فاز ثبت نام در

که در آن NG و NI به ترتیب نشان دهنده تعداد کل مقایسه‌های درون کلاسی و تعداد کل مقایسه‌های بین کلاسی می‌باشند.

۴-۲- نتایج پیاده سازی

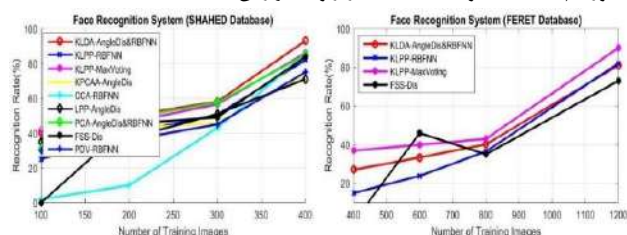
هدف مرحله آموزش، محاسبه پارامترهای لازم به منظور استخراج ویژگی از تصاویر (داده‌های خام) است بطوریکه تمایز حاصل شده از بردارهای ویژگی، تابع هدف (که می‌تواند مقدار دقت شناسایی مطلوب باشد) را محقق نماید به این ترتیب در مرحله تست ابتدا با اعمال همین پارامترها بر روی داده‌های جدید، میزان تمایز بردارهای ویژگی حاصل از آنها را تعیین می‌کنیم.

در آزمایش‌ها تعداد ۱۰۰ کلاس جهت آموزش و تست سیستم در نظر گرفته شده است. برای این منظور، چهره، عنبیه‌های چپ و راست و اثر انگشتان اشاره دست چپ و راست ۱۰۰ نفر از پایگاه داده‌های اشاره شده را انتخاب و بردار ویژگی آنها استخراج گردید. ۸۰٪ تصاویر هر نفر (کلاس) برای آموزش و ۲۰٪ تصاویر باقیمانده هر کلاس برای تست استفاده می‌شود.

در این بخش در ابتدا بهترین نتایج سیستم‌های بازشناسی تک بیومتریک چهره، عنبیه و اثر انگشت با طبقه‌بندهای مناسب با نتایج تلفیق دو اثر انگشت و دو عنبیه چپ و راست در سیستم چند نمونه و چند الگوریتم گزارش شده است و در نهایت نتایج بازشناسی سیستم مالتی مدال بیومتریک پیشنهادی که از تلفیق ویژگی‌های چهره، عنبیه و اثر انگشت بدست آمده است با همان طبقه‌بندها نشان داده می‌شود.

۴-۲-۱- سیستم تک بیومتریک چهره

شکل ۸، نتایج دقت بازشناسی سیستم تک بیومتریک چهره برای الگوریتم‌های خطی استخراج کننده ویژگی PCA، LDA، LPP، FSS، PDV و CCA و الگوریتم‌های غیر خطی مبتنی بر هسته (KPCA، KLDA، KLPP) بر روی دو پایگاه داده FERET و Shahed با محدوده تغییرات متفاوت در چهره هر نفر را نمایش می‌دهد. AUC و بالاترین درصد بازشناسی مربوط به الگوریتم KLDA و طبقه‌بند زاویه گزارش شده است.



Uni-biometric face recognition system					
AUC	دقت بازشناسی	طبقه‌بند کننده	الگوریتم استخراج ویژگی	تصاویر تست	تصاویر آموزش
0.7881	90%	Dis-Angle	Gaussian($\sigma=11$)	200	1200
0.9094	91%	Dis-Angle	Gaussian($\sigma=11$)	100	400

شکل ۸: نتایج دقت سیستم بازشناسی تک بیومتریک چهره برای پایگاه داده‌های FERET و Shahed

۴-۲-۲- سیستم‌های تک بیومتریک و چندالگوریتم عنبیه

دو الگوریتم داگمن و تبدیل هاف برای استخراج ویژگی‌های عنبیه استفاده شده و ۹۶۰۰ ویژگی برای هر یک از عنبیه‌ها استخراج می‌شود که با استفاده

پایگاه داده ذخیره شده مقایسه می‌نماید و بر اساس شباهت بیشتر (یا فاصله کمتر) الگوی جدید با الگوهای ذخیره شده، کلاس آن مشخص می‌گردد.

کارکرد صحیح ماژول طبقه‌بند از اهمیت بالایی در کارایی سیستم برخوردار است. در این مقاله نتایج خروجی نه طبقه‌بند ارزیابی شده است. این طبقه‌بندها شامل چهار طبقه‌بند با تابع فاصله (فواصل زاویه، منتهن، اقلیدسی و فاصله اقلیدسی از مرکز) و دو طبقه‌بند شبکه عصبی تابع شعاعی [38] (RBF) و شبکه عصبی احتمالی [39] (PNN) و طبقه‌بند K نزدیکترین همسایگی [40] (KNN) و طبقه‌بند ماشین بردار پشتیبان مبتنی بر هسته [41] (KSVM) و طبقه‌بند گوسین [۴۲] می‌باشد.

۴- آزمایش‌ها و نتایج پیاده سازی

۴-۱- پارامترهای ارزیابی

برای بررسی عملکرد سیستم پیشنهادی، آزمایش‌ها بوسیله پارامترهای عملکرد شامل دقت، منحنی مشخصه عملکرد (ROC)، AUC (منطقه زیر منحنی ROC)، حساسیت و... در شش پایگاه داده مختلف (دو پایگاه داده چهره FERET و Shahed (پایگاه داده مالتی بیومتریک جمع‌آوری شده در دانشگاه شاهد، تهران، ایران)، تصاویر عنبیه راست و چپ از پایگاه داده عنبیه CASIA و هر دو اثر انگشت اشاره راست و چپ از پایگاه داده Shahed ارزیابی می‌شوند.

جدول ۲: پارامترهای عملکرد

Sensitivity	$TPR \text{ (True Positive Rate)} = \frac{TP}{TP + FN} = \frac{TP}{P} = \text{Recall}$
	$= 1 - FNR \text{ (False Negative Rate)}$ measures the rate of positives that are correctly identified
Specificity	$TNR \text{ (True Negative Rate)} = \frac{TN}{TN + FP} = \frac{TN}{N} = \text{Selectivity}$
	$= 1 - FPR \text{ (False Positive Rate)}$ measures the rate of negatives that are correctly identified
positive likelihood ratio	$\frac{\text{Sensitivity}}{1 - \text{Specificity}} = \frac{TPR}{FPR}$
	PLR, likelihood ratio positive, likelihood ratio for positive results
negative likelihood ratio	$\frac{1 - \text{Sensitivity}}{\text{Specificity}} = \frac{FNR}{TNR}$
	NLR, likelihood ratio negative, likelihood ratio for negative results
Accuracy	$\frac{TP + TN}{TP + FN + TN + FP} = \frac{TP + TN}{P + N}$
	closeness of the measurements to a specific value
Precision	$\frac{TP}{TP + FP}$
	closeness of the measurements to each other

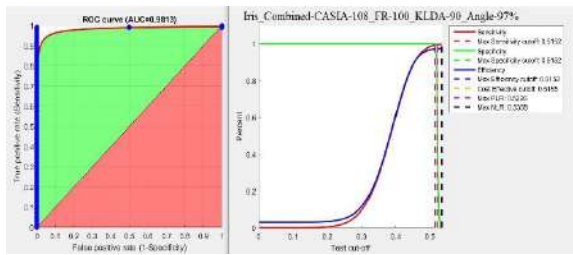
البته منحنی‌های ROC و پارامترهای دقت برای تأیید عملکرد کلی سیستم چندبیومتریک کافی نیستند. از این رو Bengio و همکاران [43] یک آزمون آماری را ارائه کرد که از نصف میزان خطای کل (HTER) و فاصله اطمینان (CI) استفاده می‌کند:

$$HTER = \frac{FPR + FNR}{2} \quad (23)$$

برای محاسبه CI حول HTER، به دنبال کران $\sigma \times z\alpha/2$ می‌گردیم. در اینجا، σ و $z\alpha/2$ به صورت زیر تعریف می‌شوند:

$$\sigma = \sqrt{\frac{FPR.TNR}{4.NI} + \frac{FNR.TPR}{4.NG}} \quad (24)$$

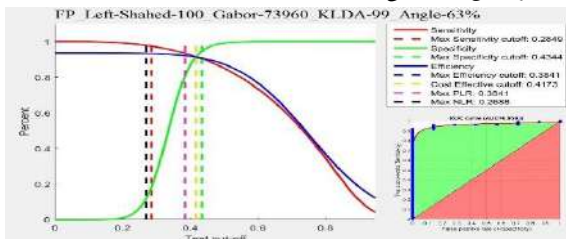
$$z\alpha/2 = \begin{cases} 1.645 & \text{for } 90\% \text{ CI} \\ 1.960 & \text{for } 95\% \text{ CI} \\ 2.576 & \text{for } 99\% \text{ CI} \end{cases} \quad (25)$$



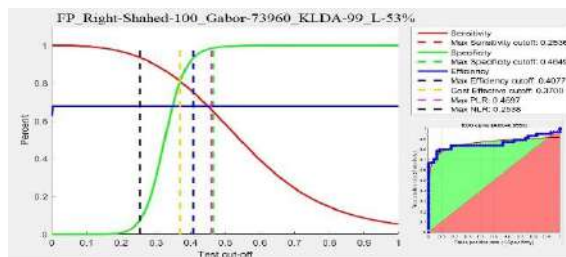
شکل ۹: منحنی ROC برای سیستم های تک بیومتریک و چندالگوریتم عنبیه بر روی پایگاه داده CASA، الف) عنبیه راست (Daugman (Alg.)، ب) عنبیه چپ (Hough Tra.)، ج) ترکیب عنبیه های راست و چپ با استراتژی کاهش ابعاد (الگوریتم KLDA)

۳-۲-۴- سیستم های تک بیومتریک و چند نمونه اثر انگشت

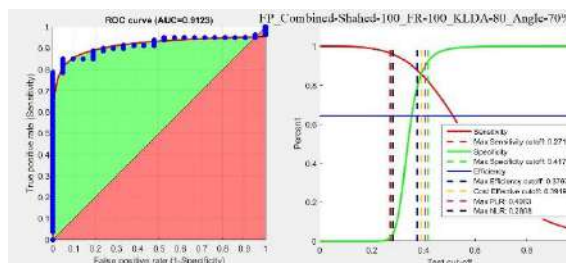
با اعمال هشت فیلتر لگاریتم-گابور در فرکانس های متفاوت، ۷۳۹۶۰ ویژگی از هر اثر انگشت استخراج می شود. سپس با اعمال آنالیز تفکیک کننده خطی (LDA) و آنالیز مولفه اصلی (PCA) در RKHS، ۷۳۹۶۰ ویژگی اثر انگشت را به ۱۵۰ ویژگی کاهش داده و پس از ترکیب سری بردارهای ویژگی اثر انگشت اشاره چپ و راست، بردار ترکیبی را به عنوان ورودی به طبقه بند می دهیم. منحنی ROC برای سیستم تک بیومتریک و سیستم چند نمونه اثر انگشت در شکل ۱۰ نمایش داده شده است.



الف)



ب)



ج)

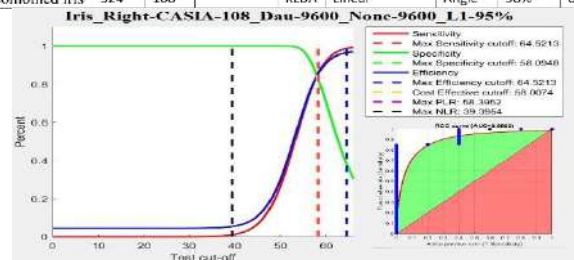
شکل ۱۰: منحنی ROC سیستم های تک بیومتریک و چند نمونه اثر انگشت بر روی پایگاه داده Shaded، الف) اثر انگشت اشاره چپ، ب) اثر انگشت اشاره راست و ج) ترکیب اثر انگشتان اشاره راست و چپ با استراتژی ترکیب قانون سریال و طبقه بند Dis-Angle

از پنج تابع هسته (gaussian, polyplus, polynomial, linear and hamming) بردارها از فضای ویژگی به فضای هیلبرت هسته بازآفرین نگاشت شده است. دو بردار ویژگی عنبیه چپ و راست در RKHS با الگوریتم KLDA (استراتژی کاهش ابعاد) تلفیق می شوند که بردار تلفیقی تنها با ۹۰ ویژگی دقت باز شناسی بالایی دارد. نتایج سیستم تک بیومتریک عنبیه چپ و راست و همچنین سیستم مالتی-الگوریتم عنبیه ترکیبی در فضای هیلبرت هسته بازآفرین به ازای هر کدام از پنج تابع هسته اشاره شده گزارش شده است. همچنین منحنی ROC سیستم برای هر یک از تک بیومتریکهای عنبیه چپ و راست و عنبیه ترکیبی مبتنی بر استراتژی کاهش ابعاد در سیستم چندالگوریتم در شکل ۹ نمایش داده شده است.

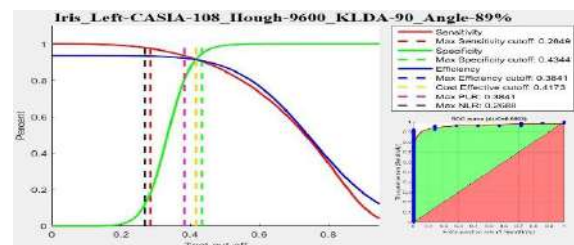
جدول ۳: مقایسه نتایج دقت سیستم های باز شناسی تک بیومتریک و چندالگوریتم عنبیه

Iris Multi-algorithm System (CASIA Database)											
Reproducing Kernel Hilbert Space											
Reproducing Kernel	Biometric	Algorithm	Data	Classifier	Accuracy	Reproducing Kernel	Biometric	Algorithm	Data	Classifier	Accuracy
NONE	Right Iris	Daugman	9600	Discrim-L1	93.37%	Gaussian kernel ($\sigma=30$)	Right Iris	Daugman	9600	Gaussian	93.37%
	Left Iris	Daugman	9600	Discrim-L1	93.37%		Left Iris	Daugman	9600	Gaussian	93.37%
	Combined Iris	Daugman	9600	Discrim-L1	93.37%		Combined Iris	Daugman	9600	Gaussian	93.37%
	Right Iris	Daugman	9600	Discrim-L1	93.37%		Left Iris	Daugman	9600	Gaussian	93.37%
	Left Iris	Daugman	9600	Discrim-L1	93.37%		Combined Iris	Daugman	9600	Gaussian	93.37%
	Combined Iris	Daugman	9600	Discrim-L1	93.37%		Right Iris	Daugman	9600	Gaussian	93.37%
Linear kernel	Right Iris	Daugman	9600	Discrim-L1	93.37%	PolyPlus kernel ($d=2$)	Right Iris	Daugman	9600	PolyPlus	93.37%
	Left Iris	Daugman	9600	Discrim-L1	93.37%		Left Iris	Daugman	9600	PolyPlus	93.37%
	Combined Iris	Daugman	9600	Discrim-L1	93.37%		Combined Iris	Daugman	9600	PolyPlus	93.37%
	Right Iris	Daugman	9600	Discrim-L1	93.37%		Left Iris	Daugman	9600	PolyPlus	93.37%
	Left Iris	Daugman	9600	Discrim-L1	93.37%		Combined Iris	Daugman	9600	PolyPlus	93.37%
	Combined Iris	Daugman	9600	Discrim-L1	93.37%		Right Iris	Daugman	9600	PolyPlus	93.37%
Hamming kernel	Right Iris	Daugman	9600	Discrim-L1	93.37%	Polynomial kernel	Right Iris	Daugman	9600	Polynomial	93.37%
	Left Iris	Daugman	9600	Discrim-L1	93.37%		Left Iris	Daugman	9600	Polynomial	93.37%
	Combined Iris	Daugman	9600	Discrim-L1	93.37%		Combined Iris	Daugman	9600	Polynomial	93.37%
	Right Iris	Daugman	9600	Discrim-L1	93.37%		Left Iris	Daugman	9600	Polynomial	93.37%
	Left Iris	Daugman	9600	Discrim-L1	93.37%		Combined Iris	Daugman	9600	Polynomial	93.37%
	Combined Iris	Daugman	9600	Discrim-L1	93.37%		Right Iris	Daugman	9600	Polynomial	93.37%

Uni-biometric & Multi-algorithm Iris recognition system								
دیتابیس	تصاویر آموزش	تصاویر تست	استخراج ویژگی	استراتژی کاهش ابعاد		طبقه بند	دقت بازشناسی	AUC
Right-CASIA	324	108	Daugman	None	Dim=9600	Dis-L1	95%	0.8892
Left-CASIA	324	108	Hough	KLDA	PolyPlus(d=2)	Angle	90%	0.9593
combined iris	324	108		KLDA	Linear	Angle	98%	0.9813



الف)



ب)

۴-۲-۴- سیستم بیومتریک ترکیبی هایبرید

مبتنی بر کوآترینیون (QPCA و QSVD) بدست آمده تنها با طول ۱۵ و ۱۲۰ ویژگی دقت بازشناسی ۱۰۰٪ را نتیجه می‌دهند که نسبت به سایر سیستم‌های مالتی بیومتریک و تک بیومتریک کارایی بیشتری دارد.

مراجع

- [1] Li, Y., Zou, B., Deng, S., & Zhou, G. (2020). Using feature fusion strategies in continuous authentication on smartphones. *IEEE Internet Computing*, 24(2), 49-56.
- [2] Jain, A. K., Ross, A. A., & Nandakumar, K. (2011). *Introduction to biometrics*. Springer Science & Business Media.
- [3] Joseph, T., Kalaiselvan, S. A., Aswathy, S. U., Radhakrishnan, R., & Shamna, A. R. (2021). A multimodal biometric authentication scheme based on feature fusion for improving security in cloud environment. *Journal of Ambient Intelligence and Humanized Computing*, 12(6), 6141-6149.
- [4] Wang, Z., Zhen, J., Li, Y., Li, G., & Han, Q. (2019). Multi - feature multimodal biometric recognition based on quaternion locality preserving projection. *Chinese Journal of Electronics*, 28(4), 789-796.
- [5] Tong, Y., Bai, J., & Chen, X. (2020, October). Research on Multi-sensor Data Fusion Technology. In *Journal of Physics: Conference Series* (Vol. 1624, No. 3, p. 032046). IOP Publishing.
- [6] Zhang, Y., Gao, C., Pan, S., Li, Z., Xu, Y., & Qiu, H. (2020). A score-level fusion of fingerprint matching with fingerprint liveness detection. *IEEE Access*, 8, 183391-183400.
- [7] Abd Ghani, M. K., Mohammed, M. A., Arunkumar, N., Mostafa, S. A., Ibrahim, D. A., Abdullah, M. K., ... & Burhanuddin, M. A. (2020). Decision-level fusion scheme for nasopharyngeal carcinoma identification using machine learning techniques. *Neural Computing and Applications*, 32(3), 625-638.
- [8] Schölkopf, B., Smola, A. J., & Bach, F. (2002). *Learning with kernels: support vector machines, regularization, optimization, and beyond*. MIT press.
- [9] Suykens, J. A., Van Gestel, T., De Brabanter, J., De Moor, B., & Vandewalle, J. P. (2002). *Least squares support vector machines*. World scientific.
- [10] Kafai, M., & Eshghi, K. (2017). CROification: accurate kernel classification with the efficiency of sparse linear SVM. *IEEE transactions on pattern analysis and machine intelligence*, 41(1), 34-48.
- [11] Zhu, J., & Hastie, T. (2005). Kernel logistic regression and the import vector machine. *Journal of Computational and Graphical Statistics*, 14(1), 185-205.
- [12] Drucker, H., Burges, C. J., Kaufman, L., Smola, A., & Vapnik, V. (1997). Support vector regression machines. *Advances in neural information processing systems*, 9, 155-161.
- [13] Scholkopf, B., Mika, S., Burges, C. J., Knirsch, P., Muller, K. R., Ratsch, G., & Smola, A. J. (1999). Input space versus feature space in kernel-based methods. *IEEE transactions on neural networks*, 10(5), 1000-1017.
- [14] Dhillon, I. S., Guan, Y., & Kulis, B. (2004, August). Kernel k-means: spectral clustering and normalized cuts. In *Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining* (pp. 551-556).
- [15] Liu, F., Huang, X., Chen, Y., & Suykens, J. A. (2020). Random Features for Kernel Approximation: A Survey on

جدول ۴ مقایسه نتایج دقت بازشناسی سیستم ترکیبی هایبرید پیشنهادی با سه استراتژی کاهش ابعاد، تلفیق سری و تلفیق موازی بوسیله کوآترینیون با سیستم‌های تک بیومتریک چهره، عنبیه و اثر انگشت و سیستم‌های مالتی نمونه اثر انگشت و مالتی الگوریتم عنبیه در فضای هیلبرت هسته بازآفرین و طبقه بند KSVM گزارش نموده است.

جدول ۴: مقایسه سه استراتژی تلفیق در فضای ویژگی RKHS ترکیب سری - استراتژی تلفیق مبتنی بر کاهش ابعاد - تلفیق موازی بوسیله کوآترینیون در سیستم ترکیبی هایبرید پیشنهادی

Face, Iris and Fingerprint Hybrid Multimodal Biometric System					
Reproducing Kernel Hilbert Space	Bio.	Algorithm	Reproducing Kernel	Dim	Kernel SVM Linear
	Face	KLDA	Hamming	90	95%
	R-Iris	Daugman	PolyPlus(d=2)	90	95.37%
	L-Iris	Hough	Polynomial(d=2)	90	92.60%
	Combined Iris	Dim Reduction	Linear	100	97.22%
	R-Index	Log-Gabor	Gaussian($\sigma=269$)	99	49%
	L-Index	Log-Gabor	Gaussian($\sigma=269$)	99	60%
	Combined Fingerprint	Dim Reduction (KLDA)	Gaussian($\sigma=49$)	99	79%
Fusion Strategy	Serial Combination	Serial rule		289	84%
	Dimensionality Reduction	KLDA	Polynomial	15	100%
	Parallel fusion	Quaternion	QPCA	120	100%
			QSVD	120	100%

در سیستم پیشنهاد شده که در فضای هیلبرت هسته بازآفرین، بردارهای ویژگی چهره، عنبیه ترکیبی و اثر انگشت ترکیبی را بر اساس استراتژی کاهش ابعاد تلفیق می‌نماید تنها با ۱۵ ویژگی به بازشناسی ۱۰۰٪ دست پیدا نموده ایم. همچنین تلفیق موازی بوسیله کوآترینیون با دو الگوریتم QPCA و QSVD به ۱۰۰٪ دقت بازشناسی دست می‌یابیم.

۵- نتیجه

از آنجا که فضای ویژگی دارای اطلاعات غنی تر (کیفیت و کمیت بالاتر) نسبت به فضاهای تصویر و تصمیم گیری است، تلفیق در سطح ویژگی نسبت به تلفیق در سایر سطوح (سُور، نمره و تصمیم) موثرتر است. در این مقاله یک سیستم بیومتریک ترکیبی هایبریدی پیشنهاد گردید که برای بدست آوردن یک الگوی ترکیبی قوی و ایمن، ویژگی‌های چهره، هر دو عنبیه و دو اثر انگشتان اشاره چپ و راست را تلفیق می‌نماید. در متولوزی پیشنهادی با هدف استخراج روابط خطی فضاهای ویژگی، از نگاشت بردارهای ویژگی به RKHS بهره برده شده است و سپس تلفیق بوسیله الگوریتم‌های کاهش ابعاد مبتنی بر هسته و روش ترکیب موازی مبتنی بر کوآترینیون هم به منظور مدیریت محاسبات و صرفه زمان و هم به دست آوردن الگویی که کاهش یافته و به طور موثرتر تفکیک کننده بین کلاسی و ایمن و مقاوم در برابر جعل باشد، ارائه شده است.

در بررسی عملکرد سیستم پیشنهادی، آزمایش‌ها بوسیله پارامترهای دقت تشخیص، ROC، AUC و حساسیت و... در شش پایگاه داده مختلف (دو پایگاه داده چهره FERET و Shahed، دو پایگاه داده عنبیه چپ و راست CASIA و هر دو اثر انگشت اشاره راست و چپ از پایگاه داده Shahed ارزیابی شدند. نتایج آزمایش‌ها به خوبی نشان می‌دهد الگوی ترکیبی هایبریدی که با استفاده از استراتژی کاهش ابعاد با الگوریتم KLDA و تلفیق موازی

- [32] Kamlaskar, C., Deshmukh, S., Gosavi, S., & Abhyankar, A. (2019). Novel canonical correlation analysis based feature level fusion algorithm for multimodal recognition in biometric sensor systems. *Sensor Letters*, 17(1), 75-86.
- [33] Zhang, H., Li, S., Shi, Y., & Yang, J. (2019). Graph fusion for finger multimodal biometrics. *IEEE Access*, 7, 28607-28615.
- [34] Kabir, W., Ahmad, M. O., & Swamy, M. N. S. (2019). A multi-biometric system based on feature and score level fusions. *IEEE Access*, 7, 59437-59450.
- [35] De Marsico, M., Nappi, M., Riccio, D., & Tortora, G. (2010). NABS: Novel approaches for biometric systems. *IEEE Transactions on Systems, Man, and Cybernetics, Part C (Applications and Reviews)*, 41(4), 481-493.
- [36] Yao, Y. F., Jing, X. Y., & Wong, H. S. (2007). Face and palmprint feature level fusion for single sample biometrics recognition. *Neurocomputing*, 70(7-9), 1582-1586.
- [37] Yang, J., Yang, J. Y., & Frangi, A. F. (2003). Combined fisherfaces framework. *Image and Vision computing*, 21(12), 1037-1044.
- [38] Roguia, S., & Mohamed, N. (2019). An optimized rbf-neural network for breast cancer classification. *International Journal of Informatics and Applied Mathematics*, 1(1), 24-34.
- [39] Tang, Z. (2020, August). Leaf image recognition and classification based on GBDT-probabilistic neural network. In *Journal of Physics: Conference Series* (Vol. 1592, No. 1, p. 012061). IOP Publishing.
- [40] Prabavathy, S., Rathikarani, V., & Dhanalakshmi, P. (2020) Classification of Musical Instruments using SVM and KNN. *International Journal of Innovative Technology and Exploring Engineering (IJITEE)* ISSN, 2278-3075.
- [41] Hekmatmanesh, A., Wu, H., Jamaloo, F., Li, M., & Handroos, H. (2020). A combination of CSP-based method with soft margin SVM classifier and generalized RBF kernel for imagery-based brain computer interface applications. *Multimedia Tools and Applications*, 79(25), 17521-17549.
- [42] Dan, C., Wei, Y., & Ravikumar, P. (2020, November). Sharp statistical guarantees for adversarially robust gaussian classification. In *International Conference on Machine Learning* (pp. 2345-2355). PMLR.
- [43] Bengio, S., & Mariéthoz, J. (2004). A statistical significance test for person authentication. In *Proceedings of Odyssey 2004: The Speaker and Language Recognition Workshop* (No. CONF). Algorithms, Theory, and Beyond. arXiv preprint arXiv:2004.11154.
- [16] Saitoh, S., & Sawano, Y. (2016). *Theory of reproducing kernels and applications*. Singapore: Springer Singapore.
- [17] Seto, M., Suda, S., & Taniguchi, T. (2014). Gram matrices of reproducing kernel Hilbert spaces over graphs. *Linear Algebra and its Applications*, 445, 56-68.
- [18] Klus, S., Schuster, I., & Muandet, K. (2020). Eigendecompositions of transfer operators in reproducing kernel Hilbert spaces. *Journal of Nonlinear Science*, 30(1), 283-315
- [19] Fister, I., Yang, X. S., Brest, J., & Fister Jr, I. (2013). Modified firefly algorithm using quaternion representation. *Expert Systems with Applications*, 40(18), 7220-7230.
- [20] Tran, T., You, D., & Lee, K. (2020, October). Quaternion-based self-attentive long short-term user preference encoding for recommendation. In *Proceedings of the 29th ACM International Conference on Information & Knowledge Management* (pp. 1455-1464).
- [21] Wang, D., Lu, H., & Yang, M. H. (2015). Kernel collaborative face recognition. *Pattern Recognition*, 48(10), 3025-3037.
- [22] Wang, A., & Zhao, W. (2021, June). Research on face recognition technology based on KPCA and SVM under convolutional filtering. In *International Conference on Signal Image Processing and Communication (ICSIPC 2021)* (Vol. 11848, p. 1184800). International Society for Optics and Photonics.
- [23] Kak, S. F., Mustafa, F. M., & Valente, P. (2018). A review of person recognition based on face model. *Eurasian Journal of Science & Engineering*, 4(1), 157-168.
- [24] Alam, M. M., Khan, M. A. R., Salehin, Z. U., Uddin, M., Soheli, S. J., & Khan, T. Z. (2020). Combined PCA-Daugman Method: An Efficient Technique for Face and Iris Recognition. *Journal of Advances in Mathematics and Computer Science*, 34-44.
- [25] Abikoye, O. C., Ojo, U. A., Awotunde, J. B., & Ogundokun, R. O. (2020). A safe and secured iris template using steganography and cryptography. *Multimedia Tools and Applications*, 79(31), 23483-23506.
- [26] Patel, R. B., Hiran, D., & Patel, J. (2021). Biometric Fingerprint Recognition Using Minutiae Score Matching. In *Data Science and Intelligent Applications* (pp. 463-478). Springer, Singapore.
- [27] Manickam, A., Devarasan, E., Manogaran, G., Priyan, M. K., Varatharajan, R., Hsu, C. H., & Krishnamoorthi, R. (2019). Score level based latent fingerprint enhancement and matching using SIFT feature. *Multimedia Tools and Applications*, 78(3), 3065-3085.
- [28] Onifade, O. F., Akinde, P., & Isinkaye, F. O. (2020). Circular Gabor wavelet algorithm for fingerprint liveness detection. *Journal of Advanced Computer Science & Technology*, 9(1),1-5.
- [29] Jain, A. K., Nandakumar, K., & Ross, A. (2016). 50 years of biometric research: Accomplishments, challenges, and opportunities. *Pattern recognition letters*, 79, 80-105.
- [30] Joseph, T., Kalaiselvan, S. A., Aswathy, S. U., Radhakrishnan, R., & Shamna, A. R. (2021). A multimodal biometric authentication scheme based on feature fusion for improving security in cloud environment. *Journal of Ambient Intelligence and Humanized Computing*, 12(6), 6141-6149.
- [31] Saini, N., & Sinha, A. (2020). Efficient fusion of face and palmprint in Gabor filtered Wigner domain. *International Journal of Biometrics*, 12(3), 301-316.



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

یک مدل ریاضی دو هدفه فازی برای مسأله زمانبندی زنجیره بحرانی با در نظر گرفتن محدودیت بودجه و عدم قطعیت

فاطمه دلوچی^۱، سید میثم موسوی^۲، احمد مینائی^۳

^۱ دانشجوی کارشناسی ارشد، گروه مهندسی صنایع، دانشگاه شاهد، تهران، fatemeh.dalouchei@shahed.ac.ir

^۲ دانشیار، گروه مهندسی صنایع، دانشگاه شاهد، تهران، sm.mousavi@shahed.ac.ir

^۳ دانشجوی دکتری، گروه مهندسی صنایع، دانشگاه شاهد، تهران، ahmad.minaei@shahed.ac.ir

چکیده: زمانبندی و ارائه‌ی برنامه‌ی زمانی پروژه با استفاده از روش زنجیره‌ی بحرانی، مسأله‌ی محدودیت منابع را مورد توجه قرار داده و با معرفی و تعیین اندازه‌ی زمان‌هایی تحت عنوان زمان ایمنی یا بافر به مقابله با تأخیرهای احتمالی در انجام فعالیت‌ها می‌پردازد. پژوهش حاضر، یک مدل ریاضی دو هدفه شامل کاهش میزان تأخیر و افزایش کیفیت پروژه را برای ارائه یک برنامه بهینه براساس اصول روش زنجیره بحرانی و محدودیت بودجه پیشنهاد می‌کند. به دلیل وجود عدم قطعیت در مسائل واقعی، ابهام در پارامتر مدت زمان فعالیت‌ها با استفاده از اعداد فازی مثلثی بیان شده سپس با توجه به یک روش از ادبیات موضوع، معادل قطعی مدل ارائه می‌شود. همچنین، برای افزودن بافر به برنامه پروژه نیاز به تعیین فعالیت‌های بحرانی است که در این پژوهش این فرآیند با مینا قرار دادن میزان شناوری فعالیت‌ها انجام می‌شود. در پایان نیز با استفاده از یک مثال کاربردی، مدل ریاضی حل شده که به دلیل چند هدفه بودن مدل فرآیند حل آن با روش برنامه‌ریزی آرمانی انجام می‌پذیرد. نتایج حاصل شده از حل مدل در حالت عدم قطعیت، بهبود در جواب‌ها را نسبت به شرایطی که مدل در حالت قطعی است، نشان می‌دهد.

کلمات کلیدی: زمانبندی پروژه، زنجیره‌ی بحرانی، محدودیت بودجه، کیفیت، عدم قطعیت فازی

احتمالی زمان‌های ایمنی یا بافر در قسمت‌های مختلفی از برنامه ارائه می‌شود

[1].

۱- مقدمه

به منظور اندازه‌گیری بافرها، نیولبد^۲ [2] روش جذر مجموع مربعات خطا را برای تعیین اندازه بافر پیشنهاد نمود. تاگل^۳ و همکاران [3] با یکپارچه نمودن ویژگی‌های پروژه به روش‌های اندازه‌گیری بافر، روش سازگار با فشردگی منابع و روش سازگار با چگالی شبکه را ارائه دادند. در مطالعه‌ی بالتا^۴ و همکاران [4] با در نظر گرفتن اهمیت اتمام پروژه‌های سد و نیروگاه‌های هیدروالکتریکی در پایان زمان تعیین شده، یک روش اندازه‌گیری بافر بر مبنای ارزیابی ریسک فازی معرفی شده است. به منظور کنترل برنامه، با در نظر گرفتن ویژگی پویای اجرای

در دنیای بزرگ و پیچیده‌ی امروز، در هر لحظه از زمان پروژه‌های بسیاری آغاز و اختتام می‌یابند که بر اساس سنجش عملکردشان می‌توان میزان موفقیت را در آنها ارزیابی کرد. معیارها و اهداف متفاوتی در ارزیابی میزان موفقیت یک پروژه اثرگذار هستند که یکی از آنها انجام پروژه‌ها در چارچوب محدودیت و برنامه زمانی تعیین شده است. یکی از روش‌های زمانبندی پروژه‌ها روش زنجیره‌ی بحرانی^۱ است که در سال ۱۹۹۷ توسط گلدنر معرفی شد. در این روش توالی از فعالیت‌ها، با در نظر گرفتن روابط وابستگی و محدودیت منابع، زنجیره‌ی بحرانی پروژه را تشکیل می‌دهند و به منظور مواجهه با تأخیرهای

³ Tukel

⁴ Balta

¹ Critical Chain

² Newbold

پروژه یک روش کنترل بافر جدید توسط هو¹ و همکاران [5] بر اساس معیار حساسیت فعالیت پیشنهاد شد. در یک مطالعه اخیر، زهرهوندی و خلیلزاده [6] با در نظر گرفتن ارزیابی حالات بالقوه‌ی خطا و اثرات آن به ارائه یک روش مؤثر برای اندازه‌گیری بافر پروژه پرداختند.

محدودیت منابع یکی از عوامل تأخیر در زمان انجام فعالیت است. محدودیت منابع در ساختار یک مسئله‌ی استاندارد برنامه‌ریزی پروژه‌ی منابع^۲ محدود قابل ارزیابی است که هدف از آن محاسبه‌ی زمان بهینه‌ی اجرای پروژه است [7]. به دلیل بهینه نبودن برنامه زمانبندی در روش زنجیره بحرانی، یک روش زنجیره‌ی بحرانی فازی برای زمانبندی پروژه‌ها تحت محدودیت‌های منابع و عدم قطعیت توسط لانگ و اوساتو^۳ [8] معرفی شده است. در مطالعه روغنیان و همکاران [9] نیز، به ارائه‌ی یک روش زنجیره‌ی بحرانی بهبود یافته شده با رویکردی فازی برای زمانبندی پروژه‌ها تحت عدم قطعیت پرداخته شده است.

بر اساس ادبیات موضوع، در نظر گرفتن اهدافی مانند کیفیت پروژه همزمان با ایجاد یک برنامه زمانی بهینه در روش زنجیره بحرانی مورد بررسی قرار نگرفته است. در نتیجه، در این پژوهش یک مدل ریاضی عدد صحیح مختلط دو هدفه بر اساس اصول روش زنجیره‌ی بحرانی و محدودیت در میزان بودجه که محدودیت در دسترسی منابع مورد نیاز را ایجاد می‌نماید، ارائه می‌شود. مدل پیشنهادی با تعریف چند حالت اجرایی برای فعالیت‌ها در ضمن ایجاد یک برنامه زمانی بهینه به اهداف حداقل سازی میزان تأخیرهای پروژه و افزایش کیفیت انجام فعالیت‌ها نیز می‌پردازد.

پژوهش حاضر برای تعیین اندازه بافر و قرار دادن آن در برنامه از میزان چگالی شبکه استفاده نموده است. به دلیل اینکه برای اندازه‌گیری بافر نیاز به مشخص کردن فعالیت‌های بحرانی است، با استفاده از تعریف دو محدودیت در مدل پیشنهادی میزان شنواری فعالیت‌ها تعیین شده و سپس فعالیت‌های بحرانی مشخص می‌شوند. این موضوع در ادبیات به ندرت مورد بررسی قرار گرفته است.

همچنین به دلیل وجود عدم قطعیت در مسائل دنیای واقعی، ابهام موجود در پارامتر مدت زمان هر یک از فعالیت‌ها با استفاده از اعداد فازی مثلثی بیان شده که برای قطعی‌سازی مدل از روش معرفی شده توسط خیمنز^۴ و همکاران [10] استفاده می‌شود. در نهایت با توجه به چند هدفه بودن مدل از روش برنامه‌ریزی آرمانی فرآیند حل آن انجام می‌پذیرد. در ادامه به بیان مسئله، حل مثال کاربردی و نتایج پرداخته شده است.

۲- بیان مسئله و مدل‌سازی

ابتدا در بخش ۲-۱، مسئله مورد مطالعه بیان شده و سپس مدل ریاضی پیشنهادی در بخش ۲-۲ معرفی می‌شود. به دلیل وجود عدم قطعیت فازی در پارامتر مدت زمان، فرآیند معادل‌سازی با استفاده از روش خیمنز و همکاران [10] در قسمت ۲-۳ ارائه می‌شود.

۲-۱ بیان مسئله

مسئله مورد مطالعه در این پژوهش عبارت است از برنامه‌ریزی به منظور انجام فعالیت‌های پروژه با در نظر داشتن اصول روش زنجیره‌ی بحرانی و محدودیت بودجه به نحوی که میزان تأخیر پروژه حداقل شده و فعالیت‌ها در بالاترین کیفیت ممکن انجام پذیرند. بر اساس روش زنجیره‌ی بحرانی زمان‌های ایمنی که با هدف اطمینان از انجام به موقع فعالیت‌ها پیش‌بینی شده‌اند از پایان آنها حذف می‌شوند؛ و فعالیت‌ها در مدت زمان میانگین برنامه‌ریزی می‌شوند. با توجه به اینکه برای هر فعالیت چندین حالت اجرایی تعریف شده است؛ بنابراین مدت زمان انجام یک فعالیت بسته به حالت اجرایی آن دارد که به منظور برنامه‌ریزی فعالیت‌ها در حالت‌های اجرایی با مدت زمان کمتر نیازمند استفاده از منابع بیشتر است که تأمین منابع مورد نیاز نیز بر اساس بودجه از پیش تعیین شده با محدودیت مواجهه است. ساختار کلی محدودیت‌های مدل به سه بخش محدودیت‌های پیشینازی، محدودیت دسترسی منابع تجدیدپذیر و نحوه تعیین فعالیت‌های بحرانی قابل تقسیم است.

۲-۲ مدل‌سازی ریاضی

مجموعه‌ها

$j = 1, 2, \dots, N$	مجموعه فعالیت‌های پروژه
$t = 1, 2, \dots, T$	مجموعه دوره‌های زمانی پروژه
$m = 1, 2, \dots, M_j$	مجموعه حالت‌های اجرای فعالیت j
$k = 1, 2, \dots, K$	مجموعه منابع تجدیدپذیر

پارامترها

dd	موعد تحویل پروژه
tb	کل بودجه منابع
cr_k	هزینه هر واحد منبع تجدیدپذیر k
Ef_j	زودترین زمان پایان فعالیت j
Lf_j	دیرترین زمان پایان فعالیت j
d_{jm}	مدت زمان فعالیت j ام در حالت m ام
r_{jmk}	منبع تجدیدپذیر k مورد نیاز فعالیت j ، زمانی که در حالت m اجرا می‌شود.
$p(i, j)$	مجموعه روابط پیشینازی فعالیت‌های پروژه
σ_i	انحراف معیار فعالیت i
Q_{jm}	کیفیت فعالیت j در حالت m
M	عدد بسیار بزرگ

متغیرها

BR_k	مقدار منبع تجدیدپذیر k تخصیص یافته به پروژه
Tc	میزان تأخیر پروژه
R_k	مقدار کل منبع تجدیدپذیر k مورد نیاز

⁴ Jiménez

¹ Hu

² Resource-constrained project scheduling problem

³ Long & Ohsato

$\forall i$

$$epp_i \leq M \cdot (1 - beta_i), \quad \forall i \quad (10)$$

$$epp_i \geq 1 - beta_i, \quad \forall i \quad (11)$$

$$(kk - 1) \cdot \sum_{i=1}^N beta_i = \sum_{i=1}^{N-1} \sum_{j=2}^N p(i, j) \cdot beta_i \cdot beta_j \quad (12)$$

رابطه (۱) تابع هدف اول مسأله است که مجموع تأخیرهای پروژه (شامل میزان تفاوت بین زمان پایان پروژه و موعد تحویل تعیین شده و اندازه بافر) را کمینه می‌سازد. رابطه (۲) تابع هدف دوم است که بیشینه‌سازی کیفیت پروژه را نمایش می‌دهد. رابطه (۳) اطمینان می‌دهد که هر فعالیت تنها در یک حالت اجرایی و در یک زمان انجام شود. رابطه (۴) نمایش روابط پیشیناری میان فعالیت‌ها است. رابطه (۵) بیشترین میزان استفاده از منابع تجدیدپذیر را تعیین می‌کند. رابطه (۶) اطمینان می‌دهد که هزینه کل منابع تجدیدپذیر از میزان بودجه در دسترس فراتر نمی‌رود. رابطه (۷) میزان تأخیر از موعد تحویل تعیین شده برای پروژه محاسبه می‌کند. رابطه (۸) و (۹) نحوه محاسبه شناوری هر فعالیت را نمایش می‌دهد. روابط (۱۰) و (۱۱) فعالیت‌های بحرانی را تعیین می‌نمایند. رابطه (۱۲) مرتبط با تعیین اندازه بافر است که میزان پیچیدگی شبکه را محاسبه می‌کند.

۲-۳- رویکرد حل فازی

مدل معرفی شده در بخش ۲-۲ یک مدل غیر خطی است. به دلیل دشواری حل مدل‌های غیر خطی نسبت به مدل‌های خطی، فرآیند خطی‌سازی انجام گرفته است. همچنین با توجه به اینکه پارامتر مدت زمان در مدل با استفاده از اعداد فازی مثلثی ارائه شده است مدل معادل قطعی بر اساس روش معرفی شده توسط خیمنز و همکاران [10] با یک مدل آلفا پارامتری جایگزین می‌شود. مدل برنامه‌ریزی ریاضی فازی رابطه (۱۳) را در نظر بگیرید که همه‌ی پارامترها در این مدل به صورت فازی بیان شده‌اند. مدل با استفاده از روش خیمنز و همکاران [10] به یک مدل قطعی مطابق با رابطه (۱۹) تبدیل شده که در آن سطح آلفا توسط تصمیم‌گیرنده مشخص می‌شود.

$$\min \tilde{c}^t x$$

Subject to:

$$\tilde{a}_i \geq \tilde{b}_i, \quad i = 1, 2, \dots, l \quad (13)$$

$$\tilde{a}_i = \tilde{b}_i, \quad i = l + 1, \dots, m$$

$alfa_i$ مرتبط با شناوری فعالیت i

epp_i میزان شناوری فعالیت i

$beta_j$ متغیر باینری که اگر فعالیت شناوری صفر داشته مقدار ۱، در غیر این صورت مقدار ۰ دارد.

x_{jmt} متغیر باینری که اگر فعالیت j در زمان t و در حالت m پایان باید مقدار ۱، در غیر این صورت مقدار ۰ دارد.

kk میزان پیچیدگی شبکه

$$z_1 = \min Tc + \sum_{i=1}^N beta_i \cdot KK \cdot \sigma_i \quad (1)$$

$$z_2 = \max \sum_{j=1}^N \sum_{m \in M_j} \sum_{t \in Ef_j}^{Lf_j} x_{jmt} \cdot Q_{jm} \quad (2)$$

S.t.

$$\sum_{m \in M_j} \sum_{t \in Ef_j}^{Lf_j} x_{jmt} = 1, \quad \forall j \quad (3)$$

$$\sum_{m \in M_j} \sum_{t \in Ef_j}^{Lf_j} (t - d_{jm}) \cdot x_{jmt} \geq \sum_{m \in M_i} \sum_{t \in Ef_i}^{Lf_i} t \cdot x_{imt}, \quad \forall (i, j) \in p \quad (4)$$

$$\sum_{j=1}^N \sum_{m \in M_j} \sum_{q=\max\{t, Ef_j\}}^{\min\{t+d_{jm}-1, Lf_j\}} r_{jkm} \cdot x_{jmq} \leq BR_k, \quad \forall k \in K, t \in T \quad (5)$$

$$\sum_{k=1}^K cr_k \cdot BR_k \leq tb \quad (6)$$

$$Tc \geq \sum_{t \in Ef_N}^{Lf_N} \sum_{m \in M_j} (t \cdot x_{Nmt} - dd) \quad (7)$$

$$\sum_{m \in M_j} \sum_{t \in Ef_j}^{Lf_j} (t - d_{jm}) \cdot x_{jmt} - \sum_{m \in M_i} \sum_{t \in Ef_i}^{Lf_i} d_{im} \cdot x_{imt} = alfa_i, \quad \forall (i, j) \in p \quad (8)$$

$$alfa_i - \sum_{m \in M_i} \sum_{t \in Ef_i}^{Lf_i} (t - d_{im}) \cdot x_{imt} = epp_i, \quad (9)$$

$$\sum_{m \in M_j} \sum_{t=Ef_j}^{Lf_j} \left(t - ((1-\alpha).E_2^{d_{jm}} + \alpha.E_1^{d_{jm}}) \right) . x_{jmt} \\ \geq \sum_{m \in M_i} \sum_{t=Ef_i}^{Lf_i} t . x_{imt}$$

$$\forall (i, j) \in p \quad (20)$$

$$\sum_{j=1}^N \sum_{m \in M_j} \sum_{q=\max\{t, Ef_j\}}^{\min\{t+(\frac{d_{jm}^p+2.d_{jm}^m+d_{jm}^o}{4})-1, Lf_j\}} r_{jkm} . x_{jmq} \leq BR_k, \quad (21)$$

$$\forall k \in K, t \in T$$

$$\sum_{m \in M_j} \sum_{t=Ef_j}^{Lf_j} \left(t - \left(\left(1 - \frac{\alpha}{2} \right) . E_2^{d_{jm}} + \left(\frac{\alpha}{2} \right) . E_1^{d_{jm}} \right) \right) . x_{jmt} \\ - \sum_{m \in M_i} \sum_{t=Ef_i}^{Lf_i} \left(\left(1 - \frac{\alpha}{2} \right) . E_2^{d_{im}} + \left(\frac{\alpha}{2} \right) . E_1^{d_{im}} \right) . x_{imt} \geq alfa_i,$$

$$\forall (i, j) \in p \quad (22)$$

$$\sum_{m \in M_j} \sum_{t=Ef_j}^{Lf_j} \left(t - \left(\left(\frac{\alpha}{2} \right) . E_2^{d_{jm}} + \left(1 - \frac{\alpha}{2} \right) . E_1^{d_{jm}} \right) \right) . x_{jmt} \\ - \sum_{m \in M_i} \sum_{t=Ef_i}^{Lf_i} \left(\left(\frac{\alpha}{2} \right) . E_2^{d_{im}} + \left(1 - \frac{\alpha}{2} \right) . E_1^{d_{im}} \right) . x_{imt} \leq alfa_i$$

$$\forall (i, j) \in p \quad (23)$$

$$alfa_i - \sum_{m \in M_i} \sum_{t=Ef_i}^{Lf_i} \left(t - \left(\left(1 - \frac{\alpha}{2} \right) . E_2^{d_{im}} + \left(\frac{\alpha}{2} \right) . E_1^{d_{im}} \right) \right) . x_{imt} \geq eppi,$$

$$\forall i \quad (24)$$

$$alfa_i - \sum_{m \in M_i} \sum_{t=Ef_i}^{Lf_i} \left(t - \left(\left(\frac{\alpha}{2} \right) . E_2^{d_{im}} + \left(1 - \frac{\alpha}{2} \right) . E_1^{d_{im}} \right) \right) . x_{imt} \leq eppi, \quad (25)$$

$$\forall i$$

$$x \geq 0$$

بنابراین مدل با استفاده از روش خیمنز و همکاران [10] به یک مدل قطعی مطابق با رابطه (۱۹) تبدیل شده که در آن سطح آلفا توسط تصمیم‌گیرنده مشخص می‌شود. همچنین نحوه محاسبه ارزش انتظاری و بازه انتظاری به ترتیب در روابط (۱۴) تا (۱۸) نمایش داده شده است.

$$EV(\tilde{c}) = \frac{c^p + 2.c^m + c^o}{4} \quad (14)$$

$$E_1^a = \frac{1}{2}(c^p + c^m) \quad (15)$$

$$E_2^a = \frac{1}{2}(c^m + c^o) \quad (16)$$

$$E_1^b = \frac{1}{2}(c^p + c^m) \quad (17)$$

$$E_2^b = \frac{1}{2}(c^m + c^o) \quad (18)$$

$$\min [EV(\tilde{c}^t)]x$$

Subject to:

$$\left((1-\alpha).E_2^{a_i} + \alpha.E_1^{a_i} \right) x \geq \alpha.E_2^{b_i} + (1-\alpha).E_1^{b_i} \quad (19)$$

$$\left(\left(1 - \frac{\alpha}{2} \right) . E_2^{a_i} + \left(\frac{\alpha}{2} \right) . E_1^{a_i} \right) x \geq \left(\frac{\alpha}{2} \right) . E_2^{b_i} + \left(1 - \frac{\alpha}{2} \right) . E_1^{b_i}$$

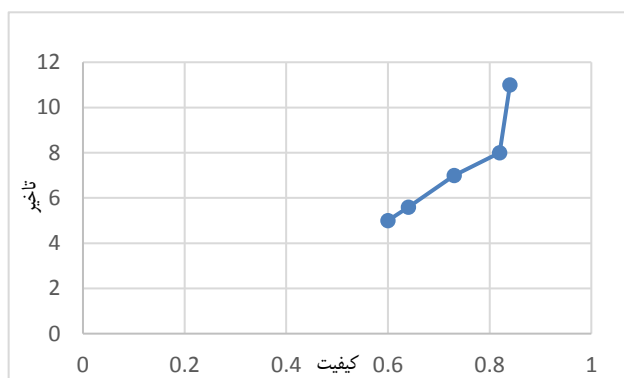
$$\left(\left(\frac{\alpha}{2} \right) . E_2^{a_i} + \left(1 - \frac{\alpha}{2} \right) . E_1^{a_i} \right) x \leq \left(1 - \frac{\alpha}{2} \right) . E_2^{b_i} + \left(\frac{\alpha}{2} \right) . E_1^{b_i}$$

$$x \geq 0$$

پارامتر مدت زمان تنها در روابط (۴)، (۵)، (۸) و (۹) وجود دارد که معادل قطعی آنها به صورت رابطه (۲۰) تا (۲۵) نمایش داده شده است.

۶	۹	۶	۳	۲
۷	۱۳	۹	۷	۴
۸	۱۴	۱۰	۱۰	۸
۹	۰	۰	۰	۰

در واقع هر چه میزان تأخیر پروژه افزایش یابد فعالیت‌ها در حالت‌های اجرایی با مدت زمان بیشتر انجام شده در نتیجه بر میزان کیفیت پروژه افزوده می‌شود. بالعکس زمانی که میزان تأخیر پروژه کاهش یابد فعالیت‌ها در حالت اجرایی با مدت زمان کمتر قرار گرفته که کاهش کیفیت فعالیت‌ها را به همراه دارند. علاوه بر این نتایج حاصل از حل مدل در حالت قطعی و زمانی که پارامترهای زمان دارای عدم قطعیت فازی هستند در جدول (۳) گزارش شده است. بر اساس نتایج بدست آمده میزان تابع هدف کیفیت پروژه در هر دو حالت قطعی و عدم قطعیت فازی یکسان است. اما تابع هدف مجموع تأخیرها مقدار کمتری را نسبت به حالت قطعی به خود اختصاص داده است. بنابراین می‌توان نتیجه گرفت که بررسی مسأله در شرایط عدم قطعیت جواب‌های بهتری نشان داده است.



