

الحمد لله
الرحمن الرحيم



دانشکده مهندسی کامپیوتر

پایان نامه کارشناسی ارشد مهندسی هوش مصنوعی

بهبود دقت قوانین استخراج شده توسط شبکه‌های عصبی چند لایه با بهره‌گیری از الگوریتم تکاملی در آموزش و هرس کردن شبکه

نگارنده: یاسر ایرانی

استاد راهنما

دکتر حمید حسن پور

بهمن ۱۳۹۷

شماره: ۹۱۳ / ۹۷

تاریخ: ۹۷/۱۲/۱۴



دانشگاه علمی کاربردی

باسمه تعالی

مدیریت تحصیلات تکمیلی

فرم شماره (۳) صورتجلسه نهایی دفاع از پایان نامه دوره کارشناسی ارشد

با نام و یاد خداوند متعال، ارزیابی جلسه دفاع از پایان نامه کارشناسی ارشد آقای یاسر ایرانی با شماره دانشجویی ۹۵۰۲۱۵۴ رشته مهندسی کامپیوتر گرایش هوش مصنوعی و رباتیکز تحت عنوان بهبود دقت قوانین استخراج شده توسط شبکه های عصبی چند لایه با بهره گیری از الگوریتم تکاملی در آموزش و هرس کردن شبکه که در تاریخ ۱۳۹۷/۱۱/۷ با حضور هیأت محترم داوران در دانشگاه صنعتی شاهرود برگزار گردید به شرح ذیل اعلام می گردد:

قبول (با درجه: <u>خیلی خوب</u>) ☒	مردود ☐
نوع تحقیق: نظری ☐	عملی ☐

عضو هیأت داوران	نام و نام خانوادگی	مرتبه علمی	امضاء
۱- استاد راهنمای اول	حمید حسن پور	استاد	
۲- نماینده تحصیلات تکمیلی	محسن فرهادی	مربی	
۳- استاد ممتحن اول	محسن بیگلری	استادیار	
۴- استاد ممتحن دوم	هادی گرایلو	استادیار	

نام و نام خانوادگی رئیس دانشکده:

تاریخ و امضاء و مهر دانشکده:

تبصره: در صورتی که کسی مردود شود حداکثر یکبار دیگر (در مدت مجاز تحصیل) می تواند از پایان نامه خود دفاع نماید (دفاع مجدد نباید زودتر از ۴ ماه برگزار شود).

تقدیم بہ

پاکسزار کسافی، ہستم کہ سر آغاز تولد من ہستند۔ از یکی زادہ می شوم و از دیگری جاودانہ۔ استاد می کہ سپیدی را بر تختہ
سیاہ زندگیم مگاشت و مادری کہ تار مویی از او پای من سیاہ نماند۔

تقدیم بہ

مقدس ترین واژہ ہاد لغت نامہ دلم، مادر مہربانم کہ زندگیم را دیون مہر و عطوفت آن می دانم۔

پدر، مہربانی مشفق، بردبار و حامی۔

برادر و خواہرا نم، ہمراہان، ہمیشگی و پشتوانہ های زندگیم۔

سپاسگزاری

سپاس خداوندگار حکیم را که با لطف بی کران خود، آدمی رازیور عقل آراست. از زحمات بی دریغ استاد راهنمای خود، جناب آقای دکتر حمید حسن پور، صمیمانه تشکر و قدردانی کنم که قطعاً بدون راهنمایی های ارزنده ایشان، این مجموعه به انجام نمی رسید.

یاسر ایرانی

بهمن ۱۳۹۷

تعهد نامه

اینجانب یاسر ایرانی دانشجوی دوره کارشناسی ارشد رشته مهندسی کامپیوتر دانشکده مهندسی کامپیوتر دانشگاه صنعتی شاهرود، نویسنده پایان نامه با عنوان **بهبود دقت قوانین استخراج شده توسط شبکه های عصبی چند لایه با بهره گیری از الگوریتم تکاملی در آموزش و هرس کردن شبکه**، تحت راهنمایی حمید حسن پور متعهد می شوم:

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش های دیگر پژوهش گران، به مرجع مورد استفاده استناد شده است.
- مطالب این پایان نامه، تا کنون توسط خود، یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ جا ارایه نشده است.
- حقوق معنوی این اثر، به دانشگاه صنعتی شاهرود تعلق دارد، و مقالات مستخرج با نام “دانشگاه صنعتی شاهرود” یا “Shahrood University of Technology” به چاپ خواهد رسید.
- حقوق معنوی تمام افرادی که در به دست آوردن نتایج اصلی پایان نامه تاثیرگذار بوده اند، در مقالات مستخرج از پایان نامه رعایت می گردد.
- در تمام مراحل انجام این پایان نامه، در مواردی که از موجود زنده (یا بافت های آنها) استفاده شده است، ضوابط و اصول اخلاقی رعایت شده است.
- در تمام مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته (یا استفاده شده است)، اصل رازداری و اصول اخلاق انسانی رعایت شده است.

یاسر ایرانی
بهمن ۱۳۹۷

مالکیت نتایج و حق نشر

- تمام حقوق معنوی این اثر و محصولات آن (مقالات مستخرج، کتاب، برنامه های رایانه ای، نرم افزارها و تجهیزات ساخته شده) متعلق به دانشگاه صنعتی شاهرود می باشد. این مطلب باید به نحو مقتضی، در تولیدات علمی مربوطه ذکر شود.
- استفاده از اطلاعات و نتایج موجود در این پایان نامه بدون ذکر منبع مجاز نمی باشد.

چکیده

امروزه شاهد جمع‌آوری داده‌های فراوانی در زمینه‌های مختلفی از جمله پزشکی، تجاری و صنعتی هستیم. این داده‌ها برای مقاصد گوناگونی گردآوری می‌شوند. یکی از اهداف مهم جمع‌آوری داده‌ها، تحقیق روی آن‌ها برای بدست آوردن نتایج و الگوهای مفید موجود در آن‌ها است. پیچیدگی و حجم زیاد داده‌ها، موجب سردرگمی در تحلیل این داده‌ها شده است. از این رو، روش‌های داده‌کاوی متعددی برای تجزیه و تحلیل داده‌ها ارائه شده است. استخراج قانون یکی از کاربردهای مهم داده‌کاوی برای کشف دانش و روابط بین داده‌ها است. استخراج قانون به کمک ابزارهای مختلفی مانند شبکه‌عصبی، درخت تصمیم و ماشین بردار پشتیبان انجام می‌شود. شبکه‌عصبی، به دلیل عدم نیاز به دانش اولیه در مورد داده‌ها و داشتن دقت بالا، بیشتر مورد توجه قرار گرفته است.

اغلب روش‌های ارائه شده برای استخراج قانون متکی بر هرس‌سازی شبکه هستند، هرس‌سازی باعث کاهش دقت و کارایی شبکه‌عصبی و در نتیجه ضعف در قوانین استخراج شده می‌شود. در این پژوهش روش نوینی ارائه می‌شود که با بهره‌گیری از الگوریتم ژنتیک، وزن‌های نرون‌های شبکه‌عصبی را به گونه‌ای تنظیم می‌کند که خسارت ناشی از مرحله هرس‌سازی را به میزان قابل توجهی تقلیل دهد. این روش در مرحله اول، به هنگام آموزش شبکه‌عصبی و به شکل سراسری و در مرحله بعد، پس از اتمام آموزش به صورت محلی با بررسی وزن‌های هر نرون از لایه میانی و خروجی، به اصلاح وزن‌های اتصالات شبکه می‌پردازد. آزمایشات انجام شده بر روی چند پایگاه داده استاندارد نشان می‌دهد که قوانین استخراج شده به کمک روش پیشنهادی از نظر سادگی بهتر از روش‌های موجود و از نظر دقت بطور متوسط دو درصد بهبود می‌یابد.

کلمات کلیدی: استخراج قانون، الگوریتم ژنتیک، شبکه‌عصبی مصنوعی، کشف روابط بین داده‌ای، هرس‌سازی

شبکه‌عصبی، کشف دانش

لیست مقالات مستخرج از پایان نامه

۱. ایرانی، ی. حسن پور، ح. (۱۳۹۷)، ”استفاده از خوشه بندی برای بهبود گسسته سازی خروجی لایه مخفی” بیست و چهارمین کنفرانس ملی سالانه انجمن کامپیوتر ایران، دانشگاه صنعتی شریف، تهران.
۲. ایرانی، ی. حسن پور، ح. (۱۳۹۸)، ”استفاده از الگوریتم ژنتیک برای بهبود هرس سازی شبکه عصبی در استخراج قانون” بیست و هفتمین کنفرانس مهندسی برق ایران، دانشگاه یزد، یزد.

فهرست مطالب

ز فهرست تصاویر

س فهرست جداول

ص فهرست اختصارات

۱ مقدمه

۱.۱ مقدمه ۱

۲.۱ اهمیت استخراج قانون از شبکه عصبی ۳

۳.۱ مراحل استخراج قانون از شبکه عصبی ۵

۴.۱ چالش‌ها و راهکارها ۷

۵.۱ ساختار پایان نامه ۸

۲ ادبیات پژوهش

۱.۲ مقدمه ۹

۲.۲ انواع قوانین قابل استخراج ۱۰

۳.۲	دسته‌بندی روش‌های استخراج قانون	۱۳
۱.۳.۲	بر اساس نوع شبکه	۱۳
۲.۳.۲	بر اساس کاربرد	۱۴
۳.۳.۲	بر اساس دیدگاه استخراج قانون	۱۵
۴.۳.۲	بر اساس نوع قوانین استخراج شده	۱۵
۵.۳.۲	بر اساس نوع ویژگی‌های موجود در پایگاه داده	۱۶
۴.۲	کارهای مرتبط	۱۶
۱.۴.۲	مروری بر کارهای انجام شده	۱۷
۲.۴.۲	معرفی روش FSRE	۲۵
۵.۲	نتیجه‌گیری	۲۶
۳	روش پیشنهادی	۲۹
۱.۳	مقدمه	۲۹
۲.۳	مراحل استخراج قانون به روش FSRE	۳۰
۱.۲.۳	پیش‌پردازش داده‌ها	۳۰
۲.۲.۳	تعیین ساختار مناسب و آموزش شبکه	۳۲
۳.۲.۳	هرس‌سازی شبکه	۳۲
۴.۲.۳	استخراج قانون	۳۳
۳.۳	چالش‌ها و راهکارها	۳۵
۴.۳	معرفی الگوریتم ژنتیک	۳۷

۳۸	کروموزوم	۱.۴.۳
۳۹	تابع برازندگی	۲.۴.۳
۴۰	عملگرهای ژنتیک	۳.۴.۳
۴۳	روش پیشنهادی	۵.۳
۴۴	غیرمتوازن کردن وزن ها به صورت سراسری	۱.۵.۳
۴۶	غیر متوازن کردن وزن ها به صورت محلی	۲.۵.۳
۴۷	معرفی الگوریتم گسسته سازی	۳.۵.۳
۴۹	نتیجه گیری	۶.۳

۴ بررسی و تحلیل نتایج

۵۱	مقدمه	۱.۴
۵۳	مجموعه داده Iris	۲.۴
۵۷	مجموعه داده Wisconsin Breast Cancer	۳.۴
۶۲	مجموعه داده Glasses	۴.۴
۶۵	تحلیل و بررسی نتایج	۵.۴
۶۷	نتیجه گیری	۶.۴

۵ جمع بندی و پیشنهادات

۶۹	مقدمه	۱.۵
۷۰	نواوری تحقیق	۲.۵

۳.۵ پیشنهادات ۷۰

مراجع ۷۵

واژه‌نامه فارسی به انگلیسی ۷۹

واژه‌نامه انگلیسی به فارسی ۸۳

فهرست تصاویر

۱.۳	فرآیند استخراج قانون با استفاده از روش پیشنهادی	۳۷
۲.۳	بررسی ورودی یک نرون به صورت جداگانه	۴۷
۱.۴	ساختار اولیه شبکه برای آموزش برای مجموعه داده Iris	۵۵
۲.۴	روند آموزش و کارایی شبکه برای مجموعه داده Iris	۵۵
۳.۴	ساختار نهایی هرس شده شبکه برای مجموعه داده Iris	۵۶
۴.۴	نتایج حاصل از گسسته‌سازی نرون‌های لایه میانی برای مجموعه داده Iris	۵۶
۵.۴	ماتریس درهم ریختگی برای قوانین استخراج شده مجموعه داده Iris	۵۸
۶.۴	ساختار اولیه شبکه برای آموزش برای مجموعه داده Wisconsin Breast Cancer	۵۹
۷.۴	روند آموزش و کارایی شبکه برای مجموعه داده Wisconsin Breast Cancer	۶۰
۸.۴	ساختار نهایی هرس شده شبکه برای مجموعه داده Wisconsin Breast Cancer	۶۱
۹.۴	ماتریس درهم ریختگی برای قوانین استخراج شده مجموعه داده Wisconsin Breast Cancer	۶۲
۱۰.۴	ساختار اولیه شبکه برای آموزش برای مجموعه داده Glasses	۶۳
۱۱.۴	روند آموزش و کارایی شبکه برای مجموعه داده Glasses	۶۴

۱۲.۴ ساختار نهایی هرس شده شبکه برای مجموعه داده Glasses ۶۵

۱۳.۴ ماتریس درهم ریختگی برای قوانین استخراج شده مجموعه داده Glasses ۶۶

فهرست جداول

۱.۴	نمونه‌ای از جدول در هم‌ریختگی	۵۲
۲.۴	پارامترهای الگوریتم ژنتیک	۵۳
۳.۴	اطلاعات کلی در مورد مجموعه داده Iris Flower	۵۴
۴.۴	اطلاعات کلی در مورد مجموعه داده Wisconsin Breast Cancer	۵۸
۵.۴	اطلاعات کلی در مورد مجموعه داده Glasses	۶۳
۶.۴	نتایج مقایسه روش پیشنهادی نسبت به سایر روش‌ها برای مجموعه داده Iris	۶۶
۷.۴	نتایج مقایسه روش پیشنهادی نسبت به سایر روش‌ها برای مجموعه داده Glasses	۶۷
۸.۴	نتایج مقایسه روش پیشنهادی نسبت به سایر روش‌ها برای مجموعه داده Wisconsin Breast Cancer	۶۷