شکل ۱: روند تغییر تابع هدف اول (تأخیر) نسبت به تابع هدف دوم (کیفیت) در نقاط پارتویی

جدول ۳: نتایج حل مدل

تابع هدف	حالت عدم قطعیت	حالت قطعی
تأخیر	۰.۸۴	۰.۸۴
کیفیت	۴	۱۱.۸۷۵

۴- نتیجه‌گیری

روش زنجیره‌ی بحرانی در مقایسه با روش‌های سنتی زمانبندی پروژه، مسأله‌ی محدودیت منابع را مورد توجه قرار داده و با معرفی و تعیین اندازه‌ی زمان‌هایی تحت عنوان زمان اِی‌م‌نِی یا بافر به مقابله با تأخیرهای احتمالی در انجام فعالیت‌ها می‌پردازد. به همین منظور در این پژوهش، بر اساس اصول روش زنجیره‌ی بحرانی و با توجه به محدودیت بودجه یک مدل ریاضی دو هدفه به منظور زمانبندی پروژه‌ها ارائه شد که در ضمن ایجاد یک برنامه بهینه

یکی از روش‌های حل مدل‌های ریاضی چند هدفه روش برنامه‌ریزی آرمانی است که نیاز به داشتن اطلاعات اولیه از تصمیم‌گیرنده دارد. این روش ابتدا در سال ۱۹۶۰ توسط چارلز^۱ و کوپر^۲ ایجاد شد. در این رویکرد برای هر یک از اهداف موجود سطح تمایل مشخص و دقیق از سوی تصمیم‌گیرنده ارائه شده و تلاش می‌شود فرآیند حل با هدف کاهش میزان انحراف از این سطوح انجام شود [11]. در نهایت، مدل دو هدفه این پژوهش نیز با استفاده رویکرد برنامه‌ریزی آرمانی با در نظر گرفتن سطح آرمان هر یک از اهداف حل می‌شود.

۳- یک مثال کاربردی

در این بخش با استفاده از یک مثال، مدل مورد بررسی قرار می‌گیرد. یک شبکه پروژه متشکل از ۹ فعالیت، که فعالیت‌های ابتدا و انتهای شبکه موهومی هستند؛ در نظر گرفته شده است. هر فعالیت دارای دو حالت اجرایی است. همچنین دو نوع منبع تجدیدپذیر برای اجرای فعالیت‌ها نیاز است که با توجه به مدت زمان و حالت اجرایی انجام شده از آنها استفاده می‌کنند. میزان بودجه تعیین شده برای پروژه ۳۰۰ واحد اعلام شده است. همچنین زمان ۳۵ روز از سوی کارفرما برای موعد تحویل پروژه در نظر گرفته شده است. مدت زمان انجام فعالیت‌ها نیز در هر یک از حالت‌های اجرایی با استفاده از اعداد فازی مثلی بیان شده است. در جدول (۱) و (۲) برخی از اطلاعات مورد نیاز نمایش داده شده است.

جدول ۱: پارامترهای مدت زمان و روابط پیشینازی

شماره فعالیت	پیشینازها	زمان حالت ۱	زمان حالت ۲
۱	-	۰	۰
۲	۱	۴	۶
۳	۲	۵	۸
۴	۲	۳	۷
۵	۲	۲	۴
۶	۳ و ۴	۴	۸
۷	۵	۸	۱۲
۸	۶ و ۷	۹	۱۵
۹	۸	۰	۰

نتایج موازنه دو تابع هدف در پنج نقطه پارتویی به منظور بررسی صحت مدل در سطح آلفا مینا برابر ۰.۷، در شکل (۱) ارائه شده است. همانطور که مشخص است با افزایش میزان تابع هدف اول (مجموع تأخیرها)، میزان تابع هدف دوم نیز (کیفیت) افزایش می‌یابد.

جدول ۲: پارامترهای منابع مورد نیاز فعالیت‌ها

شماره فعالیت	منبع نوع اول		منبع نوع دوم	
	حالت ۱	حالت ۲	حالت ۱	حالت ۲
۱	۰	۰	۰	۰
۲	۸	۵	۵	۳
۳	۷	۴	۶	۴
۴	۸	۴	۴	۲
۵	۵	۳	۸	۵

² Cooper

¹ Charnes

برای انجام فعالیت‌ها دو هدف کاهش میزان تأخیر پروژه و افزایش کیفیت را نیز در نظر دارد. در مدل پیشنهادی با استفاده از دو محدودیت میزان شناسایی هر یک از فعالیت‌ها محاسبه شده و با توجه به آن فعالیت‌های بحرانی به منظور اندازه‌گیری بافر مشخص می‌شوند. همچنین به دلیل وجود عدم قطعیت در مسائل واقعی ابهام موجود در پارامتر مدت زمان فعالیت‌ها با استفاده از اعداد فازی مثلثی بیان شده که برای قطعی‌سازی مدل از روش ارائه شده توسط خیمنز و همکاران استفاده شد. در نهایت نیز مدل با استفاده از روش برنامه‌ریزی آرمانی با توجه به چند هدفه بودن آن حل شد. با توجه به تحلیل حساسیت انجام شده در چند نقطه پارتویی، صحت مدل اثبات شده و نتایج حاصل از حل در شرایط عدم قطعیت نسبت به حالت قطعی بهبود یافته است. در نظر گرفتن عدم قطعیت در فرآیند زمانبندی کارکرد مدل ریاضی را برای بسیاری از پروژه‌ها به خصوص پروژه‌هایی که برای اولین بار انجام می‌شوند و داده‌های تاریخی برای آنها وجود ندارد؛ فراهم می‌سازد.

مراجع

- [1] Ahlemann, F., El Arbi, F., Kaiser, M. G., & Heck, A., "A process framework for theoretically grounded prescriptive research in the project management field", International Journal of Project Management, 31(1), 43-56, 2013.
- [2] Newbold, Robert C., "Project management in the fast lane: applying the theory of constraints", CRC Press, 1998.
- [3] Tukul, O. I., Rom, W. O., & Eksioğlu, S. D., "An investigation of buffer sizing techniques in critical chain scheduling", European Journal of Operational Research, 172(2), 401-416, 2006.
- [4] Balta, S., Birgonul, M. T., & Dikmen, I., "Buffer sizing model incorporating fuzzy risk assessment: case study on concrete gravity dam and hydroelectric power plant projects", ASCE-ASME Journal of Risk and Uncertainty in Engineering Systems, Part A: Civil Engineering, 4(1), 1-10, 2018.
- [5] Hu, X., Cui, N., Demeulemeester, E., & Bie, L., "Incorporation of activity sensitivity measures into buffer management to manage project schedule risk", European Journal of Operational Research, 249(2), 717-727, 2016.
- [6] Zohrehvandi, S., & Khalilzadeh, M., "APRT-FMEA buffer sizing method in scheduling of a wind farm construction project", Engineering, Construction and Architectural Management, 1-23, 2019.
- [7] Herroelen, W., De Reyck, B., & Demeulemeester, E., "Resource-constrained project scheduling: a survey of recent developments", Computers & Operations Research, 25(4), 279-302, 1998.
- [8] Long, L. D., & Ohsato, A., "Fuzzy critical chain method for project scheduling under resource constraints and uncertainty", International Journal of Project Management, 26(6), 688-698, 2008.
- [9] Roghanian, E., Alipour, M., & Rezaei, M., "An improved fuzzy critical chain approach in order to face uncertainty in project scheduling", International Journal of Construction Management, 18(1), 1-13, 2018.
- [10] Jiménez, M., Arenas, M., Bilbao, A., & Rodrı, M. V., "Linear programming with fuzzy parameters: an interactive method resolution", European journal of operational research, 177(3), 1599-1609, 2007.
- [11] Hocine, A., Guellil, M. S., Dogan, E., Ghouali, S., & Kouaissah, N., "A fuzzy goal programming with interval

مقالات پوستر کنگره



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

ارائه یک روش تلفیقی مبتنی بر سیستم‌های چندعاملی در تصاویر ماموگرافی جهت تشخیص زود هنگام بیماری سرطان پستان

حمیدرضا ناجی^۱، بهنام رضائی بزنجان^۲

^۱دانشیار، دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته کرمان، دانشکده مهندسی برق و کامپیوتر
^۲دانشجوی دکتری مهندسی کامپیوتر، گروه کامپیوتر و فناوری اطلاعات، دانشگاه آزاد اسلامی واحد کرمان، کرمان،

چکیده: سرطان سینه یکی از شایع‌ترین علل مرگ و میر در میان زنان در سراسر جهان است. از این رو، تشخیص زود هنگام به نجات جان زنان کمک می‌کند. ماموگرافی یک تست غربالگری اصلی برای سرطان سینه است. این مجموعه شامل مصنوعات بسیاری است که تأثیر منفی در تشخیص سرطان سینه دارند. بنابراین، حذف نویز و بهبود کیفیت تصویر یک فرآیند مورد نیاز در سیستم تشخیص به کمک کامپیوتر است. دقت و کارایی بینایی ماشین با ارائه ناحیه دقیق (ROI) افزایش می‌یابد. استخراج ROI یک کار چالش برانگیز در پیش‌پردازش تصویر است زیرا حضور ماهیچه‌های سینه بر تشخیص ناهنجاری تأثیر می‌گذارد. در اینجا، نشان داده می‌شود که با استفاده از سیستم‌های چند عاملی، الگوریتم خوشه بندی k-means و تکنیک‌های فیلتر وینر و تعادل هیستوگرام تطبیقی محدود کنتراست به طور موثری به افزایش کیفیت تصویر و سگمنت بندی ناحیه های مختلف تصویر کمک می‌کنند، در نتیجه پس زمینه ناخواسته و عضله سینه را نیز با استفاده از آستانه گذاری و خوشه بندی و تکنیک رشد ناحیه اصلاح شده حذف می‌کند. علاوه بر این، از ترکیب الگوریتم پیشنهادی با سیستم‌های چند عاملی جهت تشخیص تومور در تصاویر استفاده می‌کنیم. این پژوهش بر روی پایگاه داده mini - MIAS آزمایش شد؛ نتایج به دست آمده با کامل بودن و صحت حذف عضله سینه‌ای مقایسه شد و در مجموع، این نتایج نشان می‌دهد که روش پیشنهادی برای بهبود کیفیت و تشخیص سرطان در تصویر ماموگرافی برای سیستم‌های بینایی ماشین مناسب است.

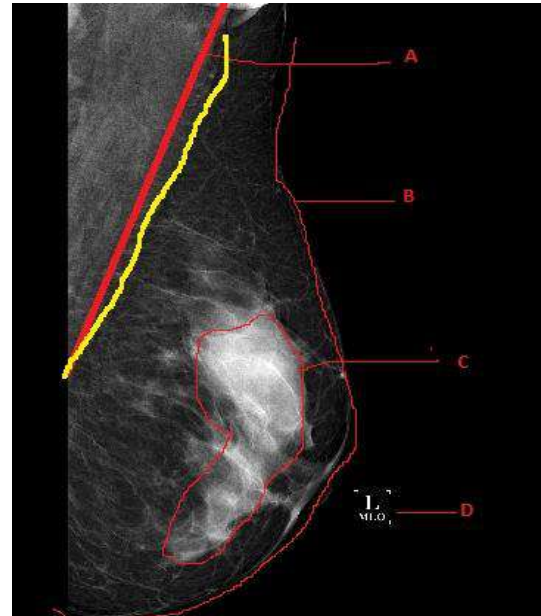
کلمات کلیدی: سرطان سینه، ماموگرافی، پیش‌پردازش، رشد منطقه، فیلتر وینر، سیستم‌های چند عاملی

۱- مقدمه

داخلی پستان را بهتر نشان می‌دهد. معمولاً تصاویر ماموگرافی شامل نویز هستند و تشخیص و درک سرطان در مراحل اولیه بسیار دشوار است [۴]. بنابراین، استاندارد سازی کیفیت تصویر و استخراج ROI برای محدود کردن جستجو برای مشکلات احتمالی ضروری است. پیش‌پردازش ماموگرافی به شناسایی نواحی غیر طبیعی کمک می‌کند که از نظر فیزیکی یا تصویری نمی‌توانند تجربه شوند، اما می‌توانند از طریق سیستم بینایی ماشین تشخیص داده شوند. سیستم بینایی ماشین به پزشکان کمک می‌کند و رادیولوژیست‌ها به راحتی و به سرعت این ناهنجاری را تشخیص می‌دهند. در اینجا، روش پیشنهادی به استاندارد سازی کیفیت تصویر و استخراج ROI نهایی کمک

سرطان سینه دومین علت مرگ زنان در سراسر جهان است. با توجه به جامعه سرطان آمریکا، در حدود ۱ تا ۸ زن در طول زندگی خود مبتلا به سرطان پستان خواهند بود و تنها ۵ تا ۱۰ درصد از سرطان‌های سینه در زنانی رخ می‌دهد که ارتباط ژنتیکی مشخصی دارند [۲]. از این رو، تشخیص زود هنگام به داشتن کیفیت زندگی بهتر، درمان اقتصادی و آرامش ذهنی بیمار و خانواده کمک خواهد کرد. ماموگرافی با دوز پایین تصویربرداری اشعه ایکس، اساسی‌ترین تست غربالگری برای سرطان سینه است و همچنین اطلاعات

می‌کند. با استخراج بخش سینه، حذف بخش ماهیچه و افزایش کیفیت ماموگرام، الگوریتم پیشنهادی به رادیولوژیست کمک می‌کند تا بیماری را دقیق‌تر تشخیص دهد و انواع نوپز مشاهده شده در تصویر ماموگرافی که در شکل ۱ نشان داده شده‌اند و ناخواسته و اضافی هستند را از تصویر ماموگرافی حذف می‌کند.



شکل ۱ ناحیه بندی سینه

چندین روش برای پیش‌پردازش تصاویر ماموگرافی از سال ۱۹۸۰ به دلیل تاثیر آن در تشخیص سرطان گزارش شده است [۳]. تکنیک‌هایی مانند فیلتر میانه تطبیقی، فیلتر میانگین، فیلتر میانگین تطبیقی، تعادل هیستوگرام، بهبود کنتراست محلی تغییر یافته هیستوگرام، ناحیه سینه و استخراج عضله سینه، تکنیک CLAE و مورفولوژی‌هایی که قبلاً مورد بحث قرار گرفته‌اند. بخش ۲ تا ۷ مجموعه داده مورد استفاده و روش‌های پیشنهادی را نشان می‌دهد. بخش ۸ به بررسی آزمایش و نتایج روش پیشنهادی می‌پردازد.

۲- پردازش تصاویر

پردازش تصاویر دیجیتال به مجموعه‌ای از تکنیک‌ها برای دستکاری و تغییرات تصاویر دیجیتال توسط کامپیوتر است. داده‌های خام دریافت شده از حسگرهای تصویربرداری دارای نقص و کمبود می‌باشد. برای غلبه بر نقیصه و کمبود به منظور رسیدن به یک تصویر خوب و پایدار به انجام چند مرحله پردازش تصویر نیاز دارد.

پردازش تصویر دیجیتال تحت سه مرحله زیر است

- پیش پردازش
- نمایش و بالا بردن و بهبود آن
- استخراج اطلاعات

۳- هدف از پردازش تصویر

- تشدید تصویر و بهبود
- بازیابی تصویر
- اندازه گیری الگو
- تشخیص تصویر

۴- رویکردهای مبتنی بر عامل

رویکردهای مبتنی بر عامل، رویکردهای الهام گرفته از زیست هستند که یک راه‌حل و چندین مزیت مانند استقلال، طرفدار فعالیت، توزیع، خود سازماندهی و انطباق را فراهم می‌کنند. عامل‌ها برای حل مسأله جمعی هم‌کاری می‌کنند. یک عامل مجزا نمی‌تواند کل مشکل را حل کند زیرا ادراک و قابلیت‌های زیادی ندارد.

عوامل مشارکتی و جایگاه‌مند را به عنوان یک الگوی برنامه‌نویسی برای اجرای مکانیزم‌های قبلاً توضیح داده شده معرفی می‌کنیم. عوامل ثبت شده از AI واکنشی مستقل قرض می‌گیرند، هر عامل دانش لازم برای رسیدن به هدف خود را به طور خودکار کسب می‌کند از شناخت جایگاه‌مند [۳]، آن‌ها اصل محلی‌سازی بافتار را استفاده می‌کنند، هر عامل در یک بافتار تکاملی محلی با محدودیت‌های همسایگی خاص قرار می‌گیرد تا ثبات جهانی فرآیند تفسیر حفظ شود. از نظریه چند عاملی، آن‌ها از اصل همکاری استفاده می‌کنند و بهره می‌گیرند، هر عامل با آشنایان خود تعامل می‌کند تا به هدف خود برسد. سیستم متشکل از عواملی است که توسط یک مدیر سیستم هدایت می‌شود، که نقش آن مدیریت ایجاد، تخریب، فعال‌سازی و غیر فعال‌سازی آن‌ها است. هر نماینده به نوبه خود دارای چندین رفتار است که انتخاب آن‌ها توسط یک مدیر رفتار هدایت می‌شود. عوامل در گروه‌های [۴] توسط یک مدیر گروه که هماهنگی را تضمین می‌کند، سازماندهی می‌شوند (شکل ۲ را ببینید).

۴-۱ تعریف عامل

سه نوع از عوامل در سیستم وجود دارند

عوامل کنترل محلی

عوامل کنترل گلوبال

عوامل اختصاصی بافت

نقش عامل کنترل گلوبال، انجام وظایف خاص اختصاص داده شده به کل تصویر و ایجاد عوامل کنترل محلی توزیع شده در حجم تصویر است. توزیع از یک شبکه ۳ بعدی منظم پیروی می‌کند که در آن گره‌ها در فاصله DI قرار می‌گیرند. نقش عوامل کنترل محلی ایجاد عوامل اختصاصی بافت، برآورد پارامترهای مدل و مقابله با مدل‌های بافت برای تصمیمات برچسب گذاری نهایی است. دو قلمرو، در مرکز یک‌گره داده شده در شبکه پارتیشن بندی، به هر عامل کنترل محلی اختصاص داده می‌شوند: قلمرو ادراک، مکعبی با طول برابر است که در آن فعالیت تخمین مدل رخ خواهد داد؛ و یک محدوده برچسب گذاری، که در آن فعالیت طبقه‌بندی رخ خواهد داد. قلمروهای درک عوامل کنترل محلی همسایه ممکن است برای تقویت انسجام بین مدل‌های مجاور همپوشانی داشته باشند. نقش عوامل اختصاصی بافت، اجرای وظایف توزیع شده توسط نوع بافت، یعنی درونیابی مدل‌های بافتی از همسایگی و برچسب زدن با استفاده از یک فرآیند رشد ناحیه است. این نوع طراحی اجازه می‌دهد تا وظایف پردازشی با توجه به سه زمینه مختلف توزیع شوند: کل حجم تصویر برای عامل کنترل جهانی؛

یک ناحیه مشخص در حجم تصویر برای عوامل کنترل محلی؛

یک ناحیه مشخص و نوع بافت برای عوامل اختصاصی بافتی.

پیشرفت در طول این زمینه‌های دقیق‌تر به تدریج کنترل و انطباق فرآیند تفسیر را بهبود می‌بخشد. [۴]

۴-۲ عامل رفتار

یک مرور کلی از نحوه توزیع و توالی یابی سیستم را ارائه می‌دهد. همانطور که قبلاً ذکر شد، فرآیند تفسیر شامل دو فاز است که هر کدام از آن‌ها از یک تخمین مدل و یک مرحله طبقه‌بندی تشکیل شده‌اند. وظیفه کنترل به عنوان یک رفتار عامل کنترل به تفصیل در زیر و در بخش‌های بعدی انجام می‌شود:

۴-۳ نماینده هماهنگ کننده و مدیریت اطلاعات

نمایندگان یک منطقه اطلاعاتی مشترک را به اشتراک می‌گذارند، که براساس نوع بافت و روابط فضایی سازماندهی شده است، که اطلاعات آماری محلی و جهانی را ذخیره می‌کند. نقشه‌های اطلاعات کیفی برای جمع‌آوری، بازیابی و اضافه کردن آسان اطلاعات تکمیلی، مانند تعداد و همسایه‌ها با یک برچسب مشخص معرفی می‌شوند. عوامل به طور ضمنی به اشتراک گذاری اطلاعات و به طور صریح از طریق تحریک رفتار و فرآیند غیر فعال‌سازی تعامل می‌کنند. هنگامی که یک عامل اطلاعات قوی تولید می‌کند، آن را در منطقه مشترک به اشتراک گذاری ذخیره می‌کند، فعالیت خود را سازگار می‌کند و آشنایان خود را تحریک می‌کند، یعنی لیست عوامل اختصاص داده شده به مناطق همسایه. سپس، آشنایانی که به بافت مشابهی اختصاص دارند، در طول رفتار بررسی مدل خود در هر دو فاز اولیه و نهایی و در طول رفتارهای رشد ریشه / منطقه خود تعامل می‌کنند. به طور مشابه، هر عامل دارای فهرستی از عوامل اختصاص یافته به یک بخش یکسان از تصویر است (کنترل محلی یا عوامل اختصاصی بافت). [۳]

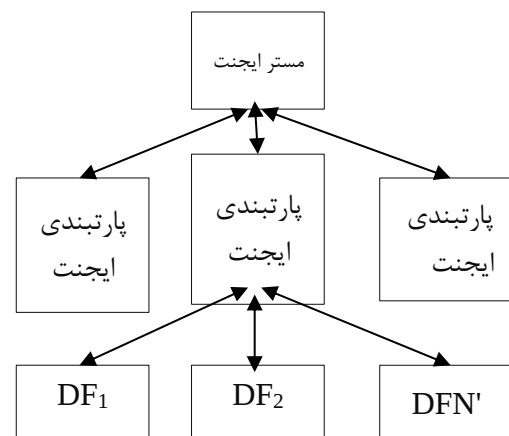
۵- سازماندهی سیستم چند عاملی

- گروه مستر ها که سطح پردازشی بالایی دارند، این گروه شامل یک عامل است
- گروه‌های پارتیشن بندی: سطح متوسط پردازشی دارند این گروه برابر است با K عامل.
- گروه دنبال کننده‌های ردیاب پردازش سطح پایین در اینجا عامل‌ها به صورت دینامیکی ایجاد و حذف می‌شوند.

۵-۱ عامل ارشد

- مامور ارشد ایجنت ها را تحت کنترل دارد
- و قوانین زیر را اجرا می کند
- دریافت تصویر دیجیتال از یک دستگاه ورودی
- پیچیدگی تصویر را به منظور ایجاد تعداد افراد
- ارزیابی می کند
- عوامل را ایجاد و راه اندازی می کند
- دریافت جداول داده از عوامل پارتیشن بندی
- بازسازی تصویر بخش بندی شده

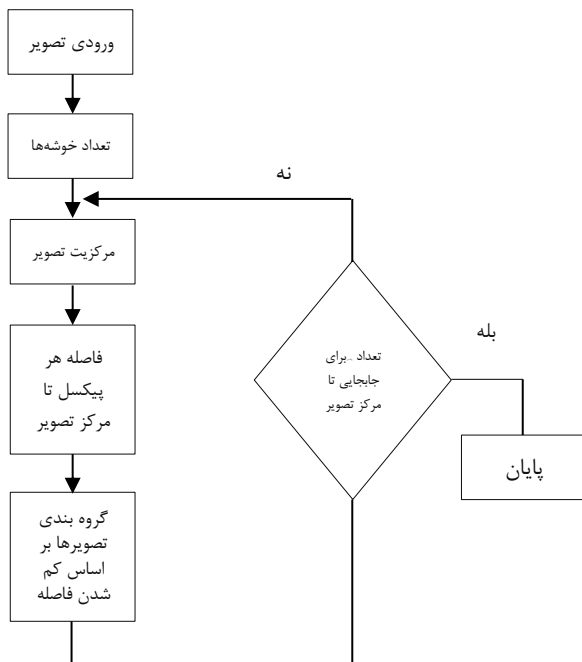
یک سیستم چند عاملی به این دلیل مورد استفاده قرار گرفت که برای آن لازم بود. جمع‌آوری مستمر تصاویر در زمان واقعی. نماینده‌ها مقدار اطلاعات در تصاویر را با اعمال فیلترهایی مانند فیلتر دو طرفه یا فیلتر گابور که لایه فیلتر را تشکیل می‌دهند، کاهش می‌دهند.



شکل ۲ نمودار سلسله ای عامل ها

۶- الگوریتم k-means

K-means نمونه‌ای از الگوریتم‌های خوشه‌بندی انحصاری است و از ستون‌های اصلی در این مقاله است [۵]. الگوریتم میانگین K خوشه‌بندی را با تعیین k نقطه مرکزی اولیه، چه به صورت تصادفی و چه با استفاده از برخی داده‌های اکتشافی آغاز می‌کند. سپس هر پیکسل تصویر را زیر نقطه مرکزی نزدیک‌ترین به آن گروه‌بندی می‌کند. سپس، نقاط مرکزی جدید را با میانگین گیری از پیکسل‌های گروه‌بندی شده زیر هر نقطه مرکزی محاسبه می‌کند. دو گام الگوریتمی قبلی به طور متناوب تکرار می‌شوند تا هم‌گرایی (نقاط مرکزی دیگر با میانگین گیری تغییر نمی‌کنند).



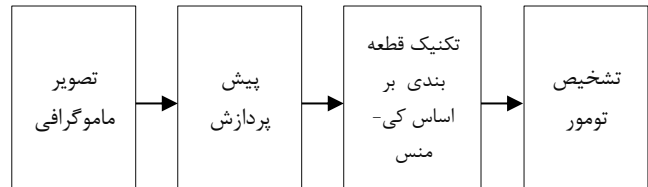
شکل ۳ فلوچارت کاربردهای

۶-۱ k-means مبتنی بر قطعه‌بندی

K به معنای خوشه‌بندی است که به طور گسترده در میان فرمول‌های خوشه‌بندی مورد استفاده و مطالعه قرار می‌گیرد که براساس به حداقل رساندن یک تابع هدف رسمی می‌باشد. اصلاحاتی در روش خوشه‌بندی K-means که آن را سریع‌تر و کارآمدتر می‌کند، پیشنهاد شده است. که در آن روش‌های مبتنی بر پیکسل مبتنی بر خوشه‌بندی K-means ساده هستند و پیچیدگی محاسباتی در مقایسه با دیگر روش‌های مبتنی بر ناحیه یا مبتنی بر لبه نسبتاً پایین است، کاربرد عملی‌تر است [۵]

۷- مجموعه داده مورد استفاده

در این مقاله ۱۰۰ تصویر از پایگاه داده جامعه تجزیه و تحلیل تصاویر گرافیکی (mini - MIAS)، انتخاب شد [۱۳]. آرشیو های گرفته شده در برنامه ملی غربالگری پستان بریتانیا (NBSP)، تمام تصاویر در اندازه ۱۰۲۴ * ۱۰۲۴ هستند. پایگاه داده شامل ۳۲۲ ماموگرام دیجیتالی شده است (که در میان آن ها ۲۰۲ تصویر طبیعی و ۱۲۰ تصویر غیر طبیعی وجود دارد). این مرکز همچنین شامل علائم رادیولوژیست ها در مکان های غیر طبیعی در صورت وجود است. تصاویر به صورت آنلاین در آرشیو پردازش تصویر آزمایشی اروپایی (PEIPA) در دانشگاه اسکس در دسترس هستند.



شکل ۴

۸-روش پیشنهادی

یک طرح CAD سرطان سینه مناطق مشکوکی را که ممکن است حاوی توده هایی از پارانشیم پس زمینه ویژگی بافت یک اندام باشند، همانطور که از بافت های پیوندی یا حمایتی مرتبط متمایز است، از هم جدا می کند [۶ - ۲]. به عبارت دیگر، چنین طرح هایی ماموگرام ها را به چند ناحیه غیر متقاطع تقسیم کرده و نواحی مورد نظر (ROI ها) و کاندیدهای جرم مشکوک را از تصویر اولترا سوند استخراج می کنند. با اینکه یک ناحیه مشکوک تاریک تر از محیط اطراف خود است، اما تراکم مشابهی دارد که یک شکل منظم با اندازه متغیر است. بنابراین، جداسازی تصویر برای حفظ حساسیت و دقت کل سیستم تشخیص جرم و طبقه بندی ضروری است. ما یک روش تطبیقی - K means را برای تشخیص میکرو کلسیفیکاسیون ها در ماموگرام های دیجیتال پیشنهاد کرده ایم. در کار حاضر، ما تلاش کرده ایم تا عملکرد روش - K means موجود را با تغییر مقادیر مختلف پارامترهای خاص مورد بحث در الگوریتم بهبود بخشیم [۱۱ - ۱۳]. الگوریتم K - means یک تکنیک تکرار شونده است که برای تقسیم یک تصویر به K خوشه استفاده می شود. در آمار و یادگیری ماشین، خوشه بندی K - means روشی از تحلیل خوشه ای است که هدف آن تقسیم n مشاهده به k خوشه است که در آن هر مشاهده به خوشه ای با نزدیک ترین میانگین تعلق دارد. الگوریتم پایه عبارت است از: مراکز خوشه K را انتخاب کنید، چه به صورت تصادفی و چه براساس برخی اقدامات اکتشافی،

هر پیکسل در تصویر را به خوشه ای اختصاص دهید که فاصله بین پیکسل و مرکز خوشه را به حداقل می رساند.

با توجه به مجموعه ای از مشاهدات (x₁, x₂, ..., x_n)، که در آن هر مشاهده یک بردار واقعی d بعدی است، خوشه بندی K - means با هدف تقسیم n مشاهده به مجموعه های k S(n) = {S₁, S₂, ..., S_k} که به منظور حداقل کردن مجموع مربعات درون خوشه ای (WCSS):

$$\argmin_S \sum_{i=1}^K \sum_{x_j \in S_i} \|x_j - \mu_i\|^2 \quad (1)$$

در مقاله پایه الگوریتم را به صورت داده شده است

نمودار هیستوگرام یک نمودار خلاصه است که تعداد نقاط داده در حال سقوط در محدوده های مختلف را نشان می دهد. این اثر تقریب سختی از توزیع فراوانی داده ها است. گروهی از داده ها، کلاس ها نامیده می شوند و در زمینه هیستوگرام، آن ها به عنوان سطل ها شناخته می شوند، زیرا می توان آن ها را به عنوان کانتینر هایی در نظر گرفت که داده ها را جمع آوری می کنند و با سرعتی برابر با فرکانس آن دسته از داده ها پر می کنند. شکل هیستوگرام گاهی اوقات به طور خاص به تعداد پیمانها حساس است. اگر سطل پیمانها خیلی عریض باشند، اطلاعات مهم ممکن است حذف شوند. با کاهش تعداد پیمانها و افزایش تعداد کلاس ها در الگوریتم K - means، دقت تشخیص افزایش می یابد. کمینه از نظر هیستوگرام های رنگی به فرآیند کاهش تعداد پیمانها با در نظر گرفتن رنگ هایی که بسیار به یکدیگر شبیه هستند و قرار دادن آن ها در یک پیمانها اشاره دارد. [۵]

ابزار جدید را به عنوان مرکز ثقل مشاهدات در خوشه محاسبه کنی هدف اصلی از سیستم بینایی ماشین در سرطان سینه کمک به رادیولوژیست و پزشکان برای تصمیم گیری سریع است. با ارائه ROI و طبقه بندی دقیق و ایجاد نواحی مختلف پیکسلی با استفاده از الگوریتم خوشه بندی K - means به شناسایی ناهنجاری کمک خواهد کرد و روش مالتی ایجنت سیستم ها به تشخیص تومور. روش پیشنهادی در پنج مرحله کار می کند که در شکل ۵ توضیح داده شده است.

به طور پیش فرض حداکثر تعداد پیمانها که می توان با استفاده از تابع هیستوگرام به دست آورد، ۲۵۶ است. به منظور صرفه جویی در زمان در هنگام تلاش برای مقایسه هیستوگرام های رنگ، می توان تعداد پیمانها را کم کرد.

$$\argmin_S \sum_{i=1}^K \sum_{x_j \in S_i} \|x_j - \mu_i\|^2 \quad (1)$$

با توجه به مجموعه اولیه k به معنی ۱m، ۱، ..., mk، که ممکن است به طور تصادفی یا توسط برخی روش های اکتشافی مشخص شود، الگوریتم به صورت متناوب بین دو مرحله پیش می رود [۱۴]. هر مشاهده را به خوشه با نزدیک ترین میانگین ارسال کنید.

(۲)

$$S_r = \{x_j: \|x_j - m_i^{(t)}\| \leq \|x_j - m_i^{(t)}\| \text{ for all } i = 1, \dots, k\}$$

جداسازی تصویر است [۱۴]، به دو روش براساس مقدار مکانی پیکسل انتخابی کار می کند و دیگری انتخاب نقطه بنیادی است. در روش پیشنهادی، نقطه بنیادی به طور خودکار با در نظر گرفتن جهت گیری ماموگرافی انتخاب می شود. این روش پیکسل های همسایه نقطه دانه را تعیین می کند و بررسی می کند که آیا پیکسل های بعدی باید به منطقه اضافه شوند یا خیر. این فرآیند تا استخراج ROI کامل تکرار می شود [۳].

```
import jade.Agent;

public class HelloWorldAgent extends Agent {
    protected void setup() {
        System.out.println("I'm an agent!");
    }
}
```

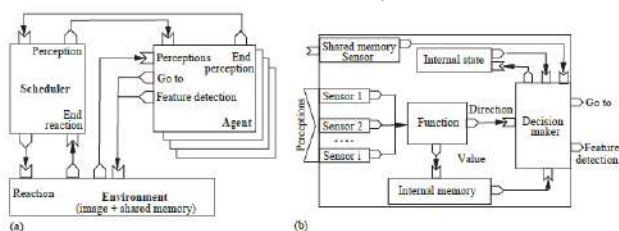
شکل ۷ دستور ساخت ایجنت

۸-۳ بهبود تصویر

مرحله سوم برای افزایش کیفیت تصویر با استفاده از فیلتر وینر و CLAH [۳] استفاده شد.

۸-۴ استفاده از ایجنت ها جهت تشخیص تومور

برنامه ریزی ایجنت ها براساس رنگ بندی و ایجاد و برنامه توسط ایجنت مستر، رنگ سفید برای بافت سالم، رنگ سیاه برای بافت داری تومور و بافت خاکستری برای بافت های مشکوک در نظر گرفته شد.

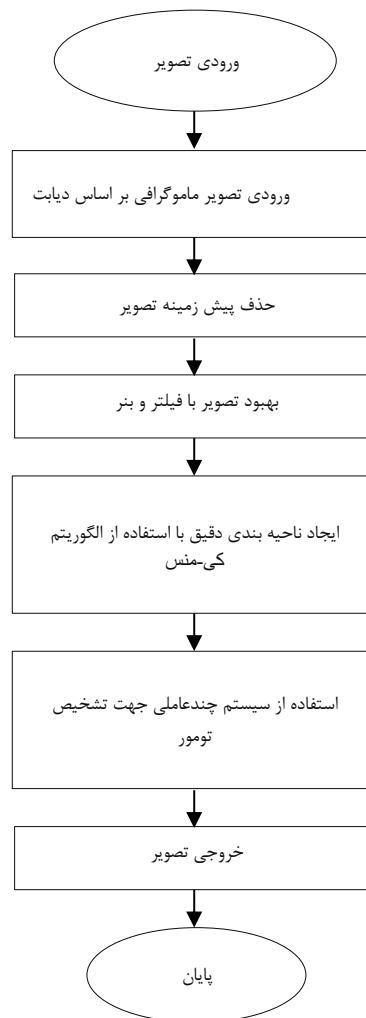


شکل ۸ [۴]

۹- شبیه سازی

شبیه سازی در محیط MATLAB و در یک سیستم با پردازنده ۷ هسته ای و پردازنده ۳،۴ گیگاهرتزی اینتل به همراه ۶ مگابایت کش و همین طور ۶ گیگابایت حافظه، شبیه سازی و اجرا گرفته شده است.

در این تحقیق از مجموعه داده MIAS به صورت تصویری و بر اساس ویژگی های آماری استفاده شده است. در این مجموعه داده، هم تصاویر با ویژگی های دارای سرطان سینه و هم عدم سرطان سینه و حالت های مشکوک وجود دارد که بر اساس داده بخش آماری آن، این عملیات تشخیص به صورت صحیح انجام خواهد گرفت.



شکل ۵

```
class Example // Declaration of the class 'Example'
{
    int age; // An 'Example' contains an integer 'age'
    void main(void); // An 'Example' has a method 'main()'
    // --> it is an agent

    void Example::main(void) // Method 'main()' definition
    {
        println("The age of ", this, " is ", ++age);
    }

    execute // Code to execute (initialization)
    {
        println("---- begin ----");
        for(int i=0; i<4; i++) new Example; // Instantiate four 'Example's
        println("---- end ----");
    }
}
```

The result of the agent-oriented program execution is as follows:

```
---- begin ----
---- end ----
The age of Example.2 is 1
The age of Example.3 is 1
The age of Example.4 is 1
The age of Example.1 is 1
The age of Example.1 is 2
The age of Example.1 is 2
The age of Example.3 is 2
The age of Example.2 is 2
The age of Example.4 is 2
```

Remark : agents are chosen in a random order.

شکل ۶ [۴]

۸-۱ حذف پس زمینه

در ابتدا تصویر با مقدار آستانه ۰،۱ دو شاخه شد سپس مولفه متصل شده به ترتیب نزولی برای استخراج بزرگترین حباب که پروفایل پستان است اما شامل عضله سینه است، سازماندهی شد.

۸-۲ مهار ماهیچه سینه ای

مرحله دوم برای کاهش بخش عضله سینه با استفاده از روش اصلاح شده رشد منطقه مورد استفاده قرار گرفت. ناحیه دانه دار در حال رشد یکی از روش های

نمونه کد های استفاده شده

```
# Subscribers to the positions of the other agents
agent_names=[]
agent_positions = {}
def agent_callback(msg, name):
    global agent_positions
    LOCK.acquire()
    agent_positions[name] = [msg.x, msg.y]
    LOCK.release()
# Handler of the service "AddAgent"
position_subscribers={}
def add_agent_handler(req):
    global agent_names
    global agent_positions
    global position_subscribers
    position_subscribers[req.name]=rp.Subscriber(
        name='/' + req.name + '/position',
        data_class=gms.Point,
        callback=agent_callback,
        callback_args=req.name,
        queue_size=1)
    LOCK.acquire()
    agent_names.append(req.name)
    agent_positions[req.name]=None
```

Code1

```
function [C,B,Cindex,Bindex] = main.m
clc;
clear all;
close all;
AgentNum=10; % Number of Agents
AgentSize=100; % Size of agents in plot
Dimension=3; % Select Dim
SizeOfEnvironmet=[15 15 15 ; -4 -4 -4]; % siz
if Dimension==2
    Dim='2';
else
    Dim='3';
end

sMat=ServerMat(AgentNum,Dimension,SizeOfEnvi
whitebg('black')
%% Make 1st Plot of agents using scatter plo

switch Dim
case '2'
    scatter(sMat(:,1),sMat(:,2),AgentSize,s
case '3'
    scatter3(sMat(:,1),sMat(:,2),sMat(:,3),.
otherwise
    error('myann:argchk' 'wrong number
```

Code 2

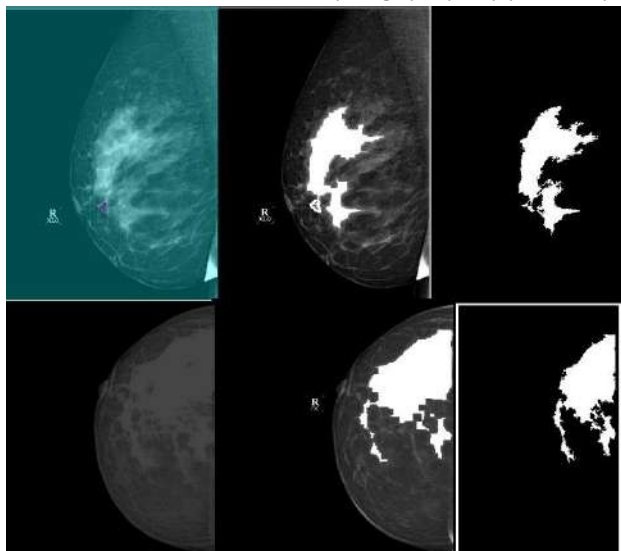
۱۰- نتایج

در این مقاله، تلاشی برای پیاده سازی الگوریتم خوشه بندی k - means برای جداسازی تصویر سینه برای تشخیص میکروکلسیفیکاسیون ها و ناحیه بندی قسمتهای مختلف سینه و همچنین یک سیستم تصمیم گیری مبتنی بر کامپیوتر برای تشخیص زودهنگام سرطان سینه به روش اصلاح شده، انجام شده است. ما یک سیستم تصمیم گیری با کمک سیستم های چند عاملی را برای تشخیص سرطان با سه رنگ سفید، سیاه و خاکستری در تصاویر ماموگرام توسعه داده ایم. این سیستم قادر به تشخیص میکروکلسیفیکاسیون ها با بازرسی بصری ماموگرام های دیجیتال است.

برای فیلتر وینر و CI is. عضله سینه کامل توسط تکنیک های بهبود یافته رشد منطقه کاهش یافت، روش پیشنهادی تست شده بر روی تصاویر پایگاه داده Mini - MIAS، ROI استخراج شده از تمام تصاویر به طور

دقیق و ثابت شده مناسب برای سیستم CAD تشخیص زودهنگام سرطان سینه است. در مجموع، این نتایج نشان می دهد که کمک موثر و مناسب برای تشخیص پزشکی. از این رو، روش پیشنهادی قطعاً می تواند برای تشخیص خودکار ناهنجاری های خوش خیم، بدخیم و میکروکلسیفیکاسیون ها در نظر گرفته شود.

نمونه های از ورودی و خروجی تصاویر



مراجع

- [1] Mambou, S.J., Maresova, P., Krejcar, O., Selamat, A. and Kuca, K., 2018. Breast cancer detection using infrared thermal imaging and a deep learning model. *Sensors*, 18(9), p.2799.
- [2] Benamrane, N. and Nassane, S., 2007, September. Medical image segmentation by a multi-agent system approach. In *German Conference on Multiagent System Technologies* (pp. 49-60). Springer, Berlin, Heidelberg.
- [3] Mohammed, R. and Said, A., 2012. An adaptative multi-agent system approach for image segmentation. *International Journal of Computer Applications*, 51(12).
- [4] Rodin, V., Benzinou, A., Guillaud, A., Ballet, P., Harrouet, F., Tisseau, J. and Le Bihan, J., 2004. An immune oriented multi-agent system for biological image processing. *Pattern Recognition*, 37(4), pp.631-645.
- [5] Vijay, J. and Subhashini, J., 2013, April. An efficient brain tumor detection methodology using K-means clustering algorithm. In *2013 International Conference on Communication and Signal Processing* (pp. 653-657). IEEE.

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بم

پیاده سازی یادگیری تقویتی عمیق برای کنترل هوشمند مبدل افزایشده

سبا عسکری نوغانی^۱، ناصر پریز^۲، محمد باقر نقیعی سیستمانی^۳

^۱ دانشجوی کارشناسی ارشد، گروه برق، دانشگاه فردوسی، مشهد، saba.askari1992@gmail.com

^۲ استاد، گروه برق، دانشگاه فردوسی، مشهد، n-pariz@um.ac.ir

^۳ دانشیار، گروه برق، دانشگاه فردوسی، مشهد، mb-naghibi@um.ac.ir

چکیده: طراحی کنترل گر برای مبدل‌های افزایشده DC یکی از چالش‌های مهم در کنترل غیرخطی است. در این مقاله، سعی شده است تا نحوه پیاده‌سازی کنترل گر غیرخطی بر اساس یادگیری تقویتی عمیق برای کنترل مبدل افزایشده DC شرح داده می‌شود. در این پژوهش، روند کنترل مبدل افزایشده DC به عنوان یک محیط برای عامل در نظر گرفته شده است. عامل‌های مورد استفاده در این پژوهش شامل روش PPO و DQN هستند. پس از بیان چگونگی پیاده‌سازی روش‌ها در محیط سیمولینک، برای پیاده‌سازی مناسب روش‌های یادگیری تقویتی عمیق، عملیات بهینه یابی ابرپارامترهای شبکه عصبی عمیق انجام شده است. مقایسه روش‌ها نشان می‌دهد که تکنیک‌های مبتنی بر یادگیری تقویتی عمیق عملکرد مناسب‌تری نسبت به روش PI در کنترل مبدل‌های افزایشده DC دارند. بررسی حساسیت روش‌ها نسبت به نویز سفید نشانگر توانمندی روش PPO در کنترل غیرخطی مبدل افزایشده DC است.

کلمات کلیدی: مبدل افزایشده، یادگیری تقویتی عمیق، یادگیری ماشین، کنترل غیرخطی، یادگیری تقویتی

یک فرایند تصمیم‌گیری متوالی بهینه در نظر گرفته شده است. باید توجه داشت که با توجه به رشد روزافزون استفاده از یادگیری تقویتی عمیق، از این روش‌ها در کنترل مسائل متنوعی استفاده شده است. فیض نژاد و همکاران [۸] روش کنترل مدل فوق محلی را بر اساس یادگیری تقویتی DDPG برای کنترل مبدل کاهنده استفاده کردند. همچنین آن‌ها در پژوهشی دیگر این روش ابداعی را برای کنترل مبدل کاهشی-افزایشی نیز استفاده کردند [۹]. حاجی حسینی و همکاران [۱۰] از روش یادگیری تقویتی PPO برای تنظیم پارامترهای کنترل گر PID در کنترل مبدل کاهشی-افزایشی استفاده کردند. البته باید توجه داشت که در زمینه‌ی کنترل مبدل افزایشده، توجه کمی به استفاده از روش‌های یادگیری عمیق شده است و هنوز بسیاری از روش‌های یادگیری تقویتی عمیق به طور مستقیم برای کنترل مبدل افزایشده مورد آزمون قرار نگرفته‌اند.

هدف اصلی این پژوهش، پیاده‌سازی یادگیری تقویتی عمیق در کنترل مبدل افزایشده است. به طور کلی می‌توان موارد زیر را به عنوان نوآوری در این پژوهش در نظر گرفت:

- مقایسه عملکرد کنترل خطی PID با نتایج حاصل از کنترل بر اساس یادگیری تقویتی عمیق.

۱- مقدمه

طراحی کنترل گر برای مبدل افزایشده یکی از چالش‌های مهم در کنترل غیرخطی محسوب می‌شود. نکته قابل توجه این است که تکنیک‌های کنترل خطی کلاسیک عملکرد ضعیفی در کنترل مبدل‌های افزایشده دارند. از طرفی باید توجه داشت که مسئله کنترل مبدل‌های افزایشده در یک فضای حالت پیوسته اتفاق می‌افتد. از این رو در این پژوهش سعی در پیاده‌سازی روش یادگیری تقویتی عمیق در کنترل مبدل افزایشده می‌گردد. باید توجه داشت که تا به امروز، روشی جامع که توانایی کنترل مبدل افزایشده در یک فضای حالت پیوسته را داشته باشد، پدید نیامده است.

استفاده از روش‌های کنترل خطی مانند کنترل کننده‌های (PID) برای کنترل مبدل افزایشده به طور گسترده‌ای در مقالات گزارش شده است. روش‌های سنتی طراحی کنترل گر با هدف تنظیم و ثابت نگه داشتن کنترل کننده PID انجام شده است. تا مبدل افزایشده، ولتاژ خروجی ثابتی داشته باشد [۱-۵]. لازم به ذکر است، تکنیک‌های کنترل خطی توصیف شده در مقالات به دلیل غیرخطی بودن سیستم مبدل افزایشده عملکرد رضایت بخشی را ارائه نمی‌دهند [۵]. نحوه عملکرد مبدل افزایشده به دو صورت کلید باز و کلید بسته است [۶]. در پژوهش پرادپ و همکاران [۷] ابتدا مسئله کنترل مبدل افزایشده به عنوان

که در این رابطه‌ها، V_{in} ولتاژ ورودی، L ضریب خودالقایی، r_L مقاومت معادل القاگر، R_{Load} مقاومت در برابر بار، C ظرفیت خازن و r_C مقاومت معادل خازن هستند.

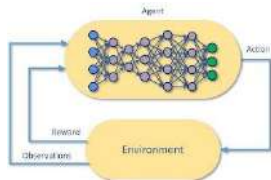
روابط ۱ تا ۴ در حقیقت معادلات دیفرانسیل خطی هستند که جداگانه بیان شده‌اند. مدل فضای حالت خطی که توسط رابطه‌های (۱) و (۲) ارائه شده است، سیستم را زمانی که کلید روشن است، توصیف می‌کند. رابطه‌های (۳) و (۴) توصیف‌کننده سیستم به هنگامی بسته بودن کلید هستند. از آنجا که مدل کلی سیستم به وضعیت المان سوئیچینگ بستگی دارد، سیستم کلی رفتاری غیرخطی دارد. در حقیقت، یک سیستم خطی طبق تعریف، سیستمی است که توسط یک معادله فضای حالت واحد به صورت

$$\frac{dx}{dt} = f(x, u) = Ax + Bu$$

توصیف می‌شود. در این رابطه تابع f باید خطی باشد. در مدل افزایشده، تابع f با یک معادله خطی واحد توصیف نمی‌شود و در واقع بین دو معادله حالت متفاوت به طور مداوم در حال تغییر است؛ بنابراین سیستم مدل افزایشده به طور کلی غیرخطی است.

۳- یادگیری تقویتی عمیق

یادگیری تقویتی عمیق زیرشاخه‌ای از یادگیری ماشین است که ترکیبی از یادگیری تقویتی و یادگیری عمیق است. در یادگیری تقویتی در واقع یادگیرنده به دنبال هدفی خاص و مشخص هست و با سعی و خطا در رسیدن به هدف تلاش دارد. در حقیقت می‌توان گفت که با استفاده از الگوریتم‌هایی موجب کمک به عامل در یافتن عمل می‌شود. در یادگیری تقویتی عمیق از شبکه‌های عصبی عمیق برای حل مسئله‌های یادگیری تقویتی استفاده می‌شود، به همین دلیل در نام آن از کلمه عمیق استفاده شده است. الگوریتم‌های یادگیری تقویتی عمیق توانایی این را دارد که در فضای پیوسته عمل کنند. همچنین تصمیم بگیرند که چه اقدام‌هایی را برای پیدا کردن جواب بهینه انجام دهند. یادگیری تقویتی عمیق برای مجموعه متنوعی از پژوهش‌ها و تحقیقات مثل رباتیک، بازی‌های ویدیویی، پردازش زبان طبیعی، بینایی رایانه‌ای، آموزش، حمل و نقل، امور مالی و مراقبت‌های بهداشتی مورد استفاده قرار گرفته است [۱۳-۱۴]. شکل ۲ نشان‌دهنده روند کار یادگیری تقویتی عمیق است. در ادامه به طور جداگانه، روش‌های یادگیری تقویتی عمیقی که در این پژوهش استفاده شده است، ذکر خواهد شد.



شکل ۲: روند یادگیری تقویتی عمیق [۱۳].

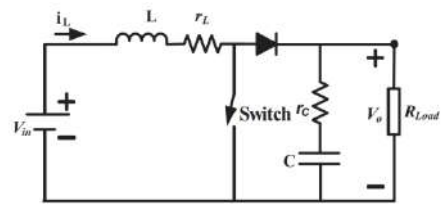
۴- روش شبکه عمیق Q

الگوریتم شبکه عمیق Q^2 یک روش یادگیری تقویتی بدون مدل و آنلاین است. عامل شبکه عمیق Q ، یک عامل یادگیری تقویتی مبتنی بر ارزش است که منتقد را برای تخمین بازده یا پاداش‌های آینده تربیت می‌کند. شبکه

- بررسی عملکرد روش‌های PPO و DQN در کنترل بهینه غیرخطی مدل‌های افزایشده.
- دستیابی به ابر پارامترهای شبکه عمیق در فرایند یادگیری تقویتی برای یافتن بهترین مشخصات شبکه عمیق در کنترل بهینه غیرخطی مدل افزایشده.

۲- مدل افزایشده

مدل افزایشده^۱ [۱۱ و ۱۲] یک مدل DC به DC است که کاربردهای زیادی در استفاده از انرژی خورشیدی، پیل‌های سوختی، خودروهای برقی و نورهای فلورسنت دارد. در حقیقت، مدل‌های افزایشده برای دستیابی به ولتاژهای بالاتر در منابع ولتاژ DC استفاده می‌شوند. همه مدل‌های افزایشده به حداقل دو دستگاه سوئیچینگ و همچنین یک عنصر ذخیره انرژی نیاز دارند. کنترل ولتاژ خروجی مدل افزایشده یک مسئله کنترل غیرخطی چالش برانگیز است. در واقع، پاسخ مدل به حالت عناصر سوئیچینگ بستگی دارد. باید توجه داشت که استراتژی‌های معمول طراحی کنترل‌کننده با استفاده از مدل‌های خطی تقریبی، عملکرد خوبی ندارند؛ بنابراین دستیابی به روش‌های جایگزین بهینه و غیرخطی بسیار مهم است. شکل ۱ مدار مدل افزایشده را به نمایش گذاشته است. رفتار مدل افزایشده را می‌توان با دو مدل خطی فضای حالت در نظر گرفت. اولین مدل در حالتی است که سوئیچ مدل افزایشده باز است و حالت دوم سوئیچ مدل افزایشده بسته است.



شکل ۱: مدار مدل افزایشده [۱۱].

مدل فضای حالت مدل افزایشده با المان سوئیچینگ در حالت باز در رابطه‌های (۱) و (۲) به نمایش گذاشته شده است.

$$\begin{bmatrix} \dot{x}_1 \\ \dot{x}_2 \end{bmatrix} = \begin{bmatrix} \frac{-r_L}{L} & 0 \\ 0 & \frac{-1}{C(R_{Load} + r_C)} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} \frac{1}{L} \\ 0 \end{bmatrix} V_{in} \quad (1)$$

$$V_o(x) = \begin{bmatrix} 0 & \frac{R_{Load}}{(R_{Load} + r_C)} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \quad (2)$$

مدل فضای حالت مدل افزایشده با المان سوئیچینگ در حالت بسته در رابطه‌های (۳) و (۴) به نمایش گذاشته شده است.

$$\begin{bmatrix} \dot{x}_1 \\ \dot{x}_2 \end{bmatrix} = \begin{bmatrix} \frac{-r_L + R_{Load} / r_C}{L} & \frac{-R_{Load}}{L(R_{Load} + r_C)} \\ \frac{-R_{Load}}{L(R_{Load} + r_C)} & \frac{-1}{C(R_{Load} + r_C)} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} \frac{1}{L} \\ 0 \end{bmatrix} V_{in} \quad (3)$$

$$V_o(x) = \begin{bmatrix} R_{Load} / r_C & \frac{R_{Load}}{(R_{Load} + r_C)} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \quad (4)$$

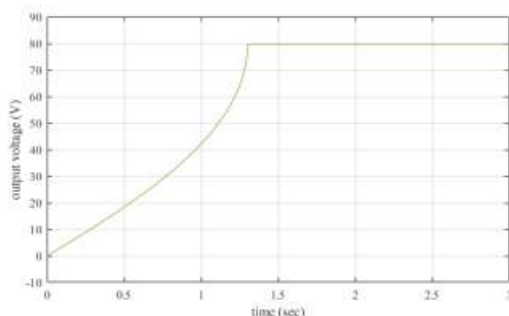
² Deep Q Network

¹ Boost Converter

میانگین مربعات خطا (MSE)^۴، میانگین خطای مطلق (MAE)^۵ و انتگرال خطای مطلق (IAE)^۶ به عنوان معیار خطا مورد استفاده قرار گرفته اند.

۷- کنترل کننده PI

با اعمال کنترل کننده PI بر سیستم نتایج کنترل سیستم برای ولتاژ مرجع مورد بررسی قرار گرفته است. هدف از استفاده از کنترل گر PI مقایسه نتایج آن با عملکرد کنترل گر یادگیری تقویتی عمیق است. شکل ۳ نتیجه خروجی از کنترل گر PI را برای ولتاژ مرجع ۸۰ ارائه می دهند. زمان کل شبیه سازی کنترل ۳ ثانیه است.



شکل ۳: خروجی کنترل گر PI برای ولتاژ مرجع ۸۰ ولت.

۷-۱- بهینه یابی ابر پارامترهای شبکه عصبی

باید توجه داشت که برای اعمال روش های یادگیری تقویتی عمیق لازم است پیش از دستیابی به نتایج نهایی با استفاده از روش های جستجو و سعی و خطا بهترین پارامترهای شبکه عصبی بهینه یابی گردند. مهم ترین ابر پارامترهای شبکه عصبی شامل تعداد لایه ها و تعداد نورون ها در هر لایه از شبکه هستند. برای دستیابی به بهترین تعداد لایه ها و نورون ها در این پژوهش به تأثیر این پارامترها بر میزان خطا و هزینه های محاسباتی توجه شده است. در حقیقت برای یافتن بهترین تعداد لایه سعی شده است تا تأثیر افزایش لایه ها بر میزان خطا و هزینه های محاسباتی سنجیده شود. برای این منظور تعداد لایه ها از تعداد ۳ لایه تا ۱۲ لایه مورد بررسی قرار گرفته اند. باید توجه داشت که برای همه محاسبات تعداد نورون ها در تمام لایه ها برابر با ۲۵۶ در نظر گرفته شده است تا فقط تأثیر افزایش تعداد لایه ها در عملکرد روش ها سنجیده شود. با توجه به نتایج موجود واضح است که افزایش تعداد لایه ها موجب افزایش چشمگیر هزینه محاسباتی در فرایند یادگیری است. همچنین با توجه به تغییرات اندک در MSE می توان نتیجه گرفت که افزایش تعداد لایه ها تأثیر چندانی در بهبود نتایج کنترل در این مسئله ندارد. بنابراین واضح است که استفاده از تعداد ۳ لایه برای شبکه عصبی عمیق در فرایند یادگیری امری منطقی است. در ادامه پژوهش کلیه نتایج بر اساس روش های یادگیری تقویتی عمیق با سه لایه ارائه می گردد.

عمیق Q یک نوع آموزش Q است. عامل های شبکه عمیق Q را می توان در محیط هایی با فضاهای مشاهده و عمل جدول ۱ آموزش داد.

جدول ۱: فضاهای مشاهده و عمل در آموزش عامل شبکه عمیق Q

فضای عمل	فضای مشاهده
گسسته	پیوسته یا گسسته

۵- روش سیاست بهینه سازی مبدأ^۳ (PPO)

الگوریتم سیاست بهینه سازی مبدأ یک روش یادگیری تقویتی بدون مدل و آنلاین است. این الگوریتم مستلزم استفاده از تجربیات کوچک از تعامل با محیط برای به روزرسانی سیاست تصمیم گیری است. پس از به روزرسانی خطامشی، تجربیات گذشته کنار گذاشته می شوند و دسته جدیدی برای به روزرسانی خطامشی ایجاد می شود. الگوریتم سیاست بهینه سازی مبدأ یک کلاس از روش های گرادیان سیاست است که سعی می کند واریانس گرادیان را نسبت به سیاست های بهتر کاهش دهد؛ که باعث پیشرفت مداوم و اطمینان از عدم تغییر شدید سیاست نسبت به سیاست قبلی و یا کاهش مسیرهای غیر قابل برگشت شود. الگوریتم سیاست بهینه سازی مبدأ دارای دو تابع تقریب، شبکه بازیگر و شبکه منتقد است.

۵-۱- شبکه بازیگر

این شبکه انتخاب های عمل را مستقیماً به حالت مشاهده شده نگاشت می کند. در هر زمان خاص بازیگر حالت مشاهده شده را به دست می آورد و احتمال انجام عمل در فضای عمل را پیش بینی می کند. باین حال، این تابع میزان عمل مناسب را در مقایسه با سایر اقدامات موجود اندازه گیری نمی کند. از این رو، شبکه منتقد برای نقد اقدامات بازگشت داده شده توسط شبکه بازیگر به کار گرفته می شود.

۵-۲- شبکه منتقد

این شبکه حالت مشاهده شده و عمل برگشت داده شده توسط شبکه بازیگر را به عنوان ورودی در نظر می گیرد و انتظار می رود پاداش طولانی مدت تخفیف داده شده را برآورده کند.

۶- روش پژوهش

همان طور که اشاره گردید، هدف اصلی این پژوهش پیاده سازی روش های یادگیری عمیق در کنترل غیرخطی مبدل افزایشده است. برای این منظور از دو روش یادگیری تقویتی عمیق که شامل DQN و PPO می شود استفاده شده است؛ همچنین این روش ها در محیط سیمولینک پیاده سازی شده است. در ادامه، عملکرد روش های کنترلی با جزئیات مورد بررسی قرار گرفته است. نتایج عددی بررسی عملکرد کنترل در محیط سیمولینک ارائه می گردد. لازم به ذکر است که در این پژوهش ولتاژ مطلوب ۸۰ ولت در نظر گرفته شده است؛ همچنین عملکرد روش های مورد بحث در برابر نویز سفید نیز مورد بررسی قرار می گیرد. برای سنجش و مقایسه روش ها از معیارهای خطا استفاده شده است.

⁴ Mean squared error

⁵ Mean absolute error

⁶ Integral absolute error

³ Proximal policy optimization

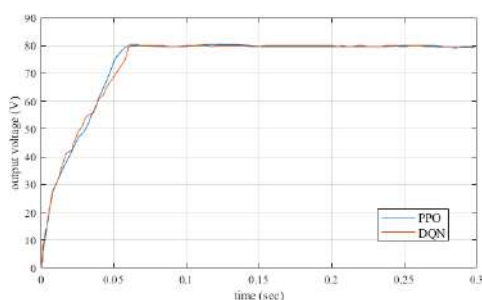
بهینه‌سازی در روند یادگیری تقویتی در کنترل غیرخطی مبدل‌های افزایشده است. روش‌های بهینه‌سازی^۷ sdgm،^۸ ame و^۹ rmsp در این پژوهش مورد ارزیابی قرار گرفتند. لازم به ذکر است که طبق جدول ۳، روش ame نسبت به سایر روش‌ها نتیجه بهتری را به نمایش گذاشت. بنابراین در ادامه پژوهش از روش ame برای بهینه‌سازی وزن‌های شبکه استفاده می‌شود.

جدول ۳: مقایسه تأثیر روش‌های مختلف بهینه‌سازی در عملکرد کنترل‌گر یادگیری تقویتی عمیق.

V _{ref}	RL method	PPO			DQN		
80 V	optimization method	rmsp	ame	sdgm	rmsp	ame	sdgm
	MSE	398.6587	391.2658	392.3697	395.3686	394.9941	401.3695

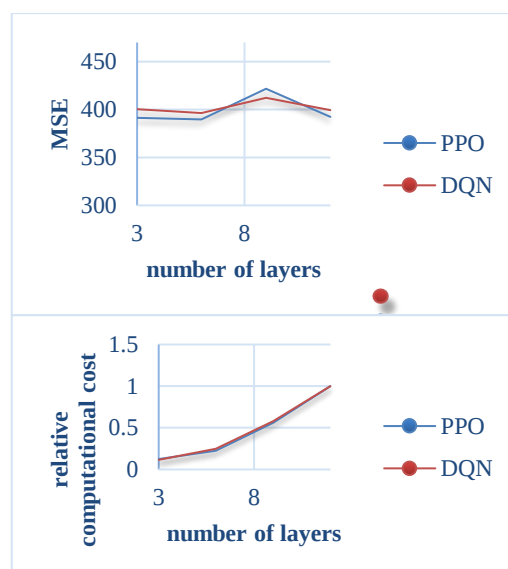
۸- پیاده‌سازی کنترل‌گر بر اساس یادگیری تقویتی عمیق

هدف اصلی این پژوهش مقایسه عملکرد کنترلی مبتنی بر یادگیری تقویتی عمیق با عملکرد کنترل‌گر PI است. خروجی کنترل‌گر بر اساس یادگیری تقویتی عمیق برای ولتاژ مرجع ۸۰ ولت در شکل ۵ نشان داده شده است. با مقایسه نتایج حاصل از کنترل بر اساس یادگیری تقویتی عمیق و داده‌های حاصل از کنترل‌گر PI واضح است که روش‌های یادگیری تقویتی عمیق عملکرد بسیار خوبی به لحاظ کاهش زمان نشست دارند. باید توجه داشت که با توجه به عملکرد روش‌های یادگیری تقویتی و کاهش چشمگیر زمان نشست، کل فرآیند شبیه‌سازی ۰.۳ ثانیه انتخاب شده است.



شکل ۵: خروجی کنترل‌گر یادگیری تقویتی عمیق برای ولتاژ مرجع ۸۰ ولت.

نتایج عددی ارائه شده در شکل ۶ و شکل ۷ شامل زمان نشست و میزان خطا است. باید توجه داشت که برای بررسی دقیق‌تر رفتار کنترل‌گرها، میزان خطا در کل زمان شبیه‌سازی گزارش شده است. نتایج ارائه شده برای روش‌های یادگیری تقویتی عمیق در حقیقت میانگین نتایج اعتبارسنجی روش‌های کنترلی در ۱۰۰ آزمایش است.



شکل ۴: میزان خطا (MSE) و هزینه محاسبات کنترل‌گر یادگیری عمیق با استفاده از لایه‌های مختلف برای ولتاژ مرجع ۸۰ ولت.

یکی دیگر از پارامترهای بسیار مهم و تأثیرگذار بر عملکرد شبکه عصبی تعداد نورون‌ها در هر لایه است. هرچند که واضح است که افزایش تعداد نورون‌ها در بیشتر حالات موجب عملکرد بهتر شبکه عصبی می‌گردد، لازم به ذکر است تعداد بالای نورون موجب افزایش چشمگیر حجم محاسبات می‌گردد. بنابراین در ادامه حالت‌های مختلف تعداد نورون موردبررسی قرار گرفته است. جدول ۲ تأثیر تعداد نورون را بر میزان MSE نشان می‌دهد.

جدول ۲: مقایسه عملکرد کنترل‌گر یادگیری تقویتی عمیق با توجه به تعداد نورون‌های مختلف.

80 V		
	neurons	MSE
PPO	32	585.3651
	64	452.6352
	128	412.253
	256	391.2658
	512	392.6850
DQN	32	623.365
	64	433.658
	128	394.9941
	256	395.6352
	512	398.8526

کمترین میزان خطا برای هر یک از روش‌ها در جدول ۲ به صورت برجسته نشان داده شده است. با توجه به نتایج موجود در جدول ۲ می‌توان نتیجه گرفته که روش PPO برای عملکرد مناسب نیازمند به شبکه‌ای عصبی با ۲۵۶ نورون است، درحالی‌که روش DQN بهترین عملکردش را با ۱۲۸ نورون نشان می‌دهد. هرچند باید توجه داشت که هنگامی که تعداد نورون‌ها از ۱۲۸ عدد در هر لایه بیشتر می‌شود، در همه حالات، هر دو روش PPO و DQN عملکرد مناسبی نشان می‌دهند. در حقیقت افزایش تعداد نورون‌ها بیشتر از تعداد ۱۲۸ تغییر شگرفی در نتایج ایجاد نمی‌کند. در ادامه این پژوهش برای اطمینان از عملکرد مناسب روش‌ها از ۲۵۶ نورون در هر لایه استفاده شده است.

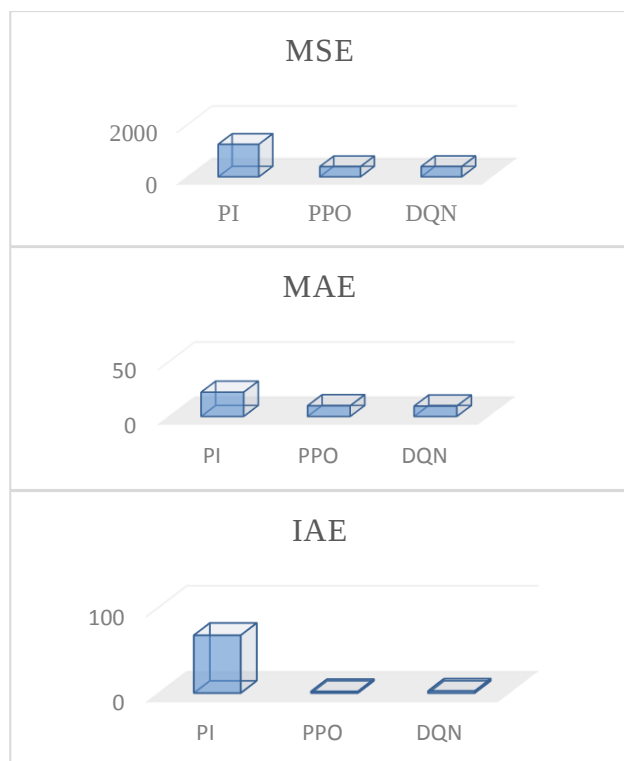
لازم به ذکر است که در حین فرایند یادگیری، لازم است تا دائماً وزن‌های شبکه بهینه گردند. در حقیقت روش‌های بهینه‌سازی تأثیر فراوانی در نتایج یادگیری دارند. یکی دیگر از اهداف این تحقیق بررسی عملکرد ۳ روش

⁷ stochastic gradient descent with momentum

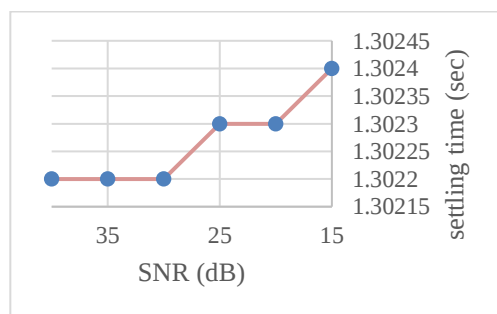
⁸ adaptive moment estimation

⁹ root mean square propagation

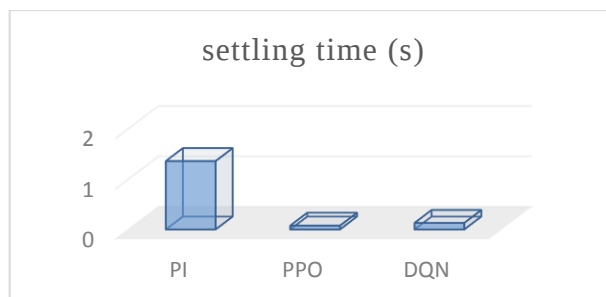
شکل ۹ میزان عملیات موفقیت آمیز با استفاده از روش‌های PPO و DQN را در برابر ۱۰۰ آزمایش دارای نویز سفید نشان می‌دهد. نتایج نشان می‌دهد که روش PPO در تمامی حالات برای SNR برابر با ۲۵ و ۳۰ دسی‌بل موفق به کنترل سیستم دارای نویز شده است. لیکن روش DQN عملکرد مطلوبی ندارد و حتی تا مقدار ۹۰ درصد آزمایشات را شکست می‌خورد. یعنی در حقیقت در بیشتر حالات دارای نویز با SNR کمتر از ۳۰ اصلاً قادر به دستیابی به ولتاژ مرجع موردنظر نیست. با توجه به نتایج در شکل ۹ روش PPO بسیار نیرومندتر از روش DQN است.



شکل ۷: مقایسه میزان خطا در تمام فرایندهای کنترل برای روش‌های کنترلی مختلف در ولتاژ مرجع برابر با ۸۰ ولت.



شکل ۸: تأثیر افزایش میزان نویز در افزایش زمان نشست در کنترل گر PI ولتاژ مرجع برابر با ۸۰ ولت.



شکل ۶: مقایسه میزان زمان نشست برای روش‌های کنترلی مختلف در ولتاژ مرجع برابر با ۸۰ ولت.

شکل ۶ مقادیر زمان نشست روش‌های مختلف را باهم مورد مقایسه قرار می‌دهد. همان‌طور که واضح است زمان نشست روش‌های یادگیری تقویتی عمیق نسبت به روش PI بسیار کمتر است. این امر عملکرد بسیار مناسب روش‌های یادگیری تقویتی عمیق را در دستیابی سریع به ولتاژ مرجع در کنترل مبدل‌های افزایش دهنده DC نشان می‌دهد. همچنین با مقایسه میزان زمان نشست روش‌های یادگیری تقویتی، می‌توان نتیجه گرفت که عملکرد روش PPO به لحاظ سرعت کنترلی بهتر از روش DQN است.

شکل ۷ نتایج مربوط به فرایند کنترل برای ولتاژ مرجع ۸۰ ولت را نشان می‌دهند. این نتایج عملکرد بسیار خوب روش‌های PPO و DQN را نسبت به روش PI نشان می‌دهند. باید توجه داشت که در بیشتر حالات و همچنین بر اساس اکثر معیارهای خطا، عملکرد روش PPO بهتر از روش DQN در فرایند کنترل غیرخطی مبدل‌های افزایش دهنده DC است.

۹- حساسیت کنترل گرها به نویز سفید

برای سنجش نیرومندی روش‌های کنترلی در این بخش عملکرد این روش‌ها در برابر اعمال نویز گوسی سفید^{۱۰} سنجیده شده است. نویز گوسی سفید یک نویز خطی است که به سیگنال ارسالی سیستم اعمال می‌گردد. این نویز دارای توزیع گوسی نسبت به زمان است. برای تعریف نویز گوسی سفید معمولاً از نسبت سیگنال به نویز (SNR)^{۱۱} استفاده می‌شود. این نسبت بر اساس رابطه-۵ تعریف می‌شود.

$$SNR = \frac{P_{signal}}{P_{noise}} \quad (5)$$

که در این رابطه P_{signal} توان اطلاعات موردنظر و P_{noise} توان سیگنال ناخواسته است.

برای سنجش اثر نویز ابتدا نویز سفید به کنترل گر PI اعمال شده است. شکل ۸ تأثیر افزایش نویز سفید را بر عملکرد کنترل گر PI به لحاظ افزایش زمان نشست نشان می‌دهد. همان‌طور که در شکل ۸ نشان داده شده است، با کاهش SNR (افزایش نویز) مقدار زمان نشست افزایش می‌یابد. با توجه به شکل می‌توان نتیجه گرفت که مقدار SNR برابر با ۲۵ و ۳۰ دسی‌بل در حقیقت مرز ایجاد تغییرات در فرایند کنترلی است. از این‌رو، رفتار روش‌های DQN و PPO نسبت به اعمال نویز با مقادیر ۲۵ و ۳۰ دسی‌بل سنجیده می‌شود.

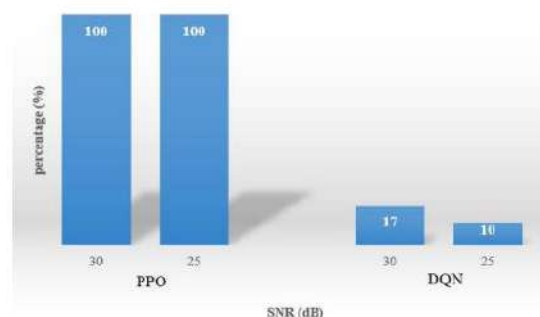
¹⁰ Additive White Gaussian Noise

¹¹ Signal to noise ratio

- مقایسه عملکرد روش‌های یادگیری تقویتی به لحاظ زمان نشست و میزان خطا نشان داد که روش PPO نسبت به روش DQN در کنترل غیرخطی مبدل افزایشده بهتر عمل می‌کند.
- بررسی اعمال نویز سفید بر کنترل‌گرهای مبتنی بر یادگیری تقویتی نشان داد که روش DQN حساسیت بالایی به اعمال نویز سفید در فرایند کنترل مبدل افزایشده DC دارد. لیکن روش PPO در تمامی حالات مورد آزمایش، عملکرد خوبی در برابر اعمال نویز داشت.

مراجع

- [1] Hung JY, Gao W, Hung JC. *Variable structure control: A survey*. IEEE transactions on industrial electronics. 1993 Feb;40(1):2-2.
- [2] Cominos P, Munro N. *PID controllers: recent tuning methods and design to specification*. IEE Proceedings-Control Theory and Applications. 2002 Jan 1;149(1):46-53.
- [3] Guo L, Hung JY, Nelms RM. *Digital controller design for buck and boost converters using root locus techniques*. In: IECON'03. 29th Annual Conference of the IEEE Industrial Electronics Society (IEEE Cat. No. 03CH37468) 2003 Nov 2 (Vol. 2, pp. 1864-1869). IEEE.
- [4] Balestrino A, Corsanini D, Landi A, Sani L. *Circle-based criteria for performance evaluation of controlled dc-dc switching converters*. IEEE Transactions on Industrial Electronics. 2006 Nov 30;53(6):1862-9.
- [5] Perry AG, Feng G, Liu YF, Sen PC. *A new design method for PI-like fuzzy logic controllers for DC-to-DC converters*. In: 2004 IEEE 35th Annual Power Electronics Specialists Conference (IEEE Cat. No. 04CH37551) 2004 Jun 20 (Vol. 5, pp. 3751-3757). IEEE.
- [6] Andalibi M, Hajihosseini M, Teymouri S, Kargar M, Gheisarnejad M. *A Time-Varying Deep Reinforcement Model Predictive Control for DC Power Converter Systems*. In: 2021 IEEE 12th International Symposium on Power Electronics for Distributed Generation Systems (PEDG) 2021 Jun 28 (pp. 1-6). IEEE.
- [7] Pradeep DJ, Noel MM, Arun N. *Nonlinear control of a boost converter using a robust regression based reinforcement learning algorithm*. Engineering Applications of Artificial Intelligence. 2016 Jun 1;52:1-9.
- [8] Gheisarnejad M, Farsizadeh H, Khooban MH. *A Novel Nonlinear Deep Reinforcement Learning Controller for DC-DC Power Buck Converters*. IEEE Transactions on Industrial Electronics. 2020 Jul 1;68(8):6849-58.
- [9] Gheisarnejad M, Farsizadeh H, Tavana MR, Khooban MH. *A Novel Deep Learning Controller for DC-DC Buck-Boost Converters in Wireless Power Transfer Feeding CPLs*. IEEE Transactions on Industrial Electronics. 2020 May 20;68(7):6379-84.
- [10] Hajihosseini M, Andalibi M, Gheisarnejad M, Farsizadeh H, Khooban MH. *DC/DC power converter control-based deep machine learning techniques: Real-time implementation*. IEEE Transactions on Power Electronics. 2020 Mar 2;35(10):9971-7.
- [11] Sundareswaran K, Sreedevi VT. *Boost converter controller design using queen-bee-assisted GA*. IEEE Transactions on industrial electronics. 2008 Oct 31;56(3):778-83.
- [12] Rashid MH. *Power electronics: circuits, devices, and applications*. Pearson Education India; 2009.
- [13] Li Y. *Deep reinforcement learning: An overview*. arXiv preprint arXiv:1701.07274. 2017 Jan 25.



شکل ۹: مقایسه عملکرد روش‌های PPO و DQN در برابر اعمال نویز سفید برای ولتاژ مرجع برابر با ۸۰ ولت.

۱۰- نتیجه‌گیری

پژوهش حاضر شامل چگونگی پیاده‌سازی روش‌های یادگیری تقویتی عمیق در کنترل غیرخطی مبدل افزایشده DC است. نتایج عددی این پیاده‌سازی و همچنین مقایسه عملکرد روش‌ها موجب دستیابی به نتایج مفیدی گردید. نتیجه‌های به‌دست‌آمده از پیاده‌سازی و مقایسه روش‌ها به‌طور خلاصه در ادامه بیان می‌گردد.

- بهینه‌یابی تعداد لایه‌های شبکه عصبی عمیق در عملکرد روش‌های یادگیری عمیق نشان داد که افزایش تعداد لایه‌ها تأثیر چندانی در میزان خطا ندارد. باید توجه داشت که این افزایش تعداد لایه‌ها هرچند ممکن است که تأثیری کم در کاهش خطای عملکرد کنترل‌گر داشته باشد، لیکن موجب افزایش چشمگیر هزینه محاسباتی می‌گردد. از این‌رو، برای کنترل غیرخطی مبدل افزایشده DC در این پژوهش از ۳ لایه بهره‌برد شد.
- بررسی تعداد نورون‌ها در میزان خطای کنترل‌گر نشان داد که پس از رسیدن تعداد نورون‌ها به عدد ۲۵۶، عامل عملکرد مناسبی از خود نشان می‌دهد. واضح است که افزایش تعداد نورون‌ها موجب افزایش هزینه محاسباتی می‌گردد. لیکن عملکرد عامل‌های یادگیری تقویتی عمیق با تعداد نورون کمتر خیلی مناسب نیست. همچنین باید توجه داشت که عامل با تعداد نورون بالاتر از ۲۵۶ نتایج مشابه ارائه می‌کرد. بنابراین در این پژوهش تعداد نورون‌ها ۲۵۶ انتخاب گردید.
- یکی دیگر از عوامل مهم در عملکرد روش‌های یادگیری عمیق، انتخاب روش بهینه‌سازی برای شبکه عصبی است. در این پژوهش عملکرد سه روش معروف بهینه‌سازی موردبررسی قرار گرفت. نتایج مقایسه‌ای بیانگر برتری روش adam نسبت به سایر روش‌ها بود. باید توجه داشت که هر سه روش بهینه‌سازی توانایی خوبی داشتند.
- بررسی مقایسه‌ای روش‌های مختلف کنترلی نشان داد که کنترل‌گر مبتنی بر یادگیری تقویتی عمیق دارای زمان نشست بسیار کمتری نسبت به کنترل‌گر PI است. این امر موجب می‌شود که خطا در فرایند کنترلی کاهش چشمگیری نسبت به کنترل‌گر PI داشته باشد.



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بم

تشخیص اخبار جعلی مبتنی بر ترکیبی از ویژگی‌های زبانی و صفات گوینده خبر

یکتا نصیری پور^۱، سعیده انبایی فریمانی^۲ و مجید وفایی جهان^۳

^۱دانشجو ارشد کامپیوتر، دانشگاه آزاد اسلامی، واحد مشهد، nasiripooryekta@mshdiau.ac.ir

^۲دانشجو دکتری کامپیوتر، دانشگاه آزاد اسلامی، واحد مشهد، anbaee@mshdiau.ac.ir

^۳دانشیار کامپیوتر، دانشگاه آزاد اسلامی، واحد مشهد، VafaeiJahan@mshdiau.ac.ir

چکیده: انتشار سریع اطلاعات نادرست یک نگرانی رو به رشد در سراسر جهان است، این امر نیاز به الگوریتم‌های موثر تشخیص اخبار جعلی که در آن نویسندگان با مقصودی خاص اطلاعات کاملاً غیرواقعی یا نادرستی را منتشر می‌کند، برانگیخته است. اکثر روش‌های پیشین برای شناسایی اخبار جعلی تنها بر استخراج ویژگی‌های زبانی از متن خبرها و یا اطلاعات مربوط به گوینده خبر مانند اسم گوینده، شغل و موقعیت سیاسی فرد، متکی هستند. در این مقاله، هدف طراحی یک مدل برای شناسایی اخبار جعلی است که در بخش بازنمایی ویژگی‌های متنی و در راستای بازنمایی اطلاعات نحوی و معنایی واژه‌ها، ابتدا تک تک واژه‌ها در متن بیانیه‌ها با استفاده از یک مدل زبانی مبتنی بر برت بازنمایی شده‌اند و سپس با بکارگیری شبکه عصبی کانولوشن زمانی که عملکرد چشم‌گیری در طبقه‌بندی جملات بدست آورده است، ویژگی‌های متن بیانیه‌ها استخراج شده‌اند و از طرفی در راستای استفاده از ویژگی‌های گوینده‌ی خبر، تعبیه متاداده‌ها به شبکه آموزش داده می‌شود و در نهایت ویژگی‌های استخراج‌شده از متن بیانیه‌ها و متاداده‌ها تلفیق شده‌اند. نتایج آزمایش‌ها نشان از کسب دقت ۹۴ درصد در قسمت تست مجموعه داده‌ی دروغگو دارد.

کلمات کلیدی: شبکه‌های کپسولی، طبقه‌بندی جملات، کشف خبرهای جعلی، پردازش زبان طبیعی، هوش مصنوعی، یادگیری عمیق

۱_ مقدمه

جدول ۱- یک نمونه خبر جعلی در مجموعه داده به همراه صفات گوینده-
ی خبر

Statement: Under the health care law, everybody will have lower rates, better quality care and better access.
Speaker: Nancy Pelosi
Context: on Meet the Press
Label: False
Justification: Even the study which Pelosis staff cited as a source of that the statement suggested that some people would pay more for health insurance. Analysis at the state level found the same thing. The general understanding of the word everybody is every person. The predictions do not back that up. We rule this statement False.

امروزه رشد فزاینده اسناد متنی در وب بر اهمیت طبقه‌بندی آنها تاکید دارد [۱-۳]. یک نوع از این اسناد متنی، اخبار جعلی است، یعنی خبری با محتوای کاملاً جعلی که توسط نویسندگان مقالات ساخته شده‌است، یا اخباری دارای درجاتی از محتوای واقعی که کاملاً دقیق نیستند [۴]. جدول ۱ یک نمونه تصادفی خبر جعلی از مجموعه داده استفاده شده در این پژوهش را نشان می‌دهد. این نوع اخبار جعلی باعث می‌شود مردم دیگر به آنچه واقعی است اعتماد نکنند [۵]. تشخیص اطلاعات نادرست در شبکه‌های اجتماعی یک مسئله فوق‌العاده مهم بلکه از نظر فنی نیز یک چالش است [۳]. از این رو مدل‌هایی با استفاده از شبکه‌های عصبی کانولوشن ارائه شده‌است [۶-۹] که اغلب مبتنی بر یکی از دو ویژگی زبانی متن اخبار یا صفات داده‌ای گوینده‌ی خبر هستند.

ویلیام یانگ و انگ [۱۰] مدل کانولوشن هیبرید را با استفاده از ویژگی‌های متنی با در نظر گرفتن مجموعه داده دروغگو^۱ پیشنهاد داد. شبکه عصبی کانولوشن عملکرد بهتری نسبت به همه مدل‌ها با دقت ۰.۲۷۰ در مجموعه تست داشت. در روش یانگ و همکاران [۱۰] هم از ویژگی‌های متنی متاداده‌ها^۲ و هم متن بیانیه‌ها استفاده شده است، اما استفاده از لایه‌ی شبکه‌های عصبی مبتنی بر حافظه کوتاه مدت^۳ در حالیکه داده‌ها ارتباط زمانی نسبت به هم ندارند، سبب شده مدل به دقت بالایی دست پیدا نکند. محمدهادی گلدانی و همکاران [۱۱] معماری مبتنی بر شبکه‌های کپسولی بر روی مجموعه داده‌ی دروغگو پیشنهاد داده است. در این روش فقط از بازنمایی متاداده‌ها استفاده شده و نتایج تجربی در این مجموعه داده با بهبود ۳.۱ درصد روی داده‌های اعتبار سنجی و ۱ درصد روی داده‌های تست است. آن‌ها با استفاده از ساختار شبکه کپسولی سه لایه مدل خود را بنا نموده‌اند و بالاترین دقت بدست آمده در این روش [۱۱] نزدیک ۴۰ درصد می‌باشد. در پژوهش پیش رو با هدف بهبود دقت تشخیص اخبار جعلی از واقعی از ترکیبی از ویژگی‌های متن بیانیه‌ها و همچنین صفات متاداده‌ای مربوط به گوینده‌ی بیانیه استفاده شده است. یکی از اجزای این سیستم شبکه عصبی کانولوشن زمانی^۴ می‌باشد که از ویژگی‌های اصلی یک شبکه عصبی کانولوشن می‌توان به اتصال محلی، تجمع، چند لایه و وزن‌های مشترک اشاره کرد. ردیابی پیوندهای محلی ویژگی‌ها از لایه‌های قبلی، نقش لایه کانولوشن است و ادغام ویژگی‌های معنایی مشابه در یکی، نقش لایه همبستگی است [۱۲]. مزیت استفاده از روش کانولوشن زمانی در حفظ ترتیب کلمات در متن بیانیه‌ها و استخراج ویژگی‌هایی با بعد کم بر اساس توالی رخداد کلمات در متن بیانیه‌ها می‌باشد. بنابراین تعداد پارامترهای کل کاهش یافته که منجر به افزایش سرعت سیستم می‌شود [۱۳]. برای شناسایی مستندات جعلی، نوآوری‌های این پژوهش شامل موارد زیر است:

- در قسمت استخراج ویژگی از متن اعلامیه‌ها ما از بازنمایی کلمات دیستیل برت^۵ [۱۴] به عنوان یک مدل زبانی قوی در انعکاس اطلاعات معنایی و نحوی واژه‌ها استفاده کرده‌ایم.
- همچنین ما با در نظر گرفتن بازنمایی دو دسته اطلاعات در مجموعه داده یعنی هم متاداده‌ها و هم متن اعلامیه‌ها، به عنوان قصد گوینده برای

تشخیص جعلی یا واقعی بودن خبر توانستیم به نتایج خوبی دست پیدا کنیم و دقتی برابر با ۹۴ درصد در مجموعه داده تست را بدست بیاوریم.

- روش پیشنهادی خود را به ازای روش‌های مختلف بازنمایی متن بیانیه‌ها و همچنین حالات مختلف استخراج ویژگی که شامل تنها متن بیانیه خبر، متاداده‌ها که اطلاعات و مشخصات گوینده است و ترکیب هر دو این‌ها مورد ارزیابی قرار می‌دهیم.

۲- مرور ادبیات

اگرچه مسئله اخبار جعلی موضوع جدیدی نیست، اما با توجه به اینکه انسان تمایل دارد اطلاعات گمراه کننده که به سرعت در رسانه‌های اجتماعی منتشر می‌شوند را باور کند، کشف خبرهای جعلی یک کار پیچیده است [۴]. اخبار جعلی در سال‌های اخیر به ویژه پس از انتخابات ۲۰۱۶ آمریکا بیشتر مورد توجه قرار گرفته است [۱۵]. طیه رسول و همکاران [۱۵] مجموعه داده‌ی دروغگو را که دارای ۶ برچسب است، در نظر گرفت و با روش استخراج ویژگی انتخاب رو به جلو^۶، مدل ماشین‌های بردار پشتیبانی^۷ با بالاترین دقت برابر ۳۹.۵٪ بود. همچنین عبدالعزیز البحر و همکاران در مقاله خود [۱۶]، مدل‌های مختلف تعبیه کلمات مانند گلاو^۸ و کلمه به بردار^۹ [۱۶] را روی مجموعه داده دروغگو ارزیابی کردند و در نهایت بالاترین دقت با استفاده از مدل زبانی گلاو و مدل شبکه عصبی کپسولی به دقت ۶۴.۴٪ رسید. جونید یونس خان و همکاران [۱۷] مدلی مبتنی بر یادگیری ماشین و یادگیری عمیق ارائه می‌دهد. آنها از فراوانی تکرار اصطلاح_فرکانس سند معکوس^{۱۰} برای استخراج ویژگی استفاده می‌کنند و سپس از روش گلاو از قبل آموزش دیده برای بازنمایی کلمات استفاده می‌کنند. سرانجام در میان مدل‌ها شبکه بیزین ساده با دقت ۶۰٪ برای شناسایی اخبار جعلی با مدل‌های مبتنی بر شبکه عصبی رقابت می‌کند. در تمامی پژوهش‌هایی که مورد استناد قرار گرفت به اغلب از مدل‌های زبانی مانند گلاو و کلمه به بردار و فرکانس اصطلاح_فرکانس سند معکوس برای بازنمایی متنی استفاده شده، که در این روش‌ها ترتیب کلمات حفظ نمی‌شوند و از چالش‌هایی نظیر ابعاد بالای ویژگی‌ها و همچنین وجود کلمات با معانی متعدد رنج می‌برند. این در صورتیست که روش برت از قبل آموزش دیده، اطلاعات نحوی در ساختار جمله‌ها را در نظر می‌گیرد و برای کلمات با معانی متفاوت بازنمایی متفاوتی تولید می‌کند [۱۸].

Forward Feature Selection^۱

SVM^۲

Glove^۸

Word2vec^۹

TF_IDF^{۱۰}

<https://www.kaggle.com/khandalavyan/liar-preprocessed-dataset>

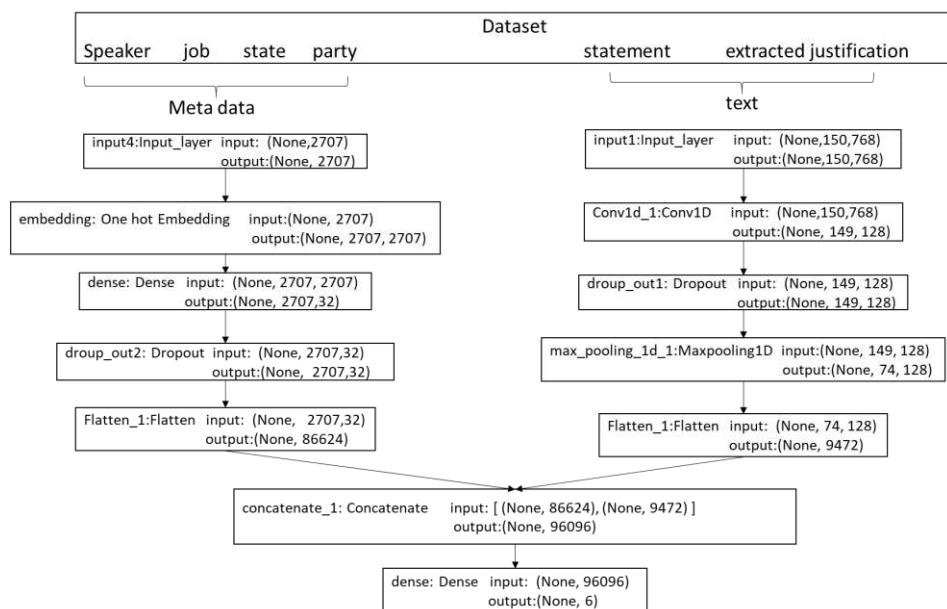
dataset

metadata^۲

Long short-term memory^۳

Temporal Convolution (Conv 1D)^۴

DistilBERT^۵



شکل ۱. معماری مدل پیشنهادی برای تشخیص اخبار جعلی در بیانیه های خبری کوتاه

از توکنایزر برت و تبدیل آن‌ها به فرم پایه است. پس از تبدیل شدن هر متن از مجموعه داده به لیستی از کلمات تمیز، می‌توانیم آن‌ها را بازنمایی کنیم. بازنمایی به این صورت است که به ازای هر لغت در متن اعلامیه‌ها بازنمایی برداری را با ابعاد (۷۶۸ تعداد لایه‌های پنهان، تعداد توکن‌های ورودی که حداکثر ۵۱۲ می‌باشد) از دیستیل برت [۲۰] استخراج می‌شود و به این دلیل که توالی کلمات در متن اعلامیه‌ها دارای اهمیت است برای استخراج ویژگی و حفظ ترتیب زمانی از یک لایه کانولوشن یک بعدی^{۱۲} استفاده می‌کنیم. در هر لایه ما یک پولینگ داریم که ماکزیمم پولینگ در نظر گرفته شده است که در بین اعداد فیلترها بزرگترین عدد را به نمایندگی در نظر بگیرد.

۲-۳ استخراج ویژگی از متاداده‌ها

متاداده‌ها در این پژوهش شامل اطلاعات سخنران خبر (اسم گوینده، شغل و موقعیت سیاسی فرد) می‌باشد. برای استخراج ویژگی‌های متاداده‌ها از بازنمایی وان‌ها^{۱۳} استفاده کردیم. این بازنمایی به صورتی است که به ازای هر کلمه در مجموعه واژگان یک بردار داریم که همگی صفر هستند و تنها در جایگاه آن کلمه مقدار یک قرار می‌گیرد. به ازای هر کدام از ۴ صفت متاداده، بردارهای متناظر در کنار هم قرار گرفته و به عنوان ورودی به یک لایه امبدینگ ارسال می‌شوند. بازنمایی تعبیه شده‌ی هر کدام از این چهار ویژگی در فرایند پسانتشار خطا یاد گرفته می‌شود. برای جلوگیری از بیش‌برازش مدل از یک لایه دراپ‌اوت^{۱۴} استفاده شده‌است. پس برای استخراج ویژگی با هدف سادگی مدل و کم شدن پارامترها از یک لایه تمام متصل استفاده شده است.

۳- مدل پیشنهادی

شکل ۱ معماری مدل پیشنهادی ما را نشان می‌دهد. در این معماری علاوه بر متن بیانیه‌ها، از متاداده‌ها به عنوان ورودی اضافی که قصد گوینده را در مورد جعل یا واقعی بودن بیانیه نشان می‌دهد، استفاده می‌کنیم پس از محاسبه‌ی بازنمایی هر دو بخش، با استفاده از دو شبکه مجزا از هم به استخراج ویژگی می‌پردازیم. پس با تخت^{۱۱} کردن هر دو ماتریس خروجی دو زیر شبکه و الحاق آن‌ها با استفاده از یک طبقه بند احتمال رخداد هر کلاس محاسبه می‌شود.

۱-۳ بازنمایی متن اعلامیه‌ها

ما از روش بازنمایی دیستیل برت [۲۰] برای تعبیه متن بیانیه‌ها استفاده می‌کنیم. جامعه پردازش زبان‌های طبیعی شاهد یک تغییر الگوی چشمگیر به سمت مدل برت از قبل آموزش دیده‌است، که به عنوان مثال می‌توان به طبقه‌بندی احساسات و مدل سازی شباهت دست یافت [۲۱-۱۹]. نوآوری فنی اصلی دیستیل برت مبتنی بر برت استفاده از آموزش دو طرفه ترانسفورمر، برای مدل سازی زبان است. نتایج نشان می‌دهد که یک مدل زبانی که به طور دو طرفه آموزش دیده‌است می‌تواند نسبت به مدل‌های تک جهت زبانی، مفهوم عمیق‌تری از متن و جریان زبان داشته باشد [۲۲]. مزیت مدل زبانی برت نسبت به بقیه مدل‌های زبانی دیگر بازتاب اطلاعات معنایی و نحوی واژه‌ها و همچنین غلبه بر چالش وجود لغات با معانی متعدد در پیکره‌ی متنی است. در ابتدا باید هر سند از مجموعه داده را پیش پردازش کرد که شامل مراحل (۱) حذف کلمات ایست از متن، (۲) حذف حروف اضافه مثل اعداد و علامت‌ها، (۳) تبدیل حروف بزرگ به کوچک، (۴) توکنایز کلمات با استفاده

^{۱۳} Drop out

^{۱۱} Flat

^{۱۲} Temporal Convolution

۳-۳. تلفیق ویژگی‌های استخراج شده

برای ادغام ویژگی‌های استخراج شده از متاداده‌ها و متن بیانی‌ها هر دو ماتریس ویژگی بدست آمده از دو زیر شبکه را تخت کرده و به بردار تبدیل می‌کنیم و در نهایت این دو خروجی را با هم تلفیق می‌کنیم و در پایان، احتمال تعلق هر نمونه داده را به هر برچسب از طریق یک لایه تمام متصل با تابع فعالیت softmax دریافت می‌کنیم. معادله تابع فعالیت استفاده شده به صورت زیر است:

$$s(Y_i) = \frac{e^{y_i}}{\sum_j e^{y_j}} \quad (1)$$

هر احتمال کلاس پیش‌بینی شده با خروجی مطلوب کلاس واقعی ۰ یا ۱ مقایسه می‌شود و یک امتیاز/ضرر محاسبه می‌شود. مقدار احتمال بدست آمده بر اساس فاصله آن با مقدار واقعی مورد انتظار جریمه می‌کند. این جریمه ماهیت لگاریتمی دارد و هنگام تنظیم وزن‌های مدل از تابع خطای آنتروپی متقاطع^{۱۴} استفاده می‌شود.

۴-۲. تنظیمات آزمایش‌ها

۴-۱. مجموعه داده

یکی از مجموعه داده‌های شناخته شده اخیر، که توسط وانگ [۱۰] ارائه شده، مجموعه داده بزرگی به نام دروغگو است که شامل بیانی‌های کوتاه با برچسب انسانی از POLITIFACT.COM API است. یک پروژه غیرانتفاعی آمریکایی است که توسط موسسه پوینتر اداره می‌شود. گزارشگران و سردبیران روزنامه و شرکای رسانه‌های خبری وابسته به آن که در مورد صحت اظهارات مقامات منتخب و افراد دیگر در سیاست ایالات متحده گزارش می‌دهند. اعتبار هر گزاره توسط وبسایت POLITIFACT.COM^{۱۵} ارزیابی شده است. این مجموعه داده به صورت سه بخش جداگانه آموزش، تست و اعتبارسنجی ارائه شده است. تاریخ بیانی‌های موجود عمدتاً مربوط به سال‌های ۲۰۰۷-۲۰۱۶ است. سخنرانان مجموعه داده ترکیبی از دمکرات‌ها و جمهوری خواهان می‌باشد و همچنین شامل تعداد قابل توجهی از پست‌های رسانه‌های اجتماعی آنلاین می‌باشد. برای هر سخنران مجموعه‌ای غنی از متاداده را در اختیار داریم که علاوه بر وابستگی به احزاب، شغل فعلی، ایالت خانه و تاریخچه اعتبار نیز ارائه شده است. یک نمونه‌ی تصادفی از مجموعه داده در جدول ۱ آورده شده است. تعداد رکوردهای این مجموعه داده در جدول ۲ آورده شده است که دارای شش برچسب با عنوان غلط^{۱۶}، درست^{۱۷}، نیمه درست^{۱۸}، به سختی درست^{۱۹}، بیشتر درست^{۲۰} و شلوار آتش^{۲۱} برای درجه صحت در نظر گرفته شده است و برای راحتی کار به هر کدام از برچسب‌ها یک عدد به ترتیب از ۰ تا ۵ اختصاص داده‌ایم.

جدول ۲. معرفی دیتاست

Liar DataSet	records
train	11522
test	1266
Validation	1283
Democrats	4150
Republicans	5687
None (e.g., FB posts)	2185

۴-۲. تنظیم پارامترهای مدل

مدل پایه برت شامل یک رمزگذار با ۱۲ بلوک ترانسفورماتور، ۱۲ هدر خودآگاهی و تعداد لایه پنهان ۷۶۸ است. ورودی توالی حداکثر ۵۱۲ توکن را می‌گیرد. دنباله کلمات ورودی دارای یک یا دو بخش است که اولین توکن از دنباله همیشه [CLS] است و یک توکن ویژه دیگر [SEP] است که بین دو جمله‌ی مجاور افزوده می‌شود [۲۳].

برای پیاده‌سازی مدل پیشنهادی، در لایه‌ی کانولوشن از تعداد ۱۲۸ فیلتر استفاده شده است و ساین kernel را برابر ۲، stride را ۱ و تعداد اپیک را برابر ۵۰ در نظر گرفته‌ایم. آزمایش‌ها در این پژوهش بر روی رایانه شخصی با پردازنده Intel Core i7-6820HQ, 2.7GHz انجام شده‌است. حافظه‌ی RAM این سیستم ۸ گیگابایت و دارای پردازنده گرافیکی AMD Radeon R7 با ویندوز ۱۰ می‌باشد. برای پیاده‌سازی این مدل از کتابخانه کراس در نرم‌افزار پایتون نسخه ۳.۸.۸ استفاده شده است. معیارهای ارزیابی مورد نظر در آزمایشات ما Accuracy, Precision, Recall و F1-score می‌باشد [۲۴].

۴-۳. انتخاب مدل زبانی

با هدف بررسی کارایی مدل‌های پیش آموزش دیده شده در بازنمایی کلمات متن بیانی‌ها، روش خود را با چند نمونه مختلف مانند گلاو، کلمه به بردار و دیستیل برت آزموده‌ایم. طبق جدول زیر برای انتخاب مدل زبانی از مدل دیستیل برت^{۲۲} استفاده می‌کنیم چون نتیجه بهتری نسبت به بقیه روش‌های تبیین کلمه دارد و بقیه آزمایش‌ها با همین مدل پیاده‌سازی خواهند شد. برای استفاده از این مدل زبانی نیاز به تنظیمات اولیه می‌باشد. همانطور که در بخش ۳-۱ گفته شد ورودی این مدل باید حداکثر ۵۱۲ عدد توکن را دریافت کند پس چالش این جاست که طول همه‌ی جملات این مجموعه داده باید برابر و کمتر از ۵۱۲ باشد. ما پس از ارزیابی مطابق شکل زیر مشاهده کردیم که تعداد ۱۳۵۵ خبر در مجموعه جمله‌هایی که بیشتر از ۱۵۰ توکن داشتند را حذف کردیم و کمتر

^{۱۹}mostly-true

^{۲۰}barely-true

^{۲۱}pants-fire

^{۲۲}<https://github.com/laboroai/Laboro-DistilBERT-Japanese>

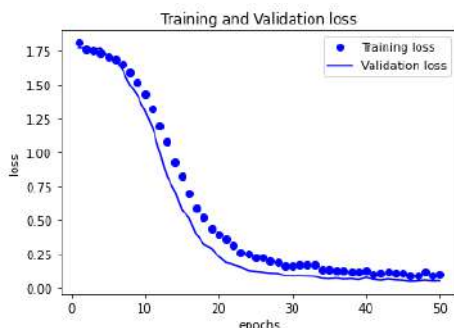
^{۱۴}Cross Entropy

^{۱۵}<https://www.politifact.com/>

^{۱۶}false

^{۱۷}true

^{۱۸}half-true

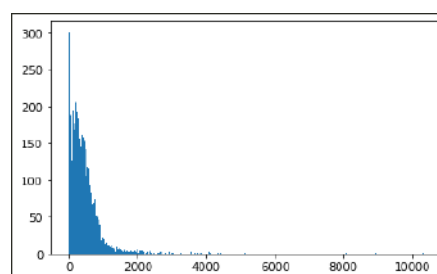


شکل ۴- نمودار خطا روش در مجموعه آموزش و اعتبارسنجی

جدول ۳- ارزیابی نتایج بر اساس مدل زبانی تعبیه‌سازی واژه‌ها

Vectorization scheme	Accuracy% (test set)
Pre-trained Glove	20.1
Pre-trained Word2Vec	43.3
Pre-trained DistilBert	94.5

از آن را با توکن <unk> پر کردیم. همچنین ما برای گرفتن بازنمایی‌ها از سرویس کولب^{۲۳} استفاده کردیم.



شکل ۲- نمودار طول جملات متن بیانیه‌ها در مجموعه داده

۱-۵. مقایسه با مقالات پایه

در این بخش با هدف ارزیابی روش پیشنهادی خودمان با استفاده از حالات و ویژگی‌های مختلف مانند فقط متن بیانیه‌ها، فقط متاداده‌ها و ترکیب هر دو همچنین مقالات پایه مورد بررسی قرار می‌دهیم.

• الگوریتم‌های یادگیری ماشین [۷، ۲۵]:

در برخی مقالات پایه از الگوریتم‌های یادگیری ماشین استفاده شده است. ما این الگوریتم‌ها را بر روی مجموعه داده خود با تعبیه کلمات فرکانس اصطلاح_فرکانس سند معکوس پیاده‌سازی کردیم و نتایج آن را در جدول ۴ بیان کردیم. بالاترین دقت مربوط به رگرسیون لجستیک با مقدار ۰,۲۳۵ می‌باشد.