فهرست اختصارات

Artificial Neural Network.....	ANN
Four Step Rule Extraction.....	FSRE
Radial Basis Function.....	RBF
Multi-layer Perceptron.....	MLP
Mean Squared Error.....	MSE
Self-Organizing Map.....	SOM
Sum of Squared Error.....	SSE
Support Vector Machine.....	SVM

فصل ۱

مقدمه

۱.۱ مقدمه

امروزه با توجه به تولید حجم عظیمی از داده‌ها در زمینه‌های مختلف از جمله پزشکی [۳]، تجاری [۲] و صنعتی [۳۸]، کشف روابط میان این داده‌ها برای اهداف مختلف از اهمیت ویژه‌ای برخوردار است. یادگیری ماشین^۱ شاخه‌ای از هوش مصنوعی^۲ است که به دنبال استقراء^۳ یا پالایش^۴ دانش است. استنتاج مفاهیم عمومی از نمونه‌های آموزشی خاص، با شناسایی ویژگی‌هایی که بتوانند نمونه‌های منفی از مثبت را تشخیص دهند، برای یادگیری ماشین اهمیت دارد. در طول دهه گذشته، یادگیری ماشین از حالت تئوری به جایی رسیده

^۱Machine learning

^۲Artificial intelligence

^۳Knowledge induction

^۴Knowledge refinement

است که ارزش تجاری قابل توجهی پیدا کرده است. الگوریتم های یادگیری ماشین با موفقیت توانسته‌اند در مسائل کاربردی گوناگونی به کار گرفته شوند [۵۳، ۶، ۳۲]. استخراج قانون^۱ یکی از کاربردهای مهم یادگیری ماشین و داده‌کاوی^۲ برای کشف روابط پنهان میان داده‌ها است. استخراج قانون برای یک مدل پیش‌بینی کننده غیرشفاف و داده‌هایی که با آن آموزش دیده است، یک توصیف قابل درک برای فرضیه های مدل پیش بینی کننده ارائه می دهد که رفتاری نزدیک به آن را دارد. استخراج قانون به کمک ابزارهای گوناگونی مانند شبکه عصبی مصنوعی (ANN)^۳ [۲۰]، درخت تصمیم^۴ [۵۲] و ماشین بردار پشتیبان (SVM)^۵ [۴۰] انجام می‌گیرد. در این میان شبکه‌های عصبی، به خصوص شبکه عصبی چند لایه پرسپترون (MLP)^۶، به دلیل عدم نیاز به دانش اولیه در مورد داده‌ها و دقت^۷ بالا، یکی از پرکاربردترین ابزارهای موجود در این حوزه است.

اگرچه شبکه‌های عصبی به طور کلی عملکرد بهتری را در مسائل طبقه‌بندی^۸ الگوها ارائه می‌دهند، اغلب از آن‌ها به عنوان یک جعبه سیاه^۹ یاد می‌شود [۱۱، ۱۳]. به عبارت دیگر، پیش‌بینی‌های آن‌ها به اندازه درخت‌های تصمیم قابل تفسیر نیستند. با استفاده از یک الگوریتم برای استخراج قانون، می‌توان پیش‌بینی‌های انجام شده توسط شبکه عصبی را درک کرد، این قوانین در عین حال که می‌توانند دقت شبکه عصبی را تا حد خوبی حفظ کنند، دقت بالاتری را در مقایسه با درخت‌های تصمیم ارائه می‌دهند. پس اگر بتوانیم قوانینی را از شبکه‌های عصبی مانند قوانین تولید شده از درخت‌های تصمیم استخراج کنیم، مطمئناً پیش‌بینی انجام شده توسط شبکه را می‌توانیم بهتر تفسیر کنیم. علاوه بر این، قوانین یک نوع دانش هستند که کارشناسان می‌توانند به راحتی آنها را تأیید، رد و یا توسعه دهند [۴۹].

سیستم‌های یادگیری ماشین نمادین^{۱۰}، ارائه‌ای صریح از دانش موجود در یک مسئله دارند که می‌تواند

¹Rule Extraction

²Data mining

³Artificial Neural Network

⁴Decision tree

⁵Support Vector Machine

⁶Multi-layer Perceptron

⁷Accuracy

⁸Classification

⁹Black Box

¹⁰Symbolic machine learning

چگونگی دستیابی به نتایج را برای کاربر شرح دهد. یک سیستم کاربر پسند^۱، به طور گسترده در برنامه‌های کاربردی^۲ مورد نیاز است. کاربرد موفق روش‌ها و سیستم‌های شبکه‌عصبی در حوزه‌هایی مانند تجارت^۳[۱۷]، مهندسی، پزشکی^۴[۳۲]، اجتماعی^۵[۵] و غیره نشان دهنده توانایی بالای این ابزار است [۴]. با این حال، شبکه‌های عصبی نیازمند یک مکانیزم توضیحی مشابه سیستم‌های یادگیری ماشین نمادین دارند. این امر کاربران را قادر می‌سازد فرآیند تصمیم‌گیری خود را درک کنند و از این طریق شبکه‌های عصبی به کاربران گسترده‌تری دست یابند.

۲.۱ اهمیت استخراج قانون از شبکه‌عصبی

برطبق تحقیقات صورت گرفته در این حوزه، مدل‌های جعبه سیاه به طور میانگین عملکرد بهتری را نسبت به مدل‌های جعبه سفید^۶ ارائه می‌دهند. در نتیجه قوانین استخراج شده از این الگوریتم‌ها کیفیت بالاتری از مدل‌های جعبه سفید خواهند داشت. از طرفی استفاده از این مدل‌ها در زمینه‌هایی مانند پزشکی و مالی که نیاز به تایید متخصصان دارد، امکان‌پذیر نمی‌باشد. همچنین استفاده از این ابزار در کاربردهایی از داده‌کاوی که هدف بدست آوردن دانش در مورد داده‌های ورودی است، ممکن نیست. استخراج قانون، استفاده از مدل‌های جعبه سیاه در این زمینه‌ها را امکان‌پذیر می‌سازد [۹].

در سال‌های اخیر شبکه‌های عصبی از یک رویکرد تئوری به یک تکنولوژی پرکاربرد در مسائل مختلف تبدیل شده‌اند. بیشتر این کاربردهای مربوط به شناسایی الگو و شبکه‌های عصبی اغلب به عنوان تخمین‌گر غیرخطی^۴ استفاده می‌شوند که در مقابل نویز مقاوم هستند. روش‌های یادگیری شبکه‌عصبی مصنوعی، یک رویکرد قدرتمند و غیر خطی برای تقریب تابع هدف برای بسیاری از مسائل طبقه‌بندی، رگرسیون^۵ و خوشه‌بندی^۶

^۱User friendly

^۲Application

^۳White Box

^۴Non-linear approximators

^۵Regression

^۶Clustering

است. شبکه‌عصبی عملکرد خوبی را در طیف وسیعی از مسائل عملی نشان داده است. با این حال، دلایل محکمی برای مناسب نبودن شبکه‌های عصبی برای ارائه دانش، وجود دارد. ناتوانی در توضیح ساختارها و قابلیت درک پایین شبکه آموزش دیده، این که چرا بعضی از داده‌ها به یک دسته خاص تعلق می‌گیرند، از جمله مهم‌ترین این دلایل هستند. بنابراین، شبکه‌های عصبی، برای افزایش میزان مقبولیت خود برای طیف گسترده‌تری از کاربران به عنوان یک سیستم تصمیم‌گیری^۱ و ارتقای میزان سودمندی خود به عنوان یک ابزار مفید و تعمیم‌پذیر یادگیری ماشین، باید بر محدودیت‌های خود غلبه کند [۱].

مسئله نحوه تصمیم‌گیری یک شبکه آموزش دیده، مشکلی دیرینه در حوزه مدل‌سازی پیوندگرا^۲ است. دانش در یک شبکه‌عصبی به صورت توزیع شده در وزن‌های عددی شبکه تعبیه شده است. یک رویکرد امید بخش برای حل این مشکل، نمایش پارامترهای وزنی به صورت یک توصیف نمادین است. تفسیر نمادین دانش تعبیه شده در شبکه آموزش دیده، رویکردی برای غلبه بر مشکل ”جعبه سیاه“ است. این نوع از اصلاح فرموله‌سازی^۳ که به عنوان استخراج قانون شناخته می‌شود، می‌تواند رفتار شبکه را به خوبی توضیح و به انتقال دانش بدست آمده، کمک کند. تکنیک‌های شبکه‌عصبی روابط نهفته میان داده‌ها را تشخیص و رفتار ویژگی‌ها را پیش‌بینی می‌کند، تکنیک‌های استخراج قانون ابزارهایی برای درک فرآیند تصمیم‌گیری شبکه‌های عصبی هستند [۳۹].

نیاز به درک رفتار شبکه‌عصبی در حوزه‌های بحرانی^۴، منجر به ارائه رویکردهای زیادی برای توضیح فرآیند تصمیم‌گیری شبکه با استفاده از انواع مختلفی از قوانین شده است. قوانین استخراج شده اغلب به شکل گزاره‌ای^۵ هستند. با این حال، محدودیت‌های ذاتی در این نوع قوانین وجود دارد. مجموعه قوانین استخراج شده از شبکه، ممکن است بیش از حد برای داده‌های طبقه‌بندی شده دقیق باشد، در حالی که مسئله واقعی

¹Decision System

²Connectionist modeling

³Reformulation

⁴Critical

⁵Propositional

به درستی توصیف نشده باشد [۵۶].

۳.۱ مراحل استخراج قانون از شبکه عصبی

اساساً، روش‌های استخراج قانون از شبکه عصبی در سه گروه آموزشی^۱، تجزیه‌ای^۲ و ترکیبی^۳ قرار می‌گیرند [۶]. در دیدگاه آموزشی، قوانین برطبق یک تحلیل تجربی از الگوهای ورودی-خروجی تولید می‌شوند که به آن رویکرد جعبه سیاه نیز گفته می‌شود. در تکنیک‌های تجزیه‌ای قوانین با بررسی وزن‌های هر نرون^۴ لایه مخفی و خروجی تعیین می‌شوند. در نهایت، دیدگاه ترکیبی هر دو دیدگاه آموزشی و تجزیه‌ای را ادغام می‌کند، این روش‌ها از دانش موجود در ساختار درونی شبکه برای بهبود دیدگاه آموزشی استفاده می‌کنند. در این پژوهش استخراج قانون از شبکه عصبی با استفاده از دیدگاه تجزیه‌ای مورد توجه قرار گرفته است. روش‌های تجزیه‌ای برای استخراج قانون از شبکه عصبی، به دلیل توجه به ساختار داخلی شبکه و ارتباطات موجود، نتایج بهتری در مقایسه با دیدگاه آموزشی ارائه می‌دهند. در دیدگاه تجزیه‌ای، استخراج قانون با تجزیه و تحلیل وزن‌ها، بایاس‌ها^۵ و مقادیر تابع فعال‌ساز^۶، به صورت مجزا برای هر نرون از لایه مخفی و خروجی انجام می‌شود. دیدگاه آموزشی، بر خلاف دیدگاه تجزیه‌ای، به ساختار داخلی شبکه توجهی ندارد و صرفاً با تحلیل خروجی‌های شبکه به عنوان یک جعبه سیاه برای نمونه‌های ورودی به استخراج قانون می‌پردازد [۶].

به طور کلی استخراج قانون با استفاده از دیدگاه تجزیه‌ای طی چند مرحله انجام می‌شود:

• پیش پردازش^۷ داده‌ها: معمولاً داده‌ها به صورت خام به شبکه اعمال نمی‌شوند. مقادیر خام گسترده‌گی

زیادی دارند، ممکن است برای یک مسئله برخی از ویژگی‌ها غیر موثر و یا اضافی باشند. بدین ترتیب

¹ Pedagogical

² Decompositional

³ Eclectic

⁴ Neuron

⁵ Bias

⁶ Activation Function

⁷ Preprocessing

مجموعه عملیاتی در قالب پیش پردازش قبل از انجام عملیات اصلی بر روی داده‌ها استفاده می‌شود. گسسته‌سازی و رقمی کردن^۱، نرمال‌سازی داده‌ها^۲، نمونه‌برداری^۳، تجمیع^۴، کاهش بعد^۵ و غیره از جمله عملیاتی است که بر روی داده‌ها انجام می‌شود.

● **تعیین ساختار مناسب شبکه و آموزش آن:** با توجه به داده‌های موجود و قوانین مورد نیاز ساختار شبکه عصبی مشخص می‌شود. تعیین ساختار مناسب و بهینه موجب تولید نتایج قابل قبولی خواهد شد. تعداد نرون‌های لایه ورودی و خروجی با توجه به داده‌ها در مرحله قبل تعیین می‌گردد، اما تعداد لایه‌های مخفی و تعداد نرون‌های موجود در هر لایه با توجه به میزان پیچیدگی و پراکندگی داده‌ها تعیین می‌شود.

● **هرس‌سازی^۶ شبکه عصبی:** در این فاز، هدف، ساده‌سازی شبکه به منظور کاهش ابهام و پیچیدگی به منظور استخراج قوانین است. در این گام برخی اتصالات و گره‌های کم اثر حذف خواهند شد. این امر ممکن است باعث کاهش دقت و کارایی^۷ شبکه شود.

● **گسسته‌سازی خروجی لایه مخفی:** استخراج قوانین با تحلیل خروجی لایه مخفی انجام می‌شود. الگوها با استفاده از دسته‌بندی این مقادیر تولید و در نهایت قوانین استخراج می‌شوند. پیچیدگی قوانین استخراج شده از شبکه، وابسته به مقادیر گسسته شده‌ی خروجی لایه مخفی است.

● **استخراج قانون:** در انتها قوانین با تجزیه و تحلیل ساختار هرس شده شبکه و مقادیر گسسته خروجی لایه مخفی تولید می‌شوند.

¹Discretization

²Normalization

³Sampling

⁴Aggregation

⁵Dimensionality reduction

⁶Pruning

⁷Performance

۴.۱ چالش‌ها و راهکارها

قوانین استخراج شده به کمک شبکه‌های عصبی معمولاً دقت پایین‌تری از خود شبکه دارند. دلیل این امر، متکی بودن این روش‌ها به هرس‌سازی شبکه و استخراج قانون از ساختار هرس شده شبکه می‌باشد. با هرس‌سازی، یال‌های کم اهمیت (وزن‌های نزدیک به صفر) حذف می‌شوند که موجب کاهش دقت و کارایی شبکه می‌شود. این عملیات بخش جداناپذیر و مهمی از روش‌های استخراج قانون به کمک شبکه‌های عصبی است و باعث شده است که تأثیر به‌سزایی بر کارایی قوانین استخراج شده داشته باشد.

ساختار هرس شده شبکه‌های عصبی، فقط شامل اتصالات موثر شبکه است. با این حال استخراج قانون، به دلیل وجود مقادیر پیوسته در خروجی لایه مخفی، به راحتی امکان‌پذیر نمی‌باشد. از این رو گسسته‌سازی این مقادیر یکی از اجزای مهم و جداناپذیر در استخراج قانون از شبکه‌های عصبی با دیدگاه تجزیه‌ای می‌باشد. تشکیل دسته‌های صحیح برای استخراج قوانین مربوط به هر کلاس، سادگی قوانین استخراج شده و کارایی مناسب آن‌ها به طور مستقیم وابسته به عملیات گسسته‌سازی هستند.

FSRE^۱ یک روش برای استخراج قانون از شبکه‌های عصبی MLP است که توانایی استخراج قوانین گزاره‌ای if-then-else برای مسائل دسته‌بندی^۲ را دارد [۵۳]. استخراج قانون در این روش با استفاده از دیدگاه تجزیه‌ای انجام می‌شود. شبکه‌های عصبی با استفاده از داده‌های آموزشی در صورت نیاز ابتدا پیش‌پردازش شده و سپس برای آموزش شبکه‌های عصبی استفاده می‌شوند. استخراج قانون در گام آخر در چهار مرحله و با استفاده از ساختار هرس شده شبکه و مقادیر گسسته شده در لایه مخفی انجام می‌شود. عملیات گسسته‌سازی خروجی لایه مخفی در روش FSRE با صرف وقت زیاد به صورت دستی و با استفاده از ترسیم خروجی هر یک از نرون لایه مخفی انجام می‌شود.

^۱Four Step Rule Extraction

^۲Classification

در این پایان نامه قصد داریم به بهبود این روش در دو فاز بپردازیم. در فاز اول با ارائه یک راهکار نوین و با استفاده از الگوریتم ژنتیک^۱ خسارات ناشی از هرس‌سازی را به میزان قابل توجهی کاهش می‌دهیم. این عملیات در دو گام به صورت سراسری^۲ و محلی^۳ به اصلاح وزن‌های شبکه می‌پردازد، در این روش قصد داریم با غیرمتوازن^۴ کردن وزن‌ها آن‌ها را در دو دسته بی‌اهمیت و بااهمیت قرار دهیم تا ساختار هرس شده بهینه‌ای برای استخراج قانون تولید شود. در فاز بعد به منظور بهبود کارایی، سرعت و سادگی قوانین استخراج شده، یک الگوریتم خوشه‌بندی ابتکاری برای گسسته‌سازی خروجی نرون‌های لایه مخفی ارائه خواهیم داد.

۵.۱ ساختار پایان نامه

در این فصل به بیان کاربرد و اهمیت استخراج قانون، شبکه‌عصبی و رابطه این دو با یکدیگر پرداختیم. در نهایت روش استخراج قانون FSRE و چالش‌های موجود در این روش بررسی و راهکارهای ارائه شده به صورت مختصر بیان شد. در فصل ۲ به معرفی انواع قوانین، برخی دسته‌بندی‌های موجود برای روش‌های استخراج قانون، معرفی روش استخراج قانون FSRE و برخی روش‌های دیگر موجود در این حوزه می‌پردازیم. در فصل ۳ روش پیشنهادی ارائه خواهد شد. در فصل ۴ به بررسی و تحلیل نتایج و تاثیر روش پیشنهادی در بهبود استخراج قانون می‌پردازیم. در نهایت در فصل ۵ نتیجه‌گیری پایان نامه و راهکارهایی برای بهبود این پژوهش در آینده بیان خواهد شد.

¹Genetic algorithm

²Global

³Local

⁴Unbalance

فصل ۲

ادبیات پژوهش

۱.۲ مقدمه

در فصل ۱ کلیات استخراج قانون با استفاده از شبکه‌عصبی مورد بررسی قرار گرفت. همچنین چالش‌های موجود در روش FSRE بررسی شد. در این فصل به معرفی انواع قوانین قابل استخراج از شبکه‌عصبی خواهیم پرداخت. همچنین دسته‌بندی روش‌های استخراج قانون از دیدگاه‌های مختلف مورد بررسی قرار خواهد گرفت. در ادامه روش استخراج قانون FSRE به عنوان روش پایه این پژوهش و سایر روش‌های مهم موجود معرفی خواهند شد.

۲.۲ انواع قوانین قابل استخراج

قوانین استخراج شده می توانند در ماهیت خود به طور قابل توجهی متفاوت باشند، از این رو خوانایی آن‌ها بسته به روش استخراج استفاده شده متفاوت است. هر نوع از قوانین، به شیوه‌ای متفاوت مفهوم و همبستگی درونی پارامترهای ورودی را توصیف می‌کند. اغلب یک مصالحه^۱ بین خوانایی و نرخ بالای طبقه‌بندی وجود دارد. در ادامه به طور خلاصه هر کدام از قوانین منطقی را شرح خواهیم داد [۹]:

• **قوانین گزاره‌ای If-Then / If-Then-Else:** این نوع قوانین که متداول‌ترین نوع قوانین هستند شامل یک جزء شرطی If و یک جزء طبقه‌بند Then هستند. یک قانون می‌تواند در صورت نیاز شامل یک بخش طبقه‌بند Else نیز باشد. جزء شرطی If در یک قانون گزاره‌ای، یک ترکیبی بولی^۲ از شرطهای منطقی بر روی ویژگی‌های ورودی است. بخش شرطی می‌تواند شامل عبارات عطفی، فصلی و نقیض باشد. نمونه‌ای از یک قانون گزاره‌ای در رابطه ۱.۲ آمده است. برای مقادیر ورودی پیوسته معمولاً شرطها به صورت محدودیت در یک بازه تعیین می‌شوند. رابطه ۲.۲ مثالی از این نوع قانون است. نمونه‌هایی از این قوانین در [۱۶، ۲۵] استفاده شده‌اند.

$$\text{If } X = x \text{ and } Y = y \text{ then } Class = A \quad (1.2)$$

$$\text{If } X \in [c_1, c_2] \text{ and } Y > c_r \text{ then } Class = A \text{ where } c_1, c_2, c_r \in \mathbb{R} \quad (2.2)$$

¹Trade off

²Boolean

در این رابطه‌ها X و Y متغیرهای ورودی، x, y و c_i مقادیر ممکن برای این متغیرها و A یک نمونه کلاس خروجی است.

- **قوانین M-of-N**: این نوع قوانین بسیار مشابه قوانین گزاره‌ای هستند. در این نوع قانون کلاس‌بندی با محقق شدن حداقل، دقیقاً و یا حداکثر M شرط از N شرط موجود انجام می‌شود. معمولاً این قوانین را می‌توان به سادگی به قوانین گزاره‌ای تبدیل کرد. برای مثال رابطه ۴.۲ نمونه قانون گزاره‌ای تبدیل شده از قانون M-of-N در رابطه ۳.۲ است. همانطور که مشاهده می‌شود در صورتی که M برابر با یک باشد می‌توان قانون را به شکل یک ترکیب فصلی^۱ از تمام شروط موجود بیان کرد و در صورتی که M برابر با N باشد، این ترکیب عطفی^۲ خواهد بود. [۴۵] نمونه‌ای از استخراج این قوانین است.

$$\text{If exactly } 2 \text{ of } X = a_1, Y = a_2, Z = a_3 \text{ then } Class = A \quad (3.2)$$

$$\text{If } ((X = a_1 \text{ and } Y = a_2) \text{ or } (X = a_1 \text{ and } Z = a_3) \text{ or } (Y = a_2 \text{ and } Z = a_3)) \text{ then } Class = A \quad (4.2)$$

- **قوانین ضمنی**^۳: در حالی که قوانین فوق نشان دهنده‌ی مناطقی از ابرفضاها^۴ هستند که محورهای مختصات را با استفاده از بازه‌هایی جدا می‌کنند، گاهی ممکن است برای تعریف یک زیرفضا^۵ نیاز به

¹Disjunction

²Conjunction

³Oblique

⁴Hyperspace

⁵Subspace

تعداد شرایط زیادی برای دستیابی به یک طبقه‌بند^۱ با کیفیت مطلوب داشته باشیم. قوانین ضمنی قادر هستند ابرفضاها را به زیرفضاهای دلخواهی تقسیم کنند. با استفاده از این قوانین می‌توان زیرفضاهای مورد نیاز را با کیفیت مطلوب و با تعداد قوانین کمتری تولید کرد. نمونه‌ای از این قوانین در رابطه ۵.۲ آمده است. نمونه‌ای از استخراج این قوانین در [۵۰] بیان شده است.

$$\text{If } (a_1X_1 + a_2X_2 + a_3X_3 > a_4) \text{ and } (a_5X_1 + a_6X_2 > a_7) \text{ then } Class = A \quad (5.2)$$

- **قوانین چند جمله‌ای:** این نوع قوانین شباهت زیادی به قوانین ضمنی دارند. قوانین تولید شده از معادلات پیچیده‌تری برای تولید بازه‌های مورد نظر خود استفاده می‌کند. به طور مثال رابطه ۶.۲ یک قانون برای تولید یک زیرفضای بیضوی است. اشکال اصلی این قوانین دشوار بودن درک آن‌ها است.

$$\text{If } (a_1X_1^2 + a_2X_2^2 + a_3X_1X_2 + a_4X_1 + a_5X_2 + a_6) \leq a_7 \text{ then } Class = B \quad (6.2)$$

- **قوانین فازی^۲:** این نوع قوانین شباهت بسیار زیادی به قوانین گزاره‌ای دارند، با این تفاوت که از معیارهای توصیفی به جای معیارهای کمی و بولی استفاده می‌کنند. در قوانین فازی زبان گفتمان متغیرها یک مجموعه فازی است. برای مثال رابطه ۷.۲ یک قانون فازی را بیان می‌کند. در این رابطه *low*، *medium* و *high* اعضای مجموعه فازی هستند. در مجموعه‌های فازی ارتباط هر عنصر با مجموعه استفاده از یک درجه تعلق^۳ مشخص می‌شود. استفاده از مفاهیم زبانی^۴ سبب شده است که قوانین فازی به خوبی قوانین گزاره‌ای قابل درک باشند. [۲۹، ۳۶] نمونه‌هایی از تولید این قوانین

¹Classifier

²Fuzzy

³Corresponding grade

⁴Linguistic concepts

هستند.

If X is high and Y is medium then $Class = A$ (۷.۲)

لازم به ذکر است که دانش موجود در شبکه عصبی را می توان به شکل های مختلفی از جمله ماشین های حالت^۱ و نمایش بصری^۲ استخراج کرد. کیفیت^۳ اصلی ترین معیار برای ارزیابی قوانین استخراج شده است. به طور خاص، کیفیت معادل میزان قابل درک بودن است که این معیار ممکن است بیشتر و یا کمتر از دقت باشد، اغلب افزایش یکی از این معیارها منجر به کاهش دیگری خواهد شد.

۳.۲ دسته بندی روش های استخراج قانون

روش های زیادی برای استخراج قانون از شبکه عصبی ارائه شده اند. هرکدام از این روش ها دارای خصوصیات متفاوتی هستند که می توان آن ها را با استفاده از این خصوصیات دسته بندی نمود. در این بخش پنج دسته بندی مختلف معرفی خواهند شد.

۱.۳.۲ بر اساس نوع شبکه

در این دسته بندی نوع شبکه استفاده شده برای استخراج قانون به عنوان پارامتر دسته بندی استفاده می شود. استخراج قانون از شبکه های عصبی مختلفی نظیر شبکه عصبی چندلایه پرسپترون (MLP) [۵۵]، شبکه های تابع پایه شعاعی (RBF)^۴ [۳۷، ۵۸] و شبکه های نگاشت خود سازمان یافته (SOM)^۵ [۳۵] صورت گرفته است.

^۱State machines

^۲Visual representation

^۳Quality

^۴Radial Basis Function

^۵Self-Organizing Map

هر کدام از این شبکه‌ها کاربرد خاص خود را دارند. شبکه‌های MLP که رایج‌ترین و پر کاربردترین نوع شبکه عصبی برای استخراج قانون می‌باشند، برای مسائل دسته‌بندی و خوشه‌بندی استفاده می‌شوند و از دقت بسیار خوبی بهره می‌برند. شبکه‌های RBF همانند شبکه‌های چند لایه عمده برای مسائل خوشه‌بندی مورد استفاده قرار می‌گیرند اما کارایی و دقت این شبکه‌ها در مقایسه با شبکه‌های MLP در حد پایین‌تری قرار دارد. شبکه‌های نگاشت خود سازمان یافته در نمایش همسایگی‌ها در مجموعه داده و یا شباهت/ تفاوت بین نمونه‌های مختلف کاربرد دارند و کاربرد چندانی در عملیات دسته‌بندی ندارد. به همین دلیل روش‌های استخراج قانون برای این شبکه کمتر می‌باشند. مزیت شبکه‌های SOM در نمایش داده‌هاست اما قوانین استخراج شده با این روش ساده بوده و تنها برای مسائل ساده کاربرد دارد.