• شبکه کپسولی کانولوشن [۱۱]:

برای تشخیص اخبار جعلی معماری را با استفاده از شبکه‌های کپسولی کانولوشن پیشنهاد دادند که دقت این مدل در این مجموعه داده ۴۰ درصد بروی داده‌های تست بوده است. در این مقاله [۱۱] از روش بازنمایی کلمه به بردار استفاده شده است.

• شبکه هیبرید کانولوشن [۱۰]:

در این روش هم از متن بیانیه و هم از متاداده‌ها استفاده شده است و نویسنده برای بازنمایی کلمات به ترتیب از روش کلمه به بردار و وان‌هاست استفاده کرده است. همچنین در این مقاله [۱۰] برای استخراج ویژگی از لایه‌ی شبکه‌های عصبی مبتنی بر حافظه کوتاه مدت استفاده کرده است و دقتی معادل ۲۷ درصد را بدست آورده است.

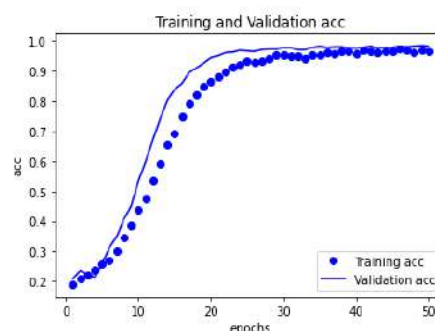
جدول ۴- ارزیابی مدل پیشنهادی با الگوریتم‌های

مختلف یادگیری ماشینی

Methods	Accuracy
Multinomial Naïve Bayes (MNB)	0.222
Support Vector Machine (SVM)	0.209
Decision Tree Classifier (DT)	0.180
K-nearest neighbors Classifier (KNN)	0.209
Logistic Regression Classifier (LR)	0.235
PassiveAggressiveClassifier (PAC)	0.183
XGBoost	0.228

۵- نتایج

شکل ۳ و ۴ به ترتیب روند دقت و خطا را روی داده‌های آموزش و اعتبارسنجی نشان می‌دهد که بیان می‌کند مدل ما به دقت خوبی در مجموعه داده اعتبارسنجی دست یافته است.



شکل ۳- نمودار دقت روش در مجموعه آموزش و اعتبارسنجی

^{۲۳} <https://colab.research.google.com/drive/1-mzaU2MXkaegTlc-wBrDdwKiHz1LwbBD>

جدول ۵- مقایسه مدل پیشنهادی با مقالات پایه و قسمت‌های مختلف دیتاست

Methods	Accuracy	Time(min)
Capsul CNN[۱۱]	0.405	45
HybridCNN[۱۰]	0.273	60
Proposed model(text)	0.764	5
Proposed model(metadata)	0.811	25
Proposed model(All)	0.945	30

مدل‌های زبانی که در مقالات پایه استفاده شده‌است مانند گلاو و کلمه به بردار، که با توجه به جدول ۳ نتایج ضعیف‌تر از برت کسب کرده‌اند. در مقاله پایه [۱۱] به علت استفاده از شبکه کپسولی چند لایه، مدل دچار پیچیدگی شده و باتوجه به جدول ۵ زمان آموزش زیاد می‌باشد در صورتیکه زمان آموزش مدل ما کمتر از این روش است. همچنین در این روش [۱۱] فقط از متاداده‌ها استفاده شده است در حالیکه ما هم از متن بیانیه‌ها و هم از متاداده‌ها به عنوان اطلاعات اضافی پروفایل گوینده استفاده کردیم و باتوجه به جدول ۵ دقت بالاتری را بدست آوردیم. در مقاله پایه دیگر [۱۰] از لایه‌ی شبکه‌های عصبی مبتنی بر حافظه کوتاه مدت استفاده شده است و چون متاداده‌ها از نظر زمانی با هم در ارتباط نیستند به دقت بالایی دست نیافته است.

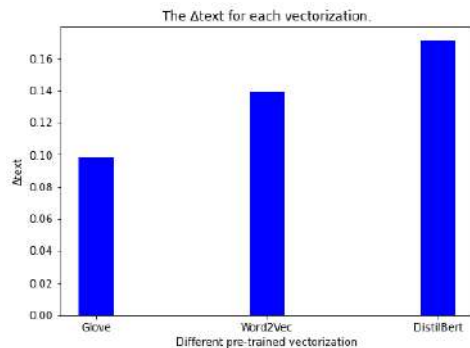
باتوجه به جدول ۵ ما در این مقاله مدل پیشنهادی را یکبار فقط با متن بیانیه، باری دیگر تنها با متادیتاها و در نهایت با هردو این‌ها ارزیابی کردیم که به ترتیب دقتی برابر با ۷۶۴٪، ۸۱۱٪ و ۹۴۵٪ را بدست آوردیم. در روش پیشنهادی ما، زمانیکه از بازنمایی متاداده‌ها و متن بیانیه‌ها باهم استفاده می‌کنیم به دقت بالاتری دست پیدا کرده‌ایم. این امر گواه بر این است که متادیتاها که شامل اطلاعات گوینده و سابقه اعتباری او می‌باشد، حاوی ویژگی‌های مهم زبانی است و می‌تواند تاثیر زیادی روی تشخیص خبر جعلی بگذارد.

۲-۵. تاثیر ویژگی‌های متن بیانیه‌ها

به منظور بررسی نقش استفاده از ویژگی‌های زبانی متن بیانیه‌ها و بهبود دقت در پیش‌بینی تشخیص اخبار جعلی آزمایشی ترتیب دادیم. ما میزان تاثیر مطلق ویژگی‌های زبانی متن بیانیه‌ها را با استفاده از معادله‌ی ۲ محاسبه می‌کنیم. در این معادله نقش ویژگی‌های متاداده به طور مطلق از میزان دقت کسب شده مبتنی بر ترکیب ویژگی‌های زبانی و متاداده حذف شده‌است. acc_{All} میزان دقت کسب شده مبتنی بر ترکیب ویژگی‌های زبانی (بازنمایی دیستیل برت) و متاداده می‌باشد. acc_{meta} میزان دقت کسب شده در آموزش مدل وقتی که تنها از ویژگی‌های متاداده‌ای استفاده کنیم، می‌باشد. Δ_{text} برابر میزان تاثیر مطلق استفاده از ویژگی‌های زبانی متن بیانیه‌هاست که ما در مدل پیشنهادی آن را به ازای هر سه روش بازنمایی مختلف گلاو، کلمه به بردار و دیستیل برت مورد بررسی قرار دادیم. مقدار رشد مطلق دقت بهبود یافته را با Δ_{text} به صورت زیر تعریف می‌کنیم:

$$\Delta_{text} = \frac{acc_{All} - acc_{meta}}{acc_{meta}} \quad (2)$$

شکل ۵ مقدار دقت بهبود یافته را در روش‌های مختلف بازنمایی بیان می‌کند، که نشان می‌دهد متن بیانیه‌ها بازنمایی شده با روش DistilBert با مقدار ۰,۱۶۵ بالاترین تاثیر را دارد و بالاترین Δ_{text} متعلق به این روش است. همچنین مقادیر acc_{All} ، acc_{meta} در جدول ۶ نشان داده شده است.



شکل ۵- بررسی بهبود دقت مدل‌های زبانی مختلف

جدول ۶- مقادیر تاثیر ویژگی‌های متن بیانیه‌ها در روش‌های مختلف زبانی

Vectorization scheme	Acc(All)	Acc(meta)	Δ_{text}
DistilBert	0.945	0.811	0.165
Word2Vec	0.201	0.382	0.139
Glove	0.433	0.183	0.098

۳-۵. تاثیر ویژگی‌های متادیتا

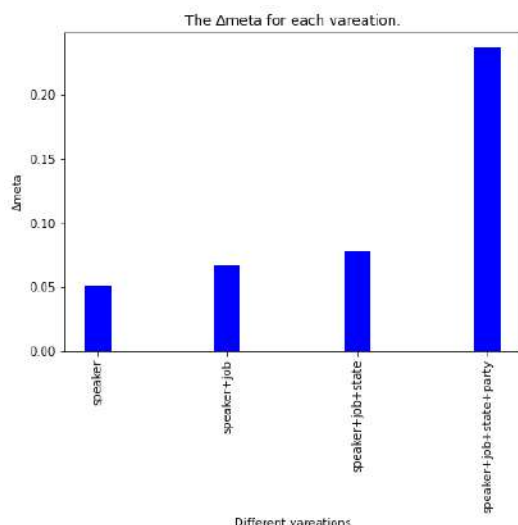
به منظور بررسی نقش استفاده از ویژگی‌های متاداده‌ای و بهبود دقت در پیش‌بینی تشخیص اخبار جعلی آزمایشی ترتیب دادیم. ما میزان تاثیر مطلق ویژگی‌های متاداده‌ای را با استفاده از معادله‌ی ۳ محاسبه می‌کنیم. در این معادله نقش ویژگی‌های زبانی متن بیانیه‌ها به طور مطلق از میزان دقت کسب شده مبتنی بر ترکیب ویژگی‌های زبانی و متاداده حذف شده است. acc_{All} میزان دقت کسب شده مبتنی بر ترکیب ویژگی‌های زبانی (بازنمایی دیستیل برت) و متاداده می‌باشد. acc_{text} میزان دقت کسب شده در آموزش مدل وقتی که تنها از ویژگی‌های زبانی متن بیانیه‌ها استفاده کنیم، می‌باشد. Δ_{meta} برابر میزان تاثیر مطلق استفاده از ویژگی‌های متاداده‌ای است که ما در مدل پیشنهادی آن را به ازای چهار حالت (۱) تنها گوینده، (۲) گوینده و شغل، (۳) گوینده و شغل و ایالت، (۴) تمام ویژگی‌ها مورد بررسی قرار داده‌ایم. مقدار رشد مطلق دقت بهبود یافته را Δ_{meta} به صورت زیر تعریف می‌کنیم:

$$\Delta_{meta} = \frac{acc_{All} - acc_{text}}{acc_{text}} \quad (3)$$

شکل ۶ تاثیر هر بخش از متادیتا را به صورت مقدار مطلق دقت بهبود یافته نشان می‌دهد به طوریکه در هر مرحله با اضافه کردن یک بخش از متادیتا، مقدار Δ_{meta}

- [2] N. Ruchansky, S. Seo, and Y. Liu, "Csi: A hybrid deep model for fake news detection," in *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management*, 2017, pp. 797-806.
- [3] S. Vosoughi, M. N. Mohsenvand, and D. Roy, "Rumor gauge: Predicting the veracity of rumors on Twitter," *ACM transactions on knowledge discovery from data (TKDD)*, vol. 11, no. 4, pp. 1-36, 2017.
- [4] A. Bondielli and F. Marcelloni, "A survey on fake news and rumour detection techniques," *Information Sciences*, vol. 497, pp. 38-55, 2019.
- [5] S.S. Jadhav and S. D. Thepade, "Fake news identification and classification using DSSM and improved recurrent neural network classifier," *Applied Artificial Intelligence*, vol. 33, no. 2, pp. 1058-1068, 2019.
- [6] P. Afshar, A. Mohammadi, and K. N. Plataniotis, "Brain tumor type classification via capsule networks," in *2018 25th IEEE International Conference on Image Processing (ICIP)*, 2018: IEEE, pp. 3129-3133.
- [7] R. Aly, S. Remus, and C. Biemann, "Hierarchical multi-label classification of text with capsule networks," in *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics: Student Research Workshop*, 2019, pp. 323-330.
- [8] S. Altalhi, M. Abulkhair, and E. Alkayal, "Capsule Network for Cyberthreat Detection."
- [9] Z. Zhou, H. Guan, M. M. Bhat, and J. Hsu, "Fake news detection via NLP is vulnerable to adversarial attacks," *arXiv preprint arXiv:1901.09657*, 2019.
- [10] W. Y. Wang, "liar, liar pants on fire": A new benchmark dataset for fake news detection," *arXiv preprint arXiv:1705.00648*, 2017.
- [11] M. H. Goldani, S. Momtazi, and R. Safabakhsh, "Detecting fake news with capsule neural networks," *Applied Soft Computing*, p. 106991, 2020.
- [12] H. Ahmed, I. Traore, and S. Saad, "Detection of online fake news using n-gram analysis and machine learning techniques," in *International conference on intelligent, secure, and dependable systems in distributed and cloud environments*, 2017: Springer, pp. 127-138.
- [13] H. W. Fentaw and T.-H. Kim, "Design and investigation of capsule networks for sentence classification," *Applied Sciences*, vol. 9, no. 11, p. 2200, 2019.
- [14] Z. Gao, A. Feng, X. Song, and X. Wu, "Target-dependent sentiment classification with BERT," *IEEE Access*, vol. 7, pp. 154290-154299, 2019.
- [15] T. Rasool, W. H. Butt, A. Shaukat, and M. U. Akram, "Multi-label fake news detection using multi-layered supervised learning," in *Proceedings of the 2019 11th International Conference on Computer and Automation Engineering*, 2019, pp. 73-77.
- [16] A. Albahr and M. Albahr, "An empirical comparison of fake news detection using different machine learning algorithms," *IJACSA*, 2020.
- [17] J. Y. Khan, M. Khondaker, T. Islam, A. Iqbal, and S. Afroz, "A benchmark study on machine learning methods for fake news detection," *arXiv preprint arXiv:1905.04749*, 2019.
- [18] C. Sun, X. Qiu, Y. Xu, and X. Huang, "How to fine-tune BERT for text classification?," in *China National Conference on Chinese Computational Linguistics*, 2019: Springer, pp. 194-206.

افزایش پیدا کرده و بالاترین میزان بهبود دقت با مقدار ۰.۲۳۶، زمانی است که از هر ۴ بخش از متادیتا استفاده کنیم.



شکل ۶- بررسی تاثیر تک تک بخش‌های متادیتا بر روی مدل

جدول ۷- مقادیر تاثیر ویژگی‌های تک تک بخش -

Different Variation	Acc(All)	Acc(text)	$\Delta meta$
Speaker	0.803	0.764	0.051
Speaker+job	0.815	0.764	0.066
Speaker+job+state	0.823	0.764	0.077
Speaker+job+state+party	0.945	0.764	0.236

۴-۵. مقایسه زمان اجرا

با توجه به مشخصات سیستم و جدول ۵ زمان آموزش مدل پیشنهادی ما خیلی کمتر از بقیه روش‌ها است و سرعت بالاتری دارد.

۶- نتیجه‌گیری

در این پژوهش ما بروی مجموعه داده دروغگو برای شناسایی خودکار اخبار جعلی کار کردیم که شامل گفته‌های کوتاه معتبر و جعلی است که در رسانه‌های اجتماعی منتشر شده‌اند. این بیانیه‌ها در زمینه‌های مختلف با سخنرانان متنوع می‌باشد. معماری پیشنهادی برای شناسایی اخبار جعلی شامل دو زیر شبکه جهت استخراج ویژگی از تعبیه کلمات متن بیانیه‌ها با استفاده از روش برت و شبکه‌ی عصبی کانولوشن و همچنین متاداده‌ها که پروفایل‌های گوینده‌اند به عنوان ورودی‌های اضافی و ارائه اطلاعات بیشتر در نظر می‌گیریم که این را می‌توان به عنوان قصد گوینده برای بیان حقیقت یا جعلی بودن آن تفسیر کرد که تا حد زیادی به مشخصات وی، به ویژه سابقه اعتباری وی بستگی دارد. در نهایت نتایج ما نشان داد که بالاترین دقت و کمترین زمان اجرا را در بین کارهای مشابه دارد.

مراجع

- [1] S. Bajaj, "The pope has a new baby! fake news detection using deep learning," *CS 224N*, pp. 1-8, 2017.

- [19] A. Hansrajh, T. T. Adeliyi, and J. Wing, "Detection of Online Fake News Using Blending Ensemble Learning," *Scientific Programming*, vol. 2021, 2021.
- [20] V. Dogra, "Banking news-events representation and classification with a novel hybrid model using DistilBERT and rule-based features," *Turkish Journal of Computer and Mathematics Education (TURCOMAT)*, vol. 12, no. 10, pp. 3039-3054, 2021.
- [21] P. Bambroo and A. Awasthi, "LegalDB: Long DistilBERT for Legal Document Classification," in *2021 International Conference on Advances in Electrical, Computing, Communication and Sustainable Technologies (ICAECT)*, 2021: IEEE, pp. 1-4.
- [22] Z. Chen, V. Badrinarayanan, C.-Y. Lee, and A. Rabinovich, "Gradnorm: Gradient normalization for adaptive loss balancing in deep multitask networks," in *International Conference on Machine Learning*, 2018: PMLR, pp. 794-803.
- [23] J.-S. Lee and J. Hsiang, "Patent classification by fine-tuning BERT language model," *World Patent Information*, vol. 61, p. 101965, 2020.
- [24] P. Ramya and B. Karthik, "A Comparative Analysis for Effective Text Document Classification Using Machine Learning Algorithms and Deep Convolution Neural Network," *Annals of the Romanian Society for Cell Biology*, pp. 16071-16086, 2021.
- [25] A. AlZubi, "Modified hierarchical method for task scheduling in grid systems," *Int J Adv Comput Sci Appl*, vol. 8, pp. 67-75, 2017.



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بوم

تنظیم کننده خودکار چندهدفه آنلاین داده‌محور برای سیستم کنترل یادگیر تکرارشونده با استفاده از الگوریتم بهینه‌یابی ازدحام ذرات

محسن جعفری^۱، ملیحه مغفوری فرسنگی^۲

^۱ گروه کنترل دانشکده مهندسی برق، دانشگاه شهید باهنر، کرمان، mohsen.jafari.mj73@gmail.com

^۲ گروه کنترل دانشکده مهندسی برق، دانشگاه شهید باهنر، کرمان، mmaghfoori@uk.ac.ir

چکیده: در این مقاله الگوریتم بهینه‌یابی ازدحام ذرات (PSO) به عنوان روشی برای تنظیم خودکار ضرایب یادگیری سیستم کنترل یادگیر تکرارشونده (ILC) به صورت آنلاین و داده‌محور معرفی شده است (PSOAT-ILC). در روش پیشنهادی، ضرایب یادگیری متغیر با تکرار هستند که پس از هر بار تکرار شدن سیستم، توسط الگوریتم پیشنهادی PSO بر اساس ارزیابی عملکرد تکرار قبلی تنظیم می‌شوند. مزیت روش پیشنهادی این است که داده‌محور بوده و قابلیت حل مسائل خطی، غیرخطی و چندمتغیره و جبران‌سازی انواع نویزها و اغتشاشات و حذف سایر محدودیت‌های روش‌های پیشین دارا است. مثال‌های شبیه‌سازی شده نشان می‌دهند که این رویکرد منجر به همگرایی یکنواخت و سریع‌تر سیستم ILC نسبت به مدل‌های پیشین با ضرایب ثابت می‌شود.

کلمات کلیدی: سیستم کنترل یادگیر تکرارشونده (Iterative Learning Control)، بهینه‌یابی ازدحام ذرات (Particle Swarm Optimization)، تنظیم کننده خودکار، بهینه‌یابی چندهدفه، کنترل داده‌محور، کنترل هوشمند، الگوریتم‌های فرا ابتکاری، الگوریتم‌های هوش جمعی

دینامیکی سیستم نمی‌کند، بلکه فقط ورودی کنترل را تغییر می‌دهد. از این رو یک سیستم داده‌محور است بنابراین نیازی به اطلاعاتی در مورد دینامیک سیستم ندارد و بار محاسباتی کمتری دارد [3].

ILC به طور گسترده در صنایعی مانند ربات‌های صنعتی، درایو هارددیسک، فرایندهای شیمیایی مانند راکتور دوره‌ای، تقطیر دوره‌ای، فرایندهای عملیات حرارتی برای محصولات فلزی، سیستم‌های چندعاملی و بسیاری دیگر از فرایندهای دوره‌ای و تکرارشونده استفاده شده است [4]. از نظر نحوه یادگیری، ILC را می‌توان به چهار طرح کلی طبقه‌بندی کرد: P-type ILC [6]، [5]، D-type ILC [8]، [7]، PD-type ILC [12]–[9] و PID-type ILC [17]–[13]. با این حال، اکثر این الگوریتم‌ها فقط برای سیستم‌های خطی کار می‌کنند. این یک محدودیت سخت است زیرا دینامیک سیستم‌های تکرارشونده می‌تواند بسیار غیرخطی باشد یا دارای اهداف ردیابی با فرکانس بالا باشد. برای حل این مشکل، بسیاری از محققین

۱- مقدمه

سیستم کنترل یادگیر تکرار شونده (Iterative Learning Control) یک سیستم کنترل هوشمند حلقه باز است که برای سیستم‌هایی طراحی شده است که یک کار ثابت را به طور مکرر در مدت زمان محدود انجام می‌دهد. در آغاز دهه ۱۹۸۰ توسط آریموتو [1] معرفی شد. هدف این کنترل کننده معمولاً ردیابی سیگنال مرجع $r(t)$ در یک بازه زمانی مشخص $[0, T]$ است. ورودی در یک سیستم ILC در هر تکرار به روز می‌شود و سپس دقت ردیابی خروجی سیستم به تدریج بهبود می‌یابد. این ماهیت تکرار شونده ویژگی کلیدی است که ILC را از کنترل کننده های بازخورد معمولی متمایز می‌کند و جزو گروه کنترل‌های هوشمند قرار می‌دهد [2]. این سیستم بر اساس روش کنترل بازگشتی است و برخلاف طرح‌های تطبیقی، ILC تلاشی برای شناسایی مدل

۲- تعریف مسئله

سیستم گسسته غیرخطی و چندمتغیره زیر را در نظر بگیرید که یک کار تعیین شده را در فاصله یک زمان محدود $[0, T]$ بارها و بارها تکرار می کند. ورودی ها و خروجی ها فرض می شود که در فواصل T_s نمونه برداری می شوند.

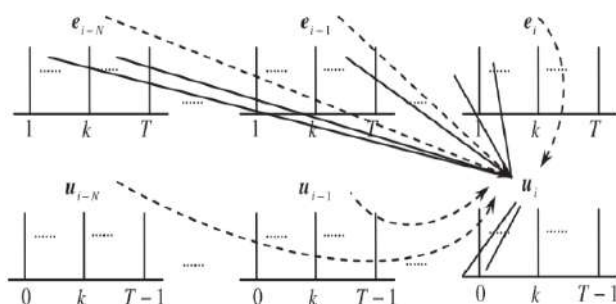
$$\begin{cases} x_k(t + Ts) = f(x_k(t), u_k(t)) \\ y_k(t) = g(x_k(t), u_k(t)) \end{cases} \quad (1)$$

که f و g توابع معین هستند، k شاخص تکرار و t شاخص زمان است. $x \in R_n$ ، $u \in R_r$ و $y \in R_q$ در رابطه (۱) به ترتیب متغیرهای حالت، ورودی و خروجی سیستم هستند. علاوه بر این، یک سیگنال مرجع $r(t)$ مشخص شده است و هدف کنترل کننده، یافتن یک سیگنال ورودی $u(t)$ است تا خروجی $y(t)$ این سیگنال مرجع را تا حد امکان با دقت دنبال کند. قانون کنترلی ILC به شرح زیر است:

$$u_{k+1} = f(u_k, u_{k-1}, \dots, u_{k-r}, e_{k+1}, e_k, \dots, e_{k-s}) \quad (2)$$

$$e_k(t) = r(t) - y_k(t), \quad \text{for } 1 \leq t \leq T \quad (3)$$

ویژگی خاص ILC این است که پس از رسیدن سیستم به نقطه زمانی $t = T$ ، حالت های سیستم به صفر برمی گردند و پس از تنظیم مجدد، قرار است سیستم دوباره همان سیگنال مرجع را دنبال کند. این ماهیت تکرار شونده باعث قابلیت یادگیری و تغییر مکرر سیگنال ورودی می شود تا با افزایش تعداد تکرارها، سیستم بتواند ردیابی کاملی را انجام دهد. روند یادگیری و استفاده از اطلاعات تکرارهای قبل در شکل (۱) قابل مشاهده است.



شکل ۱: فرآیند یادگیری از تکرارهای قبل در ILC [3]

۲-۱- الگوریتم P-type ILC

ساده ترین قانون یادگیری در ILC، یادگیری نوع تناسبی (Proportional) است که از یک بهره تناسبی برای به روزرسانی سیگنال ورودی به صورت مکرر استفاده می کند. به دلیل محاسبه ساده و کارایی آن، این نوع الگوریتم سال ها ست در جامعه ILC رایج بوده است و ساختار الگوریتم به صورت زیر است:

$$u_{k+1}(t) = u_k(t) + k_p e_k(t + T_s) \quad (4)$$

صرفاً در همه موارد از یک ضریب یادگیری کوچک استفاده می کنند، زیرا ضرایب یادگیری بزرگ باعث واگرایی می شود، همچنین یک ضریب یادگیری کوچک باعث کاهش سرعت یادگیری می شود. بنابراین، ضرایب یادگیری متفاوتی برای سیستم های تحت کنترل مختلف و شرایط کنترلی مختلف مورد نیاز است، اما از آنجایی که ما اطلاعاتی در مورد رابطه دقیق بین ضرایب یادگیری و عملکرد یادگیری نداریم، انتخاب دستی یک ضریب یادگیری مناسب تقریباً غیرممکن است. از این رو، تنظیم ضرایب یادگیری به یک مشکل عملی و چالش برانگیز تبدیل می شود.

الگوریتم های یادگیری مبتنی بر بهینه سازی را می توان به عنوان راه حل بسیار مفیدی برای مسئله فوق در نظر گرفت. در [18] یک رویکرد بهینه سازی برای ILC و در [19] یک بهینه سازی پارامتر مبتنی بر Norm برای ILC مرتبه بالا (High-order ILC) معرفی شده است. یک ILC نوع PID با ضرایب بهینه در [14] و در [20] یک ILC با فاکتور فراموشی برای مسائل خطی و در [21]، یک الگوریتم ILC با ضریب متغیر پیشنهاد شده است. در [12]، یک PD-type ILC بهینه برای سیستم های خطی با استفاده از شبکه عصبی و الگوریتم بهینه سازی ژنتیک پیشنهاد شده است.

علاوه بر این، الگوریتم بهینه یابی ازدحام ذرات (PSO) یکی از بهترین الگوریتم های اکتشافی است و برای حل مسائل بهینه یابی مورد توجه محققان زیادی قرار گرفته است. این الگوریتم از طریق شبیه سازی رفتار اجتماعی جمعی از پرندگان توسط کندی و ابرهاردت در سال ۱۹۹۵ توسعه یافت [22]. از این الگوریتم در [23]–[25] به عنوان یک تنظیم کننده مدل محور برای ILC معرفی شده است.

بالین حال، اغلب این الگوریتم ها فقط به سیستم های خطی محدود می شوند یا هنوز لازم است که مدل سیستم اصلی به طور دقیق در دسترس باشد؛ بنابراین، به دلیل عدم تطابق مدل و پیچیدگی مدل، تأثیر منفی بر پایداری و مقاومت رخ خواهد داد. به همین دلیل، یک ILC بهینه با بار محاسباتی پایین که بتواند دینامیک های غیرخطی را هم پوشش دهد و نیازی به مدل دقیق سیستم نداشته باشد، برای استفاده در دنیای واقعی و صنعت کاربردی تر است.

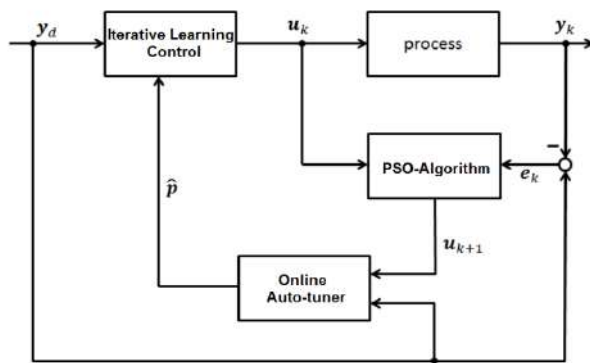
در این مقاله، یک تنظیم کننده خود کار چند هدفه داده محور مبتنی بر الگوریتم بهینه یابی ازدحام ذرات (PSO) برای سیستم کنترل یادگیر تکرار شونده (ILC) برای حل مشکل تنظیم ضرایب یادگیری به صورت آنلاین پیشنهاد شده است (PSOAT-ILC). در این طرح، ضرایب یادگیری متغیر با تکرار هستند و پس از هر تکرار با ارزیابی عملکرد تکرار قبلی تنظیم می شوند. همچنین، این یک روش مبتنی بر داده با استفاده از اندازه گیری های ورودی/خروجی است، بنابراین نیازی به یک مدل ریاضی دقیقاً شناخته شده از سیستم تحت کنترل ندارد. مثال ها می توانند نشان دهند که سرعت همگرایی این طرح سریع تر از ILC های مرسوم با ضرایب ثابت است.

ادامه مقاله به شرح زیر سازمان دهی شده است: در بخش ۲، یک تعریف مختصر از مسئله ILC مطرح شده است. در بخش ۳، توضیحاتی در مورد الگوریتم پیشنهادی PSOAT-ILC ارائه شده است. در بخش ۴، اثباتی الگوریتم پیشنهادی با مثال های عددی نشان داده شده است. در نهایت، بخش ۵ نتایج کلی و جهت گیری ها را برای تحقیقات آتی ارائه می دهد.

دلیل ماهیت غیرخطی مسئله، معمولاً بسیار دشوار یا حتی غیرممکن است. بنابراین، پیشنهاد می‌شود که از الگوریتم تکاملی و هوش جمعی PSO به عنوان یک روش محاسباتی برای یافتن ضرایب بهینه برای هر تکرار استفاده شود. زیرا PSO از قواعد جستجوی احتمالات استفاده می‌کند، نه روش‌های قطعی ریاضی و همچنین به اطلاعات مشتق یا اطلاعات کمکی دیگری نیاز ندارد، یعنی غیر خطی‌ها و اختلالات خارجی بر الگوریتم جستجو تأثیر نمی‌گذارند و فقط تابع هدف و ارزیابی شایستگی مربوطه بر جهت جستجو تأثیر می‌گذارد.

۳-۲- تنظیم کننده خودکار چندهدفه مبتنی بر الگوریتم بهینه‌یابی ازدحام ذرات (PSOAT-ILC)

در این بخش، تنظیم کننده خودکار چندهدفه آنلاین داده‌محور پیشنهادی بر اساس الگوریتم بهینه‌یابی ازدحام ذرات برای سیستم ILC ارائه خواهد شد. پیکربندی اصلی این تنظیم کننده خودکار در شکل (۲) نشان داده شده است. کل فرایند اساساً در سه مرحله انجام می‌شود:



شکل ۲: ساختار تنظیم کننده خودکار ILC

۳-۱- مرحله اول: تخمین رابطه بین ورودی و خروجی

در این مرحله ما از یک سیگنال ورودی دلخواه برای سیستم اصلی در تکرار اول استفاده می‌کنیم و سیگنال خروجی آن را اندازه‌گیری کرده سپس از این سیگنال‌ها جهت تخمین یک تابع انتقال ساده با مرتبه پایین برای سیستم توسط الگوریتم حداقل مربعات بازگشتی (Recursive Least Squares) استفاده می‌کنیم. الگوریتم استاندارد RLS استفاده شده به صورت زیر است:

$$\hat{\theta} = (\phi^T \phi)^{-1} \phi^T U \quad (8)$$

که θ و ϕ به ترتیب پارامترهای تخمین زده تابع انتقال و ماتریس کوواریانس تابع است؛ بنابراین از این رابطه تخمین زده شده در تکرارهای بعدی برای تابع هدف PSO و ارزیابی عملکرد ضرایب تولید شده استفاده می‌شود.

۳-۲- مرحله دوم: الگوریتم بهینه‌یابی ازدحام ذرات

الگوریتم بهینه‌یابی ازدحام ذرات (Particle Swarm Optimization)، رفتار یک جمعیت را به عنوان یک سیستم اجتماعی ساده شده شبیه‌سازی می‌کند. در این الگوریتم هر فرد یک "ذره" است و به عنوان نقطه‌ای در فضای

مقدار جمله $k p e_k$ یک ضریب متناسب از خطا در زمان $t+T_s$ از تکرار k م است و از این اطلاعات می‌توان برای محاسبه سیگنال ورودی جدید u_{k+1} استفاده کرد که به سیستم در طول تکرار بعدی $k+1$ اعمال می‌شود.

۲-۲- الگوریتم PID-type ILC

نوع پیشرفته‌تر قانون یادگیری در ILC، یادگیری نوع PID است که از یک بهره تناسبی، یک بهره انتگرال گیر و یک بهره مشتق گیر از خطای تکرار قبل استفاده می‌کند. الگوریتم PID-type ILC به صورت زیر است:

$$u_{k+1}(t) = u_k(t) + k_p e_k(t + T_s) + k_I \sum_{t=0}^{t+T_s} e_k(t) + k_D (e_k(t + T_s) - e_k(t)) \quad (5)$$

ضرایب k_p ، k_I و k_D به ترتیب ضرایب تناسبی، انتگرال گیر و مشتق گیر هستند.

۲-۳- همگرایی ILC

ILC همگرا است اگر و فقط اگر [12]:

$$\|e_{k+1}\| \leq \frac{1}{1 + \varepsilon} \|e_k\| \quad (6)$$

که در آن $\varepsilon > 0$ کوچک‌ترین عدد حقیقی است؛ بنابراین (۶) نشان می‌دهد که امکان دستیابی به همگرایی یکنواخت وجود دارد. در این صورت:

$$\lim_{k \rightarrow \infty} \|e_k\| \rightarrow 0 \quad (7)$$

واضح است که اگر عملگر یادگیری برابر با معکوس تابع انتقال سیستم باشد، خطا پس از یک تکرار به صفر می‌رسد، اما این اغلب یک رویکرد غیرعملی است. اولاً، دلیل استفاده از روش یادگیری این است که به دست آوردن مدل دقیق سیستم و مدل معکوس تابع انتقال به دلیل دینامیک‌های ناشناخته و غیرخطی دشوار است. علاوه بر این، معکوس دینامیک سیستم غیرخطی ممکن است دشوار باشد. عدم تطابق بین مدل ریاضی و دینامیک واقعی می‌تواند باعث عدم همگرایی شود. ثانیاً، با افزایش فرکانس، تمام سیستم‌های عملی به سرعت به صفر می‌رسند، به این معنی که ضریب یادگیری باید بهره بالایی برای فرکانس داشته باشد. این خاصیت نامطلوب به دلیل تأثیر نویز فرکانس بالا است. ثالثاً، معکوس سیستم‌های نامینیم فاز که در کاربرد عملی رایج هستند، باعث ایجاد یک یادگیری ناپایدار می‌شود. در نهایت، ضریب یادگیری معکوس کامل، سیستم را مجبور می‌کند تمام سیگنال‌های فرکانس بالا را ردیابی کند و ممکن است فشارهای زیادی بر محرک‌ها وارد کند. در واقع، پهنای باند فرکانس سیگنال مفید معمولاً محدود است و فرکانس‌های بالاتر نیازی به جبران سازی ندارند؛ بنابراین باید یک ضریب یادگیری مناسب تنظیم شود.

حال مشکل این است که ضرایب یادگیری را به نحوی تعیین کنیم تا شرط (۷) برقرار شود. متأسفانه، حل تحلیلی ضرایب بهینه برای هر تکرار به

۳-۳- مرحله سوم: تنظیم کننده خودکار چندهدفه

مبتنی بر PSO (PSOAT-ILC)

برای PSOAT-ILC، تعداد ضرایب یادگیری الگوریتم ILC به عنوان ابعاد ذرات PSO تعریف می شود، و PSO با هر بار اجرا در هر تکرار سیستم این ضرایب را به صورت بهینه تنظیم می کند و در الگوریتم ILC قرار می گیرد و ILC بهینه ترین سیگنال ورودی را برای تکرار بعدی سیستم تولید می کند. اهداف این تنظیم کننده به اختیار کاربر می تواند متفاوت باشد که در اینجا به یک نمونه جامع از آن اشاره می شود:

$$J_{k+1} = \alpha \|e_{k+1}\|^2 + \beta \|u_{k+1} - u_k\|^2 + \gamma \|u_{k+1}\|^2 \quad (14)$$

$$\forall \rightarrow 0 \leq (\alpha, \beta, \gamma) \leq 1, \quad \text{and} \quad \alpha + \beta + \gamma = 1$$

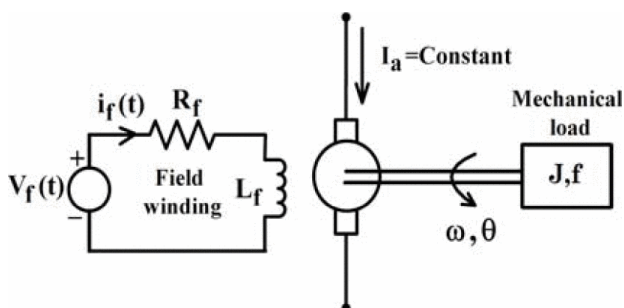
در این تابع هزینه، اولین جمله مربوط به خطای ردیابی خروجی سیستم در طول هر تکرار است، جمله دوم تضمین می کند که ورودی های کنترل بین دو تکرار زیاد تغییر نکنند و آخرین جمله انرژی ورودی کنترلی را در هر تکرار تخمین می زند. این تابع هزینه عملکرد کنترل و انرژی کنترل را متعادل می کند. رابطه (۱۴) بیان می کند که علیرغم سادگی، الگوریتم مناسبی است زیرا اندازه خطای ردیابی به طور یکنواخت در هر تکرار k کاهش می یابد و به طور همزمان ورودی (هزینه های انرژی) از اولین تکرار تا آخرین تکرار نباید به طور قابل توجهی از ورودی که در طول تکرار قبلی استفاده شده است منحرف شود و با پیشروی الگوریتم به سمت همگرایی، سرعت یادگیری کندتر می شود. به عبارت دیگر، الگوریتم منجر به همگرایی یکنواخت می شود و از این رو:

$$\|e_{k+1}\|^2 \leq J_{k+1}(u_{k+1}) \leq \|e_k\|^2 \quad (15)$$

این روش نیازی به جدا کردن حلقه کنترل برای تنظیم کردن ندارد، اما نیاز به اجرای متوالی عملیاتی تکراری دارد. برای مسائل کنترلی در واقعیت، این نوع فرایند به اصطلاح تنظیم کننده آنالاین است.

۴- مثال ها و شبیه سازی ها

در این بخش، اثربخشی تنظیم کننده خودکار پیچیده را نشان خواهیم داد. نتایج شبیه سازی عددی در این بخش ارائه شده است.



شکل ۳: موتور دی سی با آرمیچر جریان ثابت [14]

تک بعدی در نظر گرفته می شود که نشان دهنده یک راه حل بالقوه برای مشکل مربوطه است. هر ذره PSO موقعیت خود را از جمله سرعت فعلی، موقعیت و بهترین موقعیت قبلی را بر اساس تجربه خود و تجربیات همسایگان خود، تنظیم می کند. علاوه بر این، وزن اینرسی برای متعادل کردن توانایی جستجوی محلی و توانایی جستجوی جهانی اتخاذ شده است؛ بنابراین، الگوریتم PSO به صورت زیر است:

$$V_{id}(t+1) = w(t) * V_{id}(t) + c_1 * Rand * (P_{id} - X_{id}(t)) + c_2 * Rand * (P_{gd} - X_{id}(t)) \quad (9)$$

$$X_{id}(t+1) = X_{id}(t) + V_{id}(t+1) \quad (10)$$

i : از ۱ تا n که n تعداد جمعیت است.

d : ابعاد ذره است.

$V_{id}(t)$: سرعت ذره i م است.

$Rand()$: عدد تصادفی بین صفر و یک است.

X_{id} : موقعیت مکانی ذره i م است.

P_{id} : بهترین موقعیت مکانی ذره i م است.

P_{gd} : بهترین موقعیت مکانی بین تمام ذرات است.

$w(t)$: وزن اینرسی است.

c_1 و c_2 : وزن جملات شتاب تصادفی است.

پارامتر وزن اینرسی $w(t)$ برای متعادل کردن توانایی جستجوی سراسری و محلی استفاده می شود. $w(t)$ به صورت زیر تعریف می شود [24]:

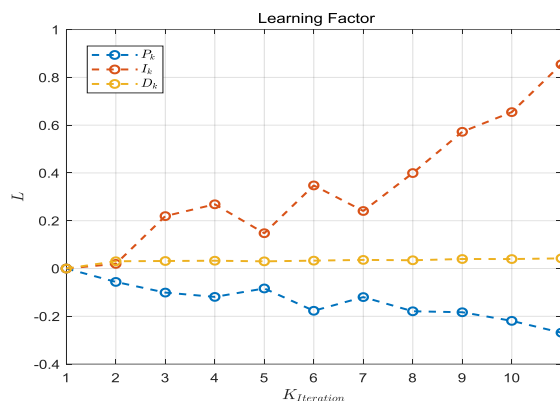
$$w(t) = w_{min} + \frac{Iter_{max} - k_{search-iteration}^{th}}{Iter_{max}} \times (w_{max}(t) - w_{min}(t)) \quad (11)$$

وزن اینرسی بزرگ جستجوی سراسری را تسهیل می کند در حالی که وزن اینرسی کوچک جستجوی محلی را تسهیل می کند. با کاهش وزن اینرسی از یک مقدار نسبتاً بزرگ به یک مقدار کوچک در طول اجرای PSO، الگوریتم تمایل دارد توانایی جستجوی جهانی بیشتری در ابتدای اجرا و توانایی جستجوی محلی بیشتر در نزدیکی پایان اجرا داشته باشد.

علاوه بر این، پیشنهاد می شود ضرایب c_1 و c_2 متغیر با زمان باشند. این ضرایب به این صورت تعریف می شوند [26]:

$$c_1 = c_{max} + \left(\frac{k_{iter}}{iter_{max}} \right) \times (c_{min} - c_{max}) \quad (12)$$

$$c_2 = c_{min} + \left(\frac{k_{iter}}{iter_{max}} \right) \times (c_{max} - c_{min}) \quad (13)$$

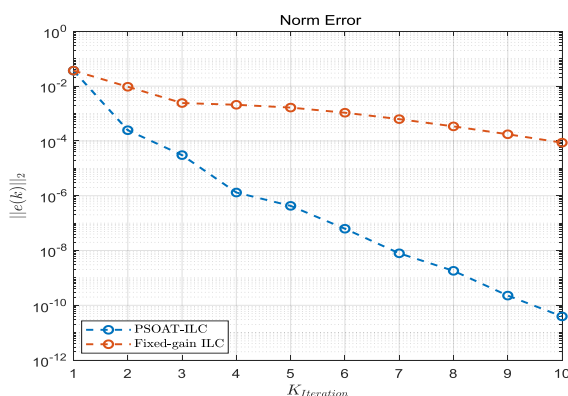


شکل ۶: نمودار تغییرات ضرایب توسط PSOAT-ILC

نتایج شبیه سازی در شکل (۴) و (۵) نشان داده شده است. همان طور که مشاهده می شود با افزایش تعداد تکرار، زاویه چرخش موتور به سرعت به مسیر خروجی مورد نظر همگرا می شود و شکل (۵) نشان می دهد که همگرایی توسط تنظیم کننده پیشنهادی بسیار سریع تر و دقیق تر از حالت قبلی است.

جدول ۲: نتایج شبیه سازی PSOAT-ILC برای ۱۰ تکرار

اندازه خطا $\ e\ $	زمان تنظیم ضرایب در هر تکرار
3.92e-11	0.06 sec



شکل ۷: نمودار مقایسه نرم خطا در مقیاس لگاریتمی

در شکل (۷) نتایج مقایسه نشان می دهد که همگرایی سیستم پیشنهادی PSOAT-ILC کاملاً یکنواخت و سریع تر و با خطای کمتری نسبت به ILC کلاسیک با ضرایب ثابت است؛ بنابراین، انتخاب بهینه ضرایب یادگیری ILC با استفاده از تنظیم کننده خودکار مبتنی بر PSO، نه تنها نرخ همگرایی را افزایش می دهد، بلکه باعث همگرایی یکنواخت می شود.

۵- نتیجه

در این مقاله، رویکرد جدیدی برای تنظیم خودکار چندهدفه آنلاین به صورت داده محور برای سیستم های کنترل یادگیر تکرار شونده مبتنی بر الگوریتم بهینه یابی ازدحام ذرات (PSOAT-ILC) به منظور بهبود عملکرد ردیابی و افزایش سرعت همگرایی پیشنهاد و توسعه داده شده است. از آنجایی که سیستم

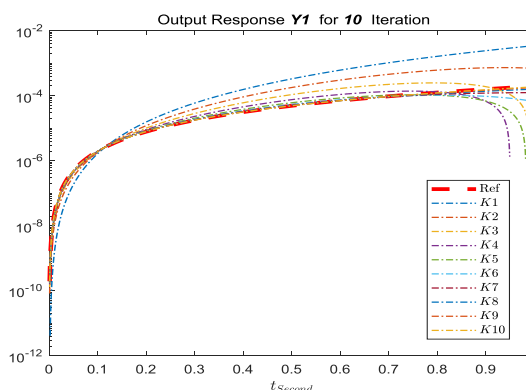
همان طور که در شکل (۳) می بینید، موتور DC را در نظر گرفتیم که آرمیچر آن توسط یک منبع جریان ثابت هدایت می شود اما جریان سیم پیچ میدان آن متغیر است به طوری که کنترل زاویه چرخش موتور با تغییر ولتاژ منبع متصل به سیم پیچ میدان انجام می شود و موتور یک بار مکانیکی را می چرخاند. در این شرایط معادلات فضای حالت موتور به صورت زیر است [14]:

$$\begin{aligned} A &= \begin{bmatrix} 0.8187 & 0 & 0 \\ 0.4526 & 0.9975 & 0 \\ 0.0023 & 0.01 & 1 \end{bmatrix}, B = \begin{bmatrix} 0 \\ 0.0197 \\ 0.0211 \end{bmatrix}, \\ C &= \begin{bmatrix} 0 & 0 & 1 \end{bmatrix} \end{aligned} \quad (16)$$

و سیگنال مرجع مطلوب به صورت زیر است:

$$y_d(t) = 10 \left(1 + \sin \left(\frac{2\pi}{t_f} t - \frac{\pi}{2} \right) \right), 0 < t \leq t_f \quad (17)$$

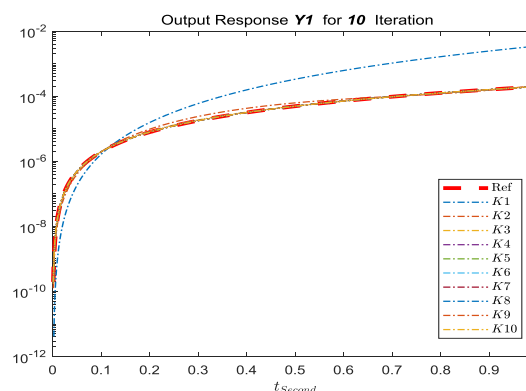
شبیه سازی در دو حالت کنترل کننده با ضرایب ثابت و کنترل کننده با تنظیم کننده خودکار انجام و مقایسه می شود.



شکل ۸: نمودار خروجی سیستم تحت ضرایب ثابت (Fixed-gain ILC)

جدول ۱: مقادیر شبیه سازی ضرایب ثابت برای ۱۰ تکرار

ضرایب	K_p	k_i	k_D	$\ e\ $
مقادیر	-0.1	0.5	0.02	8.65e-5



شکل ۹: نمودار خروجی سیستم تحت PSOAT-ILC

- pp. 470–481, 2013.
- [14] A. Madady, "PID type iterative learning control with optimal gains," *Int. J. Control. Autom. Syst.*, vol. 6, no. 2, pp. 194–203, 2008.
- [15] H. R. Reza-Alikhani, "PID type iterative learning control with optimal variable coefficients," *2010 IEEE Int. Conf. Intell. Syst. IS 2010 - Proc.*, pp. 479–484, 2010.
- [16] T. Liu, X. Z. Wang, and J. Chen, "Robust PID based indirect-type iterative learning control for batch processes with time-varying uncertainties," *J. Process Control*, vol. 24, no. 12, pp. 95–106, 2014.
- [17] Y. S. Wei and X. D. Li, "PID and EPID types of Iterative Learning Control based on Evolutionary Algorithm," *Proc. 33rd Chinese Control Conf. CCC 2014*, pp. 8889–8894, 2014.
- [18] D. H. Owens and J. Hätönen, "Iterative learning control — An optimization paradigm," *Annu. Rev. Control*, vol. 29, no. 1, pp. 57–70, Jan. 2005.
- [19] J. Hätönen, D. H. Owens, and K. Feng, "Basis functions and parameter optimisation in high-order iterative learning control," *Automatica*, vol. 42, no. 2, pp. 287–294, Feb. 2006.
- [20] X. Dai, S. Tian, W. Luo, and Y. Guo, "Iterative learning control with forgetting factor for linear distributed parameter systems with uncertainty," *J. Control Sci. Eng.*, vol. 2014, pp. 1–7, 2014.
- [21] J. Wang, C. Yu, Y. Liu, D. Shen, and Y. Q. Chen, "Variable Gain Feedback PD α -Type Iterative Learning Control for Fractional Nonlinear Systems with Time-Delay," *IEEE Access*, vol. 7, pp. 90106–90114, 2019.
- [22] J. Kennedy and R. Eberhart, "Particle swarm optimization," in *Proceedings of ICNN'95 - International Conference on Neural Networks*, 2016, vol. 4, no. 2, pp. 1942–1948.
- [23] B. Ufnalski, L. M. Grzesiak, and K. Galkowski, "Particle swarm optimization of an iterative learning controller for the single-phase inverter with sinusoidal output voltage waveform," *Bull. Polish Acad. Sci. Tech. Sci.*, vol. 61, no. 3, pp. 649–660, 2013.
- [24] Y.-C. Huang, Y.-W. Su, and P.-C. Chu, "Iterative learning control bandwidth tuning using the particle swarm optimization technique for high precision motion," *Microsyst. Technol.*, vol. 23, no. 2, pp. 361–370, Feb. 2017.
- [25] C. Peng, L. Sun, and M. Tomizuka, "Constrained Iterative Learning Control with PSO-Youla Feedback Tuning for Building Temperature Control," *IFAC-PapersOnLine*, vol. 50, no. 1, pp. 3135–3141, 2017.
- [26] Y. C. Huang and S. W. Hsu, "Particle swarm optimization on designing anticipatory iterative learning controller for positioning accuracy," in *2017 3rd International Conference on Control, Automation and Robotics, ICCAR 2017*, 2017, no. 1, pp. 483–486.

ILC نیاز به همگرایی در تکرار کمتر دارد، با بهینه‌سازی سراسری PSO، ضرایب یادگیری بهینه در هر تکرار تعیین می‌شوند. در نتیجه، روش پیشنهادی PSOAT-ILC می‌تواند به‌طور قابل توجهی تعداد تکرارها را برای ردیابی هدف موردنظر کاهش دهد. همگرایی روش ارائه شده توسط شبیه‌سازی مورد تجزیه و تحلیل قرار گرفت و نتایج اثربخشی آن به‌طور مثبت نشان داده شد. در این مقاله از هیچ‌گونه اطلاعاتی راجع به سیستم و شرایط کار آن برای کاهش اندازه فضای جستجو PSO استفاده نشده است. در نتیجه، این یک سؤال مهم برای تحقیقات آینده است که چگونه این اطلاعات را در الگوریتم PSO کدگذاری کنیم، زیرا این اطلاعات احتمال همگرایی سریع‌تر الگوریتم PSO را به یک بهینه سراسری بیشتر می‌کند.