۲.۳.۲ بر اساس کاربرد

همانطور که بیان شد استخراج قانون در مسائل گوناگونی کاربرد دارد. دسته‌بندی [۲۵] و تخمین توابع^۱ [۱۲، ۴۸] نمونه‌هایی از این کاربردها هستند. در این دسته‌بندی نوع کاربرد شبکه به عنوان معیار برای دسته‌بندی روش‌های موجود استفاده می‌شود. در مسائل دسته‌بندی هدف تعیین عضویت یک نمونه ورودی به یکی از دسته‌های خروجی موجود است. این دسته‌ها توسط مقادیر گسسته مشخص شده و بدون روی هم افتادگی^۲ هستند. این مسائل اغلب با استفاده از شبکه عصبی MLP انجام می‌شود، در این شبکه معمولاً برای هر دسته موجود یک نرون در لایه خروجی در نظر گرفته می‌شود، قوانین برای هر نرون لایه خروجی به صورت جداگانه تولید می‌شوند و کلاس مربوط به هر دسته از قوانین تعیین می‌شود. در مسائل مربوط به تخمین توابع هدف تعیین مقدار (حقیقی) خروجی برای هر نمونه بر اساس پارامترهای ورودی آن نمونه است. در این مسائل لایه خروجی تنها شامل یک نرون خواهد بود که مقدار خروجی هر نمونه را تخمین می‌زند.

¹Regression

²Overlap

۳.۳.۲ بر اساس دیدگاه استخراج قانون

همانطور که در ۳.۱ نیز اشاره شد، روش‌های موجود قوانین را بر اساس سه دیدگاه آموزشی [۴۱، ۴۲، ۱۲]، تجزیه‌ای [۲۳] و ترکیبی [۴۷] از شبکه‌عصبی استخراج می‌کنند. دیدگاه تجزیه‌ای پر کاربردترین و پربازده‌ترین روش موجود برای استخراج قانون از شبکه‌های عصبی می‌باشد. در این روش قوانین با تجزیه و تحلیل ساختار درونی شبکه، اتصالات موجود و خروجی نرون‌های لایه میانی استخراج می‌شوند. در این دیدگاه اتصالات بی‌اهمیت با استفاده از تکنیک‌های هرس‌سازی حذف و اتصالات و نرون‌های موثر شبکه شناسایی می‌شوند. ویژگی‌های موثر هر نرون لایه میانی با توجه به اتصالات باقی‌مانده مشخص می‌شود، سپس به ازای هر نمونه ورودی خروجی نرون‌های لایه میانی تعیین و نمونه‌های ورودی بر اساس این خروجی‌ها خوشه‌بندی و کلاس خروجی آن‌ها تعیین می‌شود. در نهایت، قوانین مربوط به هر خوشه با توجه به نمونه‌های ورودی و ویژگی‌های موثر هر خوشه استخراج می‌شوند.

دیدگاه آموزشی برخلاف دیدگاه تجزیه‌ای به ساختار درونی شبکه توجهی ندارد. در روش‌های مبتنی بر این دیدگاه به دلیل عدم توجه به اتصالات شبکه، معمولاً عملیات هرس‌سازی انجام نمی‌شود. قوانین با استفاده از تحلیل خروجی‌های شبکه، به عنوان یک جعبه سیاه، به ازای هر نمونه ورودی استخراج می‌شوند. استخراج قوانین در دیدگاه ترکیبی با ادغام این دو دیدگاه صورت می‌گیرد. عموماً دیدگاه ترکیبی سعی دارد تا با بهره‌گیری از برخی اطلاعات درونی شبکه به عنوان مکمل دیدگاه آموزشی، قوانین با کیفیت‌تری را تولید نماید.

۴.۳.۲ بر اساس نوع قوانین استخراج شده

قوانین می‌توانند به شکل‌های مختلفی از شبکه‌عصبی آموزش دیده استخراج شوند. در ۲.۲ برخی از شکل‌های قوانین قابل استخراج از شبکه بیان شده‌اند. همانطور که مشاهده می‌شود، هر کدام از این قوانین کاربرد خاص

خود را دارند. در این نوع دسته‌بندی، شکل قوانین استخراج شده به عنوان مبنای دسته‌بندی در نظر گرفته می‌شود. قوانین گزاره‌ای، M-of-N، و فازی از متداول‌ترین انواع این قوانین هستند.

۵.۳.۲ براساس نوع ویژگی‌های موجود در پایگاه داده

در این دسته‌بندی قابلیت اعمال روش‌های مختلف استخراج قانون بر روی انواع مجموعه‌های داده بررسی می‌شود. مجموعه داده‌های موجود ممکن است شامل ویژگی‌های پیوسته، گسسته و یا کیفی باشند. ویژگی‌های کیفی برای استفاده نیاز به تبدیل به یکی از ویژگی‌های گسسته و یا پیوسته دارند. برخی از روش‌های استخراج قانون موجود از ویژگی‌های پیوسته [۲۵]، برخی دیگر از ویژگی‌های گسسته [۱۶] و گروهی دیگر از هر دو نوع داده پیوسته و گسسته [۵۳] پشتیبانی می‌کنند. عموماً ویژگی‌های پیوسته قابل گسسته‌سازی هستند و در صورت نیاز می‌توان آن‌ها را به مقادیر گسسته تبدیل کرد، این عملیات با توجه به روش استخراج قانون و یا برای بهبود دقت قوانین انجام می‌شود. معمولاً مسائل مربوط به تخمین توابع فقط از مقادیر پیوسته نظیر دما، قد و وزن که غیر قابل شمارش هستند، پشتیبانی می‌کنند. برخی از روش‌ها فقط از ویژگی‌های گسسته مانند جنسیت، گروه خونی و مذهب که قابل شمارش هستند پشتیبانی می‌کنند. استخراج قانون با استفاده از ویژگی‌های گسسته ساده‌تر و قابل درک‌تر است اما ممکن است که نسبت به دلیل عملیات گسسته‌سازی نسبت به روش‌های مبتنی بر ویژگی‌های پیوسته از دقت پایین‌تری برخوردار باشد.

۴.۲ کارهای مرتبط

در این بخش به بررسی برخی از روش‌های مهم ارائه شده تا کنون و بررسی و تحلیل هر روش از دیدگاه‌های مختلف دسته‌بندی بیان شده در بخش قبل می‌پردازیم. در ادامه به تشریح روش FSRE می‌پردازیم که در این

پایان نامه قصد داریم به بهبود آن بپردازیم.

۱.۴.۲ مروری بر کارهای انجام شده

Setiono در [۴۳] یک روش هرس‌سازی و استخراج قانون برای تشخیص سرطان سینه بر روی پایگاه‌داده Wisconsin Breast Cancer [۸] ارائه داده است. آموزش شبکه با استفاده از یک تابع خطای اصلاح شده برای تولید وزن‌های غیرمتوان انجام شده است. در این روش، برای دستیابی به وزن‌های مناسب، هر ویژگی ورودی با ارائه یک رابطه به ۱۰ مقدار دودویی گسسته‌سازی می‌شود. البته، این عمل باعث ده برابر شدن تعداد نرون‌های لایه ورودی می‌شود. هرس‌سازی شبکه با استفاده از یک تابع مجازات^۱ انجام شده است. در پایان، خروجی تابع فعال‌ساز برای هر نرون لایه میانی گسسته‌سازی می‌شود. قوانین نهایی با ترکیب روابط بدست آمده از بررسی اتصالات باقی مانده برای هر نرون میانی، استخراج می‌شوند. روش‌های ارائه شده در این مقاله، با توسعه تکنیک‌های هرس‌سازی در [۴۴]، و تعمیم روش استخراج قانون با عنوان روش RX در [۳۴] ارائه شده‌اند.

در روش RX پس از آموزش و هرس کردن، مقادیر بدست آمده در لایه مخفی توسط عملیات خوشه‌بندی به مقادیر گسسته تبدیل می‌گردند. سپس رابطه بین مقادیر گسسته بدست آمده با خروجی‌های شبکه تعیین می‌شود. آنگاه برای هر نرون لایه مخفی تعداد اتصالات ورودی به هر نرون پس از هرس‌سازی تعیین می‌شود. در صورتیکه این تعداد کمتر از حد آستانه باشد، قوانین برای دستیابی به مقادیر گسسته شده این نرون بر اساس ورودی‌ها بدست می‌آیند. در غیر این صورت نرون لایه مخفی و ورودی‌های متصل به آن به عنوان یک شبکه جداگانه در نظر گرفته می‌شوند:

- هر مقدار گسسته بدست آمده از نرون لایه مخفی به عنوان یک خروجی شبکه عصبی جدید در نظر گرفته

^۱Penalty

می‌شود.

- هر ورودی متصل به نرون لایه مخفی به عنوان ورودی شبکه جدید در نظر گرفته می‌شود.
 - لایه مخفی جدیدی با تعداد نرون مناسب ایجاد می‌شود.
 - این شبکه توسط داده‌های موجود آموزش دیده و مراحل استخراج قانون برای آن تکرار می‌شود.
- با ترکیب قوانین بدست آمده مجموعه قوانین نهایی بدست می‌آید.

Setiono و Leow در [۴۷] یک روش ترکیبی به نام FERNN برای استخراج سریع قانون ارائه دادند. در این روش برای جلوگیری از پردازش‌های اضافه از فاز هرس‌سازی چشم پوشی می‌شود، و قوانین به شکل یک درخت تصمیم با روش C4.5 استخراج می‌شوند. این کار با استفاده از معیار بهره اطلاعات^۱ برای تعیین اولویت نرون‌های لایه مخفی برای دسته‌بندی نمونه‌ها صورت می‌پذیرد. نرون دارای بیشترین اولویت به عنوان ریشه قرار می‌گیرد و مقادیر خروجی آن به دو زیر بازه تقسیم می‌شود، این تقسیم‌بندی طوری انجام می‌شود که بتواند بیشترین کارایی را داشته باشد. این عملیات برای سایر نرون‌ها تکرار می‌شود تا در نهایت درخت تصمیم برای دسته‌بندی نمونه‌ها تکمیل گردد. این روش با استفاده از دیدگاه ترکیبی عمل می‌کند. استفاده از شبکه عصبی به عنوان یک جعبه‌سیاه برای استخراج قوانین، سبب شده است تا در فرآیند استخراج قانون نیازی به هرس‌سازی شبکه وجود نداشته باشد.

در [۲۳] روشی تحت عنوان REANN ارائه شده است. این روش به صورت تکراری مراحل استخراج قانون، خوشه‌بندی قانون، هرس کردن قانون را انجام می‌دهد تا همه نمونه‌ها پوشش داده شود. این روش برای مسائل ساده به خوبی پاسخگو بوده، اما در مقابل مجموعه داده‌های حجیم و مسائل پیچیده از پیچیدگی زمانی بالایی برخوردار است. این الگوریتم از ابتدای ایجاد شبکه عصبی اعمال شده و در انتها با تعیین قوانین مورد نیاز

¹Information Gain

برای دسته‌بندی نمونه‌های هر کلاس خاتمه می‌یابد. پس از تعیین تعداد نرون‌های ورودی و خروجی، آموزش با یک نرون در لایه مخفی آغاز می‌گردد. با افزایش تعداد نرون در لایه مخفی و آزمایش حالات مختلف آن، تعداد نرون مناسب در این لایه بدست می‌آید. سپس با استفاده از تابع پناستی شبکه آموزش دیده هرس می‌شود. پس از هرس سازی، مقادیر فعال شده هر نرون لایه مخفی توسط عملیات خوشه‌بندی به مقادیر گسسته تبدیل می‌شود. استخراج قانون با استفاده از روش REX [۲۱] انجام می‌شود. در این روش پس از گسسته‌سازی خروجی لایه مخفی، رابطه بین مقادیر گسسته لایه مخفی و دسته‌های موجود در لایه خروجی تعیین می‌گردد. سپس رابطه بین نمونه‌های ورودی و مقادیر گسسته محاسبه می‌شود. در نهایت با ترکیب روابط بدست آمده قوانین نهایی حاصل می‌شود.

Hruschka و Ebecken در [۱۹] برای بهبود روش RX، به تغییر الگوریتم خوشه‌بندی مقادیر تابع فعال ساز نرون‌های لایه میانی پرداخته‌اند. گسسته‌سازی در این روش با استفاده از الگوریتم ژنتیک صورت می‌گیرد. در نهایت قوانین توسط روش RX استخراج می‌شوند.

روش NeuroRule در [۵۱] توسط Setiono و Tanaka ارائه شده است. این روش از با استفاده از یک شبکه‌عصبی MLP با ساختار استاندارد سه لایه به استخراج قوانین برای مسائل دسته‌بندی با استفاده از دیدگاه تجزیه‌ای می‌پردازد. در این روش از یک تابع مجازات برای هرس کردن شبکه‌عصبی استفاده می‌شود. پس از گسسته‌سازی مقادیر خروجی نرون‌های لایه میانی، استخراج قانون انجام می‌شود. در این روش ابتدا روابط موجود بین لایه ورودی و مقادیر گسسته بدست می‌آید. در مرحله بعد روابط بین مقادیر گسسته و نرون‌های لایه خروجی محاسبه شده و در نهایت قوانین با ترکیب این روابط بدست می‌آید.

کارایی این روش بر روی مسئله LED ارزیابی شده است. در این مسئله هدف شناسایی رقم نشان داده شده بر اساس حالات مختلف بر روی یک دیود نوری هفت بخشی^۱ است که شامل ۱۲۸ حالت مختلف است.

^۱ Seven segment

نتایج بدست آمده قادر خواهد بود که این مسئله را با دقت ۱۰۰٪ حل نماید. پیش از این نیز در [۲۷] یک روش جهت استخراج قوانین فازی برای این مسئله ارائه شده بود که پیچیدگی بالاتری نسبت به این قوانین داشته است.

ElAlami در [۱۴] از روش تخریبی برای هرس کردن شبکه استفاده کرده است. اتصالات شبکه در روش تخریبی با یک مقدار آستانه مقایسه می‌شوند، وزن‌هایی که مقادیر کمتری نسبت به آستانه دارند، از شبکه حذف خواهند شد. هرس‌سازی برای هر کلاس خروجی به صورت جداگانه انجام می‌شود تا اتصالات موثر آن کلاس مشخص شوند. در نهایت قوانین هر کلاس، از شبکه هرس شده مربوط به همان کلاس، در دو مرحله استخراج می‌شوند. ابتدا توانایی کلاس‌بندی داده‌ها به کمک یک ویژگی بررسی می‌شود. اگر یک توانایی تشخیص یک کلاس را با دقت مطلوب داشت، قوانین مربوط به آن کلاس استخراج می‌شود. در گام بعد توانایی دسته‌بندی هر کلاس با استفاده از ترکیب‌های مختلف ویژگی‌ها بررسی می‌شود تا قوانین برای سایر کلاس‌ها نیز استخراج گردد.

روش دیگری تحت عنوان LREX در [۳۷] مطرح شده است که قادر است قوانین ساده‌ای برای خوشه‌بندی داده‌ها با استفاده از شبکه‌های عصبی RBF استخراج کند. در این روش باید هر کدام از نرون‌های لایه مخفی تنها متعلق به یک کلاس باشند تا بتوان نتایج مناسبی بدست آورد. پس از آموزش شبکه عصبی، مقادیر مختصات مرکز خوشه‌ها (نرون‌های لایه مخفی) و وزن‌های لایه خروجی استخراج می‌شوند. وزن یال‌های ورودی برای هر خوشه، مرکز خوشه را تعیین می‌کند. برای هر خوشه یک بردار حاصل می‌شود که مراکز ویژگی‌های ورودی برای این خوشه را نشان می‌دهد. کلاس خروجی برای هر خوشه (نرون مخفی) با توجه به بزرگترین وزن یال متصل به آن در لایه خروج مشخص می‌گردد. نمونه‌های موثر برای هر نرون لایه مخفی مشخص می‌گردند. این نمونه‌های موثر در هنگام اعمال به شبکه عصبی باعث تولید مقدار خروجی بیشتری

برای نرون مخفی مورد نظر، نسبت به دیگر نرون‌های مخفی می‌شوند. مقدار بدست آمده از تابع شعاع گوسی، به عنوان انحراف معیار خوشه‌ها در نظر گرفته می‌شود. حد بالا و پایین ویژگی‌های هر خوشه با استفاده از مقادیر میانگین و انحراف معیار بدست می‌آید. در نهایت با ترکیب مقادیر حد بالا و پایین یک گزاره شرطی برای یک ویژگی در خوشه مورد نظر بدست می‌آید که با ترکیب این ویژگی‌ها قانون خوشه مورد نظر بدست خواهد آمد.

روشی برای استخراج قانون با کمک شبکه عصبی SOM مبتنی بر ماتریس فاصله یکپارچه^۱ در [۳۵] ارائه شده است. این ماتریس شامل فاصله اقلیدسی وزن همه نرون‌های خروجی در شبکه عصبی SOM می‌باشد که به منظور تعیین محدوده خوشه‌ها مورد استفاده قرار می‌گیرد و حالتی مانند رابطه ۸.۲ دارد:

$$U = [U(۱), U(۱, ۲), U(۲), U(۲, ۳), U(۳), U(۳, ۴), U(۴)] \quad (۸.۲)$$

در این رابطه $U(i, j)$ فاصله بین نرون‌های i و j و $U(k)$ میانگین وزن بین نرون‌های مجاور است. در این روش ابتدا نرون‌های لایه خروجی (معمولاً در فضای دو بعدی قرار دارند) دسته‌بندی می‌شوند. بنابراین اولین نرون خروجی انتخاب می‌شود. برای نرون انتخاب شده و نرون‌های مجاور با استفاده از رابطه ۹.۲ مقدار اختلاف مرز تولید می‌شود.

$$BDV = \frac{M_L - M_O}{R_O} \quad (۹.۲)$$

در این رابطه M_L میانگین بین واحد کاندید برای مرز و واحدهای انتخاب شده قبلی، M_O میانگین بین دیگر

^۱U-matrix

واحدهای مجاور باقیمانده و R_O محدوده دیگر واحدهای مجاور باقیمانده است.

نرونی که حاوی مقدار تولید شده بزرگتری باشد به عنوان نرون بعدی انتخاب می‌شود. این عملیات تا زمانی ادامه می‌یابد که مقدار بدست آمده برای نرون‌های مجاور بیش از مقدار بدست آمده برای نرون فعلی باشد. بدین صورت یک مرز جداساز در نرون‌های لایه خروجی بدست می‌آید. این عملیات برای همه واحدهای درون ماتریس U انجام می‌شود تا مجموعه‌ای از مرزها بدست آید. سپس مقدار BDV برای همه مرزهای بوجود آمده محاسبه شده و مرزهای با مقدار BDV بزرگتر، به عنوان مرزهای جداسازی خوشه‌ها تعیین می‌گردند. تعداد مرزهای انتخاب شده به تعداد مورد نیاز برای جداسازی خوشه‌های موجود است. پس از تعیین مرزها نیاز است تا ویژگی‌های کلیدی برای دسته‌بندی/خوشه‌بندی تعیین گردد. برای تعیین این ویژگی‌ها ابتدا مرزهای جداساز برای هر ویژگی ورودی محاسبه شده، در صورتیکه این مرزهای بدست آمده با مرزهای ماتریس U مطابقت داشت، ویژگی ورودی به عنوان ویژگی مناسب انتخاب می‌شود. در نهایت میانگین مقادیر مرز در ویژگی‌های انتخاب شده به عنوان گزاره‌های شرطی انتخاب می‌شوند.

روش دیگری تحت عنوان REFANN در [۲۴] برای تخمین توابع ارائه شده است. در این روش ابتدا مراحل آموزش و هرس‌سازی با استفاده از روش N2PFA [۴۶] انجام می‌شود. در روش N2PFA پس از تقسیم مجموعه داده به سه بخش آموزش، اعتبارسنجی و تست، آموزش شبکه عصبی توسط تعداد مناسبی از نرون در لایه مخفی انجام می‌شود تا خطای آموزش و اعتبارسنجی به حداقل برسد. سپس جهت هرس کردن این شبکه عصبی، نرون‌های مخفی غیر ضروری و ورودی‌های نامرتبط به صورت آزمایش و خطا مشخص و حذف می‌گردند. این عملیات تا زمانی ادامه می‌یابد که میزان خطای حاصل از ساختار جدید بیشتر از حد آستانه نباشد. با اعمال داده‌های موجود به ساختار هرس شده، مقادیر خروجی هر نرون مخفی بدست آمده که به صورت یک منحنی می‌باشد. برای کشف معادله انجام شده در هر نرون مخفی، منحنی آنرا به سه یا پنج قسمت

تقسیم کرده و توسط توابع ریاضیاتی سه بخشی و پنج بخشی تابع مورد نظر برای منحنی تخمین زده می‌شود. در نهایت، با توجه به توابع تخمین زده شده، قوانین تولید می‌گردد.

روش ANN-DT [۴۱] یک نمونه از روش‌هایی است که عملیات استخراج قانون را با استفاده از دیدگاه آموزشی انجام می‌دهد. قوانین استخراج شده با استفاده از این الگوریتم به شکل درخت تصمیم هستند. این روش توانایی استخراج قانون بدون در نظر گرفتن ساختار شبکه و یا نوع ویژگی‌های ورودی را دارد. از این روش می‌توان در مواردی مانند رگرسیون که خروجی شبکه مقادیر پیوسته‌ای می‌باشد نیز استفاده کرد. خروجی این الگوریتم یک درخت تصمیم تک متغیری است که در هر گره تنها شرایط یک متغیر به همراه یک مقدار آستانه بررسی می‌شود. هدف این روش تعیین ویژگی‌های موثر و مقادیر آستانه برای ویژگی است. مراحل استخراج قانون:

۱. شبکه تا رسیدن به سطح قابل قبولی آموزش می‌بیند.
۲. یک مجموعه داده برای استخراج قانون تعیین می‌شود. زیر مجموعه‌ای از نمونه‌ها (با در نظر گرفتن توزیع داده‌های اصلی) نمونه‌برداری می‌شوند که سبب کاهش پیچیدگی زمانی، کاهش تعداد قوانین و افزایش کارایی می‌گردد.
۳. انتخاب ویژگی‌های مناسب و تعیین مقادیر آستانه برای تولید قوانین از روی نمونه‌های انتخاب شده در این فاز صورت می‌گیرد. به صورت بازگشتی در هر مرحله مجموعه نمونه‌های انتخاب شده را با استفاده از بهترین ویژگی به دو قسمت ایده‌آل تقسیم کرده و مقدار آستانه برای این تقسیم‌بندی را بدست می‌آورد.
۴. در این فاز زمان توقف عملیات بازگشتی تعیین می‌گردد، که برای این کار از حالتی مانند صفر شدن انحراف استاندارد یا واریانس در دسته‌ها استفاده می‌شود.

این روش بازگشتی است و می‌تواند در شبکه‌های عصبی نوع چندلایه و تابع پایه شعاعی مورد استفاده قرار می‌گیرد.

KDRuleEx [۴۲] روش دیگری برای استخراج قانون مبتنی بر دیدگاه آموزشی است. در این روش عملیات انجام شده توسط شبکه عصبی در قالب یک جدول تصمیم نمایش داده می‌شود. این الگوریتم به لطف استفاده از تکنیک آموزشی پیچیدگی زمانی کمتری دارد. این الگوریتم به دلیل عدم استفاده از حالت بازگشتی، مدیریت حافظه بهتری در مقایسه با سایر روش‌های مبتنی بر دیدگاه آموزشی دارد. این روش در مسائل ساده و نه چندان پیچیده کاربرد دارد.

مراحل استخراج قانون:

۱. تبدیل ویژگی‌های پیوسته به گسسته.
 ۲. تعیین بهترین ویژگی جداکننده. تقسیم مقادیر این ویژگی به بخش‌هایی که بتوان جداسازی بین داده‌ها را ایجاد کرد.
 ۳. تقسیم‌بندی ویژگی‌ها تا زمانی ادامه می‌یابد که هزینه محاسبات و قوانین بدست آمده از میزان تعیین شده بیشتر نباشد.
 ۴. تصمیمات حاصل شده در جدول مورد نظر وارد شده و بهترین ویژگی بعدی انتخاب می‌شود.
 ۵. این عملیات تا زمانی که به میزان درستی تعیین شده دست یابد ادامه می‌یابد.
- در روش دیگر با استفاده از دوگام زیر باعث ساده‌سازی استخراج قوانین به صورت M-of-N می‌شود. سپس با تعیین تابع فعال‌ساز مناسب برای لایه مخفی شبکه عصبی آموزش می‌بیند. در نهایت با هرس کردن شبکه قوانین مورد نظر بدست می‌آیند:

• در گام اول مقادیر نمونه‌های موجود قبل از اعمال به شبکه به مقادیر گسسته -1 و 1 تبدیل می‌شوند که باعث محدود شدن مقادیر ورودی می‌گردد.

• در گام دوم از تابع تانژانت هایپربولیک بر روی همه اتصالات بین لایه ورودی و مخفی استفاده می‌شود که موجب محدود شدن مقادیر ورودی به نرون‌های لایه مخفی می‌شود.

روش دیگری در [۲۹] ارائه شده است که قادر است قوانین فازی را برای مجموعه‌هایی شامل دو نوع ویژگی پیوسته و گسسته بدست آورد. در این روش مقادیر ویژگی‌های گسسته به مقادیر باینری به فرم 1 از n تبدیل می‌شوند. مقادیر ویژگی‌های پیوسته با توجه به کمترین و بیشترین مقدار ویژگی و طبق تابع عضویت پنج متغیری به مقادیر فازی تبدیل می‌شود. سپس شبکه با استفاده از روش DE [۵۴] آموزش می‌بیند. DE یک روش آموزش بهینه شده بر اساس جمعیت است و با استفاده از داده‌های جدید بدست آمده آموزش می‌بیند. در پایان، از روش TACO [۳۳] جهت استخراج قوانین فازی استفاده می‌شود.

۲.۴.۲ معرفی روش FSRE

در [۵۳] روش دیگری تحت عنوان FSRE برای استخراج قانون ارائه شده است. در این روش پس از آموزش شبکه، وزن‌های شبکه در دو مرحله هرس می‌شوند. در گام اول اتصالات مابین لایه مخفی و خروجی هرس می‌شود که باعث شناسایی نرون‌های موثر هر کلاس خروجی می‌گردد. ممکن است یک نرون میانی به دلیل قطع شدن تمام اتصالاتش با خروجی، از شبکه حذف شود. اگر اتصالات یک نرون دارای مقادیر مختلفی باشند، اتصالات دارای وزن‌های بزرگ‌تر حفظ، و بقیه اتصالات حذف می‌شوند. در گام بعد اتصالات مابین لایه ورودی و مخفی هرس می‌شوند تا پارامترهای ضروری هر نرون لایه مخفی شناسایی و بقیه حذف می‌شوند. در این الگوریتم پس از حذف هر اتصال کارایی شبکه سنجیده می‌شود.

استخراج قانون در روش FSRE ، پس از گسسته‌سازی نرون‌های لایه میانی، در چهار مرحله انجام می‌شود. ابتدا با اعمال داده‌های آموزشی به ساختار هرس شده شبکه، الگویی برای هر نمونه ورودی تعیین می‌شود. این الگو با کنار هم قرار گرفتن مقادیر خروجی گسسته نرون‌های لایه میانی شکل می‌گیرد. در مرحله بعد، نمونه‌های دارای الگوی مشابه در یک دسته قرار می‌گیرند. یک کلاس خروجی برای هر دسته با توجه به خروجی اکثریت نمونه‌های آن، تعیین می‌شود. در مرحله سوم، با توجه به حذف شدن برخی ویژگی‌های ورودی در فاز هرس‌سازی، ممکن است نمونه‌های موجود در یک دسته مقادیر یکسانی داشته باشند. همچنین برخی نمونه‌ها ممکن است که دارای مقادیر مشابهی باشند. در این مرحله نمونه‌های تکراری حذف و نمونه‌های مشابه ترکیب می‌شوند. در نهایت دسته‌های دارای کلاس خروجی یکسان با هم ادغام و بار دیگر عملیات حذف و ترکیب نمونه‌ها بر روی دسته‌های متعلق به یک کلاس انجام می‌شود. در انتها برای اینکه قوانین بتوانند نمونه‌های دیده نشده را نیز کلاس‌بندی کنند، یکی از کلاس‌ها به عنوان پاسخ پیش‌فرض تعیین می‌گردد. دسته‌های دارای بالاترین مقدار خطا، بیشترین تعداد نمونه و یا بیشترین قوانین تولید شده، می‌توانند کاندید انتخاب به عنوان پاسخ پیش‌فرض باشند.