مراجع

- [1] S. Arimoto, S. Kawamura, and F. Miyazaki, "Bettering operation of Robots by learning," *J. Robot. Syst.*, vol. 1, no. 2, pp. 123–140, 1984.
- [2] R. J. Li and Z. Z. Han, "A survey of iterative learning control," *IEEE Control Syst.*, vol. 26, no. 3, pp. 96–114, Jun. 2006.
- [3] Z. S. Hou and Z. Wang, "From model-based control to data-driven control: Survey, classification and perspective," *Inf. Sci. (Nij.)*, vol. 235, pp. 3–35, 2013.
- [4] H.-S. Ahn, Y. Chen, and K. L. Moore, "Iterative Learning Control: Brief Survey and Categorization," *IEEE Trans. Syst. Man Cybern. Part C (Applications Rev.)*, vol. 37, no. 6, pp. 1099–1121, Nov. 2007.
- [5] Y. Tan, S. P. Yang, and J. X. Xu, "On P-type iterative learning control for nonlinear systems without global Lipschitz continuity condition," *Proc. Am. Control Conf.*, vol. 2015-July, no. 0, pp. 3552–3557, 2015.
- [6] X. Dai, S. Tian, Y. Peng, and W. Luo, "Closed-loop P-type iterative learning control of uncertain linear distributed parameter systems," *IEEE/CAA J. Autom. Sin.*, vol. 1, no. 3, pp. 267–273, 2014.
- [7] D. Wang, "On D-type and P-type ILC designs and anticipatory approach," *Int. J. Control*, vol. 73, no. 10, pp. 890–901, 2000.
- [8] YangQuan Chen and K. L. Moore, "On D/sup α -type iterative learning control," in *Proceedings of the 40th IEEE Conference on Decision and Control (Cat. No. 01CH37228)*, 2003, vol. 5, no. December, pp. 4451–4456.
- [9] H. Li, Y. Chen, J. Zhang, and X. Wen, "A tuning algorithm of PD-type iterative learning control," *2010 Chinese Control Decis. Conf. CCDC 2010*, no. 4, pp. 1–6, 2010.
- [10] J. Zhang, B. Cui, Z. Jiang, and J. Chen, "A PD-Type Iterative Learning Algorithm for Semi-Linear Distributed Parameter Systems with Sensors/Actuators," *IEEE Access*, vol. 7, pp. 159037–159047, 2019.
- [11] Y. Q. Chen and K. L. Moore, "An optimal design of PD-type iterative learning control with monotonic convergence," *IEEE Int. Symp. Intell. Control - Proc.*, pp. 55–60, 2002.
- [12] Y. Zhang, A. Wang, T. Zhang, and Z. Huang, "PD iterative learning control based on neural network and genetic parameter optimization," *Proc. 2015 27th Chinese Control Decis. Conf. CCDC 2015*, pp. 5192–5195, 2015.
- [13] A. Madady, "An extended PID type iterative learning control," *Int. J. Control. Autom. Syst.*, vol. 11, no. 3,



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بوم

حل مسئله برنامه‌ریزی خطی کروی فازی

حسن میش مست نهی^۱، عبدالله هادی^۲

^۱دکتر، گروه ریاضی، آمار و کامپیوتر، دانشگاه سیستان و بلوچستان، زاهدان، hmnehi@hamoon.usb.ac.ir
^۲کارشناسی ارشد، گروه ریاضی، آمار و کامپیوتر، دانشگاه سیستان و بلوچستان، زابل، a.hadi8993@gmail.com

چکیده: نظریه مجموعه‌های فازی کروی برای کنترل عدم قطعیت مقید، عدم دقت در مسائل تصمیم‌گیری ویژگی چندگانه با در نظر گرفتن عضویت، عدم عضویت‌ها، و تعیین درجات نامعلوم است. در این پژوهش یک روش حل مسئله برنامه‌ریزی خطی فازی کروی شرح داده شده است. با توجه به گستردگی مفاهیم مجموعه‌های فازی و انواع آن و شباهت‌های آن‌ها با هم به منظور تشریح این روش ابتدا به بیان ادبیات این زمینه و تحقیقات انجام شده در زمینه مجموعه فازی و مجموعه فازی کروی پرداخته و سپس به شرح مراحل این روش پرداختیم.

کلمات کلیدی: مجموعه فازی کروی، برنامه‌ریزی خطی فازی کروی، مجموعه فازی، روش انتساب خطی کروی فازی، مجموعه‌های فازی، قطعیت و عدم قطعیت و ...

اطلاعات مربوط به وزن‌های معیار همبستگی هستند و مبتنی بر اندازه‌گیری λ -فازی است. بر اساس فواصل زمانی معین شده، چن^۳ [۸] روش انتساب خطی جدیدی را در چارچوب اعداد فازی دوزنقه‌ای نوع دوم برای تولید یک رتبه‌بندی ترجیحی مطلوب از جایگزینی انتخاب‌ها توسعه داد. در تحقیق دیگری، یا نگ^۴ و همکارانش [۳۹] یک روش جدید ویژگی چندگانه تصمیم‌گیری مبتنی بر مجموعه‌های نوتروزفیک فاصله‌ای و انتساب خطی تولید کردند. توسعه دادن روش تخصیص خطی گسترده برای حل مسائل تصمیم‌گیری چند معیاره (MCDM) تحت محیط فازی فیثاغورس هدف لیانگ و همکارانش بود [۲۳]. رویکرد جدیدی مبتنی بر روش تخصیص خطی توسط بشیری و دیگران [۶] برای انتخاب استراتژی نگهداری بهینه با استفاده از داده‌های کمی و کیفی از طریق تعامل با کارشناس تعمیر و نگهداری ارائه شد. با گسترش سنتی روش انتساب خطی، چن [۹] یک روش کارآمد برای حل مسائل MCDM در فضای شهودگرایانه دارای فاصله زمانی پیشنهاد کرد. سرانجام، لیانگ و همکارانش [۲۳] روش انتساب خطی را برای مجموعه‌های فازی فیثاغورس با ارزش بازه‌ای توسعه داد. نظریه مجموعه‌های فازی توسط زاده [۴۲] یک رویکرد مفید و مناسب سروکار داشتن با اطلاعات نادرست و مبهم در شرایط مبهم [۱۲] است. بعد از معرفی مجموعه‌های فازی توسط زاده [۴۲] تقریباً همه شاخه‌های علوم [۱۹]

۱- مقدمه

به عنوان یک روش کلاسیک بحث آنالیز خطی از MCDM روش انتساب خطی (LAM) در ابتدا توسط برناردو و بلین^۱ [۷]، با الهام از انتساب مسئله در برنامه‌ریزی خطی برای MADM پیشنهاد شد [۳۰]. ترکیبی از رتبه‌بندی معیارها به یک رتبه‌بندی کلی ترجیحی که یک سازش بهین میان چند رتبه‌بندی جزئی تولید می‌کند که ایده اصلی LAM است [۲۴]. در روش تخصیصی خطی کلاسیک معمولاً مقادیر قطعی تصمیم‌گیری از ماتریس استفاده می‌شود. با در نظر گرفتن عدم قطعیت اجتناب ناپذیر مسائل تصمیم‌گیری و زندگی واقعی، LAM تحت پرسوندهای مختلف فازی توسط محققان با استفاده از روش‌ها و رویکردهای مختلف حل و مدل‌سازی بسط داده می‌شود.

رضوی حاجی آقا و همکاران [۳۰] روشی را بر اساس روش تخصیصی خطی برای حل مسائل تصمیم‌گیری گروهی با استفاده از مجموعه‌های فازی زبانی مبهم ارائه کردند. نتایج با سایر روش‌ها مقایسه شده‌اند تا کارایی مدل را خلاصه کنند. یک رویکرد حل مسائل تصمیم‌گیری چند معیاره (MCDM) تحت محیط فازی مبهم توسط وی^۲ و همکاران [۳۶] توسعه داده شد.

³. Chen

⁴. Yang

¹. Bernardo and Blin

². Wei

بسیاری از محققان [۴۲، [۱۳]، [۳]، [۳۷]، [۴]، [۳۴]، [۳۵]، [۱۰]، [۳۸]، [۱۹] توسعه بسیاری از مجموعه‌های فازی معمولی را در ادبیات معرفی کرده است. محققان متعددی از این پسوندها در سال‌های اخیر در حل چند ویژگی استفاده کرده‌اند.

پسوندهای مجموعه‌های فازی معمولی تأثیر مهمی دارند و مورد استفاده قرار گرفته‌اند. غالباً در مورد پیشرفت حوزه تحقیق تصمیم‌گیری ویژگی‌های حل مسائل چندگانه (MADM) در یک محیط نامشخص [۱۱] MADM با ویژگی‌های گسسته و متغیرهای نسبت به یک برای ارزیابی وجود دارد [۱۶]. مسائل MADM شامل ویژگی‌های بی‌شماری است که ملموس یا نامشهود هستند [۱۷]. تصمیم‌گیری گروهی نوعی از مسائل MADM است که معمولاً انجام می‌شود به عنوان تجمیع ترجیحات گسسته مختلف در یک مجموعه داده شده از متغیرهای قابل درک یک اولویت جمعی واحد است. تصمیم‌گیری گروهی شامل چندین تصمیم‌گیرنده است، هر کدام با مهارت، تجربه و دانش مختلفی در رابطه با مسئله متفاوت هستند. ویژگی‌های این نوع مسائل با عنوان تصمیم‌گیری چندجانبه شناخته می‌شوند. (MAGDM) [۳۱].

چندین روش MADM به مجموعه‌های فازی گروهی و کاربردهای آن‌ها گسترش یافته است، در ادبیات مورد بررسی قرار می‌گیرند. محمود و همکاران [۲۷] مسئله تشخیص و تصمیم‌گیری در محیط فازی گروهی به عنوان یک عمل عملی کاربرد در پزشکی را بررسی کرد. در تحقیق دیگری کوتلو گوندوگدو و کهرمان [۲۱]، کوتلو گوندوگدو و همکاران [۲۲] روش کلاسیک (VIKOR) را به مدل فازی گروهی (SF-VIKOR) گسترش داد و قابلیت کاربرد و اعتبار آن را از طریق مسئله مدیریت پسماند نشان داد. کهرمان و همکاران [۱۶] مدل گروهی فازی TOPSIS را در مسئله انتخاب مکان بیمارستان توسعه داد و استفاده کرد. گسترش سلسله‌مراتب تحلیلی کلاسیک فرایند (AHP) به روش فازی گروهی (SF-AHP) و کاربرد آن برای انتخاب مکان انرژی تجدیدپذیر توسط کوتلو گوندوگدو و کهرمان [۲۰] پیشنهاد شده است. برنامه روش فرایند تحلیل سلسله‌مراتبی فازی گروهی در یک مسئله صنعتی انتخاب ربات توسط کوتلو گوندوگدو و کهرمان [۲۲] مورد بررسی قرار گرفته است. در تحقیقات کوتلو گوندوگدو و کهرمان [۲۱] روش تصمیم‌گیری چند معیاره TOPSIS به TOPSIS فازی گروهی تعمیم داده می‌شود. پسوندهای مر سوم روش‌های WASPAS و CODAS به فازی گروهی WASPAS و CODAS توسط کوتلو گوندوگدو و کهرمان [۱۹] توسعه یافتند. کوتلو گوندوگدو [۱۸] روش گروهی فازی MULTIMOORA را برای حل مؤثر مسائل پیچیده، که نیاز به ارزیابی و برآورد تحت یک محیط داده ناپایدار دارد؛ توسعه داد. مجموعه‌های فازی گروهی با مقدار کمی در توسعه پسوندهای TOPSIS تحت گنگ بودن توسط [۱۹] به کار گرفته شدند. آن‌ها از روش پیشنهادی در حل مسائل انتخاب چند معیاره در میان چاپگرهای سه‌بعدی استفاده کردند. لئو^۱ و همکاران [۲۶] رویکردی مبتنی بر مجموعه‌های فازی گروهی زبانی را برای ارزیابی عمومی اشتراک دوچرخه در چین ارائه دادند. کهرمان و همکاران [۱۵] یک روش سنجش عملکرد برای رتبه‌بندی شرکت‌ها با استفاده از یک تصمیم‌گیری چند شاخصه فازی نزدیک توسعه دادند. یک رویکرد جدید برای روش TOPSIS

فازی بر اساس معیارهای آنتروپی استفاده از تکنیک‌های تصمیم‌گیری مبتنی بر اطلاعات فازی گروهی برای مسائل MAGDAM توسط باروکاب^۲ و همکاران [۵] ارائه شد. در نهایت، برخی از مقیاس‌های مشابه جدید در محیط فازی گروهی توسط سیفی شیشاوان^۳ و همکاران [۳۳] ارائه شده است.

بسط دیگر مجموعه‌های فازی با تعریف تابع عضویت در یک سطح گروهی و به طور مستقل پارامترهای تابع عضویت را با دامنه بزرگ‌تری مشخص می‌کند.

مجموعه‌های گروهی و تصاویر فازی در یک گروه قرار دارند زیرا از تعریف توابع عضویت، مجموع مربع عضویت، عدم عضویت و درجه‌های مردد بودن در مجموعه‌های فازی گروهی برابر یا کمتر از ۱/۰ است. در حالی که برای مجموع درجه اول در مجموعه‌های فازی تصویر معتبر است. درجه عضویت مجموعه‌های فازی فیثاغورس با مجموعه‌های فازی گروهی دیگر از آنجا که مجموع مربع پارامترها حداکثر برابر با ۱ در هر دو پسوند است؛ در انطباق هستند [۲۵]. اشرف و همکاران [۱] [۲] مجموعه‌های فازی گروهی را با برخی از قوانین عملیاتی و عملیات تجمیع بر اساس معیار t-معیار و t-معیار مشترک پیشنهاد داده است.

۱-۱-۱- مقدمات

پیش از آن که وارد بحث اصلی پژوهش خود شویم ابتدا به منظور نزدیک شدن ذهن خواننده به موضوع به تعریف مفاهیم اولیه و انواع معروف مجموعه‌های فازی می‌پردازیم. این مجموعه‌ها عبارتند از:

۱-۱-۱- مجموعه‌های فازی معمولی

فرض کنید U جهان گفتمان باشد. یک مجموعه فازی معمولی یک شکل است که به صورت $\tilde{A} = \{ \langle u, u_{\tilde{A}}(u) \rangle | u \in U \}$ تشکیل می‌گردد که $0 \leq u_{\tilde{A}}(u) \leq 1$ در آن تابع درجه عضویت u به $\tilde{A}$ است. محدوده آن زیر مجموعه‌ای از اعداد حقیقی است که سوپریم آن، محدود است. زاده [42] مجموعه‌های فازی را به عنوان یک طبقه از شکل‌ها با پیوستگی درجات عضویت معرفی کرد. او مفاهیم شمول، اجتماع، اشتراک، متمم، وابستگی، محذب بودن، موانع زبانی و غیره را گسترش داد و ویژگی‌های مختلفی از این مفاهیم را در بافت مجموعه‌های فازی پایه‌ریزی کرد.

۱-۱-۲- مجموعه‌های فازی نوع ۲

زاده [۴۲] مفهوم مجموعه‌های فازی نوع ۲ را به عنوان بسطی از مجموعه‌های فازی معمولی معرفی کرد. این مجموعه‌ها مجموعه‌های فازی هستند که درجات عضویت آن‌ها مجموعه‌های فازی نوع ۱ هستند. آن‌ها شرایط بسیار مفیدی هستند در جایی که تعیین یک تابع عضویت دقیق برای یک مجموعه فازی دشوار است. این نوع مجموعه‌های فازی به پارامترهای زیادی نیاز دارد که در مدل‌سازی مسئله مورد استفاده قرار گیرند. بسیاری از محققین به منظور ساده‌سازی این مجموعه، بعد سوم را نادیده می‌گیرند و این مجموعه‌ها، مجموعه‌های فازی نوع دوم میان رده نامیده می‌شوند.

³. Seyfi Shishavan

¹. Lio

². Barukab

۱-۱-۳- مجموعه‌های فازی میان رده

مجموعه‌های فازی میان رده به طور مستقل توسط زاده [۴۲]، گرتن - گونیس [۱۳]، جان [۱۴] و سامباک [۳۲] معرفی شدند. یک مجموعه فازی میان رده، یک مورد خاص از مجموعه‌های فازی نوع ۲ است. یک مجموعه فازی میان رده (IVFS) با نگاشت F از جهان به مجموعه فاصله‌های بسته در [۰, ۱] تعریف می‌شود. فرض کنید این اجتماع، اشتراک، و متمم‌گیری از طریق گسترش سلسله مراتبی عمل‌های فازی - تئوریک بازه‌ها به دست می‌آیند.

۱-۱-۴- مجموعه‌های فازی کروی

مفهوم مجموعه‌های فازی کروی (SFS) دامنه اولویت بیشتری برای تصمیم‌گیرندگان برای تعیین درجه عضویت از مجموع مربع کروی پارامترها مجاز است حداکثر ۱ باشد را فراهم می‌کند. تعریف ۱: (نگاه کنید به [۲۱]). مجموعه‌های فازی کروی تک فازی $\tilde{A}_S$ جهان شامل U به صورت زیر داده می‌شود:

$$\tilde{A}_S = \{u, \mu_{\tilde{A}_S}(u), \nu_{\tilde{A}_S}(u), \pi_{\tilde{A}_S}(u) | u \in U\}$$

که $\mu_{\tilde{A}_S}(u), \nu_{\tilde{A}_S}(u), \pi_{\tilde{A}_S}(u): U \rightarrow [0,1]$ به ترتیب درجه‌های عضویت، غیرعضویت و نامعین از u برای $\tilde{A}_S$ هستند. و:

$$1 \leq \mu_{\tilde{A}_S}^2(u) + \nu_{\tilde{A}_S}^2(u) + \pi_{\tilde{A}_S}^2(u) \leq 1 \quad \forall u \in U$$

بنابراین، $\sqrt{1 - (\mu_{\tilde{A}_S}^2(u) + \nu_{\tilde{A}_S}^2(u) + \pi_{\tilde{A}_S}^2(u))}$ درجه امتناع از u از U تعریف شده است.

تعریف ۲: (نگاه کنید به [۲۱]). فرض کنید که $\tilde{A}_S$ و $\tilde{B}_S$ هر دو مجموعه‌های فازی کروی باشند. بنابراین عملگرهای پایه ی SFS می‌تواند به صورت زیر تعریف شده باشد:

جمع

$$\tilde{A}_S \oplus \tilde{B}_S = \left\{ \sqrt{\mu_{\tilde{A}_S}^2 + \mu_{\tilde{B}_S}^2 - \mu_{\tilde{A}_S}^2 \mu_{\tilde{B}_S}^2}, \nu_{\tilde{A}_S} \nu_{\tilde{B}_S}, \sqrt{(1 - \mu_{\tilde{B}_S}^2) \pi_{\tilde{A}_S}^2 + (1 - \mu_{\tilde{A}_S}^2) \pi_{\tilde{B}_S}^2 - \pi_{\tilde{A}_S}^2 \pi_{\tilde{B}_S}^2} \right\}.$$

ضرب

$$\tilde{A}_S \otimes \tilde{B}_S = \left\{ \mu_{\tilde{A}_S}^2 \mu_{\tilde{B}_S}^2, \sqrt{\nu_{\tilde{A}_S}^2 + \nu_{\tilde{B}_S}^2 - \nu_{\tilde{A}_S}^2 \nu_{\tilde{B}_S}^2}, \sqrt{(1 - \nu_{\tilde{B}_S}^2) \pi_{\tilde{A}_S}^2 + (1 - \nu_{\tilde{A}_S}^2) \pi_{\tilde{B}_S}^2 - \pi_{\tilde{A}_S}^2 \pi_{\tilde{B}_S}^2} \right\}.$$

ضرب در اسکالر؛ $k > 0$

$$k\tilde{A}_S = \left\{ \sqrt{1 - (1 - \mu_{\tilde{A}_S}^2)^k}, \nu_{\tilde{A}_S}^k, \sqrt{(1 - \mu_{\tilde{A}_S}^2)^k - (1 - \mu_{\tilde{A}_S}^2 - \pi_{\tilde{A}_S}^2)^k} \right\}.$$

توان $\tilde{A}_S$: $k > 0$

$$\tilde{A}_S^k = \left\{ \mu_{\tilde{A}_S}^k, \sqrt{1 - (1 - \nu_{\tilde{A}_S}^2)^k}, \sqrt{(1 - \nu_{\tilde{A}_S}^2)^k - (1 - \nu_{\tilde{A}_S}^2 - \pi_{\tilde{A}_S}^2)^k} \right\}$$

تعریف سوم (نگاه کنید به [۲۱]). برای این SFS های $\tilde{A}_S = (\mu_{\tilde{A}_S}, \nu_{\tilde{A}_S}, \pi_{\tilde{A}_S})$ و $\tilde{B}_S = (\mu_{\tilde{B}_S}, \nu_{\tilde{B}_S}, \pi_{\tilde{B}_S})$ موارد زیر معتبر است به شرط آن که $k_1, k_2 \geq 0$.

- I. $\tilde{A}_S \oplus \tilde{B}_S = \tilde{B}_S \oplus \tilde{A}_S$,
- II. $\tilde{A}_S \otimes \tilde{B}_S = \tilde{B}_S \otimes \tilde{A}_S$,
- III. $k(\tilde{A}_S \oplus \tilde{B}_S) = k\tilde{A}_S \oplus k\tilde{B}_S$,
- IV. $k_1\tilde{A}_S \oplus k_2\tilde{A}_S = (k_1 + k_2)\tilde{A}_S$,
- V. $(\tilde{A}_S \otimes \tilde{B}_S)^k = \tilde{A}_S^k \otimes \tilde{B}_S^k$,
- VI. $\tilde{A}_S^{k_1} \otimes \tilde{A}_S^{k_2} = \tilde{A}_S^{k_1+k_2}$.

پس از محاسبه تابع امتیاز و دقت، قوانین مقایسه به شرح زیر هستند:

If $SC(\tilde{A}_S) > SC(\tilde{B}_S)$, then $\tilde{A}_S > \tilde{B}_S$;
If $SC(\tilde{A}_S) = SC(\tilde{B}_S)$ and $AC(\tilde{A}_S) > AC(\tilde{B}_S)$, then $\tilde{A}_S > \tilde{B}_S$;
If $SC(\tilde{A}_S) = SC(\tilde{B}_S)$, $AC(\tilde{A}_S) < AC(\tilde{B}_S)$, then $\tilde{A}_S < \tilde{B}_S$;
If $SC(\tilde{A}_S) = SC(\tilde{B}_S)$, $AC(\tilde{A}_S) = AC(\tilde{B}_S)$, then $\tilde{A}_S = \tilde{B}_S$.

تعریف ۵ (نگاه کنید به [۲۱]). محاسبات میانگین وزنی حسابی فازی کروی

(SFWAM) با توجه به $\omega =$

$(\omega_1, \omega_2, \dots, \omega_m); \omega_i \in [0,1]; \sum_{i=1}^m \omega_i = 1$ و تعریف SWAM شده است به عنوان

$$SWAM_w(\tilde{A}_{S1}, \tilde{A}_{S2}, \dots, \tilde{A}_{Sn}) = w_1\tilde{A}_{S1} + w_2\tilde{A}_{S2} + \dots + w_n\tilde{A}_{Sn} \\ = \left\{ \sqrt{1 - \prod_{i=1}^n (1 - \mu_{\tilde{A}_S}^2)^{w_i}}, \prod_{i=1}^n \nu_{\tilde{A}_S}^{w_i}, \sqrt{\prod_{i=1}^n (1 - \mu_{\tilde{A}_S}^2)^{w_i} - \prod_{i=1}^n (1 - \mu_{\tilde{A}_S}^2 - \nu_{\tilde{A}_S}^2)^{w_i}} \right\}.$$

تعریف ۶ (نگاه کنید به [۲۱]). محاسبات میانگین وزنی هندسی فازی کروی (SFWGM) با توجه به $\omega =$

Alternatives	C_1	C_2	...	C_n
A_1	SFS_{11}^k	SFS_{12}^k	...	SFS_{1n}^k
A_2	SFS_{21}^k	SFS_{22}^k	...	SFS_{2n}^k
...	...	...	...	...
A_m	SFS_{m1}^k	SFS_{m2}^k	...	SFS_{mn}^k

Alternatives	C_1	C_2	...	C_n
A_1	SFS_{11}	SFS_{12}	...	SFS_{1n}
A_2	SFS_{21}	SFS_{22}	...	SFS_{2n}
...	...	...	...	...
A_m	SFS_{m1}	SFS_{m2}	...	SFS_{mn}

Alternatives	C_1	C_2	...	C_n
A_1	SC_{11}	SC_{12}	...	SC_{1n}
A_2	SC_{21}	SC_{22}	...	SC_{2n}
...	...	...	...	...
A_m	SC_{m1}	SC_{m2}	...	SC_{mn}

Alternatives	1st	2nd	...	mth
A_1	λ_{11}	λ_{12}	...	λ_{1m}
A_2	λ_{21}	λ_{22}	...	λ_{2m}
...	...	...	...	...
A_m	λ_{m1}	λ_{m2}	...	λ_{mm}

مرحله ۳. عنا صر ماتریس تصمیم گیری امتیازدهی را با استفاده از تابع امتیاز فازی کروی (معادله ۱۳) محاسبه کنید. ماتریس تصمیم گیری به دست آمده (امتیازگیری) در جدول ۴ ارائه شده است.

$$SC(\bar{A}_i) = \left(\mu_{\bar{A}_i} - \frac{I_{\bar{A}_i}}{2} \right)^2 - \left(\vartheta_{\bar{A}_i} - \frac{I_{\bar{A}_i}}{2} \right)^2, \quad (13)$$

مرحله ۴. ایجاد ماتریس تناوب رتبه غیرمنفی λ_{ik} با عناصری که فرکانس را نشان می دهند که A_m بر اساس رتبه بندی معیار m رتبه بندی شده است. با مقایسه ی ارزش هر ستون در ماتریس تصمیم گیری امتیازدهی شده (به جدول ۳ نگاه کنید) متغیرها می توانند با توجه به هر ضابطه ی $C_n \in C$ با توجه به ترتیب رو به کاهش SC_{mn} برای هر $A_m \in A$ رتبه بندی شوند. نتایج ماتریس فرکانس رتبه در جدول ۵ نشان داده شده است.

Alternatives	1st	2nd	...	mth
A_1	Π_{11}	Π_{12}	...	Π_{1m}
A_2	Π_{21}	Π_{22}	...	Π_{2m}
...	...	...	...	...
A_m	Π_{m1}	Π_{m2}	...	Π_{mm}

مرحله ۵. ماتریس Π فرکانس رتبه را محاسبه و مشخص کنید، جایی که Π_{ik} سهم A_m را در به رتبه بندی کلی اندازه گیری می کند. توجه داشته باشید هر ورودی Π_{ik} ماتریس فرکانس رتبه داده شده Π ، معیاری برای هماهنگی میان همه معیارها در رتبه بندی m جایگزین k ام است (جدول ۶). جایی که

$$\Pi_{ik} = w_{i1} + w_{i2} + \dots + w_{i\lambda_{mm}}.$$

مرحله ۶. ماتریس جایگشت P را به عنوان یک ماتریس مربعی $(m \times m)$ تعریف کنید و مدل تخصیص خطی زیر را با توجه به مقدار Π_{ik} مشخص کنید. مدل تخصیص خطی می تواند در قالب برنامه ریزی خطی زیر نوشته شود:

تعریف SWGM و $(\omega_1, \omega_2, \dots, \omega_n); \omega_i \in [0,1]; \sum_{i=1}^n \omega_i = 1$

شده است به عنوان

$$SWGM_w(\bar{A}_{S1}, \bar{A}_{S2}, \dots, \bar{A}_{Sn}) = \bar{A}_{S1}^{w_1} + \bar{A}_{S2}^{w_2} + \dots + \bar{A}_{Sn}^{w_n} \\ = \left[\prod_{i=1}^n \mu_{\bar{A}_i}^{w_i}, \sqrt{1 - \prod_{i=1}^n (1 - \vartheta_{\bar{A}_i}^2)^{w_i}}, \sqrt{\prod_{i=1}^n (1 - \vartheta_{\bar{A}_i}^2)^{w_i}} - \prod_{i=1}^n (1 - \vartheta_{\bar{A}_i}^2 - I_{\bar{A}_i}^2)^{w_i} \right].$$

۱-۲- روش تخصیص خطی کروی فازی (LAM (SF-

متطابق با ویژگی ها و ساختار روش انتساب خطی، روش انتساب خطی کلاسیک [۷] به روش انتساب خطی فازی کروی گسترش یافته است. SF-LAM از چند مرحله به صورت زیر تشکیل شده است که به قرار زیر است. جدول ۱ اصطلاحات زبانی و اعداد فازی کروی مربوط به آن ها را نشان می دهد.

Table 1
Linguistic terms of importance.

Linguistic terms	Spherical fuzzy numbers (μ, ϑ, I)
Absolutely Low Importance (ALI)	(0.1, 0.9, 0.0)
Very Low Importance (VLI)	(0.2, 0.8, 0.1)
Low Importance (LI)	(0.3, 0.7, 0.2)
Slightly Low Importance (SLI)	(0.4, 0.6, 0.3)
Equal Importance (EI)	(0.5, 0.4, 0.4)
Slightly More Importance (SMI)	(0.6, 0.4, 0.3)
More Importance (MI)	(0.7, 0.3, 0.2)
Very More Importance (VMI)	(0.8, 0.2, 0.1)
Absolutely More Importance (AMI)	(0.9, 0.1, 0.0)

مرحله ۱. تصمیمات تصمیم گیرندگان را با استفاده از جدول ۱ جمع آوری کنید. گروه تصمیم گیرندگان K را در نظر بگیرید،

$D = \{D_1, D_2, \dots, D_k\}$ که در گروه تصمیم گیرندگان مسئله شرکت می کنند، که در آن یک مجموعه متناهی از گزینه ها است، $A = \{A_1, A_2, \dots, A_m\}$ بر اساس مجموعه متناهی از معیارها مورد ارزیابی قرار می گیرند، $C = \{C_1, C_2, \dots, C_n\}$ ، با بردار وزن متناظر $w_i = \{w_1, w_2, \dots, w_n\}$ که $\sum_{i=1}^n w_i = 1, w_i \geq 0$ وزن هر ضابطه با استفاده از یک ماتریس مقایسه جفتی بر اساس اولویت تصمیم گیری محاسبه می شود. تصمیم گیرندگان در یک اصطلاح زبانی بر اساس جدول ۱ بیان می شود. هر تصمیم گیرنده k نظرات خود را در مورد عملکرد جایگزین A_m با توجه به ضابطه C_n استفاده شده در SFS_k^{mn} بیان می کند، بنابراین $SFS_k^{mn} = (\mu_{mn}^k, \vartheta_{mn}^k, I_{mn}^k)$ ؛ بنابراین ماتریس تصمیم گیری فردی در جدول ۲ به دست می آید.

مرحله ۲. ماتریس های تصمیم گیری فردی بر مبنای عملگرهای تجميع بدست می آیند. طبیعتاً، تصمیم گیرندگان قضاوت متفاوتی درباره عناصر ماتریس تصمیم گیری دارند. بنابراین، عملگرهای تجميع باید برای رسیدن به ماتریس یکپارچه استفاده شوند. از این رو، در این مرحله یک ماتریس تصمیم گیری جمعی تشکیل شده است، که در جدول ۳ نشان داده شده است.

- [3] Atanassov, K.T. (1986). *Intuitionistic fuzzy sets*. *Fuzzy Sets and Systems*, 20(1), 87–96. [https://doi.org/10.1016/S0165-0114\(86\)80034-3](https://doi.org/10.1016/S0165-0114(86)80034-3).
- [4] Atanassov, K.T. (1999). *Other extensions of intuitionistic fuzzy sets*. In: *Intuitionistic Fuzzy Sets*, pp. 179–198. https://doi.org/10.1007/978-3-7908-1870-3_3.
- [5] Barukab, O., Abdullah, S., Ashraf, S., Arif, M., Khan, S.A. (2019). A new approach to fuzzy TOPSIS method based on entropy measure under spherical fuzzy information. *Entropy*, 21(12), 1231. <https://doi.org/10.3390/e21121231>.
- [6] Bashiri, M., Badri, H., Hejazi, T.H. (2011). *Selecting optimum maintenance strategy by fuzzy interactive linear assignment method*. *Applied Mathematical Modelling*, 35(1), 152–164. <https://doi.org/10.1016/j.apm.2010.05.014>.
- [7] Bernardo, J.J., Blin, J.M. (1977). A programming model of consumer choice among multi-attributed brands. *Journal of Consumer Research*, 4(2), 111. <https://doi.org/10.1086/208686>.
- [8] Chen, T.Y. (2013). A linear assignment method for multiple-criteria decision analysis with interval type-2 fuzzy sets. *Applied Soft Computing Journal*, 13(5), 2735–2748. <https://doi.org/10.1016/j.asoc.2012.11.013>.
- [9] Chen, T.Y. (2014). The extended linear assignment method for multiple criteria decision analysis based on interval-valued intuitionistic fuzzy sets. *Applied Mathematical Modelling*, 38(7–8), 2101–2117. <https://doi.org/10.1016/j.apm.2013.10.017>.
- [10] Cu'o'ng, B.C. (2014). Picture fuzzy sets. *Journal of Computer Science and Cybernetics*, 30(4), 409–420. <https://doi.org/10.15625/1813-9663/30/4/5032>.
- [11] Donyatalab, Y., Farrokhzadeh, E., Garmroodi Se D, S., Seyfi Shishavan, S.A. (2019). Harmonic mean aggregation operators in spherical fuzzy environment and their group decision making applications. *Journal of Multiple-Valued Logic and Soft Computing*, 33(6), 565–592.
- [12] Farrokhzadeh, E., Seyfi Shishavan, S.A., Donyatalab, Y., Kutlu Gündoğdu, F., Kahraman, C. (2021). Spherical fuzzy bonferroni mean aggregation operators and their applications to multiple-attribute decision making. In: *Kutlu Gündoğdu, F., Kahraman, C. (Eds.), Decision making with Spherical Fuzzy Sets – Theory and Applications*. Springer, pp. 111–134. https://doi.org/10.1007/978-3-030-45461-6_5.
- [13] Grattan-Guinness, I. (1976). Fuzzy membership mapped onto intervals and many-valued quantities. *Zeitschrift für Mathematische Logik und Grundlagen der Mathematik*, 22(1), 149–160. <https://doi.org/10.1002/malq.19760220120>.
- [14] Jahn KU (1975) Intervall wertige Mengen. *Mathematische Nachrichten* 68(1):115–132.
- [15] Kahraman, C., Onar, S.C., Oztaysi, B. (2020). Performance measurement of debt collection firms using spherical fuzzy aggregation operators. In: *Kahraman, C., Cebi, S., Cevik Onar, S., Oztaysi, B., Tolga, A., Sari, I. (Eds.), Intelligent and Fuzzy Techniques in Big Data Analytics and Decision Making*. INFUS 2019. *Advances in Intelligent Systems and Computing*, Vol. 1029. Springer, Cham, 506–514. https://doi.org/10.1007/978-3-030-23756-1_63.
- [16] Kahraman, C., Öztaysi, B., Çevik Onar, S. (2019b). An integrated intuitionistic fuzzy AHP and TOPSIS approach to evaluation of outsource manufacturers. *Journal of Intelligent Systems*, 29(1), 283–297. <https://doi.org/10.1515/jisys-2017-0363>.

$$\begin{aligned} \text{Max} \quad & \sum_{i=1}^m \sum_{k=1}^m \Pi_{ik} \cdot P_{ik} \\ \text{s.t.} \quad & \sum_{k=1}^m P_{ik} = 1, \quad \forall i = 1, 2, \dots, m; \\ & \sum_{i=1}^m P_{ik} = 1, \quad \forall k = 1, 2, \dots, m; \\ & P_{ik} = 0 \text{ or } 1 \quad \text{for all } i \text{ and } k. \end{aligned}$$

مرحله ۷. مدل تخصیص خطی را حل کنید و ماتریس جایگزشت بهیبه P^* را برای تمامی i و k به دست آورید.

مرحله ۸. ضرب داخلی ماتریس $A = P^* \cdot A$ را محاسبه کنید و نظم $\begin{bmatrix} A_1 \\ A_2 \\ \vdots \\ A_m \end{bmatrix}$

بهینه را به دست آورید.

Table 7
Decision matrix based on comments of DM_1 .

Alternatives	C_1	C_2	C_3	C_4
A_1	(0.6, 0.4, 0.3)	(0.5, 0.4, 0.4)	(0.3, 0.7, 0.2)	(0.5, 0.4, 0.4)
A_2	(0.5, 0.4, 0.4)	(0.3, 0.7, 0.2)	(0.3, 0.7, 0.2)	(0.4, 0.6, 0.3)
A_3	(0.6, 0.4, 0.3)	(0.2, 0.8, 0.1)	(0.5, 0.4, 0.4)	(0.5, 0.4, 0.4)
A_4	(0.7, 0.3, 0.2)	(0.4, 0.6, 0.3)	(0.5, 0.4, 0.4)	(0.7, 0.3, 0.2)

Table 8
Decision matrix based on comments of DM_2 .

Alternatives	C_1	C_2	C_3	C_4
A_1	(0.5, 0.4, 0.4)	(0.7, 0.3, 0.2)	(0.3, 0.7, 0.2)	(0.6, 0.4, 0.3)
A_2	(0.4, 0.6, 0.3)	(0.5, 0.4, 0.4)	(0.3, 0.7, 0.2)	(0.5, 0.4, 0.4)
A_3	(0.6, 0.4, 0.3)	(0.3, 0.7, 0.2)	(0.6, 0.4, 0.3)	(0.5, 0.4, 0.4)
A_4	(0.5, 0.4, 0.4)	(0.4, 0.6, 0.3)	(0.4, 0.6, 0.3)	(0.8, 0.2, 0.1)

Table 9
Decision matrix based on comments of DM_3 .

Alternatives	C_1	C_2	C_3	C_4
A_1	(0.7, 0.3, 0.2)	(0.6, 0.4, 0.3)	(0.5, 0.4, 0.4)	(0.5, 0.4, 0.4)
A_2	(0.5, 0.4, 0.4)	(0.4, 0.6, 0.3)	(0.4, 0.6, 0.3)	(0.6, 0.4, 0.3)
A_3	(0.6, 0.4, 0.3)	(0.3, 0.7, 0.2)	(0.7, 0.3, 0.2)	(0.5, 0.4, 0.4)
A_4	(0.7, 0.3, 0.2)	(0.7, 0.3, 0.2)	(0.7, 0.3, 0.2)	(0.9, 0.1, 0.0)

۱- نتیجه

الگوریتم پیشنهادی دارای مزایای زیر است: اول اینکه، مقادیر ارزیابی به عنوان مجموعه‌های فازی کروی ارائه می‌شوند که می‌توانند میزان نامعلومی از نظر تصمیم گیرندگان را در مورد گزینه‌های دیگر در نظر بگیرند. دوم، از روش تخصیص خطی برای رتبه‌بندی جایگزین‌ها برای اجتناب از تأثیر ذهنیت استفاده شده‌است.

مراجع

- [1] Ashraf, S., Abdullah, S., Aslam, M., Qiyas, M., Kutbi, M.A. (2019a). Spherical fuzzy sets and its representation of spherical fuzzy t-norms and t-conorms. *Journal of Intelligent & Fuzzy Systems*, 36(6), 6089–6102. <https://doi.org/10.3233/JIFS-181941>.
- [2] Ashraf, S., Abdullah, S., Mahmood, T., Ghani, F., Mahmood, T. (2019b). Spherical fuzzy sets and their applications in multi-attribute decision making problems. *Journal of Intelligent & Fuzzy Systems*, 36(3), 2829–2844. <https://doi.org/10.3233/JIFS-172009>.

- 24(3), 1125–1148. <https://doi.org/10.3846/20294913.2016.1275878>.
- [31] Robinson, J.P., Amirtharaj, E.C.H. (2015). MAGDM problems with correlation coefficient of triangular fuzzy IFS. *International Journal of Fuzzy System Applications*, 4(1), 1–32. <https://doi.org/10.4018/IJFSA.2015010101>.
- [32] Sambuc R (1975) *Function U-Flous, Application a l'aide au Diagnostic en Pathologie Thyroïdienne*. University of Marseille.
- [33] Seyfi Shishavan, S.A., Kutlu Gündoğdu, F., Farrokhzadeh, E., Donyatalab, Y., Kahraman, C. (2020). Novel similarity measures in spherical fuzzy environment and their applications. *Engineering Applications of Artificial Intelligence*, 94, 103837. <https://doi.org/10.1016/j.engappai.2020.103837>.
- [34] Smarandache, F. (1998). *Neutrosophic Logic: Neutrosophy, Neutrosophic Set, Neutrosophic Probability*. American Research Press, Rehoboth.
- [35] Torra, V. (2010). Hesitant fuzzy sets. *International Journal of Intelligent Systems*, 25(6), 529–539. <https://doi.org/10.1002/int.20418>.
- [36] Wei, G., Alsaadi, F.E., Hayat, T., Alsaedi, A. (2017). A linear assignment method for multiple criteria decision analysis with hesitant fuzzy sets based on fuzzy measure. *International Journal of Fuzzy Systems*, 19(3), 607–614. <https://doi.org/10.1007/s40815-016-0177-x>.
- [37] Yager, R.R. (1986). On the theory of bags. *International Journal of General Systems*, 13(1), 23–37. <https://doi.org/10.1080/03081078608934952>.
- [38] Yager, R.R. (2017). Generalized orthopair fuzzy sets. *IEEE Transactions on Fuzzy Systems*, 25(5), 1222–1230. <https://doi.org/10.1109/TFUZZ.2016.2604005>.
- [39] Yang, W., Shi, J., Pang, Y., Zheng, X. (2018). Linear assignment method for interval neutrosophic sets. *Neural Computing and Applications*, 29(9), 553–564. <https://doi.org/10.1007/s00521-016-2575-2>.
- [40] Zachman, John A., "A Framework for Information Systems Architecture", IBM Systems Journal, Vol. 26, No. 3, 1987.
- [41] Zadeh, L.A. (1965). Fuzzy sets. *Information and Control*, 8(3), 338–353. [https://doi.org/10.1016/S0019-9958\(65\)90241-X](https://doi.org/10.1016/S0019-9958(65)90241-X).
- [42] Zadeh, L.A. (1975). The concept of a linguistic variable and its application to approximate reasoning—I. *Information Sciences*, 8(3), 199–249. [https://doi.org/10.1016/0020-0255\(75\)90036-5](https://doi.org/10.1016/0020-0255(75)90036-5).
- [17] Karasan, A., Ilbahar, E., Kahraman, C. (2019). A novel pythagorean fuzzy AHP and its application to landfill site selection problem. *Soft Computing*, 23(21), 10953–10968. <https://doi.org/10.1007/s00500-018-3649-0>.
- [18] Kutlu Gündoğdu, F. (2020). A spherical fuzzy extension of MULTIMOORA method. *Journal of Intelligent & Fuzzy Systems*, 38(1), 963–978. <https://doi.org/10.3233/JIFS-179462>.
- [19] Kutlu Gündoğdu, F., Kahraman, C. (2019a). A novel VIKOR method using spherical fuzzy sets and its application to warehouse site selection. *Journal of Intelligent & Fuzzy Systems*, 37(1), 1197–1211. <https://doi.org/10.3233/JIFS-182651>.
- [20] Kutlu Gündoğdu, F., Kahraman, C. (2019b). A novel spherical fuzzy analytic hierarchy process and its renewable energy application. *Soft Computing*, 24, 4607–4621. <https://doi.org/10.1007/s00500-019-04222-w>.
- [21] Kutlu Gündoğdu, F., Kahraman, C. (2019c). A novel fuzzy TOPSIS method using emerging interval-valued spherical fuzzy sets. *Engineering Applications of Artificial Intelligence*, 85, 307–323. <https://doi.org/10.1016/j.engappai.2019.06.003>.
- [22] Kutlu Gündoğdu, F., Kahraman, C. (2020). Spherical fuzzy Analytic Hierarchy Process (AHP) and its application to industrial robot selection. In: Kahraman, C., Cebi, S., Cevik Onar, S., Oztaysi, B., Tolga, A., Sari, I. (Eds.), *Intelligent and Fuzzy Techniques in Big Data Analytics and Decision Making*. INFUS 2019. *Advances in Intelligent Systems and Computing*, Vol. 1029. Springer, Cham, 988–996. https://doi.org/10.1007/978-3-030-23756-1_117.
- [23] Liang, D., Darko, A.P., Xu, Z., Quan, W. (2018). The linear assignment method for multicriteria group decision making based on interval-valued Pythagorean fuzzy Bonferroni mean. *International Journal of Intelligent Systems*, 33(11), 2101–2138. <https://doi.org/10.1002/int.22006>.
- [24] Liang, D., Darko, A.P., Xu, Z., Zhang, Y. (2019). Partitioned fuzzy measure-based linear assignment method for Pythagorean fuzzy multi-criteria decision-making with a new likelihood. *Journal of the Operational Research Society*, 71(5), 831–845. <https://doi.org/10.1080/01605682.2019.1590133>.
- [25] Liu, P., Zhu, B., Wang, P. (2019). A multi-attribute decision-making approach based on spherical fuzzy sets for Yunnan Baiyao's R&D project selection problem. *International Journal of Fuzzy Systems*, 21(7), 2168–2191.
- [26] Liu, P., Zhu, B., Wang, P., Shen, M. (2020). An approach based on linguistic spherical fuzzy sets for public evaluation of shared bicycles in China. *Engineering Applications of Artificial Intelligence*, 87, 103295. <https://doi.org/10.1016/j.engappai.2019.103295>.
- [27] Mahmood, T., Ullah, K., Khan, Q., Jan, N. (2019). An approach toward decision-making and medical diagnosis problems using the concept of spherical fuzzy sets. *Neural Computing and Applications*, 31(11), 7041–7053. <https://doi.org/10.1007/s00521-018-3521-2>.
- [28] Object Management Group. *Unified Modeling Language: Superstructure*, Version 2.0, ptc/03-07-06, July 2003, <http://www.omg.org/cgi-bin/doc?ptc/2003-08-02>.
- [29] Plamondon, R., Lorette, G., "Automatic Signature Verification and Writer Identification - The State of the Art", *Pattern Recognition*, Vol. 22, pp. 107-131, 1989.
- [30] Razavi Hajiagha, S.H., Shahbazi, M., Amoozad Mahdiraji, H., Panahian, H. (2018). A Bi-objective scorevariance based linear assignment method for group decision making with hesitant fuzzy linguistic term sets. *Technological and Economic Development of Economy*,



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

شناسایی بی‌درنگ حالات عاطفی چهره با روش یادگیری عمیق

معصومه صادق‌پور^۱، افشین ابراهیمی^۲

^۱ دانش آموخته‌ی کارشناسی ارشد، دانشکده برق، دانشگاه صنعتی سهند، تبریز، Sadeghpur1995@gmail.com

^۲ استاد، دانشکده برق، دانشگاه صنعتی سهند، تبریز، aebrahimi@sut.ac.ir

چکیده: هدف این پروژه طراحی شبکه عصبی کانولوشن بی‌درنگ برای شناسایی حالات عاطفی چهره می‌باشد. در ابتدا از روش تشخیص صورت ویولا- جونز استفاده می‌شود، سپس بخش استخراج شده وارد یک شبکه عصبی عمیق کانولوشنال می‌شود. برای این بخش ابتدا از شبکه‌ی از پیش‌آموزش داده شده‌ی **Xception** استفاده می‌کنیم، سپس با الهام از این شبکه، معماری جدیدی ارائه می‌دهیم که این شبکه‌ها از طریق مجموعه داده‌ی **FER-2013** آموزش داده شده است. خروجی این شبکه یکی از کلاس‌های هفت-گانه شناخته شده برای حالات عاطفی چهره را نشان می‌دهد. سیستم طراحی شده فوق، توانست به دقت حدود ۸۹/۸۱٪ دست یابد. این مقدار دقت در مقایسه با بهترین نتایج حاصل شده در روش‌های مشابه بسیار قابل توجه است.

کلمات کلیدی: یادگیری عمیق^۱، تشخیص چهره^۲، تشخیص حالت چهره^۳، شبکه عصبی کانولوشنال^۴

۱- مقدمه

مشابهی است که شامل شش احساس پایه خشم، تنفر، ترس، شادی، غم و تعجب می‌باشد که همگی قابل تفکیک از حالت خنثی در حالات چهره انسان است. ایشان همچنین به این نکته دست یافتند که صورت انسان می‌تواند حاوی حالات ترکیبی از احساسات باشد. برای مثال، شخص ممکن است همزمان خوشحال و متعجب باشد و یا همزمان ناراحت و عصبانی باشد. به بیان دیگر این حالات دارای همپوشانی هستند.

حالت چهره جلوه آشکاری از حالت احساسی، نیت، شخصیت و ویژگی‌های روانی یک فرد است که نقش مهمی در ارتباطات میان افراد دارد. حالات چهره و دیگر حرکات بدن در هنگام سخن گفتن، انتقال دهنده‌ی نشانه‌های ارتباط غیرکلامی در تعاملات چهره به چهره می‌باشند. همچنین ممکن است این نشانه‌ها کامل‌کننده‌ی صحبت باشند، به این طریق که به شنونده کمک کنند تا معانی واقعی کلمات گفته شده را دریابد. بر طبق تحقیقات انجام شده

مدل‌های ارائه شده درمورد احساس را می‌توان در دو گروه کلی احساسات گسسته و احساسات پیوسته دسته‌بندی کرد. منظور از حالات عاطفی چهره همان بعد گسسته‌ی احساسات است که شامل بخش‌های مجزا است که هریک در اثر مواجهه با موقعیت معینی ایجاد شده و مشخصه مشترکی ندارند. رویکردهای گسستگی احساسات بر این عقیده هستند که فضای احساسی را می‌توان به بخش‌هایی مجزا تفکیک نمود که هر بخش بیانگر احساس خاصی است.

با توجه به آزمون بازشناسی بیان چهره‌های هیجانی که توسط اکمن^۵ و فرایسن^۶ [۱] ارائه شد، حالات عاطفی چهره در سراسر جهان دارای معانی

^۱ Deep Learning

^۲ Face Detection

^۳ Facial Expression Recognition (FER)

^۴ Convolutional Neural Network

^۵ Ekman

^۶ Friesen

حالت چهره گوینده ۵۵ درصد، تن صدا ۳۸ درصد و کلمات گفته شده تنها ۷ درصد تأثیر بر شنونده را شامل می‌شود [۲]. این تحقیقات نشان می‌دهد که حالات چهره بعد اصلی را در ارتباط انسان تشکیل می‌دهند.

روش‌های پیشرفته در کارهای مرتبط با تصویر مانند طبقه‌بندی تصویر [۳] و تشخیص اشیاء، همگی بر اساس شبکه‌های عصبی کانولوشنال هستند. این وظایف به معماری‌های CNN با میلیون‌ها پارامتر نیاز دارند. بنابراین، استقرار آن‌ها در سیستم‌های ربات و سیستم‌های بی‌درنگ^۱ غیرممکن می‌شود. در این مقاله ما یک چارچوب کلی برای طراحی CNN‌های بی‌درنگ با هدف ایجاد بهترین دقت تشخیص نسبت به تعداد پارامترها پیشنهاد می‌دهیم.

۲- مدل پیشنهادی

تصویر ورودی که می‌تواند به صورت بی‌درنگ از دوربین دریافت شود یا فایلی ذخیره شده باشد، وارد بلوک تشخیص صورت می‌شود. در اینجا وجود یک یا چند صورت انسان یا عدم وجود آن تشخیص داده می‌شود. سپس این پنجره به منظور استفاده به عنوان ورودی بلوک شناسایی حالات چهره، پیش‌پردازش بر روی آن اعمال می‌شود که نهایتاً تبدیل به یک تصویر حاوی صورت می‌گردد. پس از آن تصویر حاوی صورت به عنوان ورودی، به بلوک شناسایی حالات صورت داده می‌شود که در این مرحله، پس از اعمال معماری انتخاب شده برای شناسایی حالات چهره، یکی از هفت حالت چهره (شش احساس پایه و خنثی) به عنوان خروجی شبکه استخراج می‌گردد.

۲-۱ سیستم تشخیص چهره

اولین قدم در سیستم‌های شناسایی احساسات، تشخیص چهره می‌باشد و نیازمند این است که استخراج ویژگی را به نواحی از تصویر که شامل چهره می‌باشند محدود شود. چرا که توجه سیستم بر قسمت‌های پس‌زمینه و یا بقیه قسمت‌های بدن باعث سردرگمی سیستم خواهد شد. امروزه سیستم‌های زیادی وجود دارند که از الگوریتم تشخیص چهره آشاری^۲ و یولا و جونز^۳ استفاده می‌کنند. دلایل این استفاده، سرعت بالا و قابل اعتماد بودن آن می‌باشد. و یولا و جونز در سال ۲۰۰۱ تشخیص با استفاده از طبقه‌بندی ویژگی‌های Haar که یک روش موثر در تشخیص چهره است را ارائه دادند [۵]. این یک روش مبتنی

بر یادگیری ماشین و دارای تصاویر بسیار برای آموزش و ساختار آشاری از طبقه‌بند هاست و از آن برای تشخیص اشیاء در تصاویر دیگر استفاده می‌شود. در این پروژه، به علت توانایی‌های بالقوه الگوریتم و یولا و جونز به منظور تشخیص چهره در تصاویری که از مجموعه داده‌ی FER-2013 می‌باشد، استفاده شده است.

۲-۲ سیستم شناسایی حالات چهره

پس از استخراج تصاویر حاوی صورت توسط سیستم تشخیص صورت این تصاویر به عنوان داده‌های ورودی شبکه عصبی عمیق مورد استفاده قرار می‌گیرند. در این پروژه برای این مرحله ابتدا از معماری xception برای شناسایی حالات صورت استفاده شده سپس یک الگوریتم جدید با ایده از معماری xception، طراحی شده است.