در فصل بعد مراحل استخراج قانون به کمک روش FSRE بررسی و رویکردی برای حل چالش‌های موجود ارائه خواهد شد.

۵.۲ نتیجه‌گیری

در این فصل با انواع قوانین قابل استخراج از شبکه‌های عصبی آشنا شدیم. در ادامه به بررسی دسته‌بندی‌های موجود برای روش‌های استخراج قانون ارائه شده پرداختیم. هر روش استخراج قانون دارای خصوصیتی می‌باشد که سبب می‌شود در دسته‌های مختلفی قرار بگیرد. در ادامه برخی از مهم‌ترین روش‌هایی که تا کنون ارائه

شده‌اند را مرور کردیم. روش FSRE یکی از جدیدترین و کاراترین روش‌های موجود برای استخراج قانون است که توانایی حل مسائل پیچیده را دارد. در فصل بعد به تشریح این روش خواهیم پرداخت.

فصل ۳

روش پیشنهادی

۱.۳ مقدمه

در این فصل ابتدا به تشریح مراحل استخراج قانون توسط روش FSRE می‌پردازیم. سپس چالش‌های موجود و در این روش و راهکارهای پیشنهادی و اهمیت اعمال آن‌ها بیان خواهد شد. در انتها روش پیشنهادی ارائه خواهد شد.

۲.۳ مراحل استخراج قانون به روش FSRE

FSRE یک روش برای استخراج قانون از شبکه‌عصبی چند لایه پرسپترون است که توانایی استخراج قوانین گزاره‌ای if-then-else برای مسائل دسته‌بندی را دارد. استخراج قانون در این روش با استفاده از دیدگاه تجزیه‌ای انجام می‌شود. داده‌های آموزشی در صورت نیاز ابتدا پیش‌پردازش شده و سپس برای آموزش به شبکه‌عصبی ارسال می‌شوند. استخراج قانون در گام آخر در چهار مرحله و با استفاده از ساختار هرس شده شبکه انجام می‌شود. در ادامه به بررسی هر کدام از این مراحل خواهیم پرداخت.

۱.۲.۳ پیش‌پردازش داده‌ها

برای دستیابی به کارایی مناسب، اغلب نیاز است قبل از آموزش شبکه، پیش‌پردازش‌هایی بر روی داده‌های خام با هدف تبدیل مجموعه داده‌های خام به داده‌های کارآمد و بهینه، اعمال شود. نرمال‌سازی، گسسته‌سازی، کاهش بعد و نمونه‌برداری از جمله پیش‌پردازش‌های رایجی هستند که می‌توانند این داده‌های خام را به داده‌های کارآمد در جهت افزایش کارایی شبکه تبدیل کنند. این پیش‌پردازش‌ها متناسب با مجموعه داده مورد نظر انتخاب می‌شوند. روش FSRE با هر دو نوع داده پیوسته و گسسته سازگاری دارد و محدودیتی برای استفاده از هر کدام از این داده‌ها وجود ندارد.

پراکندگی و اندازه‌های متفاوت طول بردار ویژگی در پایگاه‌داده، سبب می‌شود که یک ویژگی با مقادیر بزرگ نسبت به سایر ویژگی‌ها، تاثیر بیشتری در آموزش شبکه داشته باشد. برای رفع این مشکل، مقادیر هر ویژگی

ورودی با استفاده از رابطه ۱.۳ به کمک میانگین و انحراف معیار نرمال سازی می‌شوند.

$$x'_i = \frac{x_i - \mu}{\sigma} \quad (۱.۳)$$

در این رابطه x_i مقدار اولیه یک ویژگی برای نمونه i ام، μ میانگین مقادیر آن ویژگی برای تمام نمونه‌ها، σ انحراف معیار استاندارد ویژگی مورد بررسی و x'_i مقدار نرمال شده برای نمونه i ام می‌باشد.

از روش بهره اطلاعاتی^۱ برای ارزیابی ویژگی‌ها از نظر میزان کارایی و اهمیت استفاده می‌شود. این روش، یک خصوصیت آماری مبتنی بر بی‌نظمی^۲ می‌باشد که عبارت است از میزان کاهش بی‌نظمی که به واسطه جداسازی نمونه‌ها از طریق یک ویژگی حاصل می‌شود و برای هر ویژگی به صورت مجزا بدست می‌آید.

در صورت نیاز می‌توان در مرحله پیش پردازش مقادیر پیوسته را گسسته سازی کرد. روش‌های CAIM [۳۱]، Fast-CAIM [۳۰]، Modified CAIM [۵۷] و ur-CAIM [۱۰] رایج‌ترین روش‌های موجود برای گسسته سازی هستند.

پس از انجام عملیات پیش پردازش باید با استفاده از یک روش مناسب داده‌ها کدگذاری شوند، تا بهترین نتیجه حاصل شود. روش‌های باینری موجب افزایش ورودی‌های شبکه (بخصوص ۱ از N) و ساده تر شدن مسئله می‌شوند. در حالت‌های نرمال شده فضای تغییرات هر ویژگی در محدوده کوچک تری قرار می‌گیرد که از نوسان در آموزش جلوگیری می‌کند. همچنین حالت فازی برای مسائل فازی و تولید قوانین فازی استفاده می‌شود. در انتهای عملیات انجام شده در این فاز، داده‌ها به سه دسته مجزای آموزش، ارزیابی و تست تقسیم می‌شوند.

¹Information gain

²Entropy

۲.۲.۳ تعیین ساختار مناسب و آموزش شبکه

آموزش شبکه با استفاده از یک شبکه استاندارد سه لایه MLP انجام می‌شود. تعداد نرون‌های لایه ورودی متناسب با تعداد ویژگی‌های ورودی پایگاه داده تعیین می‌شود. تعداد نرون‌های لایه خروجی نیز متناسب با تعداد کلاس‌های موجود انتخاب می‌شوند. تعداد نرون‌های لایه میانی در ابتدای آموزش یک در نظر گرفته می‌شود. در صورت عدم دستیابی به کارایی مناسب یک واحد به تعداد آن‌ها افزوده می‌شود، این کار تا زمان دستیابی شبکه به کارایی مطلوب ادامه می‌یابد. آموزش با داده‌های آموزشی و کارایی با مجموعه تست محاسبه می‌گردد. تابع MSE^1 برای محاسبه خطا استفاده می‌شود. جهت اندازه‌گیری میزان کارایی از پارامترهای دقت، حساسیت^۲ و اختصاصی بودن^۳ استفاده می‌شود. حساسیت و اختصاصی بودن تنها برای حالت دو کلاسه استفاده می‌شود. توابع فعال‌ساز نرون‌های لایه میانی و خروجی و همچنین وجود بایاس در هر لایه متناسب با شبکه انتخاب می‌شود. در پایان آموزش شبکه با استفاده از روش گرادیان نزولی^۴ و الگوریتم یادگیری پس انتشار خطا^۵ صورت می‌گیرد.

۳.۲.۳ هرس‌سازی شبکه

وجود اتصالات زیاد در شبکه باعث می‌شود که استخراج قانون بسیار پیچیده یا غیرممکن شود. به همین منظور برخی اتصالات و گره‌های کم اهمیت در این مرحله حذف می‌شوند. هرس‌سازی باعث ساده‌تر شدن ساختار شبکه، شناسایی گره‌ها و اتصالات موثر در شبکه و در نهایت تعیین ویژگی‌های پراهمیت برای هر کلاس می‌شود. این عملیات علیرغم کاهش ابهام در شبکه، باعث کاهش کارایی شبکه نیز می‌شود که سبب می‌شود قوانین

¹Mean Squared Error

²Sensitivity

³Specificity

⁴Gradient descent

⁵Backpropagation

استخراج شده از این ساختار دقت کمتری نسبت به شبکه اصلی داشته باشد. در نتیجه هرس سازی باید به گونه ای صورت بگیرد که علاوه بر کاهش پیچیدگی، دقت شبکه تا حد مطلوبی حفظ شود.

هرس سازی در روش FSRE در دو مرحله صورت می پذیرد:

- در گام اول اتصالات مابین لایه مخفی و خروجی هرس شده که باعث شناسایی نرون های موثر هر کلاس خروجی می شود. در این الگوریتم ممکن است تمام اتصالات یک نرون با خروجی حذف گردند. در این حالت این نرون و همه اتصالات آن از شبکه حذف می شود. اگر اتصالات یک نرون دارای مقادیر مختلفی باشند، اتصالاتی که دارای بزرگترین وزن هستند حفظ و باقی حذف می شوند.

- در گام بعد اتصالات مابین لایه ورودی و مخفی هرس می شوند تا پارامترهای ورودی موثرتر شناسایی شوند. پارامترهای ضروری هر نرون لایه مخفی شناسایی و باقی حذف می شوند. در صورتیکه تمام اتصالات یک ورودی حذف گردند، این ویژگی ضروری نبوده و حذف می شود. در این الگوریتم یک اتصال حذف می شود و سپس کارایی شبکه سنجیده می شود.

۴.۲.۳ استخراج قانون

عملیات نهایی استخراج قانون با استفاده از ساختار هرس شده شبکه انجام می شود. قوانین گزاره ای در چهار مرحله برای مسائل دسته بندی استخراج می شوند. در این روش از الگوی نمونه های ورودی در لایه مخفی جهت تولید قوانین استفاده می گردد، تشکیل این الگوها با توجه به خروجی تابع فعال ساز نرون های لایه میانی انجام می پذیرد.

وجود مقادیر پیوسته در خروجی لایه مخفی موجب تشکیل الگوهای بسیار زیادی می شود. از این رو استخراج قانون با استفاده از این الگوها بسیار پیچیده و حتی غیرممکن خواهد بود. به همین منظور نیاز است

تا خروجی تابع فعال ساز برای هر نرون لایه میانی گسسته سازی شود. روش های گسسته سازی مختلفی برای اینکار ارائه شده است. در این روش، با توجه به توزیع مقادیر خروجی هر نرون به صورت مجزا، گسسته سازی با تعیین مقادیر آستانه و به صورت دستی انجام می شود.

پس از اتمام فاز گسسته سازی، عملیات استخراج قانون با استفاده از مقادیر گسسته در لایه میانی انجام می شود. این عملیات در چهار مرحله صورت می گیرد:

۱. ابتدا تمام نمونه ها به شبکه هرس شده اعمال می شوند تا الگوی هر نمونه تعیین شود. الگوی یک نمونه با کنار هم قرار دادن مقادیر گسسته شده خروجی تابع فعال ساز نرون های لایه میانی، تشکیل می شود.

۲. برخی نمونه های ورودی مقادیر گسسته مشابهی در لایه مخفی تولید می کنند، که باعث تولید شدن الگوهای یکسانی برای برخی نمونه ها می شود. هر الگوی بدست آمده متعلق به یک کلاس خروجی مشخص است که با توجه به خروجی شبکه تعیین می گردد. در این گام برای هر الگوی بدست آمده یک دسته جدا در نظر گرفته می شود. در هر دسته نمونه های ورودی مشابه قرار می گیرند.

۳. با توجه به حذف شدن برخی ویژگی های ورودی در فاز هرس سازی، ممکن است نمونه های موجود در یک دسته مقادیر یکسانی داشته باشند. همچنین ممکن است برخی نمونه ها مقادیر مشابهی داشته باشند. در این مرحله نمونه های تکراری حذف و نمونه های مشابه ترکیب می شوند. در نهایت دسته های دارای کلاس خروجی یکسان با هم ادغام می شوند. بار دیگر عملیات حذف نمونه های تکراری و ترکیب نمونه های مشابه بر روی دسته های متعلق به یک کلاس انجام می شود.

۴. در انتها برای اینکه قوانین بتوانند نمونه های دیده نشده را نیز کلاس بندی کنند، یکی کلاس ها به عنوان پاسخ پیش فرض تعیین می گردد. دسته های دارای بالاترین مقدار خطا، بیشترین تعداد نمونه و

یا بیشترین قوانین تولید شده می‌توانند کاندید انتخاب به عنوان پاسخ پیش‌فرض باشند.

۳.۳ چالش‌ها و راهکارها

وزن‌های بدست آمده با استفاده از روش پس انتشار خطا و گرادیان نزولی، تنها با هدف کاهش خطا و افزایش کارایی شبکه تولید می‌شوند. برخی از این وزن‌ها مقادیر بزرگی دارند که باعث می‌شود در تعیین خروجی تاثیر قابل توجهی داشته باشند، گروهی دیگر از اتصالات دارای مقادیر کوچک و نزدیک به صفر هستند، در نتیجه اغلب تاثیر چندانی در خروجی شبکه ندارند و در مرحله هرس‌سازی حذف می‌شوند. گروه سوم اتصالاتی هستند که تاثیر آن‌ها در کارایی شبکه کم است اما بی‌تاثیر نیستند. الگوریتم‌های هرس‌سازی در جهت حفظ اتصالات موثر و حذف اتصالات بی‌اثر عمل می‌کنند، اما چالش اصلی موجود مربوط به اتصالاتی است که تاثیر کمی در خروجی شبکه دارند. رویکرد الگوریتم‌های هرس‌سازی، در رابطه با این اتصالات، حذف آن‌ها تا حد امکان، به منظور استخراج قوانین ساده‌تر می‌باشد.

قوانین استخراج شده به کمک شبکه‌های عصبی معمولاً دقت پایین‌تری از خود شبکه دارند. دلیل این امر، متکی بودن این روش‌ها به هرس‌سازی شبکه و استخراج قانون از ساختار هرس شده شبکه می‌باشد. با هرس‌سازی، یال‌های کم اهمیت (وزن‌های نزدیک به صفر) حذف می‌شوند که موجب کاهش دقت و کارایی شبکه می‌شود. از آنجایی که این عملیات بخش جداناپذیر و مهمی از روش‌های استخراج قانون به کمک شبکه‌عصبی است، با ارائه یک راهکار نوین در این تحقیق خسارات ناشی از آن را به میزان قابل توجهی کاهش می‌دهیم.

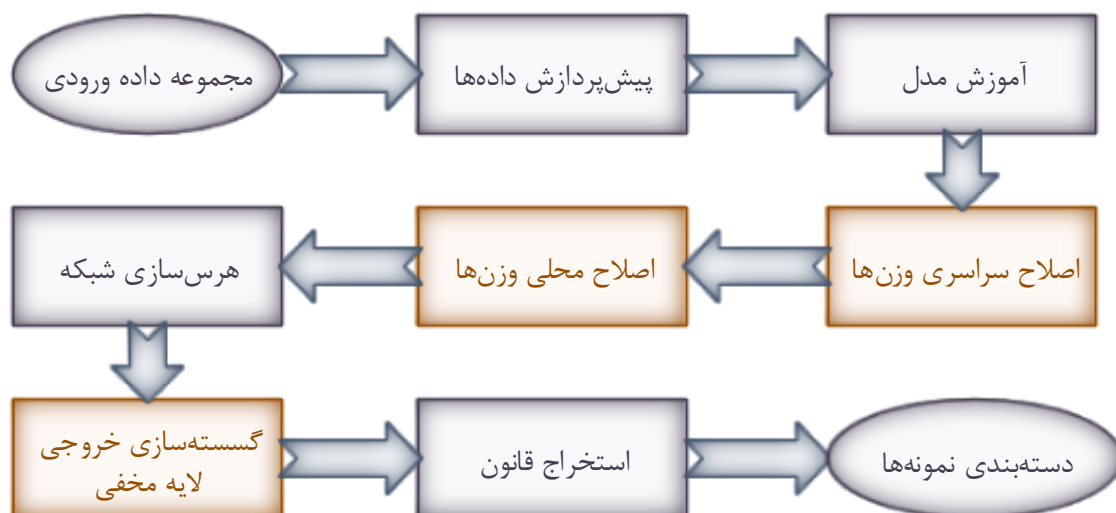
وزن‌های یک شبکه عصبی آموزش دیده را می‌توان به سه دسته کم اهمیت، اهمیت متوسط، و پراهمیت تقسیم بندی کرد. وزنهای کم اهمیت، به صفر نزدیک هستند و با حذف آنها (هرس یالهای مربوطه) تاثیر

چندانی بر عملکرد شبه عصبی نخواهد گذاشت. در مرحله هرس، یالهای با اهمیت حفظ می‌شوند. چنانچه با هرس یالهای کم اهمیت نتوان مجموعه قوانین مناسبی را استخراج نمود، یالهای با اهمیت متوسط نیز حذف خواهند شد. البته حذف یالهای با اهمیت متوسط به میزان قابل توجهی بر دقت شبکه و در نتیجه دقت مجموعه قوانین استخراج شده تاثیر خواهد داشت. در روش پیشنهادی، با استفاده از الگوریتم ژنتیک، آموزش به گونه‌ای انجام می‌شود که وزن‌های تولید شده غیرمتوازن باشند. به عبارت دیگر، شبکه فاقد یال‌های با اهمیت متوسط باشد. براین اساس، با حفظ وزن‌های پر اهمیت و حذف وزن‌های کم اهمیت خسارت ناشی از مرحله هرس‌سازی تقلیل می‌یابد. این امر باعث می‌شود که شبکه عصبی پس از هرس‌سازی، و در نتیجه مجموعه قوانین استخراج شده از آن با کاهش دقت همراه باشد. روش پیشنهادی در این تحقیق، با هدف تولید هدفمند وزن‌های غیر متوازن با استفاده از الگوریتم ژنتیک، کاهش دقت شبکه عصبی هرس شده را کنترل نماید. این کار به گونه‌ای در مرحله آموزش شبکه عصبی انجام می‌شود که ضمن حفظ دقت شبکه، مجموعه وزنهای آن به دو گروه موثر و کم تاثیر (نزدیک به صفر) قابل دسته بندی باشند.

ساختار هرس شده شبکه عصبی، فقط شامل اتصالات موثر شبکه است. با این حال استخراج قانون، به دلیل وجود مقادیر پیوسته در خروجی لایه میانی، به راحتی امکان پذیر نمی‌باشد. از این رو گسسته‌سازی این مقادیر یکی از اجزای مهم و جدا ناپذیر در استخراج قانون از شبکه عصبی با دیدگاه تجزیه‌ای می‌باشد. تشکیل دسته‌های صحیح برای استخراج قوانین مربوط به هر کلاس، سادگی قوانین استخراج شده و کارایی مناسب آن‌ها به طور مستقیم وابسته به عملیات گسسته‌سازی هستند.

عملیات گسسته‌سازی خروجی لایه میانی در روش FSRE با صرف وقت زیاد به صورت دستی و با استفاده از ترسیم خروجی هر یک از نرون لایه میانی انجام می‌شود. از این رو گسسته‌سازی نیز علاوه بر هرس‌سازی یکی از چالش‌هایی است که قصد داریم در این پژوهش به آن بپردازیم. برای رفع این مشکل قصد داریم با استفاده از

الگوریتم خوشه‌بندی ارائه شده در [۴۳]، دقت و کارایی قوانین استخراج شده با استفاده از این روش را بهبود دهیم. فرآیند استخراج قانون با استفاده از روش پیشنهادی در شکل ۱.۳ نشان داده شده است.



شکل ۱.۳: فرآیند استخراج قانون با استفاده از روش پیشنهادی

۴.۳ معرفی الگوریتم ژنتیک

غیرخطی بودن توابع هدف یا محدودیت‌ها، گسسته بودن فضای حل مسئله و اندازه مسئله از جمله عوامل تاثیرگذار در پیچیدگی محاسباتی مدل‌های تصمیم‌گیری هستند. روش‌های دقیق در این شرایط قادر به پیدا کردن جواب بهینه در زمان قابل قبولی نیستند. هدف روش‌های فراابتکاری، پیدا کردن بهترین جواب ممکن (نزدیک به بهینه) در زمان قابل قبول است.

در سی سال گذشته، نوع جدیدی از الگوریتم‌های تقریب‌ظهور یافته‌اند که اساساً هدف از آنها ترکیب روش‌های ابتکاری در چارچوب‌های کلان‌تر به منظور کاوش کارا و اثربخش فضای جستجو می‌باشد. امروزه از این روش‌ها با عنوان روش‌های فرا ابتکاری^۱ نام برده می‌شود. یک روش فرا ابتکاری، در حقیقت، یک روش

^۱ Meta-heuristic

ابتکاری برای حل یک طبقه بسیار عمومی از مسائل است و ترکیبی کارآمد از توابع هدف یا روش‌های ابتکاری مبتنی بر توابع هدف می باشد.

الگوریتم ژنتیک، الهامی از علم ژنتیک و نظریه تکامل داروین است و بر اساس بقای برترین‌ها یا انتخاب طبیعی استوار است. یک کاربرد متداول الگوریتم ژنتیک، استفاده از آن بعنوان تابع بهینه کننده است. الگوریتم ژنتیک ابزار سودمندی در بازشناسی الگو، انتخاب ویژگی، درک تصویر و یادگیری ماشینی است. در الگوریتم‌های ژنتیکی، نحوه تکامل ژنتیکی موجودات زنده شبیه‌سازی می‌شود. این الگوریتم‌ها با الهام از روند تکاملی طبیعت مسائل را حل می‌نمایند. یعنی مانند طبیعت یک جمعیت از موجودات را تشکیل می‌دهند و با اعمالی بر روی این مجموعه به یک مجموعه بهینه و یا موجود بهینه دست می‌یابند. با توجه به خصوصیات خاص خودشان به خوبی از عهده حل مسائلی که نیاز به بهینه‌سازی دارند بر می‌آیند.

الگوریتم ژنتیک نوع خاصی از الگوریتم‌های تکاملی است که از تکنیک‌های زیست‌شناسی مانند وراثت و جهش استفاده می‌کند. در این الگوریتم ابتدا یک جمعیت اولیه تصادفی از کروموزوم‌ها^۱ تولید شده و با استفاده از تابع برازندگی^۲ مناسب، ارزیابی می‌شوند. در هر مرحله تعدادی از کروموزوم‌های برتر به عنوان کروموزوم نخبه به نسل بعد منتقل می‌شوند. سایر کروموزوم‌های نسل بعد با استفاده از عملگرهای ژنتیک تولید می‌شوند. این عملیات تا زمانی ادامه می‌یابد که شرط خاتمه الگوریتم برقرار شود. در نهایت کروموزوم برتر به عنوان جواب مسئله به خروجی منتقل می‌شود.

۱.۴.۳ کروموزوم

برای اینکه بتوانیم یک مسئله را بوسیله الگوریتم‌های ژنتیک حل کنیم، بایستی آن را به فرم مخصوص مورد نیاز این الگوریتم‌ها تبدیل کنیم. در این روند بایستی راه حل مورد نیاز مسئله را به گونه‌ای تعریف کنیم که

^۱Chromosome

^۲Fitness Function

قابل نمایش بوسیله یک کروموزوم باشد. اینکه چه نوع بازنمایی^۱ برای مسئله استفاده شود، به شخص طراح و فرم مسئله بستگی دارد. چند نمونه از بازنمایی‌هایی را که معمولاً استفاده می‌شوند:

- اعداد صحیح
- رشته‌های بیتی
- اعداد حقیقی در فرم نقطه شناور
- اعداد حقیقی به فرم رشته‌های بیتی
- یک مجموعه از اعداد حقیقی یا صحیح
- ماشین‌های حالت محدود
- هر فرم دیگری که بتوانیم عملگرهای ژنتیک را بر روی آنها تعریف کنیم.

۲.۴.۳ تابع برازندگی

برای اینکه بتوانیم موجودیت‌های^۲ بهتر را درون جمعیت تشخیص بدهیم بایستی معیاری را تعریف کنیم که بر اساس آن موجودیت‌های بهتر را تشخیص دهیم. به این کار، یعنی تعیین میزان خوبی یک موجود، ارزیابی آن موجود می‌گویند. ارزیابی، اینگونه است که بر حسب اینکه موجود چقدر خوب است یک عدد به آن نسبت می‌دهیم، این عدد که برای موجودات بهتر بزرگتر (یا کوچکتر) است را شایستگی آن موجود می‌نامیم. به عنوان مثال در صورتی که به دنبال مینیمم یک تابع هستیم، مقدار شایستگی را می‌توانیم ورودیهایی که مقادیر تابع

^۱Representation

^۲Entity

برای آنها کمتر است در نظر بگیریم که ورودیهای بهتری هستند. بسته به نوع مساله ما می خواهیم شایستگی را بیشینه و یا کمینه کنیم.

۳.۴.۳ عملگرهای ژنتیک

برای پیدا کردن يك نقطه در فضای جستجو باید از عملگرهای ژنتیک استفاده کرد. در ادامه عملگرهای انتخاب^۱، ترکیب^۲ و جهش^۳ توضیح داده خواهند شد.

عملگر انتخاب

در مرحله انتخاب، يك جفت از کروموزومها برگزیده می شوند تا با هم ترکیب شوند. عملگر انتخاب رابط بین دو نسل است و بعضی از اعضای نسل کنونی را به نسل آینده منتقل می کند. بعد از انتخاب، عملگرهای ژنتیک روی دو عضو برگزیده اعمال می شوند. معیار در انتخاب اعضاء ارزش تطابق آنها می باشد اما روند انتخاب حالتی تصادفی دارد.

شاید انتخاب مستقیم و ترتیبی به این شکل که بهترین اعضا دو به دو انتخاب شوند در نگاه اول روش مناسبی به نظر برسد اما باید به نکته ای توجه داشت. در الگوریتم ژنتیک ما با ژنها روبرو هستیم. يك عضو با تطابق پایین اگرچه در نسل خودش عضو مناسبی نمی باشد اما ممکن است شامل ژنهایی خوب باشد و اگر شانس انتخاب شدنش ° باشد، این ژنهای خوب نمی توانند به نسل های بعد منتقل شوند. پس روش انتخاب باید به گونه ای باشد که به این عضو نیز شانس انتخاب شدن بدهد. راه حل مناسب، طراحی روش انتخاب به گونه ای است که احتمال انتخاب شدن اعضای با تطابق بالاتر بیشتر باشد. انتخاب باید به گونه ای صورت بگیرد که تا جایی که ممکن است هر نسل جدید نسبت به نسل قبلی اش تطابق میانگین بهتری داشته باشد.

¹Selection

³Mutation

²CrossOver

روش‌های متداول انتخاب عبارتند از:

● انتخاب چرخ رولت^۱

● انتخاب ترتیبی^۲

● انتخاب بولتزمن^۳

● انتخاب حالت پایدار^۴

● نخبه سالاری^۵

● انتخاب رقابتی^۶

عملگر ترکیب

مهمترین عملگر در الگوریتم ژنتیک، عملگر ترکیب است. ترکیب، فرآیندی است که در آن نسل قدیمی کروموزوم‌ها با یکدیگر مخلوط و ترکیب می‌شوند تا نسل تازه‌ای از کروموزوم‌ها بوجود بیاید. جفت‌هایی که در قسمت انتخاب، به عنوان والد در نظر گرفته شدند، در این قسمت ژنهایشان را با هم مبادله می‌کنند و اعضای جدید بوجود می‌آورند.

بر اساس مباحث تئوری ترکیب اعضا با تطابق بالا باعث بوجود آمدن اعضای می‌شود که از تطابق میانگین، تطابق بیشتری دارند. ترکیب در الگوریتم ژنتیک باعث از بین رفتن پراکندگی یا تنوع ژنتیکی جمعیت می‌شود. زیرا اجازه می‌دهد ژن‌های خوب یکدیگر را بیابند. شکل‌های مختلفی از عملگر ترکیب وجود دارد که مهمترین‌های آنها عبارتند از ترکیب تک نقطه‌ای، ترکیب دو نقطه‌ای، ترکیب n نقطه‌ای، ترکیب یکنواخت و ترکیب حسابی.