۲-۱-۱ معماری Xception

شبکه xception [۳] توسط فرانسوا چالت^۴ در گوگل برین^۵ مطرح شد که مدل توسعه یافته‌ای از معماری Inception-V3 [۹] است. این معماری در میان معماری‌های مرسوم مانند ResNet، VGGNet و AlexNet تعداد پارامترهای مورد نیاز را به شدت کاهش می‌دهد که این امر موجب افزایش سرعت آموزش، کاهش چشمگیر میزان پردازش و حتی بهبود دقت نتایج خروجی می‌شود. نوآوری چالت برای کاهش تعداد پارامترهای شبکه کانولوشنی استفاده از کانولوشن تفکیک پذیر عمقی^۶ به جای کانولوشن معمولی و استفاده از اتصالات میانبر^۷ یا بلوک باقی مانده^۸ به جای روش سنتی است. کانولوشن تفکیک پذیر عمقی از دو لایه‌ی مختلف تشکیل شده است: کانولوشن نقطه‌ای^۹ و کانولوشن عمقی^{۱۰}. هدف اصلی این لایه‌ها تفکیک همبستگی فاصله‌ها^{۱۱} از همبستگی‌های کانال^{۱۲} است [۳]. کانولوشن‌های تفکیک پذیر عمقی

^۱ Real-time

^۲ cascade

^۳ Viola and Jones

^۴ Francois Chollet

^۵ Google Brain

^۶ Depth-wise separable convolution

^۷ Skip Connections

^۸ Residual Connections

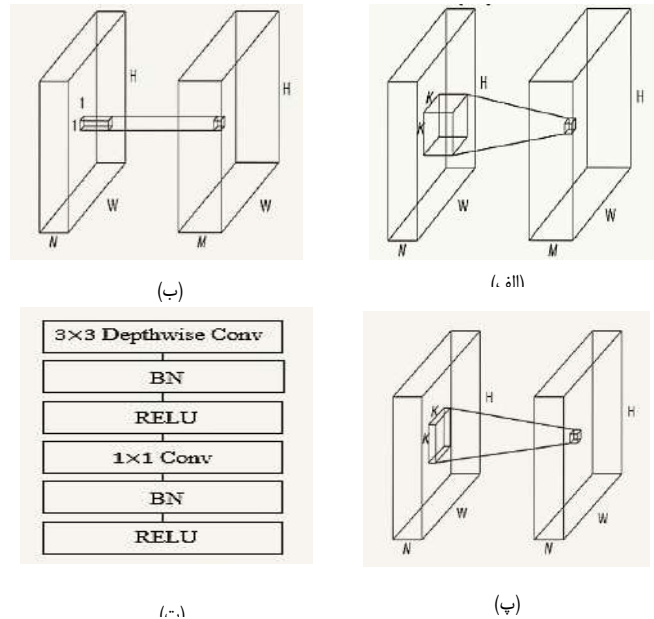
^۹ point-wise convolution

^{۱۰} Depth-wise convolution

^{۱۱} spatial cross-correlations

^{۱۲} channel cross correlations

محاسبات را نسبت به کانولوشن‌های استاندارد با ضریب $\frac{1}{M} + \frac{1}{K^2}$ کاهش می‌دهند (شکل ۱) که در حین کاهش قابل توجه پارامترها، دقت را نیز بسیار



مطلوب حفظ می‌نماید [۶].

شکل ۱: (الف) کانولوشن استاندارد (ب) کانولوشن نقطه‌ای (پ) کانولوشن عمقی (ت) کانولوشن تفکیک‌پذیر عمقی

۲-۱-۲- معماری پیشنهادی

کاهش تعداد پارامترها و استفاده از معماری‌های کانولوشنی کوچک (سبک) باعث می‌شود سیستم‌های سخت‌افزاری مانند سیستم عامل‌های ربات عملکرد سریع‌تری داشته باشند. همچنین کاهش پارامترها تعمیم بهتری را در چارچوب اصل او کام^۱ فراهم می‌کند. شکل ۲ مدل پیشنهادی (با الهام از شبکه‌ی xception) را نشان می‌دهد که با هدف ایجاد بهترین دقت نسبت به تعداد پارامترها، طراحی شده است و برای تحقق این هدف روش‌های زیر را برای طراحی شبکه در نظر گرفته ایم:

- استفاده از کانولوشن تفکیک‌پذیر عمقی [۶]
- استفاده از بلوک‌های باقی‌مانده [۷]
- حذف لایه‌های کاملاً متصل^۲ [۸]

شکل ۲: معماری پیشنهادی

در سال‌های اخیر، در حوزه شبکه‌های عصبی کانولوشنی، افزایش عمق به عنوان راهکاری برای افزایش دقت این شبکه‌ها در نظر گرفته می‌شد. اما مسئله‌ای که در خصوص افزایش عمق ما را دچار مشکل می‌سازد، کوچک شدن گرادینت‌ها و به‌روز نشدن آن‌ها طی مراحل آموزش است، که به آن محوشدگی گرادینت^۳ می‌گویند. هی و همکاران [۷] با تعریف اتصالات میانبر یا بلوک باقی‌مانده این مشکل را برطرف کردند و امکان استفاده از لایه‌های بسیار زیاد را تسهیل کردند. راهکار ارائه شده استفاده از اتصال میانبر و ساختن شبکه‌ای مبتنی بر بازنمایی باقی‌مانده می‌باشد که نمونه‌ای از بلوک این شبکه در شکل ۳ نمایش داده شده است. بلوک‌های باقی‌مانده، نقشه‌برداری مد نظر را بین دو لایه بعدی تغییر می‌دهند، به طوری که ویژگی‌های آموخته

^۱ Occam's razor

^۲ Fully Connected layer

^۳ Vanishing Gradient

آن تبدیل می‌شود سپس سیستم‌های تشخیص چهره و شناسایی حالات بر روی آن‌ها اعمال می‌شوند. شناسایی حالات چهره در هر فریم بطور متوالی انجام می‌شود و حالات مربوط به هر فریم استخراج می‌شود.



شکل ۴: نمونه ای از شناسایی حالات چهره در تصویر

۳ نتیجه

این پروژه به صورت کاملاً متن باز^۲ طراحی شده است که با استفاده از زبان برنامه نویسی پایتون و در محیط «Google Colab» پیاده سازی شده است. در این پیاده سازی از یک سیستم با پردازنده Intel Core i7 به کار رفته است. در مرحله تشخیص صورت از چهارچوب OpenCV و الگوریتم ویولا-جونز، در مرحله شناسایی حالات عاطفی چهره انسان از چهارچوب Tensorflow-Keras و برای آموزش شبکه از مجموعه داده‌ی FER-2013 استفاده گردیده که دارای ۳۵۸۸۸ تصویر رنگی با اندازه ۴۸×۴۸ است و تقسیم بندی برچسب‌های اصلی آن به ۷ کلاس (صفر: عصبانی، یک: انزجار، دو: ترس، سه: خوشحال، چهار: غمگین، پنج: شگفت‌زدگی، شش: خنثی) می‌باشد. پوشه مربوط به حالت انزجار با ۶۰۰ تصویر دارای حداقل نمونه است، در حالی که برچسب‌های دیگر تقریباً ۵۰۰۰ نمونه دارند. نمونه‌هایی از تصاویر این مجموعه داده در شکل شماره ۵ نشان داده شده است.



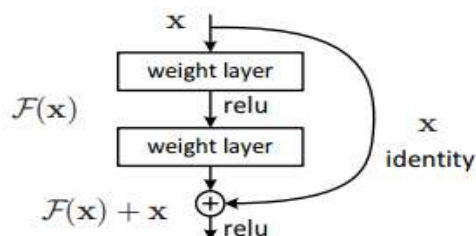
شکل ۵: نمونه تصاویر دیتاست FER-2013

به منظور ارزیابی روش پیشنهادی، عملکرد آن را با شبکه‌های از پیش آموزش داده‌ی Inception-V3 و روش پیشنهادی توسط اوکناویو و همکاران [۱۰] با استفاده از معیار ارزیابی Training Accuracy مقایسه

شده به تفاوت نقشه ویژگی^۱ اصلی و ویژگی‌های مورد نظر تبدیل می‌شود. در نتیجه، ویژگی‌های مورد نظر $H(x)$ به منظور حل مشکل یادگیری $F(x)$

به صورت زیر اصلاح می‌شوند:

$$H(x) = F(x) + x \quad (1)$$



شکل ۳: بلوک باقی‌مانده [۸]

شبکه‌های کانولوشنی معمولاً برای استخراج ویژگی‌ها شامل مجموعه‌ای از لایه‌های کاملاً متصل در انتها است، به این صورت که برای طبقه‌بندی، نقشه‌های ویژگی آخرین لایه کانولوشن به بردار تبدیل شده و وارد لایه‌های کاملاً متصلی که به دنبال آن لایه‌ی Softmax وجود دارد، می‌شوند. در اینجا لایه‌های کانولوشن به عنوان استخراج کننده‌های ویژگی در نظر گرفته و ویژگی حاصله به روش سنتی طبقه‌بندی می‌گردد.

در معماری پیشنهادی، بکارگیری استراتژی به نام Global Average Pooling (GAP) برای جایگزینی لایه‌های کاملاً متصل سنتی در CNN [۸] پیشنهاد شده است که در آن به جای افزودن لایه‌های کاملاً متصل به یکدیگر در بالای نقشه‌های ویژگی، از لایه‌ی GAP استفاده شده که با گرفتن میانگین از همه عناصر موجود در نقشه ویژگی، هر نقشه را به مقدار اسکالر کاهش می‌دهد و بردار حاصله مستقیماً در لایه‌ی Softmax تغذیه می‌شود. عملکرد میانگین، شبکه را مجبور به استخراج ویژگی‌های Global از تصویر ورودی می‌کند.

۲-۳ خروجی

پس از اینکه صورت موجود در تصویر تشخیص داده شد، احتمال مربوط به هر هفت حالت فرد را مشخص می‌کنیم. این احتمال توسط تابع Softmax که در انتهای شبکه به کار بردیم مشخص می‌شود و در نهایت حالت دارای بیشترین احتمال، به عنوان نتیجه نهایی گزارش می‌گردد. همانطور که قبلاً اشاره شد؛ یک فرد می‌تواند هم ترسیده باشد و هم ناراحت باشد پس لازم است به این نکته هم توجه شود. شبکه‌ی ارائه شده قادر است درصد هر کدام از این احساسات را نیز بیان کند. شکل ۴ نمونه‌ای از شناسایی حالت چهره بر روی تصویر را نشان می‌دهد؛ احتمال ممکن برای هر حالت مشخص می‌شود و حالتی که بیشترین احتمال را دارد به عنوان حالت (احساس) نهایی مطرح می‌شود. در این پروژه علاوه بر شناسایی احساسات در تصویر، شناسایی احساسات روی فیلم هم انجام شده است، به این صورت که ابتدا فیلم به فریم‌های تشکیل دهنده‌ی

^۱ Feature map

^۲ Open Source

- [5] Viola, P., & Jones, M. (2001, December). Rapid object detection using a boosted cascade of simple features. In *Proceedings of the 2001 IEEE computer society conference on computer vision and pattern recognition. CVPR 2001* (Vol. 1, pp. I-I). Ieee.
- [6] Howard, A. G., Zhu, M., Chen, B., Kalenichenko, D., Wang, W., Weyand, T., ... & Adam, H. (2017). Mobilenets: Efficient convolutional neural networks for mobile vision applications. *arXiv preprint arXiv:1704.04861*.
- [7] He, K., Zhang, X., Ren, S., & Sun, J. (2016). Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 770-778).
- [8] Lin, M., Chen, Q., & Yan, S. (2013). Network in network. *arXiv preprint arXiv:1312.4400*.
- [9] Szegedy, C., Vanhoucke, V., Ioffe, S., Shlens, J., & Wojna, Z. (2016). Rethinking the inception architecture for computer vision. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 2818-2826).
- [10] Arriaga, O., Valdenegro-Toro, M., & Plöger, P. (2017). Real-time convolutional neural networks for emotion and gender classification. *arXiv preprint arXiv:1710.07557*.

می‌کنیم. شرایط آموزش از جمله مجموعه داده، تعداد تکرارهای آموزشی و اندازه دسته و ... برای همه یکسان بوده‌اند تا مقایسه‌ی عادلانه‌ای را داشته باشیم.

شبکه‌ی مورد استفاده	Accuracy(%)
الگوریتم پیشنهادی توسط اوکتاویو و همکاران	68/8
Inception-V3	71/99
Xception	76/9
الگوریتم پیشنهادی	81/89

همان‌طور که مشاهده می‌شود الگوریتم پیشنهادی در مقایسه با سایر روش‌ها به صحت قابل ملاحظه‌ای دست یافته است، همچنین مدل پیشنهادی تقریباً ۹۰۰۰ پارامتر دارد که این برابر است با کاهش چشمگیری در مقایسه با تعداد پارامترهای شبکه‌ی Inception-V3 و Xception. همان‌طور که قبلاً گفته شد هدف از این پروژه طراحی شبکه‌ای با بهترین دقت نسبت به تعداد پارامترها بود؛ معماری پیشنهادی که می‌توانیم از آن به نام mini-Xception یاد کنیم هم کاهش چشمگیری در تعداد پارامترها داشت هم دقت قابل قبولی را کسب کرد.

۴ پیشنهاد

یکی از چالش‌های مهم در شناسایی حالت چهره این است که الگوریتم ارائه شده بتواند در صورت استفاده از قسمت‌های خاصی از چهره به جای تمام چهره به درستی عمل کند. از مزایای این امر می‌توان به کارایی الگوریتم در شرایط خاصی همچون وجود انسداد در نواحی مختلف چهره اشاره کرد. همچنین با انتخاب برخی از نواحی مهم چهره به جای کل آن، می‌توان هزینه محاسباتی را نیز کاهش داد. حذف برخی از نواحی چهره که در شناسایی حالت چهره تأثیر زیادی ندارند، می‌تواند باعث افزایش دقت الگوریتم نیز بشود. بنابراین پیشنهادها برای ادامه‌ی این پروژه در ابتدا، استفاده از قسمت‌های خاصی از چهره به جای تمام چهره، جهت تشخیص حالات چهره است و پس از آن اضافه کردن امکان شناسایی جنسیت و سن می‌باشد.

مراجع

- [1] Ekman, P., & Friesen, W. V. (1971). Constants across cultures in the face and emotion. *Journal of personality and social psychology*, 17(2), 124.
- [2] Khademi, M., Kiapour, M. H., Manzuri-Shalmani, M. T., & Kiaei, A. A. (2010, April). Analysis, interpretation, and recognition of facial action units and expressions using neuro-fuzzy modeling. In *IAPR Workshop on Artificial Neural Networks in Pattern Recognition* (pp. 161-172). Springer, Berlin, Heidelberg.
- [3] Chollet, F. (2017). Xception: Deep learning with depthwise separable convolutions. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 1251-1258).
- [4] Viola, Paul, and Michael J. Jones. "Robust real-time face detection." *International journal of computer vision* 57.2 (2004): 137-154.

نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بزم

شناسایی بیماری COVID-19 با استفاده از شبکه عصبی پیچشی

سید محمد موسوی^۱، سوده حسینی^۲

^۱ دانشجوی کارشناسی ارشد، دانشکده ریاضی و کامپیوتر، بخش علوم کامپیوتر، دانشگاه شهید باهنر کرمان، m.mousavi@math.uk.ac.ir

^۲ دانشیار، دانشکده ریاضی و کامپیوتر، بخش علوم کامپیوتر، دانشگاه شهید باهنر کرمان، So_hosseini@uk.ac.ir

چکیده: تصاویر اشعه ایکس قفسه سینه برای نظارت بر بیماری‌های مختلف ریه مفید بوده و اخیراً برای نظارت بر بیماری COVID-19 که باعث عفونت در دستگاه تنفسی فوقانی و ریه‌ها می‌شود استفاده شده است. در این مقاله به منظور خودکارسازی روند تشخیص بیماری COVID-19 با استفاده از تصاویر اشعه ایکس یک شبکه عصبی پیچشی سبک پیشنهاد شده است که شامل ۹ لایه و ۳۷۷،۷۳۱ پارامتر بوده که تمام آنها قابل آموزش است. نوآوری این پژوهش استفاده از یک شبکه عصبی پیچشی سبک می‌باشد که علاوه بر کاهش تعداد لایه‌ها و پارامترها، شبکه همچنان از دقت بالایی برای تشخیص موارد مبتلا برخوردار است. در مرحله اول تصاویر به سائز ۲۵۶*۲۵۶ تغییر سائز داده شده‌اند و به منظور یادگیری بهتر شبکه از روش تقویت داده استفاده شده است. در مرحله دوم با استفاده از ۳۶۵۹ تصویر به منظور آموزش و ۵۲۲ تصویر برای ارزیابی در ۲۰ مرحله و با استفاده از روش توقف زودهنگام براساس کمترین میزان خطای اعتبارسنجی شبکه آموزش داده شده است. در مرحله آخر پس از بررسی بر روی ۱۰۴۷ تصویر مجموعه آزمایش دقت ۰.۹۶۲ بدست آمده است. نتایج نشان می‌دهد شبکه‌های عصبی پیچشی به منظور شناسایی موارد مبتلا به بیماری COVID-19 با استفاده از تصاویر رادیولوژی منجر به دقت بالایی می‌شود.

کلمات کلیدی: طبقه‌بندی تصاویر پزشکی، شبکه‌های عصبی پیچشی، COVID-19

یک گام مهم در کنترل COVID-19 غربالگری موثر و دقیق بیماران است تا موارد مثبت درمان به موقع دریافت کرده و به طور مناسب از مردم جدا شوند. اقدامی که برای مهار گسترش عفونت بسیار مهم تلقی می‌شود. آزمایش^۲ RT-PCR، روش غربالگری خوبی برای تشخیص موارد COVID-19 است. این آزمایش به مواد خاصی نیاز دارد که در همه جا در دسترس نیست. به علاوه این آزمون زمان‌بر است و حساسیت نسبتاً پایینی دارد (نرخ مثبت واقعی). بنابراین، نیاز فزاینده‌ای به استفاده از تکنیک‌های غربالگری سریع‌تر و قابل اعتمادتر وجود دارد که می‌تواند با آزمایش RT-PCR تایید نهایی شود.

برخی مطالعات استفاده از تکنیک‌های تصویربرداری مانند تصاویر رادیولوژی قفسه سینه و تصاویر CT scan را برای یافتن شاخص‌های بصری مرتبط با عفونت ویروسی SARS-CoV-2 پیشنهاد کرده‌اند [۱-۳]. از طرفی دوز تابش یکی از نگرانی‌های مهم است که برای CT scan در مقایسه با تصویربرداری اشعه ایکس بسیار بالاتر است. این امر به ویژه برای زنان باردار و کودکان بسیار مهم است. امکانات تصویربرداری از قفسه سینه در سیستم‌های

۱- مقدمه

سندرم حاد تنفسی ویروس کرونا-۲ (SARS-CoV-2)، ویروس عامل COVID-19، که یک بیماری ویروسی و تنفسی است که ابتدا در اواخر سال ۲۰۱۹ در ووهان چین گزارش شد و کمی بعد در ۱۱ مارس ۲۰۲۰ توسط سازمان بهداشت جهانی پاندمی اعلام شد. از آنجا که COVID-19 یک بیماری واگیردار و به راحتی قابل انتقال است زندگی میلیاردها نفر در سراسر جهان را تحت تاثیر قرار داده است. از خطرات ناشی از آن می‌توان به شدت و میزان مرگ و میر نسبتاً بالا و فشار ناشی از آن بر سیستم درمانی کشورهای جهان اشاره کرد. به منظور کاهش ریسک ابتلا افراد جدید و حفظ امنیت افراد جامعه مواردی مانند فاصله گذاری اجتماعی و استفاده از ماسک در میان مردم سراسر جهان رواج یافت. همچنین تشخیص زودهنگام و دقیق COVID-19 برای کنترل شیوع بیماری و کاهش مرگ و میر آن از اهمیت زیادی برخوردار است.

^۲ Reverse-transcription polymerase chain reaction

^۱ Severe acute respiratory syndrome coronavirus 2

نوبین مراقبت های بهداشتی به آسانی موجود است و معاینه رادیوگرافی را مکمل خوبی برای آزمایش RT-PCR و در برخی موارد حتا شاخص حساسیت بالاتری را نشان می دهد. باتوجه به ظرافت شاخص های بصری اشعه ایکس متخصصان رادیولوژی در تشخیص این تغییرات ظریف و تفسیر دقیق آن ها با یک چالش بزرگ مواجه هستند. به همین دلیل، داشتن سیستم های تشخیصی محاسبه شده که بتواند به آن ها در تفسیر دقیق تر عکس های اشعه ایکس که از ویژگی های COVID-19 کمک کند، بسیار مطلوب و ضروری است.

هدف اصلی این مقاله پیشنهاد یک سیستم پشتیبانی تصمیم پزشکی با استفاده از یک مدل شبکه عصبی پیچشی سبک در تشخیص COVID-19 بر اساس تصاویر اشعه ایکس قفسه سینه است. مدل پیشنهادی در مقایسه با دیگر رویکردهای شبکه عصبی پیچشی کاهش چشم گیر تعداد پارامترهای شبکه باتوجه به دقت بالا را ارائه کرد. همچنین می توان به کاهش زمان و منابع مورد نیاز برای بکارگیری در مراکز درمانی به صورت بلادرنگ اشاره کرد. مدل پیشنهادی شامل دو بخش یادگیری ویژگی و طبقه بندی می باشد. در بخش اول پس از پیش پردازش که شامل تغییر سائز تصاویر مجموعه داده و استفاده از عملیات های هندسی متفاوتی به منظور تقویت داده شبکه آموزش دیده و سپس در مرحله دوم دقت مدل مورد ارزیابی قرار می گیرد.

مجموعه داده استفاده شده از ۵۲۲۸ تصویر اشعه ایکس قفسه سینه جمع آوری شده از چند مجموعه داده دیگر شامل ۱۶۲۶ عکس بیمار (کرونا) و ۱۸۰۲ عکس سالم و ۱۸۰۰ عکس بیمار (غیر کرونا) تشکیل شده است.

سایر بخش های مقاله به ترتیب زیر می باشد: در بخش بعدی برخی از کارهای مرتبط اخیر مورد بررسی قرار گرفته است. بخش ۳ بر مجموعه داده های مورد استفاده تمرکز دارد. بخش ۴ به مدل پیشنهادی می پردازد، نتایج بدست آمده در بخش ۵ مورد بررسی قرار می گیرد و در نهایت بخش ۶ نتیجه گیری را نشان می دهد.

۲- کارهای مرتبط

تکور و کومار [۲] از یک شبکه عصبی پیچشی پیشنهاد شده برای طبقه بندی تصاویر CXR^۲ و CT^۳ استفاده کردند. آنها از دو سناریو طبقه بندی دو کلاسه و چند کلاسه استفاده کردند. به این منظور از ۱۱۰۹۵ تصویر استفاده شد. در نهایت بهترین دقت در طبقه بندی دو کلاسه ۹۹٫۶٪ و در طبقه بندی چند کلاسه ۹۸٫۲٪ گزارش شد. در الپتاجی و سلام [۳] از تصاویر اشعه ایکس و CT scan برای پیش بینی موارد مبتلا به COVID-19 بر اساس مدل های DensNet, AlexNet, GoogleNet, ResNet, InceptionV3, VGG^۵،

بالاترین دقت بدست آمده با استفاده از تصاویر اشعه ایکس ۹۷٫۷٪ و با استفاده از تصاویر CT ۹۷٫۱٪ گزارش کردند. در لورنسن و همکاران [۴] از تصاویر اشعه ایکس قفسه سینه برای طبقه بندی افراد آلوده به COVID-19 به وسیله یادگیری عمیق و با استفاده از یادگیری انتقالی^۶ استفاده کرده اند. این مقاله به بررسی پنج مدل مبتنی بر شبکه های عصبی پیچشی از پیش آموزش دیده (AlexNet, VGG16, ResNet50, ResNet101, ResNet152) بر روی ۱۸۵ تصویر اشعه ایکس شامل چهار کلاس پرداخت. باتوجه به تعداد کم تصاویر در مجموعه داده، از فرایند تقویت اطلاعات^۷ از جمله چرخش ۹۰، ۱۸۰ و ۲۷۰ درجه در محورهای مختلف استفاده شد. بهترین نتایج در صورتی حاصل می شود که معماری از قبل آموزش دیده شده ResNet152 با استفاده از دسته های بزرگ تر داده در تعداد متوسط دوره های آموزش با استفاده از تابع بهینه ساز Nadam آموزش داده شود. در حسن تبار و همکاران [۷] برای تشخیص بیماران COVID-19 از دو الگوریتم شبکه عصبی عمیق به همراه مدل مقطعی^۸ برای استخراج ویژگی و یک روش شبکه عصبی پیچشی استفاده شد. در نتیجه معماری شبکه عصبی پیچشی ارائه شده با دقت ۹۳٫۲٪ و حساسیت ۹۶٫۱٪ بهتر از روش شبکه عصبی عمیق با دقت ۸۳٫۴٪ و حساسیت ۸۶٪ عملکرد کرده است. در اسماعیلی و سنگور [۵] از رویکردهای مبتنی بر یادگیری عمیق، یعنی استخراج ویژگی های عمیق^۹ و تنظیم دقیق^{۱۰} شبکه های عصبی پیچشی از پیش آموزش دیده و آموزش انتها به انتها^{۱۱} یک مدل توسعه یافته شبکه عصبی پیچشی متشکل از ۲۱ لایه شامل لایه های کانولوشن، ادغام حداکثر و لایه های کاملاً متصل و لایه طبقه بندی نهایی به همراه نرمال سازی دسته ای و لایه های ReLU به منظور طبقه بندی تصاویر اشعه ایکس قفسه سینه COVID-19 و سالم استفاده کردند. در این کار از مجموعه داده حاوی ۱۸۰ تصویر COVID-19 و ۲۰۰ تصویر سالم استفاده شد و به منظور استخراج ویژگی های عمیق از مدل های شبکه عصبی پیچشی از پیش آموزش دیده (VGG16, VGG19, ResNet101, ResNet50, ResNet18) بکار رفته است. در نهایت ویژگی های عمیق استخراج شده از مدل ResNet50 به کمک SVM دقت ۹۴٫۷٪ را بدست آوردند که بالاترین امتیاز در بین تمام نتایج بدست آمده بود. مارکوس و همکاران [۶] روشی برای تشخیص COVID-19 با استفاده از شبکه عصبی پیچشی عمیق بر اساس مدل EfficientNetB4 به منظور طبقه بندی تصاویر اشعه ایکس قفسه سینه ساختند. بهترین دقت در طبقه بندی دو کلاسه با دقت ۹۹٫۵۱٪ گزارش شد.

در ادامه در جدول ۱ کارهای انجام شده در زمینه تشخیص موارد COVID-19 با استفاده از شبکه های عصبی پیچشی آورده شده است.

^۵ Visual Geometry Group

^۶ Transfer learning

^۷ Data Augmentation

^۸ Fractal

^۹ Deep feature extraction

^{۱۰} Fine-tuning

^{۱۱} End-to-end

^۲ Chest X-Ray

^۳ Computed Tomography

COVID-19 جدول ۱: کارهای مرتبط در زمینه شناسایی بیماری

ردیف	مرجع	رویکرد	مجموعه داده	نوع داده	بهترین نتایج
۱	دب و همکاران (۲۰۲۲) [۱]	VGGNet, GoogleNet, DensNet, NASNet, ResNet, ResNext, Ensemble model	مجموعه داده اول: J.P. Cohen [13] مجموعه داده دوم: chest-xray-covid19- pneumonia dataset مجموعه داده سوم: private dataset from MGM medical college and hospital	CXR	کلاس بندی دو کلاسه: ۹۵٫۶٪ کلاس بندی چند کلاسه: ۹۳٫۴٪
۲	ثکوری و کومار (۲۰۲۱) [۲]	شبکه عصبی پیچشی	مجموع ۱۱۰۹۵ عکس اشعه ایکس و CT scan	CXR and CT scan	کلاس بندی دودویی با ۹۹٫۶٪ دقت و کلاس بندی چند کلاسه ۹۸٫۲٪
۳	الپتاجی و سلام (۲۰۲۱) [۳]	ResNet50, ResNet53, GoogleNet, AlexNet, DensNet201, VGG16, VGG19, Inception V3, ResNet50+SVM ، شبکه عصبی پیچشی	مجموعه داده جمع آوری شده [۹]	CXR and CT scan	دقت ۹۷٫۷٪ با استفاده از CXR و دقت ۹۷٫۱٪ با استفاده از CT
۴	لورنسن و همکاران (۲۰۲۱) [۴]	AlexNet, VGG16, ResNet50, ResNet101, ResNet152	مجموعه داده بدست آمده از وبسایت مرکز بالینی Kragujevac	CXR	ResNet152 با دقت ۹۵٪
۵	اسماعیل و سنگور (۲۰۲۱) [۵]	ResNet18, ResNet50, ResNet101, VGG16, VGG19 + SVM	مجموعه داده اول: COVID-19 image data collection [10] مجموعه داده دوم: مجموعه داده ذات الریه جمع آوری شده توسط paul mooney در وبسایت [11] kaggle مجموعه داده سوم [۱۲]	CXR	ResNet50 + Linear kernel SVM classifier با دقت ۹۴٫۷٪
۶	مارکوس و همکاران (۲۰۲۰) [۶]	EfficientNetB4	مجموعه داده اول: مجموعه داده ذات الریه جمع آوری شده توسط paul mooney در وبسایت [11] kaggle مجموعه داده دوم: J.P. Cohen [13]	CXR	برای کلاس بندی دو کلاسه: ۹۹٫۵۱٪ و برای کلاس بندی چند کلاسه: ۹۶٫۷٪
۷	حسن تبار، احمدی و شریفی (۲۰۲۰) [۷]	شبکه عصبی عمیق، شبکه عصبی پیچشی	مجموعه داده اول: ۶۸۲ عکس MRI از مجموعه Covid-10 مجموعه داده دوم: ۱۰۰ عکس CT از مجموعه CovidSemiseg	MRI CT scan	CNN پیشنهادی با دقت ۹۳٫۲٪
۸	پولسینلی و همکاران (۲۰۲۰) [۸]	SqueezeNet, Proposed CNN	1. a CT scan dataset about COVID-19 [14] 2. Sirm dataset of COVID-19 chest CT scan [15]	CT scan	۸۵٪ برای proposed CNN

۳- مدل پیشنهادی

در این مقاله یک مدل شبکه عصبی پیچشی سبک برای تشخیص COVID-19 از تصاویر اشعه ایکس قفسه سینه طراحی شده است. در یک طبقه بندی چند کلاسه هدف اصلی طبقه بندی و تشخیص صحیح تصاویر به سه کلاس بیمار (کرونا) و بیمار (غیر کرونا) و سالم است. دلیل استفاده از این سه مورد

۳-۱- ساختار داده

یکی از پیش نیازهای اصلی در شناسایی تصویر^{۱۲} و بینایی کامپیوتر^{۱۳} یافتن مجموعه داده مناسب است. در این بخش شرح مختصری از مجموعه داده های

^{۱۳} Computer vision

^{۱۲} Image recognition

مورد استفاده همراه با نمونه‌هایی از آن ارائه می‌شود. علاوه بر این تکنیک تقویت داده به‌منظور عملکرد بهتر شبکه ارائه خواهد شد.

۳-۱-۱- مجموعه داده

از زمان شیوع بیماری COVID-19 تاکنون چندین مجموعه داده از تصاویر پزشکی تهیه شده است. در این مقاله ما از تصاویر اشعه ایکس قفسه سینه به‌منظور تشخیص بیماری COVID-19 استفاده کردیم. داده‌ها توسط کومار [۱۶] از چندین منبع دیگر جمع‌آوری شده است. تصاویر به سه دسته بیماران COVID-19، بیماران غیر COVID-19 و افراد عادی طبقه‌بندی شده و تمامی عکس‌ها در ساینز ۲۵۶*۲۵۶ تغییر سایز داده شده‌اند.

در مجموع ۵۲۲۸ تصویر اشعه ایکس قفسه سینه استفاده شده و توزیع داده‌ها براساس کلاس‌ها در جدول ۲ ارائه شده است. این مجموعه داده شامل پوشه‌هایی با همین نام‌ها بوده و هیچ اطلاعات دیگری در دسترس نیست.

جدول ۲: مجموعه داده

کلاس داده	تعداد داده	درصد از کل
بیمار (کرونا)	۱۶۲۶	۳۱
بیمار (غیر کرونا)	۱۸۰۲	۳۵
سالم	۱۸۰۰	۳۴
مجموع	۵۲۲۸	۱۰۰

۳-۱-۲- تقویت داده

حجم عظیم تصاویر آموزشی یکی از الزامات یادگیری عمیق است در همین راستا فرایند تقویت داده به‌منظور افزایش عملکرد طبقه‌بندی استفاده شده است. این فرایند با هدف افزایش مصنوعی مجموعه داده‌های آموزشی انجام شده [۱۷] و مجموعه داده آزمایش ثابت باقی می‌ماند. مجموعه‌ای از عملیات‌های هندسی متفاوتی به‌منظور افزایش مجموعه داده قابل استفاده است. ما در این مقاله از

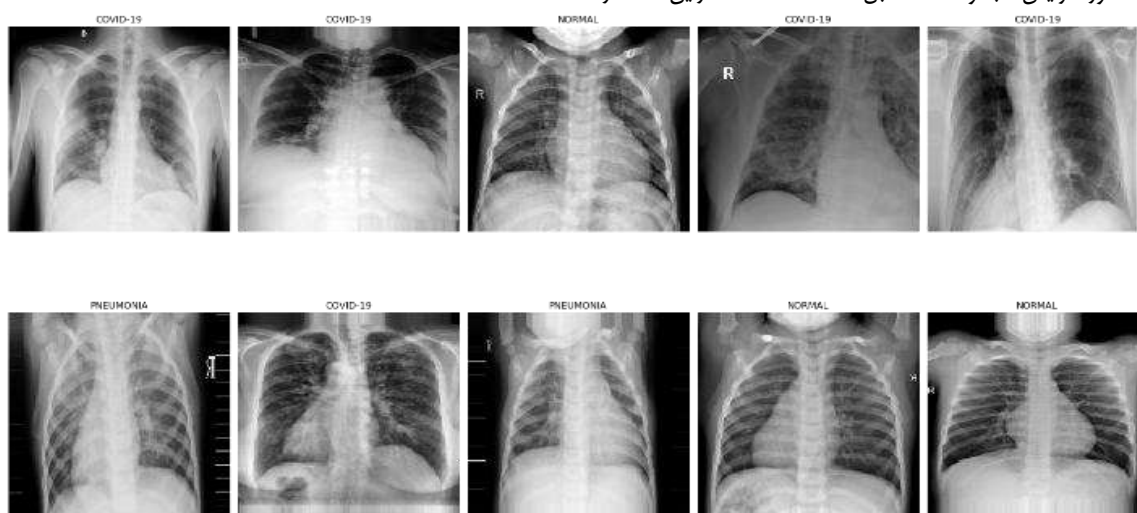
پشت و رو کردن افقی تصاویر و همچنین بزرگنمایی استفاده کرده‌ایم. شکل ۱ برخی از تصاویر هر سه کلاس را نمایش می‌دهد.

۳-۲- شبکه عصبی پیچشی

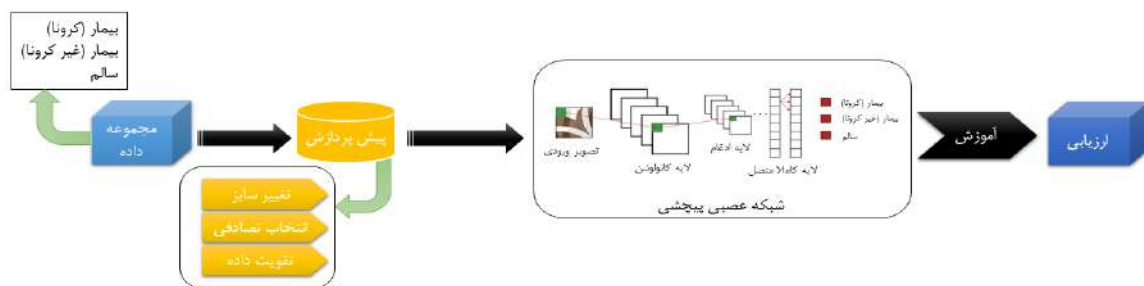
شبکه عصبی پیچشی معتبرترین الگوریتم یادگیری عمیق است [۵] و از آن‌ها برای پردازش حجم عظیمی از داده‌ها استفاده می‌شوند و نیازی به استخراج دستی ویژگی‌ها ندارند. معماری یک شبکه عصبی پیچشی به دو بخش یادگیری ویژگی و طبقه‌بندی تقسیم می‌شود. به‌طور کلی این شبکه‌ها به‌صورت سلسله‌مراتبی از سه نوع لایه تشکیل می‌شوند: لایه کانولوشن و لایه ادغام به‌منظور استخراج ویژگی‌ها و لایه کاملاً متصل به منظور طبقه‌بندی آن‌ها.

۳-۲-۱- معماری پیشنهادی

معماری پیشنهادی شامل سه مرحله می‌باشد. ابتدا تصاویر مجموعه داده پیش‌پردازش شده و سپس توسط شبکه عصبی پیچشی فرایند آموزش صورت گرفته‌است. شبکه عصبی پیچشی پیشنهادی از ۹ لایه تشکیل شده که در آن ۵ لایه کانولوشن و بلافاصله لایه ادغام حداکثر، سپس یک لایه مسطح و سه لایه کاملاً متصل استفاده شده است. این شبکه در مجموع شامل ۳۷۷,۷۳۱ پارامتر که تمامی آن‌ها قابل آموزش بوده می‌باشد. سپس در مرحله آخر شبکه عصبی پیچشی پیشنهادی توسط ۱۰۴۷ تصویر مجموعه آزمایش مورد بررسی قرار خواهد گرفت. علاوه بر این از توابع فعالسازی ReLU و سیگموئید در این شبکه استفاده شده است و وزن‌ها با استفاده از بهینه‌ساز adam تولید می‌شوند. در طول فرایند آموزش نقطه‌ای وجود دارد که خروجی مدل بهبود نمی‌یابد به همین منظور از تکنیک توقف زودهنگام بر اساس کمترین میزان خطای اعتبار سنجی استفاده شده است.



شکل ۱: نمونه‌هایی از تصاویر مجموعه داده بعد از اعمال فرایند تقویت داده

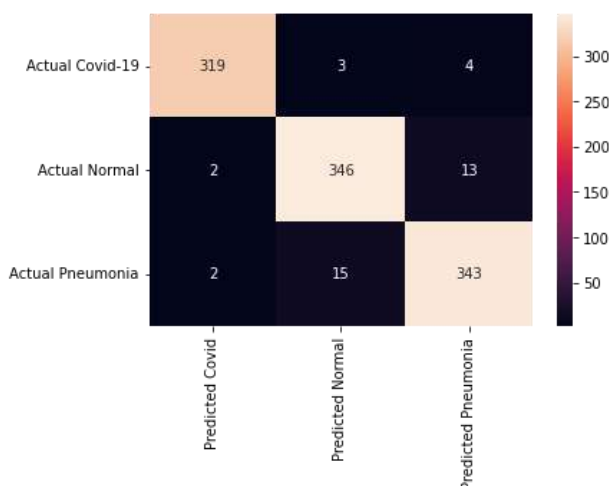


شکل ۲: روند کار مدل پیشنهادی

۴- نتایج بدست آمده

جدول ۳: نتایج بدست آمده از طبقه‌بندی چند کلاسه

کلاس	Precision	Recall	F-score
بیمار (کرونا)	۰,۹۹	۰,۹۸	۰,۹۸
بیمار (غیر کرونا)	۰,۹۵	۰,۹۶	۰,۹۵
سالم	۰,۹۵	۰,۹۵	۰,۹۵
مجموع	۰,۹۶۲	۰,۹۶۲	۰,۹۵۹



شکل ۳: ماتریس درهم‌ریختگی مجموعه آزمایش برای شناسایی افراد مبتلا به کرونا در طبقه‌بندی سه کلاسه (افراد سالم، بیمار (کرونا) و بیمار (غیر کرونا))

جدول ۴ به مقایسه دقت شبکه پیشنهادی با برخی از شبکه‌های بکار رفته در مقالات جدول ۱ می‌پردازد. سایر شبکه‌ها با داده‌های مورد استفاده در این مقاله پیاده‌سازی شده و نتایج در جدول زیر قابل مشاهده است.

در این بخش عملکرد روش پیشنهادی خود را که در بخش قبل توضیح داده شد ارزیابی می‌کنیم. در این طبقه‌بندی سه نوع داده وجود دارد: بیمار مبتلا به کرونا، بیمار غیر کرونا و افراد سالم. مدل پیشنهادی خود را روی مجموعه داده [۱۶] آموزش داده و معیارهای ارزیابی مختلفی برای تجزیه و تحلیل مدل در نظر گرفتیم. تمامی تصاویر مجموعه داده به منظور بهبود عملکرد شبکه در ابعاد 256×256 تغییر سائز داده شده و همچنین از تکنیک تقویت داده استفاده شده است. ۷۰٪ از مجموعه داده‌ها برای آموزش شبکه، ۲۰٪ برای آزمایش و ۱۰٪ برای اعتبارسنجی^{۱۴} استفاده می‌شود. آموزش در ۲۰ مرحله و با استفاده از توقف زودهنگام براساس کمترین میزان خطای اعتبارسنجی انجام شده است. تابع هزینه بکار رفته در فرمول (۱) به‌نمایش گذاشته شده است جایی که $\mathcal{Y}_k$ نتیجه درست طبقه‌بندی و $\hat{\mathcal{Y}}_k$ نتیجه پیش‌بینی شده می‌باشد. دقت^{۱۵} طبق فرمول (۲)، بازخوانی^{۱۶} طبق فرمول (۳) و ارزیابی^{۱۷} طبق فرمول (۴) به ازای هر کلاس در جدول ۳ نمایش داده شده است. شکل ۳ نتایج را در قالب ماتریس آشفتگی نشان می‌دهد. همانطور که در شکل ۳ مشاهده می‌شود مدل به‌طور دقیق ۱۰۰۸ مورد از ۱۰۴۷ مورد مجموعه آزمایش را درست پیش‌بینی کرده است و دقت مجموعه آزمایش ۹۶,۲٪ است. براساس نتایج عملکرد تشخیص COVID-19 با کمک هوش مصنوعی میتوان نتیجه گرفت شبکه عصبی پیچشی ارائه شده با استفاده از تصویر اشعه ایکس قفسه سینه از دقت و حساسیت نسبتاً بالایی برخوردار است.

$$CCE^{18} = - \sum_{k=0}^{m-1} y_k \log(\hat{y}_k) \quad (1)$$

$$Precision = \frac{TP}{(TP+FP)} \quad (2)$$

$$Recall = \frac{TP}{(TP+FN)} \quad (3)$$

$$F - score = \frac{2 * (precision * recall)}{(precision + recall)} \quad (4)$$

True Positive (TP): پیش‌بینی درست (کلاس واقعی مثبت)

True Negative (TN): پیش‌بینی درست (کلاس واقعی منفی)

False Positive (FP): پیش‌بینی غلط (کلاس واقعی منفی)

False Negative (FN): پیش‌بینی غلط (کلاس واقعی مثبت)

^{۱۴} Validation

^{۱۵} Precision

^{۱۶} Recall

^{۱۷} F1-Score

^{۱۸} Categorical Cross Entropy

[4] I. Lorencin et al., "Automatic Evaluation of the Lung Condition of COVID-19 Patients Using X-ray Images and Convolutional Neural Networks," *Journal of Personalized Medicine*, vol. 11, no. 1, p. 28, 2021.

[5] A. M. Ismael and A. Şengür, "Deep learning approaches for COVID-19 detection based on chest X-ray images," *Expert Systems with Applications*, vol. 164, p. 114054, 2021.

[6] G. Marques, D. Agarwal, and I. de la Torre Díez, "Automated medical diagnosis of COVID-19 through EfficientNet convolutional neural network," *Applied soft computing*, vol. 96, p. 106691, 2020.

[7] S. Hassantabar, M. Ahmadi, and A. Sharifi, "Diagnosis and detection of infected tissue of COVID-19 patients based on lung X-ray image using convolutional neural network approaches," *Chaos, Solitons & Fractals*, vol. 140, p. 110170, 2020.

[8] M. Polsinelli, L. Cinque, and G. Placidi, "A light CNN for detecting COVID-19 from CT scans of the chest," *Pattern recognition letters*, vol. 140, pp. 95-100, 2020.

[9] Mendeley Data. Extensive COVID-19 X-Ray and CT Chest Images Dataset. <https://doi.org/10.17632/8h65ywd2jr.3>

[10] COVID-19 image data collection. <https://github.com/ieee8023/covid-chestxray-dataset/tree/master/images>

[11] chest-xray-pneumonia. <https://www.kaggle.com/paultimothymooney/chest-xray-pneumonia>

[12] Covid-19 Imaging findings. <https://radiologyassistant.nl/chest/lk-jg-1>

[13] J. P. Cohen, P. Morrison, L. Dao, K. Roth, T. Q. Duong, and M. Ghassemi, "Covid-19 image data collection: Prospective predictions are the future," *arXiv preprint arXiv:2006.11988*, 2020.

[14] J. Zhao, Y. Zhang, X. He, P. Xie, COVID-CT-dataset: a CT scan dataset about COVID-19, *arXiv preprint arXiv:2003.13865* (2020)

[15] S. italiana di Radiologia Medica e Interventistica, Sirm dataset of COVID-19 chest CT scan, (<https://www.sirm.org/category/senza-categoria/covid-19/>). Accessed: 2020-04-05.

[16] S. Kumar. "COVID19+PNEUMONIA+NORMAL Chest X-Ray Images." <https://www.kaggle.com/sachinkumar413/covid-pneumonia-normal-chest-xray-images> (accessed 11/8/2021, 2021).

[17] N. A. Farda, J.-Y. Lai, J.-C. Wang, P.-Y. Lee, J.-W. Liu, and I.-H. Hsieh, "Sanders classification of calcaneal fractures in CT images with deep learning and differential data augmentation techniques," *Injury*, vol. 52, no. 3, pp. 616-624, 2021.

جدول ۴: مقایسه مدل پیشنهادی با کارهای مرتبط در حوزه شناسایی COVID-19

نتیجه	تعداد پارامتر	تعداد لایه	رویکرد
۰,۹۵۳	۱۳۸ (میلیون)	۱۶	VGG16 [۱]
۰,۹۴۴	۱۴۳ (میلیون)	۱۹	VGG19 [۳]
۰,۹۱۶	۲۵ (میلیون)	۵۰	ResNet50 [۵]
۰,۹۴۵	۶۰ (میلیون)	۱۵۲	ResNet152 [۴]
۰,۹۵۳	۱۹ (میلیون)	---	EfficientNetB4 [۶]
۰,۹۶۲	۳۷۷,۷۳۱	۹	مدل پیشنهادی

۵- نتیجه گیری

این مقاله یک سیستم خودکار برای پشتیبانی از تشخیص بیماران COVID-19 با استفاده از یک مدل شبکه عصبی پیچشی جدید از ابتدا تا انتها آموزش دیده^{۱۹} با استفاده از تصاویر اشعه ایکس قفسه سینه ارائه کرده است. توجه به این نکته مهم است که معماری‌های شبکه عصبی پیچشی عمیق زمانی که با تعداد دوره‌های بیشتری آموزش داده می‌شوند تمایل به بیش‌برازش^{۲۰} دارند. به‌منظور جلوگیری از آن ما از روش‌هایی مانند توقف زودهنگام و تقویت داده استفاده کرده‌ایم هرچند استفاده از روش حذف تصادفی^{۲۱} در اغلب اوقات راه حل مناسبی است. مدل پیشنهادی دقت ۹۶,۲٪ را برای طبقه‌بندی چند کلاسه ارائه کرد هرچند مجموعه داده‌های موجود هنوز قوی نیستند اما دانش تجربی درمورد برنامه‌های کاربردی شبکه‌های عصبی پیچشی بیان می‌کند افزایش تعداد نمونه‌ها و کیفیت مجموعه داده تاثیر مستقیم بر دقت بدست آمده دارد. در کارهای آینده تصاویر بیشتری جمع‌آوری خواهد شد و مدل‌های عمیق‌تری برای تشخیص COVID-19 مورد بررسی قرار خواهد گرفت. علاوه بر این سایر بیماری‌های ریوی نیز در مطالعات آینده گنجانده خواهد شد همچنین توسعه یک رابط گرافیکی برای کمک به متخصصان رادیولوژی در شناسایی COVID-19 میتواند هدف مطالعات آینده باشد.

مراجع

- [1] S. D. Deb, R. K. Jha, K. Jha, and P. S. Tripathi, "A multi model ensemble based deep convolution neural network structure for detection of COVID19," *Biomedical Signal Processing and Control*, vol. 71, p. 103126, 2022.
- [2] S. Thakur and A. Kumar, "X-ray and CT-scan-based automated detection and classification of covid-19 using convolutional neural networks (CNN)," *Biomedical Signal Processing and Control*, vol. 69, p. 102920, 2021.
- [3] M. Elpelatgy and H. Sallam, "Automatic prediction of COVID-19 from chest images using modified ResNet50," *Multimedia Tools and Applications*, pp. 1-13, 2021.

^{۱۹}End-to-end trained

^{۲۰}Overfitting

^{۲۱}Dropout

کاربرد از شبیه‌سازی مونت کارلو در آزمون کیفیت شیشه خودرو

حمیده ایرانمنش^۱، عباس پرچمی^۲، مهدی جباری نوقابی^۱ و بهرام صادق‌پور گیلده^۱

^۱ گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد، مشهد، iranmanesh.hamideh@mail.um.ac.ir

^۱ گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد، مشهد، jabbarinm@um.ac.ir

^۱ گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد، مشهد، sadeghpour@um.ac.ir

^۲ گروه آمار، دانشکده ریاضی و رایانه، دانشگاه شهید باهنر کرمان، کرمان، parchami@uk.ac.ir

چکیده: آزمون فرضیه آماری یک روش موثر برای تصمیم‌گیری در مورد کارایی یک فرایند تولیدی است. با در نظر گرفتن کیفیت فازی به جای حدود مشخصات فنی دقیق، می‌توانیم تصمیمات مطمئن‌تری برای بررسی توانایی کارایی فرایندهای تولیدی بگیریم. در این مقاله یک مطالعه کاربردی بر اساس کیفیت فازی با استفاده از شاخص یانگتینگ ارائه شده است. با توجه به پیچیدگی فرمول‌های شاخص‌های کارایی حتی تحت شرایط نرمال بودن داده‌ها، ممکن است با چالش عدم توانایی پیدا کردن توزیع آماری برآوردگر کارایی فرایند روبرو شویم. همچنین این چالش نیز برای آزمون کارایی فرایند بر اساس کیفیت فازی دیده می‌شود. رویکرد پیشنهادی به کار برده شده در این مطالعه کاربردی، یک تکنیک برای آزمون کارایی یک محصول تولیدی بر اساس کیفیت فازی و مبتنی بر تکنیک‌های نمونه‌گیری تصادفی به روش مونت کارلو است و قابلیت تعمیم برای انواع کیفیت‌های فازی را دارد. این مطالعه در صنعت خودروسازی مبتنی بر کیفیت فازی از نوع مثلثی ارائه شده است. محاسبات عددی در این مطالعه برای نشان دادن عملکرد روش مونت کارلو برای تصمیم‌گیری‌های مطمئن در آزمون شاخص کارایی یانگتینگ ارائه شده‌اند.

کلمات کلیدی: آزمون فرضیه‌ها، کیفیت فازی، کارایی فرایند، شبیه‌سازی مونت کارلو.

۱ مقدمه

در [۹، ۱] مورد بحث قرار گرفت. بر پایه مفهوم کیفیت فازی کران‌های فازی و فاصله بین آنها در [۲] معرفی شد. صادق‌پور گیلده [۱۰] شاخص‌های کارایی C_p ، C_{pk} و شاخص یانگتینگ را با لحاظ وقوع خطای اندازه‌گیری مقایسه کرد. نسل دیگری از شاخص‌های کارایی فرایند توسط پرچمی و ماشین‌چی [۷] برای اندازه‌گیری کارایی کیفیت فازی توسعه داده شد. یکی از مسائل اساسی در استنباط آماری، آزمون فرضیه‌های آماری است و آزمون کارایی یک روش رایج برای بررسی عملکرد فرایندهای تولیدی صنعتی است. ممکن است در آزمون فرضیه‌ها با مواردی روبرو شویم که داده‌ها به صورت مبهم/فازی ثبت شده باشند. در چنین شرایطی، روش‌های کلاسیک آزمون فرضیه‌ها قادر به حل این مسئله جدید نیستند و نیاز به تعمیم دارند. زاده [۱۳] مفهوم فازی نوع دوم را معرفی کرد و پرچمی و همکاران [۸] مفهوم فازی نوع دوم را برای شاخص‌های کارایی C_p ، C_{pk} و C_{pm} بر اساس داده‌های

در کنترل کیفیت، مانند سایر مسائل آماری ممکن است با مفاهیم نادقیق و مبهم روبرو شویم. یک مورد کاربردی در تجزیه و تحلیل کیفیت وضعیتی است که در آن حدود مشخصات فنی یک مجموعه فازی باشد. در روش‌های متداول کنترل کیفیت، یک کالا را "باکیفیت" و یا "بی‌کیفیت" می‌نامند، اما به کمک مفهوم فازی، به هر کالا درجه‌ای (بین صفر و یک) به عنوان میزان کیفیت کالا داده می‌شود. با این دیدگاه می‌توان قضاوت و تصمیم‌گیری منصفانه‌تری در خصوص فرایند تولید کالا اتخاذ نمود. برای نخستین بار یانگتینگ [۱۱] در سال ۱۹۹۶ مفهوم "کیفیت فازی" را در شرایطی که حدود مشخصات فنی USL و LSL هستند، به کمک جایگزین کردن تابع نشانگر $I_{\{x: x \in [LSL, USL]\}}$ با تابع عضویت مجموعه فازی $\tilde{Q}$ ارائه داد. انگیزه و مزایای استفاده از این رویکرد کیفیت فازی به جای استفاده از رویکرد کیفیت کلاسیک

معمولی به کار بردند. چن و هانگ [۳] شاخص ناکارایی فرایند $\widehat{C}_{pp}'$ را با در نظر گرفتن حدود مشخصات فازی نوع دوم معرفی کردند. چن و چانگ [۴] یک روش آزمون فرضیه آماری فازی برای داده‌های فازی با در نظر گرفتن حدود مشخصات فنی یک طرفه توسعه دادند. در حل مسئله آزمون فرضیه‌ها بر اساس داده‌های فازی، ابهام موجود در داده‌ها منجر به ایجاد ابهام در p -مقدار می‌شود. پرچمی [۵] به محاسبه p -مقدار فازی بر اساس داده‌های فازی مبتنی بر اصل گسترش می‌پردازد. همچنین، با توجه به اینکه روش p -مقدار متداول‌ترین روش آزمون فرضیه‌ها در بین کاربران علوم مختلف است، آنها دو مطالعه موردی مبتنی بر p -مقدار فازی نیز ارائه دادند.

فرض کنید انجام آزمون فرضیه صفر $C_{\bar{Q}} \leq c_0$: H_0 (معادل با "فرایند ناکارا است") در مقابل فرضیه یک $C_{\bar{Q}} > c_0$: H_1 (معادل با "فرایند کارا است") مبتنی بر شاخص کارایی یانگتینگ $C_{\bar{Q}}$ مورد نظر باشد. با توجه به پیچیدگی فرمول شاخص کارایی یانگتینگ حتی تحت شرایط نرمال بودن داده‌ها ممکن است با چالش عدم توانایی پیدا کردن توزیع آماری برآوردگر کارایی فرایند روبرو شویم. پرچمی و همکاران [۶] برای حل این مشکل یک الگوریتم ساده و کاربردی برای انجام آزمون فرض آماری پیشنهاد دادند. آنها این الگوریتم را مبتنی بر روش مونت کارلو برای آزمون شاخص کارایی یانگتینگ طراحی کردند. انگیزه اصلی از نوشتن این مقاله، معرفی یک مسئله کاربردی در صنعت مبتنی بر کیفیت فازی مثلی است که به روش مونت کارلو مورد بحث قرار می‌گیرد.

ادامه مقاله به شرح ذیل سازمان داده شده است. در بخش دوم از این مقاله به ذکر مقدمه‌ای بر آزمون کیفیت فازی و شاخص کارایی یانگتینگ برای اندازه‌گیری کیفیت فازی می‌پردازیم. در بخش سوم یک مطالعه کاربردی جدید بر اساس کیفیت فازی مثلی ارائه می‌شود. در انتها مقاله با نتیجه‌گیری پایان می‌یابد.

۲ اندازه‌گیری کیفیت فازی

۱-۲ شاخص کارایی یانگتینگ

فرض کنید f تابع چگالی احتمال یک مشخصه کیفی یک بعدی X باشد. یانگتینگ (۱۹۹۶) شاخص کارایی زیر را برای اندازه‌گیری کیفیت فازی بر اساس داده‌های دقیق ارائه داد:

$$C_{\bar{Q}} = \begin{cases} \int_{-\infty}^{+\infty} \bar{Q}(x) f(x) dx, & \text{برای مشخصه کیفیت فازی پیوسته} \\ \sum_{i=1}^n \bar{Q}(x_i) f(x_i), & \text{برای مشخصه کیفیت فازی گسسته} \end{cases} \quad (۱)$$

که در آن $\bar{Q}$ تابع عضویت کیفیت فازی برای نمایش درجه انطباق با کیفیت استاندارد فازی است. توجه داشته باشید که با در نظر گرفتن x به عنوان مشخصه کیفیت اندازه‌گیری شده یک محصول، $\bar{Q}(x)$ نشان‌دهنده میزان مطابقت با کیفیت استاندارد است [۹]. لازم به ذکر است که

شاخص کارایی فرایند معرفی شده در رابطه (۱) برابر با احتمال رخداد مشخصه کیفی X در کیفیت فازی با توجه به تعریف احتمالاتی زاده [۱۲] است، به بیان دیگر $C_{\bar{Q}} = P(X \in \bar{Q})$.

معمولاً میانگین (μ) و انحراف استاندارد (σ) فرایند نامعلوم هستند که بر اساس نمونه تصادفی جمع‌آوری شده از فرایند تحت کنترل، می‌توان آنها را با میانگین نمونه ($\hat{\mu} = \sum_{i=1}^n X_i/n$) و انحراف استاندارد نمونه ($\hat{\sigma} = \sqrt{\sum_{i=1}^n (X_i - \bar{X})^2/(n-1)}$) برآورد کرد. سرانجام تحت شرایط نرمال بودن مشخصه کیفیت یک بعدی X ، برآوردگر طبیعی شاخص یانگتینگ $C_{\bar{Q}}$ با لحاظ کیفیت فازی $\bar{Q}$ و بر اساس جایگزین کردن برآوردهای $\hat{\mu}$ و $\hat{\sigma}$ به صورت زیر به دست می‌آید:

$$\widehat{C_{\bar{Q}}} = \int_{-\infty}^{+\infty} \bar{Q}(x) \widehat{f_{\mu, \sigma}}(x) dx = \int_{-\infty}^{+\infty} \bar{Q}(x) f_{\hat{\mu}, \hat{\sigma}}(x) dx \quad (۲)$$

که در آن $f_{\mu, \sigma}(x)$ تابع چگالی احتمال توزیع نرمال با پارامترهای μ و σ^2 است.

۲-۲ آزمون کیفیت فازی بر اساس شاخص یانگتینگ

هدف اصلی برای تجزیه و تحلیل کیفیت فازی تحت شرایط نرمال بودن مشخصه کیفیت یک بعدی X ، تعیین کردن مقدار بحرانی مناسب بر اساس شاخص یانگتینگ $C_{\bar{Q}}$ برای تصمیم‌گیری در مورد توانا بودن فرایند تولیدی به تولید محصولات در حدود مشخصات فنی فازی از پیش تعیین شده است. بنابراین برای ارائه رویکرد آماری برای یک تصمیم‌گیری مطمئن، آزمون فرضیه‌های زیر را بر اساس نمونه تصادفی $X_1, \dots, X_n$ در نظر می‌گیریم:

$$\begin{cases} H_0 : C_{\bar{Q}} \leq c_0 & \text{(فرایند ناکارا است)}, \\ H_1 : C_{\bar{Q}} > c_0 & \text{(فرایند کارا است)}, \end{cases} \quad (۳)$$

که در آن $c_0 \in (0, 1)$ حداقل معیار استاندارد برای شاخص یانگتینگ $C_{\bar{Q}}$ است. اگر $\widehat{C_{\bar{Q}}}$ آماره آزمون و مقادیر بزرگ آن موجب رد فرضیه صفر باشند و تعریف کنیم:

$$\phi(X) = \begin{cases} 1 & \widehat{C_{\bar{Q}}} > c, \\ 0 & \text{در غیر این صورت} \end{cases} \quad (۴)$$

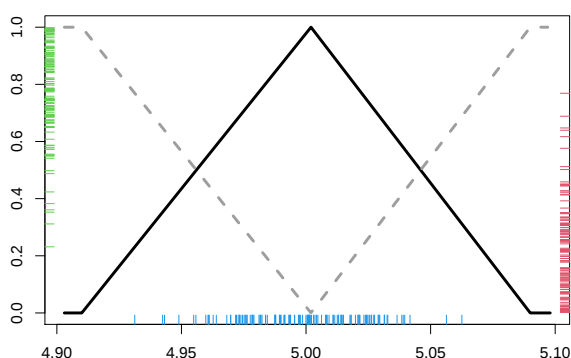
که در آن c مقدار بحرانی آزمون کیفیت فازی است، بنابراین آزمون $\phi(X)$ فرضیه صفر را در صورتی رد می‌کند که نامساوی $\widehat{C_{\bar{Q}}} > c$ برقرار باشد، لذا احتمال خطای نوع اول به صورت زیر به دست می‌آید:

$$\alpha = \sup_{H_0} Pr(\widehat{C_{\bar{Q}}} > c) = P(\widehat{C_{\bar{Q}}} > c | C_{\bar{Q}} = c_0). \quad (۵)$$

در نتیجه

$$P(\widehat{C_{\bar{Q}}} \leq c | C_{\bar{Q}} = c_0) = 1 - \alpha. \quad (۶)$$

به بیانی دیگر مقدار بحرانی c برابر با چندک $(1 - \alpha)$ ام توزیع $\widehat{C_{\bar{Q}}}$ تحت شرط $C_{\bar{Q}} = c_0$ است. نحوه تعیین مقدار بحرانی c با استفاده از روش مونت کارلو طراحی شده در الگوریتم ۱ [۶] است. از طرف دیگر، p -



شکل ۱: تابع عضویت مثلثی کیفیت فازی $\bar{Q}$ و درجه انطباق و عدم انطباق مشاهدات

کیفیت فازی مثلثی زیر آزمون کنیم:

$$\bar{Q}(x) = \begin{cases} \frac{x-4.910}{0.092} & \text{if } 4.910 \leq x < 5.002, \\ \frac{5.090-x}{0.088} & \text{if } 5.002 \leq x < 5.090, \\ 0 & \text{سایر نقاط.} \end{cases} \quad (9)$$

تابع عضویت فوق با کیفیت فازی مثلثی در شکل ۱ ترسیم شده است، و همچنین درجه انطباق و عدم انطباق برای هر مشاهده به ترتیب در سمت چپ و راست شکل ۱ نشان داده شده است.

روش مونت کارلو در این مطالعه برای آزمون فرضیه $H_0: C_{\bar{Q}} \leq 0.76$ در مقابل $H_1: C_{\bar{Q}} > 0.76$ ، بر اساس نمونه تصادفی مشاهده شده $x_1, \dots, x_{135}$ در نظر گرفته شده است.

همچنین آماره $W = 0.993$ به همراه p -مقدار $p\text{-value} = 0.711$ که مربوط به آزمون شاپیرو-ویلک می شوند، گواهی بر تایید نرمال بودن داده ها هستند. برآورد شاخص کارایی یانگتینگ بر اساس مشاهدات $x_1, \dots, x_{135}$ با استفاده از رابطه (۲) برابر است با:

$$\widehat{c}_{\bar{Q}} = \int_{4.910}^{5.090} \bar{Q}(x) f_{\hat{\mu}, \hat{\sigma}}(x) dx = 0.780,$$

که در آن $f_{\hat{\mu}, \hat{\sigma}}$ برآورد تابع چگالی احتمال توزیع نرمال است، که میانگین (μ) و انحراف استاندارد فرایند (σ) بر اساس نمونه تصادفی جمع آوری شده، با میانگین نمونه ($\hat{\mu} = 4.999$) و انحراف استاندارد نمونه ($\hat{\sigma} = 0.0250$) برآورد می شوند. در نتیجه مقدار $\widehat{c}_{\bar{Q}}$ بزرگتر از مقدار بحرانی مونت کارلو c در سطح معنی داری $\alpha = 0.01$ نیست. ($c = 0.780$) لازم به ذکر است که مقدار بحرانی مونت کارلو $c = 0.791$ با استفاده از الگوریتم ۱ طراحی شده در مرجع [۶] قابل محاسبه است و نحوه به کارگیری این الگوریتم در ادامه توضیح داده می شود.

با در نظر گرفتن ۱۲ تکرار برای اجرای شبیه سازی در گام دوم الگوریتم [۶]، دنباله زیر را برای پوشش مقادیر میانگین μ_j در نظر می گیریم:

۵/۰۰۴، ۵/۰۰۱، ۴/۹۹۷، ۴/۹۹۴، ۴/۹۹۰، ۴/۹۸۷، ۴/۹۸۴، ۴/۹۸۰، ۵/۰۱۸، ۵/۰۱۵، ۵/۰۱۱، ۵/۰۰۸

با توجه به تابع عضویت کیفیت فازی در این مطالعه، ریشه نامعلوم σ_j با

مقدار آزمون کیفیت به صورت زیر به دست می آید:

$$p\text{-value} = P(\widehat{C}_{\bar{Q}} > \widehat{c}_{\bar{Q}} | C_{\bar{Q}} = c_0) \\ = E[I(\widehat{C}_{\bar{Q}} > \widehat{c}_{\bar{Q}} | C_{\bar{Q}} = c_0)], \quad (7)$$

که در آن $\widehat{c}_{\bar{Q}}$ مقدار مشاهده شده شاخص کارایی یانگتینگ بر اساس مشاهدات $x_1, \dots, x_n$ با استفاده از رابطه (۲) است و $I(A)$ تابع نشانگر پیشامد/مجموعه A می باشد. همچنین تابع توان آزمون کیفیت فازی با استفاده از شاخص یانگتینگ طبق فرمول زیر تعریف می شود:

$$\Pi(C_{\bar{Q}}) = P(\widehat{C}_{\bar{Q}} > c | C_{\bar{Q}}), \quad (8)$$

که در آن c مقدار بحرانی آزمون است.

۳ آزمون کارایی فرایند تولید شیشه های خودرو

عمدتاً شیشه های خودرو به عنوان بخشی از بدنه و نمای آن دارای دو کارکرد اصلی زیر می باشند:

۱. تامین دید مناسب و ایمن از محیط پیرامونی خودرو برای سرنشینان

۲. تامین ایمنی سرنشینان خودرو هنگام بروز تصادفات و وقوع موارد پیش بینی نشده مانند برخورد اجسام خارجی به شیشه

در حوزه ایمنی، شیشه های خودرو به دو دسته شیشه های سکوریت و شیشه های لمینت تقسیم می شوند. عموماً شیشه های عقب و جانبی خودروها از نوع سکوریت هستند. شیشه سکوریت شیشه ای است که با اعمال یک فرایند عملیات حرارتی، استحکام بالایی یافته و در برابر ضربه ها و تنش های مکانیکی و حرارتی تا پنج برابر شیشه آیل مستحکم است. همچنین، در صورت بروز ضربات سنگین و تصادفات شدید که منجر به شکست شیشه می شود، قطعات شیشه شکسته شده شیشه های سکوریت، بسیار کوچک و فاقد لبه های برنده هستند و مانع صدمات جدی به سرنشینان خودرو می شود. در مقابل، شیشه جلوی خودرو از نوع لایه دار و لمینت است. شیشه لمینت در واقع متشکل از دو لایه شیشه و یک لایه طلق بین آنهاست. این لایه طلق، نقش اصلی در ممانعت از ورود اجسام خارجی به داخل اتومبیل و جلوگیری از پاشش خرده های شیشه لایه داخلی به روی سرنشینان در صورت شکست شیشه، دارد. در فرایند تولید کلیه شیشه های خودروئی مشخصه کیفیت نظیر ضخامت باید به شکلی بهینه باشند که ضمن حصول اطمینان از ایمنی محصول، مطابق الزامات استاندارد، فرم و شکل و انحنای آن دقیق و مطابق بدنه خودرو باشد، طوری که هیچگونه موج غیراستاندارد در سطح آن مشاهده نشود. در این مطالعه موردی به بررسی یک نمونه تصادفی به حجم ۱۳۵ از ضخامت شیشه جلوی خودرو در فرایند تولید شیشه های خودروئی پرداخته می شود. در ادامه، قصد داریم کارایی فرایند ضخامت شیشه جلوی خودرو (بر حسب میلی متر) را بر اساس شاخص یانگتینگ در سطح معنی داری $\alpha = 0.01$ با در نظر گرفتن

استفاده از گام سوم الگوریتم ۱، برای هر دوازده حالت ممکن با استفاده از روش نیوتن رافسون به دست می‌آید. سپس با تولید ۱۰۰۰ نمونه مستقل از توزیع نرمال $N(\mu_j, \sigma_j^2)$ ، برآورد شاخص کارایی یانگتینگ را برای هر نمونه شبیه‌سازی شده با استفاده از رابطه (۲) به دست می‌آوریم. بعد از مرتب کردن این ۱۰۰۰ شاخص کارایی برآورد شده، ۹۹۰ امین شاخص کارایی به عنوان مقدار بحرانی برای هر دوازده حالت ممکن در نظر گرفته می‌شود. میانگین دوازده مقدار بحرانی به دست آمده برابر با $c = 0.791$ است، که به عنوان مقدار بحرانی مونت کارلو در این مطالعه در نظر گرفته می‌شود. در نتیجه با مقایسه $\widehat{c_Q} = 0.780$ و c ، فرضیه صفر رد نمی‌شود و بنابراین فرایند مربوط به بررسی ضخامت شیشه خودرو در این مطالعه، یک فرایند ناکارا در سطح معنی داری $\alpha = 0.01$ قلمداد می‌شود. نتایج و تصمیم‌گیری مربوط به آزمون کیفیت فازی مشاهدات ضخامت شیشه جلوی خودرو بر اساس حداقل معیار استاندارد $c_0 = 0.76$ برای سطوح معنی داری مختلف ۰/۰۱، ۰/۰۲۵، ۰/۰۵۰ و ۰/۱۰۰ در جدول ۱ خلاصه شده‌اند.

۴ نتیجه‌گیری

در مقاله حاضر، مطالعه‌ای کاربردی بر اساس کیفیت فازی مثلی برای ارزیابی و آزمون کارایی فرایندهای تولیدی با استفاده از شاخص کارایی یانگتینگ ارائه شد. در این مطالعه با استفاده از روش طراحی شده الگوریتم ۱ [۶] به برآورد مقدار بحرانی و p -مقدار بر اساس رویکرد آماری آزمون کیفیت فازی مثلی پرداخته شد. آزمایشگر می‌تواند به جای حدود مشخصات فنی دقیق این مطالعه را برای انواع کیفیت‌های فازی برای تعیین مطابقت فرایند مورد بررسی با الزامات کارایی از پیش تعیین شده به کار گیرد و در صورت در نظر گرفتن کیفیت فازی مورد نظر تحقیقات خود، تصمیمات مطمئن‌تری بگیرد. هدف اصلی این مقاله، معرفی یک مسئله کاربردی در صنعت مبتنی بر کیفیت فازی است که به روش مونت کارلو مورد بحث قرار می‌گیرد. تحقیقات آتی در همین راستا، در ارتباط با آزمون شاخص‌های کارایی چندمتغیره بر اساس انواع کیفیت‌های فازی خواهد بود.

مراجع

- [۱] پرچمی، ع.، ماشین چی، م. (۱۳۹۱). کنترل کیفیت آماری، انتشارات دانشگاه شهید باهنر کرمان.
- [۲] پرچمی، ع. و ماشین چی، م. (۱۳۸۶). کیفیت فازی و نسل جدیدی از شاخص‌های کارایی. اندیشه آماری، شماره اول سال دوازدهم، صص. ۶۸ تا ۷۶.

- [3] S. M. Chen, T. M. Hung, (2021). What can fuzziness do for capability analysis based on fuzzy data. Scientia Iranica, 28(2), 1049–1064.

جدول ۱: نتایج آزمون کیفیت فازی برای نمونه‌ای به حجم ۱۳۵ در مطالعه مربوط به ضخامت شیشه جلوی خودرو

سطح معنی داری	مقدار بحرانی مونت کارلو برای C_Q	قاعده تصمیم‌گیری	وضع فرایند بر اساس کیفیت مثلی	p -مقدار مونت کارلو
۰/۰۱۰	۰/۷۹۱	عدم رد فرضیه صفر	ناکارا	۰/۰۶۸
۰/۰۲۵	۰/۷۸۶	عدم رد فرضیه صفر	ناکارا	۰/۰۶۸
۰/۰۵۰	۰/۷۸۲	عدم رد فرضیه صفر	ناکارا	۰/۰۶۸
۰/۱۰۰	۰/۷۷۸	رد فرضیه صفر	کارا	۰/۰۶۸

- [4] K. S. Chen, T. C. Chang, (2020). A fuzzy approach to determine process quality for one-sided specification with imprecise data. *Proceedings of the Institution of Mechanical Engineers, Part B: Journal of Engineering Manufacture*, 234(9), 1198–1206.
- [5] A. Parchami, (2020). Fuzzy decision in testing hypotheses by fuzzy data: Two case studies. *Iranian Journal of Fuzzy Systems*, 17(5), 127–136.
- [6] A. Parchami, H. Iranmanesh, B. Sadeghpour Gildeh (2021). Monte Carlo statistical test for fuzzy quality, *Iranian Journal of Fuzzy Systems*, 2103–6531.
- [7] A. Parchami, M. Mashinchi (2010). A new generation of process capability indices, *Journal of Applied Statistics*, 37 (1), 77–89.
- [8] A. Parchami, S. Ç. Onar, B. Öztayşi, C. Kahraman, (2017). Process capability analysis using interval type-2 fuzzy sets. *International Journal of Computational Intelligence Systems*, 10(1), 721–733.
- [9] A. Parchami, B. Sadeghpour Gildeh, M. Mashinchi (2016). Why Fuzzy Quality?. *International Journal for Quality Research*, 10 (3), 457–470.
- [10] B. Sadeghpour Gildeh, (2003). Comparison of C_p , C_{pk} and $C_{p\text{-tilde}}$ process capability indices in the case of measurement error occurrence. *IFSA World Congress, Istanbul, Turkey*, 563–567.
- [11] C. Yongting (1996). Fuzzy quality and analysis on fuzzy probability. *Fuzzy Sets and Systems*, 83, 283–290.
- [12] L.A. Zadeh, (1968). Probability measures of fuzzy events. *Journal of Mathematical Analysis and Applications*, 23 (2), 421–427.
- [13] L. A. Zadeh, (1975). The concept of a linguistic variable and its application to approximate reasoning-III. *Information sciences*, 9(1), 43–80.

گامی فراتر در پیشگویی پیوند: یک مرور سیستماتیک بر پیشگویی پیوند چندلایه

ساناز جراحی^۱، صادق سلیمانی^۲*

^۱ دانشجوی کارشناسی ارشد، گروه مهندسی کامپیوتر، دانشگاه کردستان، سنندج، s.jarahi@uok.ac.ir

^۲ استادیار، گروه مهندسی کامپیوتر، دانشگاه کردستان، سنندج، s.sulaimany@uok.ac.ir

چکیده: پیشگویی پیوند در حیطه شبکه‌های چندلایه به تازگی توجه زیادی را به خود جلب کرده است به طوری که بیش از هفتاد درصد مقالات آن متعلق به سه سال اخیر است. این شیوه پیشگویی پیوند، نکات و چالش‌های مختلف پژوهشی را با خود به همراه دارد و مستعد نوآوری‌های متعدد است. در این مقاله با تمرکز بر پیشگویی پیوند چندلایه، کل مقالات منتشر شده در این زمینه که تعداد آن‌ها ۵۲ عدد است و اولین آن‌ها در سال ۲۰۱۵ منتشر شده است، از پایگاه‌داده‌های علمی معتبر جمع‌آوری شده و از نظر سال انتشار، یکسان بودن یا نبودن تعداد گره‌ها در همه لایه‌ها، نظارتی بودن یا نبودن محاسبات، تعداد لایه‌های مورد بررسی در مقالات و پوشش یک یا هر دو زمینه بین‌لایه‌ای یا درون لایه‌ای مورد بررسی قرار گرفته‌اند. در پایان به حیطه‌های بررسی نشده و مستعد برای محاسبات آتی نیز اشاره شده است که از آن جمله می‌توان به پیشگویی پیوند چندلایه منفی، پیشگویی پیوند چند لایه برای گراف‌های دوبخشی و علامت‌دار، بهبود روش‌های پیشگویی پیوند چندبخشی از طریق الگوریتم‌های مقایسه شبکه‌ها و لحاظ کردن اثرات ویژگی‌های توپولوژیکی شبکه بر کارایی الگوریتم‌های پیشگویی پیوند درون لایه‌ای اشاره کرد.

کلمات کلیدی: پیشگویی پیوند، پیش‌بینی پیوند، چند لایه، شبکه چند لایه

۱- مقدمه

کرد [2]. این محاسبات می‌توانند، همانند سایر محاسبات بر شبکه‌های پیچیده، با ناظر یا بدون ناظر اعمال شوند. همچنین طیف وسیعی از گراف‌ها اعم از گراف‌های ساده، وزن‌دار، جهت‌دار، دوبخشی و ... را می‌توان برای پیشگویی پیوند، مورد بررسی قرار داد و برای هر کدام از این نوع گراف‌ها، تاکنون روش‌های مختلف پیشگویی پیوند متناسب پیشنهاد شده‌اند [3]. علاوه بر آن، پیشگویی پیوند در کاربردهای مختلف مورد استفاده قرار گرفته است و مفید واقع شده است که از آن جمله می‌توان به شبکه‌های اجتماعی، شبکه‌های زیستی و پزشکی، شبکه‌های علمی مانند هم‌تألیفی و ... اشاره کرد [4].

به فراخور مدل‌های چندلایه از شبکه‌های پیچیده، پیشگویی پیوند چندلایه^۲ نیز هشت سال پس از ظهور، پیشگویی پیوند، در سال ۲۰۱۵ تعریف و عرضه

پیشگویی پیوند زیر حوزه‌ای از گراف کاوی^۱ است که پژوهش‌های زیادی در ارتباط با آن انجام شده و در حال انجام است. در پایه‌ای‌ترین حالت، تعریف پیشگویی پیوند آن است که با داشتن یک گراف $G=(V,E)$ در زمان t_0 ، بخواهیم ارتباطات تشکیل شونده در آن در زمان t_1 را از نظر محاسباتی پیش‌بینی کنیم [1]. این تعریف، بعدها به شدت توسعه یافت و از جنبه‌های مختلف بررسی و تطابق پیدا کرد. به عنوان مثال، می‌توان پیشگویی پیوند را برای پیش‌بینی ارتباطات آتی حذف شونده یا حتی ترکیبی از روابط که در آینده، مراحل بعدی تحول، یک گراف حذف و اضافه خواهند شد، نیز تعریف

^۱ <http://www.sciencedirect.com>

^۸ <http://scholar.google.com>

^۹ <http://academic.microsoft.com>

^۱ Graph Mining

^۲ Multilayer Link Prediction

^۳ Intralayer

^۴ Interlayer

^۵ Multiplex

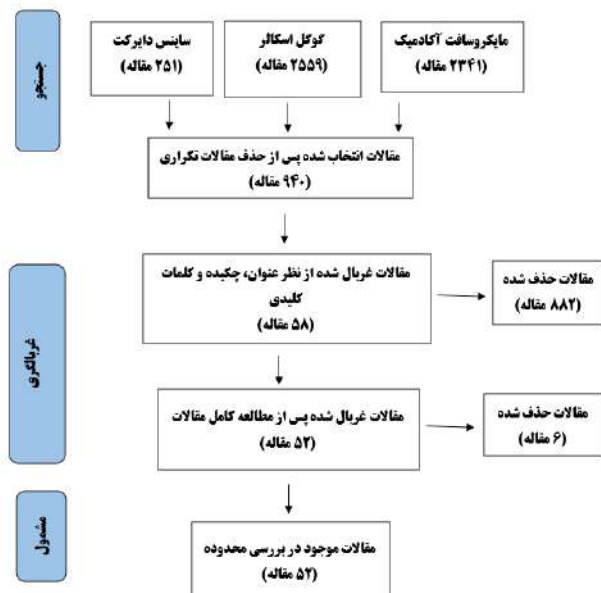
^۶ Multilayer

آنگاه V مجموعه رئوس و E مجموعه یال های بین آن ها خواهند بود. در این صورت یک شبکه با p لایه به شکل $Gm = (Gm1, Gm2, Gm3, \dots, Gmp)$ نمایش داده می شود که $Gmi = (Vi, Ei)$ بیانگر گراف متناظر با لایه i است. اگر $\exists Gmi, Gmj \quad |Vi| \neq |Vj|$ یعنی شبکه از نوع چند لایه است، در غیر این صورت شبکه را چندگانه گوئیم.

۳- روش پژوهش

هدف از این پژوهش، کسب نگرش نسبت به وضعیت کنونی پیشگویی پیوند چند لایه است. بدین منظور بررسی کتابخانه ای جامع در پایگاه داده های معتبر مقالات علمی، با سبک های مختلف انجام گرفت. سه پایگاه داده زیر به دلیل تفاوت های نحوه جستجو و نمایش خروجی و نیز اعتبار آن ها، انتخاب شدند:

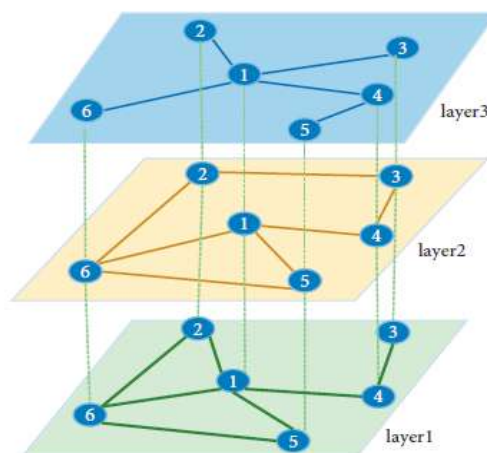
- ساینس دایرکت^۷
- گوگل اسکالر^۸
- مایکروسافت آکادمیک^۹



شکل ۲: دیاگرام جریان پریسما برای فرایند انتخاب مقالات از سه پایگاه داده معرفی شده

در ابتدا برای هر پایگاه داده عبارت "multilayer" + "link prediction" و سپس با جستجوی پیشرفته، محدود به جستجو در عنوان و چکیده مقالات، عبارت "multi-layer" + "link prediction" بدون هیچ محدودیتی از نظر زمان انتشار مقالات مورد جستجو قرار گرفت و به این صورت تمام مقالات از این سه پایگاه داده گردآوری شدند. سپس مقالات بدست آمده از نظر تکراری بودن مورد بررسی قرار گرفته شدند و بر اساس این غربالگری ۹۴۰ عدد مقاله انتخاب شد. بعد از آن این تعداد مقاله را از نظر ارتباط دقیق تر با موضوع مورد مطالعه ما، با بررسی عنوان، چکیده و کلمات کلیدی بررسی کردیم و به ۵۸ مقاله دست یافتیم و در نهایت با مطالعه کامل و دقیق متن

شد [5]. این نوع از پیشگویی پیوند، با در نظر گرفتن شبکه در چند لایه، به بررسی دو مسأله پیشگویی پیوند درون لایه ای^۳ با در نظر گرفتن شبکه های سایر لایه ها و پیشگویی پیوند بین لایه ای^۴ با در نظر گرفتن شبکه های هر لایه می پردازد و کاربردها و چالش های جدید فرارو قرار می دهد. شکل ۱، مثالی از یک شبکه یا گراف چند لایه است که گره ها در هر لایه، نسبت به سایر لایه ها یکسان هستند. برای مثال هایی از شبکه های چند لایه می توان به انواع شبکه های زیستی و پزشکی مانند دارو-هدف-ژن [6]، یا شبکه های علمی مانند مقاله-نویسنده-مکان-واژه [7] را نام برد.



شکل ۱: مثالی از یک شبکه چند لایه که پیشگویی پیوند را می توان درون لایه ای یا بین لایه ای استفاده کرد (برگرفته از [8])

علیرغم مطالعات مروری متعدد در زمینه پیشگویی پیوند [4], [9], [10]، بر اساس بررسی ها و مطالعات مؤلفین این مقاله، تاکنون پژوهشی به صورت مجزا به مرور مقالات مرتبط با پیشگویی پیوند چند لایه نپرداخته است. با توجه به اهمیت و کاربردهای متعدد شبکه های چند لایه، یک مقاله مروری از دیدگاه های متنوع می تواند در شناخت پتانسیل های مختلف پیشگویی پیوند در شبکه های چند لایه، راهگشا باشد.

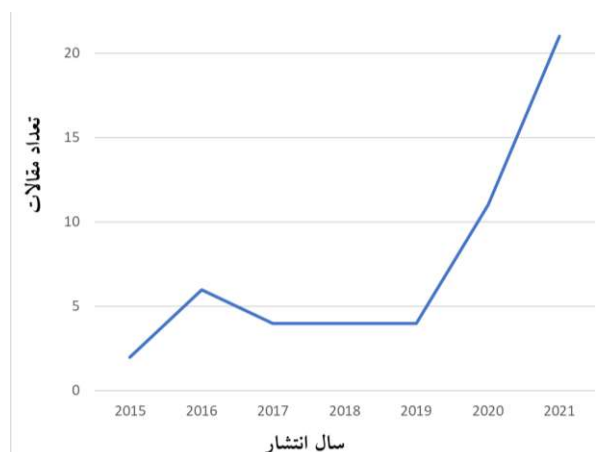
۲- پیشگویی پیوند چند لایه

دو نوع شبکه چند لایه می توان در نظر گرفت. شبکه هایی که تعداد گره ها در تمام لایه ها یکسان هستند که به آن ها شبکه های چندگانه^۵ گفته می شود و شبکه هایی که تعداد گره ها در لایه ها متفاوت است که شبکه چند لایه^۶ نامیده می شود. در این مقاله به منظور سهولت ممکن است از هر دو شبکه با نام چند لایه یاد شود، اما در صورتی که نیاز به تصریح باشد، اصطلاح شبکه چندگانه نیز ذکر خواهد شد. همچنین ممکن است کلمات پیشگویی و پیش بینی به جای هم به کار بروند. علاوه بر آن کلمات شبکه و گراف نیز ممکن است به جای هم استفاده شوند. به منظور بررسی دقیق تر این شبکه ها، تعریف زیر، را در نظر می گیریم. اگر یک شبکه ساده بدون جهت را با $G_s(V, E)$ نمایش دهیم که G_s بیانگر گراف مورد استفاده برای نمایش شبکه است.

مقالات به ۵۲ عدد مقاله مرتبط با حیطه کار خود دست یافتیم و در ادامه بر اساس معیارهای مختلفی، این مقالات را دسته‌بندی کرده‌ایم.

۴- نتایج و بحث

جمع‌بندی تعداد مقالات بر حسب سال انتشار، نشان از یک شیب صعودی در سه سال اخیر و مورد توجه واقع شدن آن از سوی پژوهشگران است که بیش از ۷۰ درصد مقالات را تشکیل می‌دهد که در شکل ۳ نشان داده شده است و آمار صعودی با شیب زیاد انتشار مقالات در سال‌های اخیر در این حیطه، حاکی از خوش آتیه بودن آن به عنوان یک زمینه پژوهشی مرتبط با پیش‌بینی پیوند است.



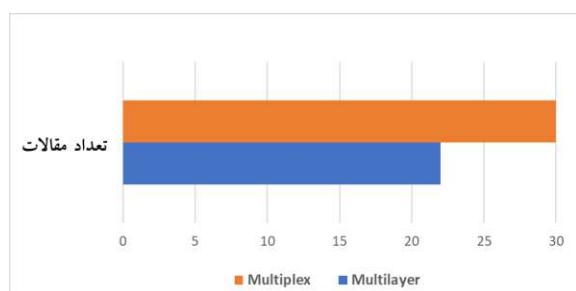
شکل ۳: تعداد مقالات منتشر شده بر حسب سال

عمده مقالات منتشر شده در مجلات معتبر هستند. آمار کلی به شرح جدول ۱ است.

جدول ۱- تنوع انتشار مقالات از نظر منبع انتشار

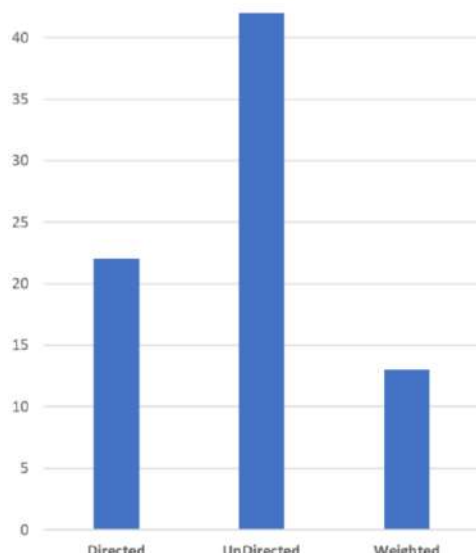
نوع انتشار	مجله	کنفرانس	فصل کتاب	پیش‌نویس
تعداد	۴۰	۵	۵	۲

در شبکه‌های چند لایه، تعداد گره‌ها در همه لایه‌ها لزوماً برابر نیستند و هر گره در هر لایه می‌تواند به تعداد مختلفی گره در لایه دیگر متصل باشد. این در حالی است که در شبکه‌های چندگانه تعداد گره‌ها در همه لایه‌ها یکسان است و گره‌های هر لایه با لایه دیگر ارتباط یک به یک دارد، یعنی یک گره در یک لایه نمی‌تواند به بیش از یک گره در لایه دیگر متصل باشد. از کل مقالات منتشر شده، ۳۰ مورد به شبکه چندگانه و ۲۲ مورد به شبکه چند لایه پرداخته‌اند که در شکل ۴ نشان داده شده است و به نظر می‌رسد علاقه پژوهش‌گران بیشتر به پژوهش در زمینه شبکه‌های چندگانه بوده است.



شکل ۴: تعداد مقالات بر اساس شبکه‌های چندگانه نسبت به شبکه‌های چند لایه

نوع گراف مورد استفاده در پیشگویی پیوند چندگانه، ساده، جهت‌دار و وزن‌دار بوده‌اند که توزیع تعداد آن‌ها حسب شکل ۵ است و با توجه به اینکه در زمینه گراف‌های علامت‌دار^{۱۰} یا دوبخشی یا هایپرگراف^{۱۱}، کاری مرتبط با شبکه‌های چندلایه انجام نشده است، پس این زمینه‌ها مستعد بررسی بیشتر می‌باشند.



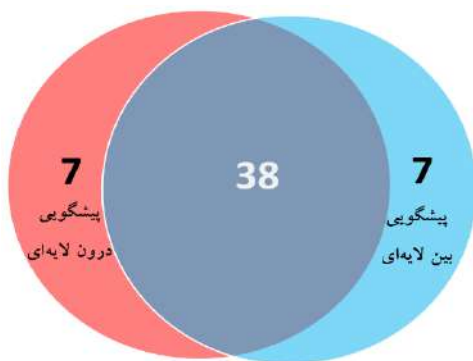
شکل ۵: دسته‌بندی مقالات از نظر نوع گراف مورد استفاده

روش‌های انجام پیشگویی پیوند، در یک دسته‌بندی کلی، به باناظر یا بدون ناظر تقسیم می‌شوند، که روش‌های باناظر، الگوریتمی را با استفاده از یک مدل یادگیری خاص، بر یک مجموعه داده موجود آموزش می‌دهند، تا نسبت به داده‌های جدید، برا ساس یادگیری پیشین، محاسبه و پیش‌بینی انجام دهد. اما روش‌های بدون ناظر، تنها بر اساس اطلاعات کنونی توپولوژیکی شبکه، نسبت به انجام پیش‌بینی اقدام می‌کنند. عمده مقالات منتشر شده در رابطه با پیشگویی پیوند چندلایه، از نوع باناظر است و این آمار نشان می‌دهد که پژوهشگران علاقه زیادی به استفاده از روش پیشگویی پیوند باناظر داشته‌اند.

^{۱۰} Signed

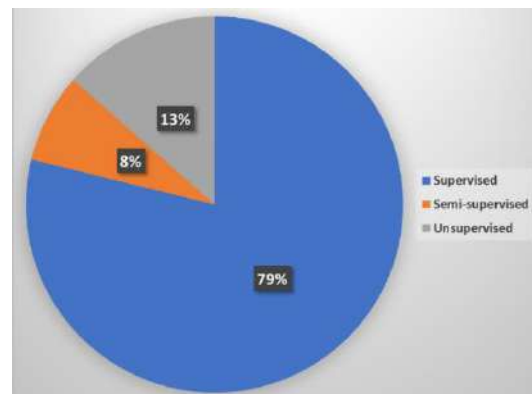
^{۱۱} Hypergraph

یک معیار مهم دسته‌بندی پیشگویی پیوند چندلایه، مکان انجام پیشگویی است. از نظر درون لایه‌ای و بین لایه‌ای بودن آن است. در پیشگویی درون لایه‌ای، هدف، یافتن پیوند جدید در یک لایه خاص است. اما در پیشگویی پیوند بین لایه‌ای، هدف، پیش‌بینی پیوندهای جدید بین لایه‌ها است. این در حالی است که مقالات متعددی بر پیشگویی همزمان پیوندهای درون لایه‌ای و بین لایه‌ای محاسبات انجام داده‌اند و تمایل پژوهشگران برای انجام پیشگویی، همزمان در هر دو موقعیت درون لایه و بین لایه را نشان می‌دهد.



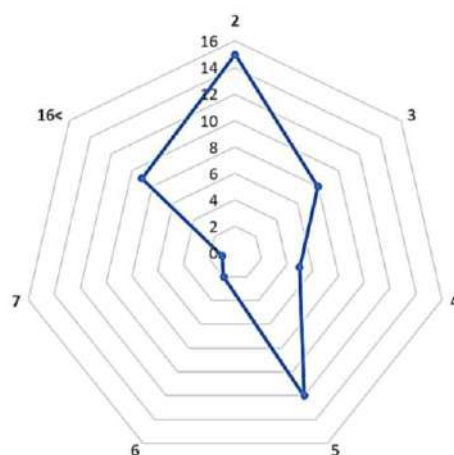
شکل ۸: تعداد مقالات مشترک نسبت به تعداد مطالعات تک زمینه‌ای

روش‌های ارزیابی مقالات عمدتاً بر حسب با ناظر یا بدون ناظر بودن روش محاسباتی مورد استفاده، AUC، Precision، منحنی ROC، ماتریس به‌هم‌ریختگی^{۱۲} که معیارهای Precision-Recall و F-Measure را در اختیار قرار می‌دهد بوده است و از این نظر با روش‌های پیشگویی پیوند تک لایه، شباهت زیادی دارد.



شکل ۶: پراکندگی پژوهش‌ها از نظر نظارتی بودن محاسبات

یک موضوع مهم در پیشگویی پیوند چندلایه، میزان توسعه‌پذیری روش‌ها از نظر پوشش تعداد لایه‌های مورد محاسبه است. شکل ۷ نشان می‌دهد که تعداد زیادی از پژوهش‌ها، بر شبکه‌های دو لایه متمرکز بوده‌اند، در حالی که پژوهش‌های متمرکز بر تعداد لایه‌های زیاد نیز قابل توجه است.



شکل ۷: نمودار راداری تعداد مقالات بر حسب تعداد لایه‌های تحت بررسی در پیشگویی پیوند. اعداد روی هفت ضلعی تعداد لایه و اعداد درون هفت ضلعی، شمارش مقالات با این تعداد لایه است

جدول ۲- لیست کل مقالات منتشر شده در رابطه با پیشگویی پیوند چند لایه

سال	عنوان	مرجع
2021	Layer reconstruction and missing link prediction of a multilayer network with maximum a posteriori estimation	[11]
2021	Layer Information Similarity Concerned Network Embedding	[8]
2021	A deep learning approach to predict inter-omics interactions in multi-layer networks	[12]
2021	Heterogeneous Types of miRNA-Disease Associations Stratified by Multi-Layer Network Embedding and Prediction	[13]
2021	Multiplex network embedding for implicit sentiment analysis	[14]
2021	HITS centrality based on inter-layer similarity for multilayer temporal networks	[15]
2021	Link prediction for multilayer networks using interlayer structural information	[16]
2021	Interlayer Link Prediction in Multiplex Social Networks Based on Multiple Types of Consistency Between Embedding Vectors	[17]
2021	Prediction of new scientific collaborations through multiplex networks	[18]

^{۱۲} Confusion matrix

2021	An information theoretic approach to link prediction in multiplex networks	[19]
2021	Effective link prediction in multiplex networks: A TOPSIS method	[20]
2021	Multilayer network analysis for improved credit risk prediction	[21]
2021	Link prediction in multiplex networks using a novel multiple-attribute decision-making approach	[22]
2021	Prediction of drug-target interactions based on multi-layer network representation learning	[6]
2021	Prediction of drug response in multilayer networks based on fusion of multiomics data	[23]
2021	Link Prediction with Multiple Structural Attentions in Multiplex Networks	[24]
2021	Multilayer Network Analysis: The Identification of Key Actors in a Sicilian Mafia Operation	[25]
2021	A hybrid approach for pair-wise layer similarity in a multiplex network	[26]
2021	Community-guided link prediction in multiplex networks	[27]
2021	Supervised-learning link prediction in single layer and multiplex networks	[28]
2021	A new link prediction in multiplex networks using topologically biased random walks	[29]
2020	A classification approach to link prediction in multiplex online ego-social networks	[30]
2020	Link prediction in multilayer networks	[31]
2020	Interlayer link prediction in multiplex social networks: An iterative degree penalty algorithm	[32]
2020	Link prediction in real-world multiplex networks via layer reconstruction method	[33]
2020	Overlapping communities and the prediction of missing links in multiplex networks	[34]
2020	Heterogeneous Multi-Layered Network Model for Omics Data Integration and Analysis	[35]
2020	Analysis of the Influence of Interlayer Structure Information of Multi-Layer Networks on Link Prediction	[36]
2020	Supervised link prediction in multiplex networks	[37]
2020	Multiplex Graph Association Rules for Link Prediction	[38]
2020	Link prediction in multiplex networks via triadic closure	[39]
2020	Link prediction via community detection in bipartite multi-layer graphs	[40]
2019	Link prediction in multiplex networks based on interlayer similarity	[41]
2019	Information spread link prediction through multi-layer of social network based on trusted central nodes	[42]
2019	Tensorial and bipartite block models for link prediction in layered networks and temporal networks	[43]
2019	Application of hyperbolic geometry in link prediction of multiplex networks	[44]
2018	Link Prediction in Multi-layer Networks and Its Application to Drug Design	[45]
2018	Multilayer Link Prediction in Online Social Networks	[46]
2018	A novel multilayer model for missing link prediction and future link forecasting in dynamic complex networks	[47]
2018	Quantifying layer similarity in multiplex networks: a systematic study	[48]
2017	Community detection, link prediction, and layer interdependence in multilayer networks	[49]
2017	Link prediction in multiplex online social networks	[50]
2017	Link prediction via layer relevance of multiplex networks	[51]
2017	Layer-Wise Model Stacking for Link Prediction in Multilayer Networks. Case of Scientific Collaboration Networks	[52]
2016	TeleLink: Link Prediction in Social Network Based on Multiplex Cohesive Structures	[53]
2016	A multilayer approach to multiplexity and link prediction in online geo-social networks	[54]
2016	A holistic approach for predicting links in coevolving multiplex networks	[55]
2016	An efficient method for link prediction in weighted multiplex networks	[56]
2016	A multilayered approach for link prediction in heterogeneous complex networks	[7]
2016	A Holistic Approach for Link Prediction in Multiplex Networks	[57]

2015	Link prediction in multiplex networks	[5]
2015	An Efficient Method for Link Prediction in Complex Multiplex Networks	[58]

and N. B. Anuar, "Applications of link prediction in social networks: A review," *J. Netw. Comput. Appl.*, vol. 166, p. 102716, 2020.

- [5] M. Pujari and R. Kanawati, "Link prediction in multiplex networks," *Networks Heterog. Media*, vol. 10, no. 1, p. 17, 2015.
- [6] Y. Shang, L. Gao, Q. Zou, and L. Yu, "Prediction of drug-target interactions based on multi-layer network representation learning," *Neurocomputing*, vol. 434, pp. 80–89, 2021.
- [7] H. Shakibian, N. M. Charkari, and S. Jalili, "A multilayered approach for link prediction in heterogeneous complex networks," *J. Comput. Sci.*, vol. 17, pp. 73–82, 2016.
- [8] R. Lu, P. Jiao, Y. Wang, H. Wu, and X. Chen, "Layer Information Similarity Concerned Network Embedding," *Complexity*, vol. 2021, 2021.
- [9] V. Martinez, F. Berzal, and J. C. Cubero, "A survey of link prediction in complex networks," *ACM Comput. Surv.*, vol. 49, no. 4, 2016.
- [10] A. Kumar, S. S. Singh, K. Singh, and B. Biswas, "Link Prediction Techniques, Applications, and Performance: A Survey," *Phys. A Stat. Mech. its Appl.*, vol. 553, p. 124289, 2020.
- [11] J. Kuang and C. Scoglio, "Layer reconstruction and missing link prediction of multilayer network with Maximum A Posteriori estimation," *arXiv Prepr. arXiv2101.02733*, 2021.
- [12] N. Borhani, J. Ghaisari, M. Abedi, M. Kamali, and Y. Gheisari, "A deep learning approach to predict inter-omics interactions in multi-layer networks," 2021.
- [13] D.-L. Yu, Z.-G. Yu, G.-S. Han, J. Li, and V. Anh, "Heterogeneous Types of miRNA-Disease Associations Stratified by Multi-Layer Network Embedding and Prediction," *Biomedicines*, vol. 9, no. 9, p. 1152, 2021.
- [14] X. Zhao, Y. Liu, and Z. Jin, "Multiplex network embedding for implicit sentiment analysis," *Complex Intell. Syst.*, pp. 1–15, 2021.
- [15] L. Lv, K. Zhang, D. Bardou, X. Li, T. Zhang, and W. Xue, "HITS centrality based on inter-layer similarity for multilayer temporal networks," *Neurocomputing*, vol. 423, pp. 220–235, 2021.
- [16] F. Tang, C. Wang, Y. Wang, and J. Su, "Link prediction for multilayer networks using interlayer structural information," *Int. J. Mod. Phys. C*, p. 2250003, 2021.
- [17] R. Tang, Z. Miao, S. Jiang, X. Chen, H. Wang, and W. Wang, "Interlayer Link Prediction in Multiplex Social Networks Based on Multiple Types of Consistency between Embedding Vectors," *IEEE Trans. Cybern.*, 2021.
- [18] M. Tuninetti, A. Aleta, D. Paolotti, Y. Moreno, and M. Starnini, "Prediction of new scientific collaborations through multiplex networks," *EPJ Data Sci.*, vol. 10, no. 1, p. 25, 2021.
- [19] S. H. Jafari, A. M. Abdolhosseini-Qomi, M. Asadpour, M. Rahgozar, and N. Yazdani, "An information theoretic approach to link prediction in multiplex networks," *Sci. Rep.*, vol. 11, no. 1, pp. 1–21, 2021.

۵- نتیجه گیری و کارهای آتی

در مقایسه با پیشگویی پیوند معمول در شبکه‌های یک لایه، نسبت به شبکه‌های چند لایه، موارد متعددی وجود دارد که هنوز لحاظ نشده است و مستعد بررسی و آزمون است. به عنوان مثال، انواع گراف‌های مورد بررسی در پیشگویی پیوند چندلایه، ساده، وزن دار و جهت دار بوده است در حالی که سایر انواع گراف مانند گراف دوبخشی، گراف علامت دار و هایپرگراف نیز می‌توانند مدل سازی کننده شبکه‌ها در هر لایه باشند. همچنین استراتژی پیشگویی پیوند صرفاً یافتن یال‌های تشکیل شونده در آینده، پیشگویی پیوند مثبت، نیست، بلکه می‌توان از پیشگویی پیوند برای یافتن ارتباطات ضعیف یا غیرواقعی یا اشتباهی نیز بهره برد، پیشگویی پیوند منفی. حتی امکان استفاده از پیشگویی پیوند به صورت ترکیبی یعنی لحاظ کردن همزمان حذف و اضافه ارتباطات نیز در برخی مقالات بررسی شده و قابلیت تعمیم و استفاده به شبکه‌های چند لایه را دارد [3]، [2]. علاوه بر آن، هر شبکه‌ای که بتوان آن را در دو لایه یا بیشتر مدل نمود می‌توان برای کاربردهای پیشگویی ارتباطات جدید، درون و بین لایه‌ای مورد بهره‌برداری قرار داد. به عنوان مثال در حیطه بیوانفورماتیک و پزشکی، به سادگی شبکه‌های چندلایه مختلف از قبیل فنوتایپ-ژنوتایپ^{۱۳} در نظر گرفت و پتانسیل سنجی استفاده از پیشگویی پیوند در بهبود یا ارتقای آن‌ها را مورد بررسی قرار داد.

الگوریتم‌هایی وجود دارد که شبکه‌های مختلف را بتوان با همدیگر مقایسه نمود. این الگوریتم‌ها از امکان مقایسه ساده دو شبکه با تعداد گره همسان، تا امکان مقایسه دو شبکه با تعداد گره‌های مختلف، متنوع هستند و به این ترتیب می‌توان برای کمک به بهبود نتایج محاسباتی پیشگویی پیوند در شبکه‌های چند لایه، از آن‌ها کمک گرفت. به عنوان مثال می‌توان به مراجع [60]، [59] اشاره کرد. علاوه بر آن تاثیر توپولوژی شبکه بر نتایج پیشگویی پیوند بین لایه‌ای، موضوعی است که تاکنون مورد بررسی قرار نگرفته است و با الهام از مقالات مشابه، قابل انجام است [61].

در نهایت، جالب آنکه درصد قابل توجهی از پژوهشگران حیطه شبکه‌های چندلایه را ایرانیان در دانشگاه‌های داخل و خارج از کشور تشکیل می‌دهند تا جایی که ۳۴ نفر از مجموع کل ۱۵۳ نویسنده، یعنی بیش از بیست درصد نویسندگان مقالات پیشگویی پیوند چندلایه، ایرانی هستند.

مراجع

- [1] D. Liben-nowell and J. Kleinberg, "The Link Prediction Problem for Social Networks," *J. Am. Soc. Inf. Sci. Technol.*, vol. 58, no. 7, pp. 1019–1031, 2007.
- [2] S. Sulaimany *et al.*, "Predicting brain network changes in Alzheimer's disease with link prediction algorithms," *Mol. Biosyst.*, vol. 13, no. 4, 2017.
- [3] S. Sulaimany, M. Khansari, and A. Masoudi-Nejad, "Link prediction potentials for biological networks," *Int. J. Data Min. Bioinform.*, vol. 20, no. 2, p. 24, 2018.
- [4] N. N. Daud, S. H. Ab Hamid, M. Saadoon, F. Sahran,

^{۱۳} Phenotype-Genotype

- [37] N. Shan, L. Li, Y. Zhang, S. Bai, and X. Chen, "Supervised link prediction in multiplex networks," *Knowledge-Based Syst.*, vol. 203, p. 106168, 2020.
- [38] M. Coscia and M. Szell, "Multiplex Graph Association Rules for Link Prediction," *arXiv Prepr. arXiv2008.08351*, 2020.
- [39] A. Aleta, M. Tuninetti, D. Paolotti, Y. Moreno, and M. Starnini, "Link prediction in multiplex networks via triadic closure," *Phys. Rev. Res.*, vol. 2, no. 4, p. 42029, 2020.
- [40] M. Koptelov, A. Zimmermann, B. Crémilleux, and L. Soualmia, "Link prediction via community detection in bipartite multi-layer graphs," in *Proceedings of the 35th Annual ACM Symposium on Applied Computing*, 2020, pp. 430–439.
- [41] S. Najari, M. Salehi, V. Ranjbar, and M. Jalili, "Link prediction in multiplex networks based on interlayer similarity," *Phys. A Stat. Mech. its Appl.*, vol. 536, p. 120978, 2019.
- [42] T. Fan, S. Xiong, W. Zhao, and T. Yu, "Information spread link prediction through multi-layer of social network based on trusted central nodes," *Peer-to-Peer Netw. Appl.*, vol. 12, no. 5, pp. 1028–1040, 2019.
- [43] M. Tarrés-Deulofeu, A. Godoy-Lorite, R. Guimera, and M. Sales-Pardo, "Tensorial and bipartite block models for link prediction in layered networks and temporal networks," *Phys. Rev. E*, vol. 99, no. 3, p. 32307, 2019.
- [44] Z. Samei and M. Jalili, "Application of hyperbolic geometry in link prediction of multiplex networks," *Sci. Rep.*, vol. 9, no. 1, pp. 1–11, 2019.
- [45] M. Koptelov, A. Zimmermann, and B. Crémilleux, "Link prediction in multi-layer networks and its application to drug design," in *International Symposium on Intelligent Data Analysis*, 2018, pp. 175–187.
- [46] H. Mandal, M. Mirchev, S. Gramatikov, and I. Mishkovski, "Multilayer link prediction in online social networks," in *2018 26th Telecommunications Forum (TELFOR)*, 2018, pp. 1–4.
- [47] Y. Yasami and F. Safaei, "A novel multilayer model for missing link prediction and future link forecasting in dynamic complex networks," *Phys. A Stat. Mech. its Appl.*, vol. 492, pp. 2166–2197, 2018.
- [48] P. Bródka, A. Chmiel, M. Magnani, and G. Ragozini, "Quantifying layer similarity in multiplex networks: a systematic study," *R. Soc. open Sci.*, vol. 5, no. 8, p. 171747, 2018.
- [49] C. De Bacco, E. A. Power, D. B. Larremore, and C. Moore, "Community detection, link prediction, and layer interdependence in multilayer networks," *Phys. Rev. E*, vol. 95, no. 4, p. 42317, 2017.
- [50] M. Jalili, Y. Orouskhani, M. Asgari, N. Alipourfard, and M. Perc, "Link prediction in multiplex online social networks," *R. Soc. open Sci.*, vol. 4, no. 2, p. 160863, 2017.
- [51] Y. Yao *et al.*, "Link prediction via layer relevance of multiplex networks," *Int. J. Mod. Phys. C*, vol. 28, no. 08, p. 1750101, 2017.
- [52] G.-K. Mbogo, A. Visheratin, and S. Rakitin, "Layer-Wise Model Stacking for Link Prediction in Multilayer Networks. Case of Scientific Collaboration Networks," in *International Conference on Complex Networks and their Applications*, 2017, pp. 117–126.
- [53] D. Jin, M. Wang, and Y.-R. Lin, "TeleLink: Link Prediction in Social Network Based on Multiplex Cohesive Structures," in *International Conference on*
- [20] S. Bai, Y. Zhang, L. Li, N. Shan, and X. Chen, "Effective link prediction in multiplex networks: A TOPSIS method," *Expert Syst. Appl.*, vol. 177, p. 114973, 2021.
- [21] M. Óskarsdóttir and C. Bravo, "Multilayer network analysis for improved credit risk prediction," *Omega*, vol. 105, p. 102520, 2021.
- [22] H. Luo, L. Li, Y. Zhang, S. Fang, and X. Chen, "Link prediction in multiplex networks using a novel multiple-attribute decision-making approach," *Knowledge-Based Syst.*, vol. 219, p. 106904, 2021.
- [23] L. Yu, D. Zhou, L. Gao, and Y. Zha, "Prediction of drug response in multilayer networks based on fusion of multiomics data," *Methods*, vol. 192, pp. 85–92, 2021.
- [24] S. Huang, Q. Ma, C. Yang, and Y. Yao, "Link Prediction with Multiple Structural Attentions in Multiplex Networks," in *2021 International Joint Conference on Neural Networks (IJCNN)*, 2021, pp. 1–9.
- [25] A. Ficara, G. Fiumara, P. De Meo, and S. Catanese, "Multilayer Network Analysis: The Identification of Key Actors in a Sicilian Mafia Operation," in *International Conference on Future Access Enablers of Ubiquitous and Intelligent Infrastructures*, 2021, pp. 120–134.
- [26] D. Mohapatra, "A hybrid approach for pair-wise layer similarity in a multiplex network," *Soc. Netw. Anal. Min.*, vol. 11, no. 1, pp. 1–10, 2021.
- [27] F. Karimi, S. Lotfi, and H. Izadkhah, "Community-guided link prediction in multiplex networks," *J. Informetr.*, vol. 15, no. 4, p. 101178, 2021.
- [28] D. Malhotra and R. Goyal, "Supervised-learning link prediction in single layer and multiplex networks," *Mach. Learn. with Appl.*, vol. 6, p. 100086, 2021.
- [29] E. Nasiri, K. Berahmand, and Y. Li, "A new link prediction in multiplex networks using topologically biased random walks," *Chaos, Solitons & Fractals*, vol. 151, p. 111230, 2021.
- [30] A. Rezaeipanaah, G. Ahmadi, and S. S. Matoori, "A classification approach to link prediction in multiplex online ego-social networks," *Soc. Netw. Anal. Min.*, vol. 10, no. 1, pp. 1–16, 2020.
- [31] D. Malik and A. Singh, "Link prediction in multilayer networks," *Int. J. Bus. Intell. Data Min.*, vol. 16, no. 4, pp. 490–505, 2020.
- [32] R. Tang, S. Jiang, X. Chen, H. Wang, W. Wang, and W. Wang, "Interlayer link prediction in multiplex social networks: an iterative degree penalty algorithm," *Knowledge-Based Syst.*, vol. 194, p. 105598, 2020.
- [33] A. M. Abdolhosseini-Qomi, S. H. Jafari, A. Taghizadeh, N. Yazdani, M. Asadpour, and M. Rahgozar, "Link prediction in real-world multiplex networks via layer reconstruction method," *R. Soc. open Sci.*, vol. 7, no. 7, p. 191928, 2020.
- [34] A. M. Abdolhosseini-Qomi, N. Yazdani, and M. Asadpour, "Overlapping communities and the prediction of missing links in multiplex networks," *Phys. A Stat. Mech. its Appl.*, vol. 554, p. 124650, 2020.
- [35] B. Lee, S. Zhang, A. Poleksic, and L. Xie, "Heterogeneous multi-layered network model for omics data integration and analysis," *Front. Genet.*, vol. 10, p. 1381, 2020.
- [36] F. Tang, "Analysis of the Influence of Interlayer Structure Information of Multi-Layer Networks on Link Prediction," 2020.

Social Computing, Behavioral-Cultural Modeling and Prediction and Behavior Representation in Modeling and Simulation, 2016, pp. 174–185.

- [54] D. Hristova, A. Noulas, C. Brown, M. Musolesi, and C. Mascolo, “A multilayer approach to multiplexity and link prediction in online geo-social networks,” *EPJ Data Sci.*, vol. 5, no. 1, p. 24, 2016.
- [55] A. Hajibagheri, G. Sukthankar, and K. Lakkaraju, “A holistic approach for predicting links in coevolving multiplex networks,” in *2016 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining (ASONAM)*, 2016, pp. 1079–1086.
- [56] S. Sharma and A. Singh, “An efficient method for link prediction in weighted multiplex networks,” *Comput. Soc. networks*, vol. 3, no. 1, pp. 1–17, 2016.
- [57] A. Hajibagheri, G. Sukthankar, and K. Lakkaraju, “A holistic approach for link prediction in multiplex networks,” in *International conference on social informatics*, 2016, pp. 55–70.
- [58] S. Sharma and A. Singh, “An efficient method for link prediction in complex multiplex networks,” in *2015 11th International Conference on Signal-Image Technology & Internet-Based Systems (SITIS)*, 2015, pp. 453–459.
- [59] P. Wills and F. G. Meyer, “Metrics for graph comparison: a practitioner’s guide,” *PLoS One*, vol. 15, no. 2, p. e0228728, 2020.
- [60] M. Tantardini, F. Ieva, L. Tajoli, and C. Piccardi, “Comparing methods for comparing networks,” *Sci. Rep.*, vol. 9, no. 1, p. 17557, 2019.
- [61] M. K. Khouzani and S. Sulaimany, “Identification of the effects of the existing network properties on the performance of current community detection methods,” *J. King Saud Univ. - Comput. Inf. Sci.*, Apr. 2020.



نهمین کنگره مشترک سیستم‌های فازی و هوشمند ایران

۱۱-۱۳ اسفندماه ۱۴۰۰، مجتمع آموزش عالی بوم

نقش پایه‌ای اینترنت اشیا در مدیریت هوشمند فرآیندهای صنعت کشاورزی

فرید شریف مقدم^۱، سید علیرضا بشیری موسوی^۲

^۱ دانشگاه بین‌المللی امام خمینی (ره) - مرکز آموزش عالی فنی و مهندسی بوئین زهرا، farid.shmoghaddam@gmail.com

^۲ دانشگاه بین‌المللی امام خمینی (ره) - مرکز آموزش عالی فنی و مهندسی بوئین زهرا، a.bashiri@bzeng.ikiu.ac.ir

چکیده: با توجه به افزایش جمعیت در جهان و افزایش نیاز انسان‌ها به محصولات کشاورزی، استفاده از روش سنتی در صنعت کشاورزی قادر به بهبود و افزایش تولید محصولات نیست و نمی‌تواند مشکل عرضه و تقاضا را با توجه به افزایش نرخ جمعیتی موجود حل کند. با توجه به افزایش روزانه‌ی استفاده از اینترنت اشیا و اثبات کارایی این فناوری در صنایع مختلف و توسعه‌ی هر روزه آن، می‌توان از این فناوری استفاده کرد تا با کمترین مشارکت انسانی بتوان تولید محصولات کشاورزی را بالا برد در حالی که منابع مورد نیاز این صنایع را کاهش داد، هزینه‌های مربوطه را تا حد ممکن پایین آورد و نظارت را بسیار ساده‌تر کرد. اما در این بین موانعی مانند هزینه‌ی بالای پیاده‌سازی و عدم آشنایی افراد با فناوری و کامپیوتر وجود دارد که مانع استفاده از این فناوری در صنعت کشاورزی می‌شود. در این مقاله هدف ما این است که راهکارهای مختلف اینترنت اشیا در کشاورزی را ارائه دهیم. علاوه بر این به سیستم‌های پیاده‌سازی شده مبتنی بر اینترنت اشیا در این صنعت می‌پردازیم و پروتکل‌های مختلف استفاده شده در این فناوری را مورد بررسی قرار می‌دهیم.

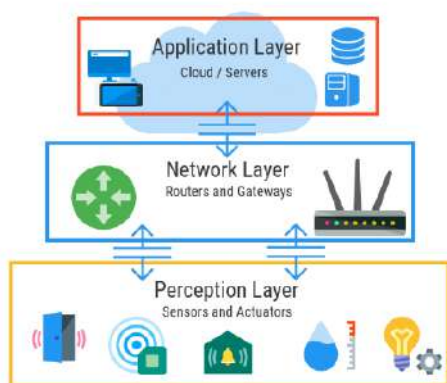
کلمات کلیدی: اینترنت اشیا، حسگر، کشاورزی هوشمند، کشاورزی دقیق

۱- مقدمه

بخش اتوماسیون کاری استفاده می‌شود و توانسته نقش موثر خود را در صنایع مختلف نشان دهد فناوری اینترنت اشیا (IoT) است. اینترنت اشیا با استفاده از اتوماسیون و خودکارسازی امور می‌تواند نقش بسیار موثری در کاهش هزینه‌ها و جایگزینی ماشین به جای افراد داشته باشد در حالی که اطلاعات بسیار بیشتری درباره‌ی محیط مزارع می‌تواند به کشاورزان ارائه دهد. علاوه بر این می‌توان گفت که پیاده‌سازی این فناوری در صنعت کشاورزی به طرز چشمگیری شدت کار و استفاده‌ی بی‌مورد از منابع را کاهش می‌دهد [1]. در صنعت کشاورزی سه ویژگی برای پیاده‌سازی فناوری‌های مختلف مهم است: ۱- نظارت دقیق ۲- راحتی در استفاده ۳- هزینه‌ی پیاده‌سازی پایین، که با استفاده از اینترنت اشیا دستیابی به این ویژگی‌ها با آموزش و پیاده‌سازی درست قابل دستیابی است. علاوه بر این، اینترنت اشیا دلیل پایش مداوم شرایط، رطوبت هوا، خاک و غیره نقش بسزایی در مصرف بهینه‌ی آب دارد که با توجه به کمبود منابع آبی می‌تواند کمک شایانی در این زمینه کند. اینترنت اشیا با استفاده از حسگرهای مختلف، فناوری‌های متفاوت مثل فناوری

کشاورزی یکی از اساسی‌ترین ارکان توسعه‌ی اقتصاد یک کشور است و نقش بسزایی در تولید مواد خام، غذا، گیاهان دارویی و ... دارد. با توجه به نقش موثر این صنعت در هر کشوری، همواره تلاش شده تا کارایی کشاورزی و محصولات آن بیشتر شود و از ضررهای موجود جلوگیری به عمل آید. در بیشتر کشورها کشاورزان به این علم پی برده‌اند که کشاورزی سنتی نمی‌تواند کارایی آنها را بالا برده و محصولات بیشتر و بهتری را برای عرضه به بازار تولید کند. علاوه بر این با توجه به چالش‌های موجود در این صنعت مثل: کمبود آب، حمله‌ی حشرات، آفات و عدم عرضه‌ی مناسب، کشاورزان باید به فکر روش‌های خودکار و اتوماسیون در صنعت کشاورزی باشند و بتوانند با استفاده از فناوری‌های مختلف موجود، علاوه بر بهره‌وری بیشتر و عرضه‌ی بهتر محصولات، هزینه‌های مربوط در همه مراحل را تا حد ممکن کاهش دهند و استفاده‌ی بصره‌ای از منابع موجود با توجه به محدودیت‌های ممکن داشته باشند. یکی از بهترین فناوری‌های موجود که در صنایع مختلف در

را چنین تعریف می‌کند: اینترنت اشیا محیطی است که در آن اشیا، حیوانات یا افراد هرکدام شناسه‌های منحصر به فردی دارند که قادر به انتقال داده‌ها از طریق اینترنت بدون نیاز به تعامل انسان با انسان و یا رایانه است. این تعاریف نشان می‌دهند که اینترنت اشیا با استفاده از فناوری‌های مختلف، می‌تواند اشیا پیرامون ما را روشن کند و خودکار کند تا بتواند با استفاده از اینترنت براساس تغییر شرایط محیطی و یا کنترل دستی از راه دور نسبت به اشیا مختلف و مولفه‌های محیطی آگاهی داشته و نسبت به این مولفه‌ها واکنش مناسبی داشته باشد. معماری اینترنت اشیا انواع گوناگونی دارد از جمله: معماری ۳ لایه، ۴ لایه و ۵ لایه که بدلیل اختصار تنها لایه‌های معماری ۳ لایه را بیان می‌کنیم که از ۳ جز: ۱- لایه ادراک (جمع‌آوری اطلاعات از حسگرها) ۲- لایه شبکه (انتقال اطلاعات جمع‌آوری شده) ۳- لایه کاربردی (ذخیره‌سازی اطلاعات و استفاده از آنها) تشکیل شده است [۱۰].



شکل ۱: معماری ۳ لایه‌ی اینترنت اشیا [۱۱]

۳- چالش‌ها و فرصت‌های اینترنت اشیا در صنعت کشاورزی

همانطور که اشاره کردیم استفاده از فناوری‌های گوناگون مثل اینترنت اشیا در صنعت کشاورزی مزایا و معایب‌های خاص خود را نیز دارد که در این بخش به آنها می‌پردازیم. اصلی‌ترین مشکل و چالش پیاده‌سازی اینترنت اشیا در صنعت کشاورزی هزینه‌های مربوط به پیاده‌سازی این فناوری در مزارع است. برای پیاده‌سازی اولیه‌ی این فناوری کشاورز باید هزینه‌های زیادی از جمله: حسگرهای مورد نیاز مثل حسگر رطوبت‌سنج هوا، رطوبت‌سنج خاک، دماسنج، حسگر بررسی pH خاک و ...، هزینه‌ی اینترنت، هزینه‌ی سرور و فضای ابری برای ذخیره‌سازی و کنترل از راه دور، پمپ آبرسانی، پمپ تزریق کود، مخزن آب، مخزن کود و ... را پرداخت کند تا بتواند از این فناوری در مزرعه‌ی خود استفاده کند که با توجه به عدم اطلاع کشاورز از مزایای پیاده‌سازی این فناوری و کاهش هزینه‌ها در آینده، تصمیم به استفاده از این فناوری نمی‌گیرد. مشکل دیگر استفاده از این فناوری در این صنعت، استفاده‌ی اجباری از نیروی کار متخصص است زیرا نگهداری سیستم مدیریت و تحلیل داده‌ها نیاز به تعمیر و

شناسایی مبتنی بر فرکانس رادیویی^۱، بلوتوث، موقعیت یاب، اینترنت و ... می‌تواند آمار و اطلاعات مفیدی مثل پارامترهای محیطی در مزارع را ذخیره کرده، به کشاورز اطلاع دهد و یا با توجه به شرایط، اقدامات لازم را در مزرعه انجام دهد که این موضوع باعث کنترل و نظارت دقیق و آنی، بهبود تولید محصولات و مدیریت بهتر مزارع می‌شود. به‌طور مثال این فناوری این امکان را به کشاورزان می‌دهد تا بتوانند از راه دور مزارع خود را آبیاری کنند [۲]. علاوه بر این، مزارع هوشمند با استفاده از فناوری‌های مختلف می‌توانند در بخش پیش‌بینی، عرضه، تقاضای محصولات و مدیریت زمان برای انبارداری و نگهداری آنها کمک بسزایی کند [۳]. در این مقاله مقدمه‌ای بر اینترنت اشیا در صنعت کشاورزی، کارهای انجام شده در بحث پیاده‌سازی اینترنت اشیا در این صنعت و توپولوژی‌های مختلف در بحث پیاده‌سازی این فناوری را بیان خواهیم کرد.

اینترنت اشیا امروزه بدلیل اثبات کارآمدی آن در صنایع مختلف تقریباً در همه‌ی صنایع نفوذ کرده و در حال پیشرفت است و باعث می‌شود که کارهایی که در گذشته پیچیده و مدیریت آنها سخت بوده را آسان کرده و علاوه بر این جزئیات همه‌ی مولفه‌هایی که برای هر سیستمی مهم است را نظارت و ضبط کنیم. به‌طور مثال، در مقاله‌ی [۴] الشهري^۲ و همکارانش، با استفاده از اینترنت اشیا و رایانش ابری به پیاده‌سازی سیستم مدیریت تصفیه خانه‌های آب پرداختند که با استفاده از حسگرهای متفاوت داده‌های تصفیه خانه‌ها را جمع‌آوری کرده و با استفاده از این دو فناوری، پایش مداوم و دقیقی روی داده‌ها انجام داده و براساس پایش داده‌ها بهینه‌ترین و کارآمدترین رویکرد را برای عملیات نمک‌زدایی ارائه می‌کند. در تحقیق [۵] یان ایی^۳ و همکارش معماری سیستم مدیریت اطلاعات^۴ را مورد بررسی قرار داده و بر اساس آن یک نمونه از سیستم‌های مدیریت و تحلیل داده‌های کشاورزی را پیاده‌سازی کردند. در پژوهش [۶] استوسس^۵ به همراه یک گروه تحقیقاتی، درباره‌ی تجهیزات اینترنت اشیا، پلتفرم‌های قابل اجرا، استانداردهای موجود در اینترنت اشیا و توپولوژی‌های مختلف استفاده شده در این فناوری بخصوص در بخش صنعت کشاورزی در کشور چک، برای پیاده‌سازی سیستم‌های مزرعه هوشمند تحقیق کردند و در قالب یک مقاله ارائه دادند. در تحقیق [۷] راجو^۶ و همکارش مقدمه‌ای بر اینترنت اشیا و کاربردهای آن در کشاورزی ارائه کردند که در آن فناوری‌های بیسیم در اینترنت اشیا و معماری‌های استفاده شده در آن بیان شده است. در مقاله‌ی [۸] محمود عباسی به همراه جمعی از محققین به بررسی مولفه‌های اساسی اینترنت اشیا، چالش‌ها و کاربردهای آن در کشاورزی و مقایسه فناوری‌های استفاده شده در آن پرداختند.

۲- مفهوم اینترنت اشیا

به‌طور کلی تعریف دقیقی از اینترنت اشیا که مورد تأیید همگان باشد وجود ندارد. در مقاله [۹] اینترنت اشیا چنین تعریف شده است: شبکه‌ای باز و جامع از اشیا هوشمند که ظرفیت سازماندهی خودکار، اشتراک‌گذاری اطلاعات، داده‌ها و منابع را در مواجهه با موقعیت‌ها و تغییرات محیطی دارد. گلوهاک^۷ نیز اینترنت اشیا

^۵ Stočes

^۶ Raju

^۷ Gluhak

^۱ RFID

^۲ Alshehri

^۳ Yan-e

^۴ MIS

نگهداری مداوم دارد تا بتوان بالاترین کارایی را از سیستم انتظار داشت که این نیز خود علاوه بر هزینه‌بر بودن، نیاز به نیروی کار متخصص نیز دارد که به دلیل کمبود نیروی کار متخصص در این زمینه بخصوص در برخی از کشورها، استفاده از این فناوری را با مشکل مواجه کرده است. اما در مقابل این فناوری مزایای بسیار زیادی دارد مثل: دقت بیشتر سیستم مبتنی بر اینترنت اشیا نسبت به کشاورزی سنتی، جمع‌آوری اطلاعات بسیار مفید و دقیق بدون نیاز به نیروی انسانی، پایش مداوم داده‌ها و عملکرد خودکار سیستم نسبت به بررسی صورت گرفته، دسترسی به سیستم از راه دور، بهبود کیفیت محصولات و افزایش تولید، امنیت بیشتر نیروی انسانی در شرایط خطرناک و پر ریسک مثل سم‌پاشی، امکان مقایسه و خروجی گرفتن از داده‌های مزرعه بصورت خودکار و آبی، امکان پیش‌بینی تولید، عرضه و تقاضای محصولات، امکان مدیریت بهینه‌تر در انبارداری و سردخانه‌های موجود در مزارع و ... همه‌ی اینها نشان می‌دهد که معایب موجود در استفاده از این فناوری در مقایسه با مزایای آن قابل چشم پوشی است و نسبت سود به زیان بالایی دارد.

۴- فناوری‌های مورد استفاده در اینترنت اشیا

فناوری‌ها و پروتکل‌های مورد استفاده در اینترنت اشیا همگی باید دارای یک ویژگی واحد باشند که آن ویژگی مصرف پایین انرژی است. همه مولفه‌ها، فناوری‌ها، توپولوژی‌های مورد استفاده‌ی شبکه در اینترنت اشیا همگی باید طوری طراحی شوند که در پایین‌ترین حالت مصرف انرژی قرار داشته باشند. در گذشته برای انتقال اطلاعات در اینترنت اشیا از فناوری‌های متفاوتی مثل: وای فای، بلوتوث، GSM، LTE و ... استفاده می‌شد که هدف از ایجاد آنها استفاده در فناوری اینترنت اشیا نبود و هدف‌های متفاوتی داشتند. اما بدلیل مصرف بسیار زیاد انرژی این فناوری‌ها در بخش انتقال اطلاعات و روی کار آمدن و استفاده از اینترنت اشیا در صنایع مختلف، طولی نکشید تا فناوری‌ها و پروتکل‌های جدید با هدف انتقال اطلاعات در اینترنت اشیا با مصرف خیلی پائین‌تر از فناوری‌های گذشته، خود را جایگزین فناوری‌های قبل در اینترنت اشیا کنند. به‌طور مثال برای ارتباطات دوربرد، استفاده از شبکه‌های سلولی بسیار مناسب است اما این نوع شبکه‌ها با اینکه فضای بسیار مناسبی از انتقال اطلاعات و ارتباطات را برای ما مهیا می‌کنند، اما مصرف انرژی آنها برای استفاده در فناوری اینترنت اشیا بسیار زیاد است [۱۲]. مصرف انرژی در فناوری‌های جدید بسیار پایین است به‌طوری که می‌توان گفت در آینده‌ای نزدیک تجهیزات و مولفه‌های استفاده شده در سیستم‌های مبتنی بر اینترنت اشیا با استفاده از یک باتری ساده تا سال‌ها و دهه‌ها به کار خود ادامه می‌دهند و بالاترین بخش مصرف انرژی در این سیستم‌ها برای بخش حسگرها خواهد بود و نه بحث جمع‌آوری و انتقال اطلاعات [۶]. این در حالیست که فناوری‌های گذشته بیشترین مصرف انرژی را در بخش شبکه و انتقال اطلاعات از حسگرها به سرورها داشتند که برای انتقال داده‌ها بصورت شبانه‌روزی و پایش مداوم تعداد زیاد حسگرها این اتفاق هزینه‌ی تجهیزات و نگهداری بالایی را به کشاورز تحمیل می‌کند. علاوه بر مصرف پایین انرژی، امنیت داده‌ها، تضمین

ارسال و دریافت داده بصورت صحیح، پهنای باند و دامنه‌ی مناسب نسبت به سیستم مورد نیاز نیز در اینترنت اشیا اهمیت دارد [۱۳]. دستگاه‌های مورد استفاده در اینترنت اشیا می‌توانند از طریق شبکه‌های IP و شبکه‌های غیر IP به سیستم متصل شوند که هر کدام مزایا و معایب خود را دارند. به‌طور مثال اتصال دستگاه‌ها از طریق شبکه‌های IP دشوار و نیازمند حافظه و قدرت زیادی هستند که به همین دلیل مصرف خیلی زیادی دارند، اما محدودیتی بابت دامنه‌ی ارسال و دریافت اطلاعات ندارند در حالی که اتصال از طریق شبکه‌های غیر IP مثل بلوتوث، مصرف انرژی پائین‌تری دارند در حالیکه محدودیت دامنه برای ارسال و دریافت اطلاعات هنگام استفاده از این روش داریم.

از پروتکل‌های امروزی مورد استفاده در اینترنت اشیا با مصرف کم می‌توان به: شبکه‌ی گسترده‌ی دوربرد^۸، بلوتوث کم مصرف^۹، Zigbee، SigFox، IEEE P802.11ah، ارتباط از نزدیک^{۱۰}، OPC-UA، حمل و نقل تله‌متری MQ^{۱۱}، Dash 7 Alliance Protocol 1.0، اینترنت اشیا باند باریک^{۱۲} و موج Z اشاره کرد که هر کدام برای اهداف مشخصی استفاده می‌شود. در ادامه به برخی از پروتکل‌هایی که بیشتر در فناوری اینترنت اشیا در صنعت کشاورزی مورد استفاده قرار می‌گیرند، می‌پردازیم.

۱-LoRaWAN-۴

همانطور که اشاره شد پروتکل‌هایی که در گذشته برای اینترنت اشیا استفاده می‌شد مصرف انرژی بالایی داشتند. بنابراین برای رفع این مشکل پروتکل‌های متفاوتی ارائه شد. یکی از این پروتکل‌ها پروتکل شبکه‌ی گسترده‌ی دوربرد است که بر پایه‌ی لایه فیزیکی دوربرد^{۱۳} قابل پیاده‌سازی است که به عنوان مناسب‌ترین پروتکل در بین پروتکل‌های شبکه گسترده کم توان^{۱۴} برای استفاده از اینترنت اشیا در صنعت کشاورزی شناخته می‌شود. این پروتکل اطلاعات را بصورت نامحدود با حداکثر ۲۴۳ بایت برای هر پیام ارسال می‌کند و برد آن در مناطق روستایی به ۲۰ کیلومتر و در مناطق شهری به ۵ کیلومتر می‌رسد [۱۴]. این پروتکل با در نظر گرفتن تعداد حسگرهای استفاده شده در سیستم، یک لایه‌ی پیوند داده با برد طولانی، مصرف انرژی پایین و نرخ بیت پایین را برای سیستم فراهم می‌آورد. استفاده از این پروتکل در استفاده از اینترنت اشیا راه حلی بسیار امیدوار کننده برای استفاده بهینه با مصرف انرژی خیلی پایین ارائه داده که می‌توان با استفاده از آن تعداد بسیار زیادی دستگاه یا حسگر را به شبکه متصل کرد و مشکلی بابت مصرف انرژی و مشکلات انتقال اطلاعات نداشت [۱۳].

این فناوری به دلیل دوربرد بودن، مصرف بسیار پایین انرژی، قدرت مناسب برای ارسال و دریافت اطلاعات در اعماق زمین مثل دستورات مربوط به خطوط لوله‌ی آب برای تنظیم رطوبت خاک، پایین بودن هزینه‌ی سخت افزار مورد نیاز بدلیل عدم نیاز به سخت افزار قوی و اتصال تعداد زیادی حسگر در شبکه در صنعت کشاورزی بخصوص در مزارع با وسعت بالا بسیار کاربرد دارد.

^{۱۲} NB-IoT

^{۱۳} LoRa

^{۱۴} LPWAN

^۸ LoRaWAN

^۹ BLE

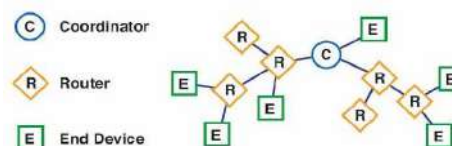
^{۱۰} NFC

^{۱۱} MQTT

۲-۴- Zigbee

Zigbee یک روش اتصال بیسیم با برد کوتاه (۱۰۰ متر)، مصرف بسیار پایین انرژی و نرخ انتقال پایین تا ۲۵۰ کیلوبیت بر ثانیه است که مبتنی بر استاندارد IEEE 802.15.4 و فرکانس ۲,۴ گیگاهرتز عمل می‌کند [۱۵]. این پروتکل یکی از بهترین و کم مصرف‌ترین روش‌های اتصال مولفه‌ها و حسگرها برای پیاده‌سازی سیستم‌های مبتنی بر اینترنت اشیا است. این فناوری برای ارتباط محلی بین دستگاه‌های اینترنت اشیا بسیار مناسب است که می‌توان با ترکیب آن با توپولوژی مش^{۱۵} پیاده‌سازی بهینه با مصرف پایین انرژی به همراه انتقال اطلاعات بصورت امن و سریع را تجربه کرد. علاوه بر این، این فناوری می‌تواند با تراکم بالای گره‌های نهایی و حسگرهای زیاد با مقیاس‌پذیری بالایی به کار خود ادامه دهد و هیچ تداخلی در پیاده‌سازی سیستم بوجود نیارد. با توجه به اینکه این فناوری از توپولوژی مش پشتیبانی می‌کند می‌توانیم با استفاده از این توپولوژی، شبکه گسترده‌تری با گره‌ها و حسگرهای بسیار زیاد برای سیستم مورد نیاز در مزارع نسبتاً بزرگ پیاده‌سازی کنیم که بدلیل مصرف بسیار پایین انرژی این نوع اتصال در اینترنت اشیا نیاز به شارژ مداوم و هزینه‌های اضافی بابت نگهداری مولفه‌ها در این نوع سیستم‌ها نیست. بنابراین با توجه به ویژگی‌های گفته شده این شبکه و چالش‌هایی که در استفاده از این فناوری گفته شد می‌توان گفت که بطرز چشمگیری این چالش‌ها و مشکلات را حل کرده است که باعث افزایش کارایی و استفاده‌ی بیشتر از این فناوری در صنایع مختلف بخصوص صنعت کشاورزی می‌شود.

این پروتکل بدلیل برد کوتاهی که دارد برای مزارع با وسعت زیاد مناسب نیست اما بدلیل پشتیبانی از توپولوژی مش می‌تواند تعداد زیادی حسگر را بصورت یکپارچه به سیستم متصل کرد و سرعت بالا برای ارسال و دریافت اطلاعات و تراکم بالای حسگرها را برای پیاده‌سازی این سیستم در مزارع متصور شد.



شکل ۲: ساختار کلی شبکه Zigbee با توپولوژی مش [۱۶]

۳-۴- NB-IoT

پروتکل اینترنت اشیا باند باریک یکی از فناوری‌های سلولی دوبرد و نسبتاً کم مصرف است که از مدولاسیون تقسیم فرکانس عمود برهم^{۱۶} استفاده می‌کند و به منظور استفاده در فناوری اینترنت اشیا ساخته شده است. استفاده از این شبکه مزایای زیادی علاوه بر مصرف پایین انرژی نسبت به سایر شبکه‌های سلولی دارد. علاوه بر این بدلیل ساختار متفاوت این شبکه و نرخ بیت پایین آن، نیاز به پیاده‌سازی گیتوی نداشته که خود علاوه بر هزینه‌ی پایین‌تر، پیاده‌سازی این شبکه را بدلیل اتصال مستقیم حسگرها به ایستگاه اصلی شبکه راحت‌تر می‌کند. فرکانس کاری این شبکه برای ارتباط حداقل ۱۸۰ کیلوهرتز است که این امکان را می‌دهد تا بتوان با GSM با فرکانس ۲۰۰ کیلوهرتز به راحتی ارتباط برقرار کرد. علاوه بر این می‌توان با استفاده از روش‌های مختلف مثل تطبیق نرخ بیت و تقسیم فرکانس و ... از فناوری LTE نیز به‌طور مناسبی استفاده کرد که این

امکان کمک می‌کند تا علاوه بر پایین آمدن هزینه‌ی پیاده‌سازی سیستم‌های مبتنی بر اینترنت اشیا بخاطر عدم سازگاری فرکانس کاری ارتباطات مولفه‌ها، زمان مورد نیاز برای تطبیق و پیاده‌سازی ارتباطات نیز به طرز چشمگیری کاهش یافته و امنیت و یکپارچگی اطلاعات بدلیل وجود پروتکل‌های امنیتی در شبکه تلفن همراه افزایش یابد.

این پروتکل بدلیل فرکانس ارتباطی خود و راحتی ارتباط از طریق تلفن همراه با سیستم‌های مبتنی بر این شبکه می‌تواند کاربرد زیادی در مزارع بزرگ داشته باشد، به‌طوری که توسعه‌ی مولفه‌های استفاده شده در این نوع شبکه سریع‌تر و آسان‌تر است و ارسال و دریافت اطلاعات می‌تواند به واسطه‌ی تلفن همراه انجام شود که این امکان راحتی استفاده در سیستم‌های پیاده‌سازی شده توسط این پروتکل را به همراه خواهد داشت.

۴-۴- MQTT

حمل و نقل تله‌متری MQ یکی از معروف‌ترین پروتکل‌های مورد استفاده در فناوری اینترنت اشیا است که در زمینه‌های مختلف کاربردهای متفاوتی دارد. این پروتکل از پروتکل‌های سبک و ساده در پیاده‌سازی است که می‌توان در زمینه‌های مختلف مثل ذخیره‌سازی داده‌ها با اندازه کم و زیاد و ارتباط دو طرفه (با از بین نرفتن داده‌ها در طی زمان) استفاده کرد. این پروتکل که نوعی پروتکل ماشین به ماشین^{۱۷} است معمولاً به روی پروتکل TCP/IP پیاده‌سازی و استفاده می‌شود و براساس ساختار مشتری-سرور عمل می‌کند. علاوه بر این در این پروتکل مفهومی به اسم انتشار/اشتراک^{۱۸} وجود دارد که ساختار ارسال و دریافت اطلاعات در سیستم‌های مبتنی بر اینترنت اشیا با استفاده از این پروتکل را تعیین می‌کند. حداقل اندازه‌ی داده در این پروتکل برای ارسال ۲ بایت و حداکثر ۲۵۶ مگابایت است. سرآیند فایل‌های ارسالی در این پروتکل بسیار سبک بوده که باعث کاهش ترافیک موجود در شبکه و سادگی در تبادل اطلاعات می‌شود. این پروتکل معمولاً زمانی استفاده می‌شود که دستگاه‌های ما حافظه و باتری کمی دارند و می‌خواهیم از فناوری اینترنت اشیا در محیطی با وسعت کم با احتمال تاخیر زیاد در ارسال داده‌ها استفاده کنیم. از ویژگی‌های این پروتکل می‌توانیم به مصرف پایین انرژی، پهنای باند پایین و کوتاه برد بودن آن اشاره کرد. لازم به ذکر است که ارسال و دریافت اطلاعات در این پروتکل به روش سوکت انجام می‌شود.

این پروتکل در صنایع کشاورزی برای مزارع با وسعت کم مناسب است که حریم امنیتی پایینی دارد و امکان دسترسی غیرمجاز به سیستم زیاد است و فضای ذخیره سازی و باتری دستگاه‌های استفاده شده در سیستم از حجم پایینی برخوردارند و نیاز به تشکیل صف اطلاعات در ورودی دستگاه‌ها برای از بین نرفتن داده‌ها و ارتباط دو طرفه است.

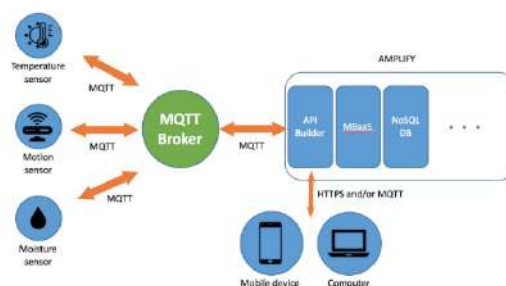
^{۱۷} M2M

^{۱۸} Publish/Subscribe

^{۱۵} Mesh

^{۱۶} OFDM

نرم افزار موبایلی ارائه شده که دو حالت را برای دسترسی کشاورز برای آبیاری مزارع ارائه می‌دهد. ۱- حالت خودکار برای آبیاری محصولات با توجه به میزان رطوبت و دمای محیط ۲- حالت دستی برای آبیاری محصولات با توجه به نظر کشاورز. این برنامه به صورتی عمل می‌کند که اگر کشاورز برنامه‌ای برای دستی براساس برنامه کشاورز عمل می‌کند و اگر بخش داده‌های ورودی خالی باشد سیستم بصورت خودکار برنامه آبیاری را به اجرا در می‌آورد. در پژوهش [۲۰] داگار^{۲۱} و همکارانش مدل پیشنهادی برای ساخت مزارع ارائه کردند که کیفیت محصولات و عملکرد مزارع را بالا می‌برد. در این مدل توصیه شده است که از PolyHouse برای گلخانه‌ها و مزارع استفاده شود. PolyHouse یک اتاقک فولادی است که از پلی‌اتیلن برای لایه حفاظتی روی آن استفاده می‌شود. پلی‌اتیلن موجود در لایه محافظتی سازه باعث جلوگیری از صدمه دیدن محصولات در برابر اشعه‌های مضر خورشید می‌شود و فولادی بودن سازه هنگام بارش باران و تگرگ و طوفان و ... محصولات مزارع را از صدمات احتمالی حفظ می‌کند. علاوه بر این استفاده از حسگرهای دیگر برای استفاده از آنها در سیستم اینترنت اشیا را پیشنهاد داده و الگوریتم‌های ساده‌ای برای پیاده‌سازی حسگرها ارائه کرده‌اند. سیستم پیشنهادی شامل حسگرهای متفاوتی با اهداف گوناگون است. از جمله آنها می‌توان به: حسگر حجم آب درون لوله برای کنترل و پایش میزان آب موجود در لوله برای جلوگیری از هدر رفت آب و جلوگیری از آلوده شدن آن به دلیل باقی ماندن آب به مدت طولانی در خطوط آبیاری محصولات استفاده می‌شود. با استفاده از این حسگر می‌توانیم حجم آب مورد نیاز محصولات و حجم آب ارسالی از پمپ‌های آب را متناسب یکدیگر در نظر گرفته و حجم بیشتری از آب درون لوله‌ها باقی نماند که هم باعث هدر رفت آب نشود و هم کیفیت محصولات را پایین نیاورد. حسگر دیگر پیشنهادی در این سیستم حسگر کنترل Ph خاک برای کنترل میزان اسیدی بودن خاک و کنترل میزان کود مورد نیاز برای محصولات مزارع به تفکیک نوع محصول کشت شده است. با توجه به حساسیت برخی از محصولات به اسیدی بودن خاک، این حسگر کمک بسزایی در زمینه کنترل کیفیت محصول و تعریف دقیق میزان کود مورد نیاز محصولات انجام می‌دهد. حسگر رطوبت خاک نیز با بررسی میزان رطوبت خاک و میزان رطوبت استاندارد مورد نیاز محصول با توجه به اطلاعات تعریف شده در سیستم باعث راه اندازی دستگاه‌های آب‌پاش و یا سیستم‌های گرماساز هوا می‌شود. حسگر دماسنج نیز مانند حسگر رطوبت خاک عمل کرده و با بررسی دمای محیط و دمای استاندارد مورد نیاز محصولات، با دستور به دستگاه‌های تهویه هوا دما را پایین و یا بالا می‌برد. حسگر متفاوتی که در این سیستم پیشنهاد شده حسگر تشخیص حرکت است که برای این منظور در نظر گرفته شده است تا با بررسی حرکات غیرعادی اطراف مزارع با ارسال دستور به بلندگوهای تعبیه شده، باعث ترس حیوانات و فراری دادن آنها شود. در مقاله‌ی [۲۱] بو^{۲۲} و همکارش، سیستمی مبتنی بر اینترنت اشیا و یادگیری تقویتی عمیق ارائه کردند که از ۴ لایه تشکیل شده است. لایه‌های تشکیل دهنده این سیستم عبارتند از: ۱- لایه‌ی جمع‌آوری اطلاعات حسگرها ۲- لایه‌ی محاسبات لایه‌ی ۳- لایه‌ی انتقال اطلاعات جمع‌آوری شده ۴- لایه‌ی محاسبات ابری.



شکل ۳: ساختار کلی پروتکل MQTT [۱۷]

۵- سیستم‌های پیاده‌سازی شده مبتنی بر اینترنت اشیا در مزارع کشاورزی

همانطور که اشاره شد با توجه به گسترش اینترنت اشیا در صنایع مختلف، صنعت کشاورزی نیز از این قاعده مستثنی نیست و سیستم‌های مختلفی در زمینه‌های مختلف در صنعت کشاورزی پیاده‌سازی شده است. به‌طور کلی سیستم‌های پیاده‌سازی شده در این زمینه از ۳ جز اصلی: سخت‌افزار، صفحات وب و برنامه‌ی موبایل تشکیل می‌شوند که براساس نیاز و هدف از استفاده از این فناوری طراحی متفاوتی دارند. به‌طور مثال در مقاله‌ی [۱۸] سرکار^{۱۹} و همکارش با توجه به اهمیت ۳ پارامتر مهم در مزارع کشاورزی یعنی ۱- دمای محیط ۲- رطوبت خاک ۳- رطوبت هوا سیستمی ارائه دادند که با استفاده از میکروکنترلر، داده‌ها را جمع‌آوری کرده و توسط مولفه‌ی esp8266 و پروتکل UART اطلاعات را به سرور منتقل می‌کند و برای نظارت بر داده‌ها نیز توسط صفحات مبتنی بر وب در دستگاه‌های مختلف مثل کامپیوتر و PDA، داده‌ها قابل رهگیری و نظارت هستند. این سیستم این قابلیت را به کشاورز می‌دهد تا در هر لحظه بصورت آنی و لحظه‌ای اطلاعات مهمی از قبیل پارامترهای گفته شده را بررسی کند و در صورت بروز مشکل اقدام مناسب را انجام دهد. در تحقیقی دیگر [۱۹] مویانگ پرات هاب^{۲۰} به همراه گروه تحقیقاتی خود، سیستمی بهینه برای آبیاری محصولات کشاورزی براساس شبکه حسگرهای بیسیم ارائه دادند که با استفاده از داده‌کاوی اطلاعات دریافتی از حسگرها را به‌طور کامل بررسی کرده و اقدام به آبیاری محصولات می‌کند. بخش‌های بیان شده در این سیستم عبارتند از: مولفه‌های سخت‌افزاری استفاده شده در سیستم، صفحات وب مربوط به پایش اطلاعات و تغییر تنظیمات سیستم آبیاری محصولات و برنامه‌ی موبایلی مورد استفاده برای آبیاری محصولات مزارع به منظور کنترل و آبیاری بهینه در زمین‌های کشاورزی. برای پیاده‌سازی این سیستم در بخش سخت‌افزار از جعبه‌ی کنترل ضدآب برای جلوگیری از صدمات احتمالی به دستگاه‌های الکترونیکی سیستم، حسگر DHT22 برای بررسی دما و رطوبت محیط، شیر برقی برای آبیاری محصولات مزارع، حسگر التراسونیک برای کنترل حجم آب درون مخازن و برد NodeMCU برای ارسال و دریافت داده‌ها از سرور از طریق اتصال به وای‌فای استفاده شده است. برای بخش صفحات مبتنی بر وب نیز صفحات مختلفی پیاده‌سازی شده است که می‌توان به صفحه‌ی تنظیم و تغییر میزان آبیاری محصولات بر اساس پارامترهای کنترلی مثل دما و رطوبت خاک اشاره کرد. در قسمت انتهایی نیز

^{۲۱} Daggar

^{۲۲} Bu

^{۱۹} Sarkar

^{۲۰} Muangprathub

- 2011 Fourth International Conference on Intelligent Computation Technology and Automation, 2011, vol. 1: IEEE, pp. 1045-1049 .
- [۶] M. Stočes, J. Vaněk, J. Masner, and J. Pavlík, "Internet of things (iot) in agriculture-selected aspects," *Agris on-line Papers in Economics and Informatics*, vol. 8, no. 665-2016-45107, pp. 83-88, 2016.
- [۷] K. L. Raju and V. Vijayaraghavan, "IoT technologies in agricultural environment: a survey," *Wireless Personal Communications*, vol. 113, no. 4, pp. 2415-2446, 2020.
- [۸] M. Abbasi, M. H. Yaghmaee, and F. Rahnema, "Internet of Things in agriculture: A survey," in 2019 3rd International Conference on Internet of Things and Applications (IoT), 2019: IEEE, pp. 1-12 .
- [۹] S. Madakam, V. Lake, V. Lake, and V. Lake, "Internet of Things (IoT): A literature review," *Journal of Computer and Communications*, vol. 3, no. 05, p. 164, 2015.
- [۱۰] C. Verdouw, H. Sundmaeker, B. Tekinerdogan, D. Conzon, and T. Montanaro, "Architecture framework of IoT-based food and farm systems: A multiple case study," *Computers and Electronics in Agriculture*, vol. 165, p. 104939, 2019.
- [۱۱] A. Calihman, "Architectural frameworks in the IoT civilization," 2019. [Online]. Available: <https://www.netburner.com/learn/architectural-frameworks-in-the-iot-civilization/>.
- [۱۲] K. Mekki, E. Bajic, F. Chaxel, and F. Meyer, "A comparative study of LPWAN technologies for large-scale IoT deployment," *ICT express*, vol. 5, no. 1, pp. 1-7, 2019.
- [۱۳] J. de Carvalho Silva, J. J. Rodrigues, A. M. Alberti, P. Solic, and A. L. Aquino, "LoRaWAN—A low power WAN protocol for Internet of Things: A review and opportunities," in 2017 2nd International Multidisciplinary Conference on Computer and Energy Science (SpliTech), 2017: IEEE, pp. 1-6 .
- [۱۴] B. Miles, E.-B. Bourennane, S. Boucherkha, and S. Chikhi, "A study of LoRaWAN protocol performance for IoT applications in smart agriculture," *Computer Communications*, vol. 164, pp. 148-157, 2020.
- [۱۵] I. Kuzminykh, A. Snihurov, and A. Carlsson, "Testing of communication range in ZigBee technology," in 2017 14th International Conference The Experience of Designing and Application of CAD Systems in Microelectronics (CADSM), 2017: IEEE, pp. 133-136 .
- [۱۶] S. Dhar, M. R. Nayeem, and H. A. Rahman, "Maximum Sun-Light Tracking and Optimum Battery Charge Monitoring for Solar Panel Using ZIGBEE Wireless Sensor Network".
- [۱۷] A. KS, "MQTT With PYTHON — Part 1," 2018. [Online]. Available: <https://medium.com/@ashiqgiga07/mqtt-with-python-part-1-a38e64308c76>
- [۱۸] P. J. Sarkar and S. Chanagala, "A Survey on IOT based Digital Agriculture Monitoring System and Their impact on optimal utilization of Resources," *Journal of Electronics and Communication Engineering (IOSR-JECE)*, vol. 11, no. 1, pp. 1-4, 2016.
- [۱۹] J. Muangprathub, N. Boonnam, S. Kajornkasirat, N. Lekbangpong, A. Wanichsombat, and P. Nillaor, "IoT and agriculture data analysis for smart farm," *Computers and electronics in agriculture*, vol. 156, pp. 467-474, 2019.

همانطور که پیداست این سیستم با ترکیب هوش مصنوعی، محاسبات ابری و اینترنت اشیا سیستمی ارائه می‌دهد که تصمیمات مهم در مزارع را بصورت آنی و لحظه‌ای با استفاده از داده‌های بدست آمده در گذشته و داده‌های ورودی کشاورز می‌گیرد. در این سیستم ابتدا داده‌ها از حسگرها جمع‌آوری می‌شود، سپس اطلاعات در لایه‌ی محاسبات لایه‌ی قرار گرفته و پیش‌پردازش اطلاعات صورت می‌گیرد. هدف از ارسال داده‌ها به لایه‌ی محاسبات لایه‌ی ۱- ناقص نبودن اطلاعات جمع‌آوری شده و کنترل وضعیت هر حسگر در سیستم ۲- فشرده سازی اطلاعات و حذف داده‌های تکراری است. این اتفاق باعث می‌شود تا وضعیت حسگرها و سالم بودن آنها در شبکه بصورت مداوم بررسی شود و علاوه بر این حجم ارسال داده‌ها بطرز چشمگیری کاهش خواهد یافت که باعث کاهش پهنای باند مورد استفاده‌ی شبکه نیز می‌شود. پس از اعمال تغییرات روی داده‌ها در لایه‌ی محاسبات لایه‌ی، اطلاعات از طریق لایه‌ی انتقال به سرورها و فضای ابری تعریف شده در سیستم منتقل می‌شود. در این لایه با استفاده از ترکیب یادگیری عمیق و یادگیری تقویتی به پیاده‌سازی الگوریتم یادگیری عمیق تقویتی پرداخته تا بتوان با استفاده از آن و داده‌های جمع‌آوری شده از حسگرها و لایه‌ی محاسبات لایه‌ی، بصورت هوشمند بهترین تصمیم با سرعت بالا اتخاذ شود و با ارسال دستور مناسب اقدام مناسب صورت گیرد.

۶- نتیجه

در این مقاله، نقش سیستم‌های مبتنی بر فناوری اینترنت اشیا بیان شده است. علاوه بر این با پروتکل‌های پرکاربرد در اینترنت اشیا بخصوص در سیستم‌های پیاده‌سازی شده در صنعت کشاورزی آشنا شدیم و در انتها برخی از سیستم‌های پیاده‌سازی شده برای استفاده در مزارع را بررسی کردیم و به سیستم‌های بهینه در کنترل، مدیریت و اتوماسیون کاری در مزارع مثل آبیاری، تهویه مطبوع و ... در مزارع پرداختیم. با توجه به رشد این فناوری در سراسر جهان می‌توان آینده‌ی خوبی برای استفاده‌ی بهینه و درست از این فناوری در صنعت کشاورزی متصور شد. امیدواریم که در آینده‌ای نه چندان دور با استفاده از این فناوری در صنعت کشاورزی، تولید محصولات بیشتر و با کیفیت‌تری در عین مصرف بهینه انرژی داشته باشیم.

مراجع

- [۱] W. Zhang, "Study about IOT's application in" digital Agriculture" construction," in 2011 International Conference on Electrical and Control Engineering, 2011: IEEE, pp. 2578-2581 .
- [۲] C. Goumopoulos, B. O'Flynn, and A. Kameas, "Automated zone-specific irrigation with wireless sensor/actuator network and adaptable decision support," *Computers and electronics in agriculture*, vol. 105, pp. 20-33, 2014.
- [۳] M. Lee, J. Hwang, and H. Yoe, "Agricultural production system based on IoT," in 2013 IEEE 16th international conference on computational science and engineering, 2013: IEEE, pp. 833-837 .
- [۴] M. Alshehri, A. Bhardwaj, M. Kumar, S. Mishra, and J. Gyani, "Cloud and IoT based smart architecture for desalination water treatment," *Environmental Research*, vol. 195, p. 110812, 2021.
- [۵] D. Yan-e, "Design of intelligent agriculture management information system based on IoT," in

- [४०] R. Dagar, S. Som, and S. K. Khatri, "*Smart farming–IoT in agriculture*," in 2018 International Conference on Inventive Research in Computing Applications (ICIRCA), 2018: IEEE, pp. 1052-1056 .
- [४१] F. Bu and X. Wang, "*A smart agriculture IoT system based on deep reinforcement learning*," *Future Generation Computer Systems*, vol. 99, pp. 500-507, 2019.