¹ Roulette wheel Selection

² Rank Selection

³ Boltzmann Selection

⁴ Steady-State Selection

⁵ Elitism

⁶ Tournament Selection

ترکیب لازم نیست در هر نسل اتفاق بیفتد. در واقع ممکن است نسل‌هایی بدون عملگر ترکیب به نسل‌های جدید، تبدیل شوند. برای تعیین رخ دادن یا ندادن ترکیب از پارامتری به نام احتمال ترکیب P_c ، استفاده می‌شود که این پارامتر مقدار بین ۰ و ۱ است. از آنجا که ترکیب نقشی اساسی در رشد مقدار میانگین تطابق جمعیت دارد، مقدار P_c بین ۰,۵ تا ۰,۸ و بیشتر بین ۰,۷ تا ۰,۸ در نظر گرفته می‌شود.

عملگر جهش

بعد از اینکه يك عضو در جمعیت جدید بوجود آمد، هر ژن آن با احتمال جهش، جهش می‌یابد. در جهش ممکن است ژنی از مجموعه ژن‌های جمعیت حذف شود یا ژنی که تا به حال در جمعیت وجود نداشته است به آن اضافه شود. جهش يك ژن به معنای تغییر آن ژن است و وابسته به نوع کدگذاری، روش‌های متفاوت جهش استفاده می‌شود.

همانطور که گفته شد، هر عضو وابسته به احتمال جهش، جهش می‌یابد. احتمال جهش، P_m مقداری است که توسط کاربر تعیین می‌شود. در الگوریتم استاندارد ژنتیک، مقدار این پارامتر، بسیار کوچک، مثل $P_m = 0,01$ یا حتی $P_m = 0,001$ در نظر گرفته می‌شود.

P_m مقداری بین ۰ و ۱ است و معمولاً عددی کوچک انتخاب می‌شود. در فرزندى که بوجود آمده است (توسط ترکیب)، به ترتیب مقداری تصادفی بین ۰ و ۱ به هر ژن اختصاص می‌یابد. اگر این مقدار اختصاص داده شده از P_m کمتر باشد، ژن جهش می‌یابد و اگر بیشتر باشد، ژن تغییر نمی‌کند. نرخ بالای جهش، باعث تنوع و پراکندگی ژنتیکی در جمعیت می‌شود که این پراکندگی ممکن است همگرایی را به تاخیر بیندازد. به همین دلیل، پیشنهاد می‌شود برای جمعیت‌های بزرگ یا در نسل‌های آخر، از P_m های کوچکتر و برای جمعیت‌های کوچک یا در نسل‌های ابتدایی از P_m های بزرگتر استفاده شود. وارونه‌سازی بیت، تغییر ترتیب قرارگیری و تغییر مقدار نمونه‌هایی از انواع جهش هستند.

۵.۳ روش پیشنهادی

همانطور که بیان شد چالش‌های اساسی مورد بررسی در روش FSRE در مرحله هرس‌سازی و گسسته‌سازی می‌باشد. در این بخش قصد داریم به معرفی راهکارهای ارائه شده برای بهبود این الگوریتم در این دو بخش بپردازیم. همانطور که در بخش ۳.۳ اشاره شد چالش اصلی، حذف برخی وزن‌های میانی است که باعث می‌شود برخی از دانش موجود در شبکه از دست برود. برای رفع این مشکل رویکردی مبتنی بر الگوریتم ژنتیک ارائه شده است که قادر است وزن‌های شبکه عصبی آموزش دیده را در دو مرحله به صورت سراسری و محلی به شکلی اصلاح کند که وزن‌های حاصل به صورت غیرمتوازن باشند. به صورت ساده‌تر قصد داریم وزن‌ها فقط در دو دسته موثر و بی‌اثر قرار بگیرند و دسته میانی تا جای ممکن حذف شود. این امر باعث می‌شود هرس‌سازی به شکل بهینه و کاراتر انجام شود. الگوریتم سراسری، وزن‌های شبکه را با استفاده از الگوریتم ژنتیک به هنگام آموزش شبکه اصلاح می‌کند. الگوریتم محلی نیز به اصلاح وزن‌های شبکه آموزش دیده برای هر نرون لایه میانی به صورت جداگانه می‌پردازد. این الگوریتم‌ها در ادامه تشریح خواهند شد.

اهمیت فاز گسسته‌سازی در استخراج قانون برای استخراج قوانین دقیق تر و کاراتر بیان شد. در گام بعد سعی داریم که با ارائه یک الگوریتم خوشه‌بندی برای گسسته‌سازی مقادیر خروجی تابع فعال‌ساز نرون‌های لایه میانی به اهداف زیر دست یابیم.

- این الگوریتم عملیات گسسته‌سازی را به صورت پویا و خودکار انجام دهد.

- گسسته‌سازی به شکل کاملاً بهینه و با حداقل خطای ممکن انجام شود.

۱.۵.۳ غیرمتوازن کردن وزن‌ها به صورت سراسری

همانطور که بیان شد، هدف روش پیشنهادی تولید وزن‌های شبکه در دو دسته موثر و کم تاثیر می‌باشد. بدین منظور آموزش شبکه با استفاده از الگوریتم ژنتیک انجام می‌شود. حل مسئله با استفاده از الگوریتم ژنتیک، نیازمند تعریف کروموزوم و تابع برازندگی مناسب است. با توجه به اینکه جواب مسئله ما در الگوریتم ژنتیک وزن‌های شبکه هستند، در نتیجه ساختار کروموزوم را این وزن‌ها تشکیل خواهد داد. بدین صورت که وزن‌های لایه ورودی به مخفی و لایه مخفی به خروجی را با ترتیب مشخصی به صورت متوالی کنار هم قرار می‌دهیم. هر وزن شبکه‌عصبی به عنوان یک ژن در کروموزوم قرار خواهد گرفت.

چالش اصلی این بخش تعریف یک تابع برازندگی مناسب جهت امتیازدهی به این کروموزوم‌ها می‌باشد. تابع برازندگی باید در عین حال که کارایی شبکه را حفظ می‌کند، به انتخاب کروموزوم‌هایی بپردازد که شامل وزن‌های غیرمتوازن بهتری باشد. در نتیجه تابع برازندگی مورد نیاز در این مسئله از دو بخش تشکیل می‌شود.

- قسمت اول مربوط به کارایی شبکه می‌شود، و می‌توان از معیارهای مجموع مربعات خطا (SSE^1)، میانگین مربعات خطا (MSE)، یا دقت ($Accuracy$) شبکه استفاده کرد.

- در قسمت دوم باید یک معیار جهت امتیازدهی به میزان دو دسته شدن وزن‌ها در هر کروموزوم تعریف کنیم. رابطه ۲.۳ به صورت تجربی برای امتیازدهی به غیر متوازن بودن وزن‌ها تعریف شده است.

$$score = \frac{1}{\sqrt{var(|w|) * R}} \quad (2.3)$$

در این رابطه، $score$ امتیاز مربوط به یک کروموزوم است که با توجه به میزان دو دسته شدن وزن‌های

¹Sum of Squared Error

آن تعیین می‌شود. w بردار وزن‌های موجود در یک کروموزوم، R نسبت تعداد وزن‌های کم‌تاثیر یک کروموزوم به کل وزن‌ها، var واریانس بردار ورودی و $||$ بیانگر قدرمطلق بردار وزن‌ها است.

استفاده از واریانس در این رابطه به منظور ایجاد پراکندگی در وزن‌های شبکه می‌باشد. استفاده از این معیار به تنهایی ممکن است تعداد وزن‌های موثر شبکه را به میزان قابل توجهی افزایش دهد، در این صورت این وزن‌ها جهت هرس‌سازی شبکه مناسب نخواهند بود. پارامتر R برای جلوگیری از انتخاب این دسته از کروموزوم‌ها تعریف شده است. با انتخاب یک مقدار آستانه مناسب می‌توان تخمینی از تعداد وزن‌های غیر موثر در هر کروموزوم را بدست آورد. پارامتر R حاصل تقسیم این عدد بر تعداد کل اتصالات است. بنابراین مقدار این پارامتر همواره عددی بین ۰ و ۱ خواهد بود. در صورتیکه تعداد وزن‌های بی‌اثر کم باشد، R عددی نزدیک به صفر خواهد بود که باعث می‌شود مقدار $score$ عددی بزرگ باشد. با توجه به اینکه در نهایت الگوریتم ژنتیک سعی در کمینه‌سازی $score$ را دارد، با این کار احتمال انتخاب این کروموزوم کاهش خواهد یافت.

با ادغام این دو معیار می‌توانیم وزن‌های مورد نظر خود را تولید کنیم. با ترکیب معیار $score$ با یک معیار ارزیابی شبکه، که در اینجا از SSE استفاده کرده‌ایم، تابع برازندگی نهایی در رابطه ۳.۳ بیان شده است که نتایج قابل قبولی را ارائه می‌دهد.

$$fitness_value = \frac{SSE + score}{2} \quad (3.3)$$

در این رابطه، $score$ همانطور که بیان شد امتیاز غیرمتوازن بودن وزن‌ها، SSE معیار ارزیابی شبکه و $fitness_value$ مقدار تابع برازندگی می‌باشد.

عملیات آموزش شبکه با استفاده از الگوریتم ژنتیک در جهت اصلاح وزن‌ها انجام خواهد شد. الگوریتم ژنتیک سعی در کمینه‌سازی تابع برازندگی معرفی شده دارد. با کمینه کردن این تابع علاوه بر دستیابی به وزن‌های مناسب برای هرس‌سازی، کارایی شبکه را تا حد مطلوبی حفظ خواهد شد. برای اینکه کارایی شبکه پس از آموزش به روش گرادیان نزولی و الگوریتم پس انتشار خطا حفظ شود، وزن‌های تولید شده را به عنوان یک کروموزوم در جمعیت اولیه قرار می‌دهیم. این کروموزوم به عنوان یک کروموزوم نخبه باعث سرعت بخشیدن به الگوریتم و حفظ حداقل کارایی شبکه خواهد شد.

۲.۵.۳ غیر متوازن کردن وزن‌ها به صورت محلی

در این الگوریتم بهبود وزن‌ها با استفاده از الگوریتم ژنتیک با بررسی وزن‌های ورودی به هر نرون به صورت جداگانه صورت می‌گیرد. در این روش ورودی هر نرون لایه خروجی و مخفی را به صورت جداگانه با استفاده از الگوریتم ژنتیک اصلاح می‌کنیم. این روش می‌تواند به دو صورت انجام پذیرد:

- در روش اول سعی می‌شود با اصلاح وزن‌های ورودی به شکل دلخواه، خروجی هر نرون به ازای نمونه‌های ورودی آموزش بدون تغییر باقی بماند.

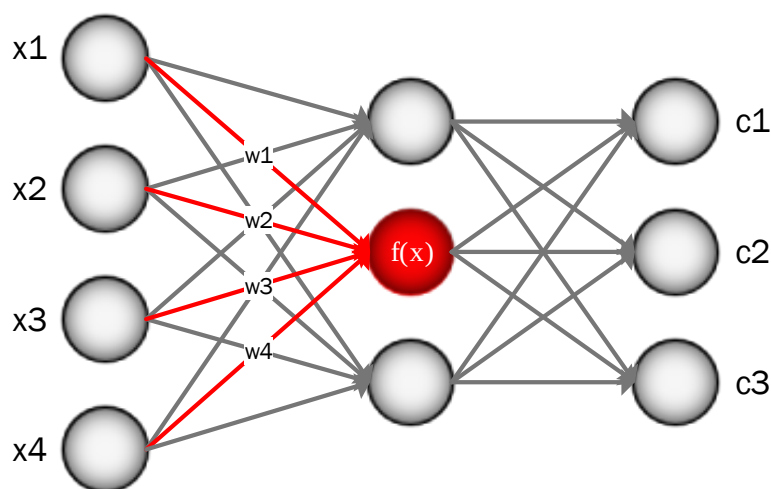
- روش دوم به عنوان یک روش ساده‌تر، استفاده از یک تابع برازندگی برای بررسی عملکرد کلی شبکه به ازای تغییر ورودی‌های یک نرون است.

تابع برازندگی در روش اول با تشکیل یک چند جمله‌ای به ازای هر نرون لایه میانی تعریف می‌شود. برای نمونه رابطه ۴.۳ مطابق شکل ۲.۳ تابع برازندگی مربوط به یک نرون لایه میانی می‌باشد.

$$diffrence_value = |f(x) - [(w_{\text{ان}} * x_{\text{ان}}) + (w_{\text{ر}} * x_{\text{ر}}) + (w_{\text{ر}} * x_{\text{ر}}) + (w_{\text{ر}} * x_{\text{ر}})]| \quad (۴.۳)$$

در این رابطه w_i ها وزن های ورودی برای یک نرون و x_i ها مقدار ویژگی نمونه های ورودی به شبکه هستند. $f(x)$ خروجی نرون به ازای این ورودی ها است. در این رابطه مقدار $f(x)$ باید همواره به ازای وزن های اولیه شبکه ثابت بماند و با خروجی حاصل از وزن های جدید مقایسه شود.

$diffrence_value$ مقدار تابع برازندگی خواهد بود که همواره مثبت است. با کمینه سازی این تابع خروجی نرون حفظ خواهد شد. همانند روش قبل برای غیرمتوازن کردن وزن ها، $diffrence_value$ می بایست با مقدار $score$ از رابطه ۲.۳ ترکیب شود تا با حفظ کارایی شبکه، وزن های شبکه غیرمتوازن شوند. ساختار کروموزوم در این روش تنها شامل وزن های ورودی به یک نرون خواهد بود و الگوریتم ژنتیک به ازای هر نرون به صورت جداگانه اجرا خواهد شد.



شکل ۲.۳: بررسی ورودی یک نرون به صورت جداگانه

۳.۵.۳ معرفی الگوریتم گسسته سازی

ایده اصلی این الگوریتم، انتخاب مقادیر گسسته با حداقل فاصله $\epsilon \in (0, 1)$ از یکدیگر است. ϵ های خیلی کوچک تضمین می کنند که همواره دقت شبکه با استفاده از این مقادیر گسسته حفظ شود. هر چه ϵ کوچک

باشد، تعداد مقادیر گسسته تولید شده بیشتر خواهد بود. اما وجود مقادیر گسسته زیاد سبب تولید الگوهای بیشتر و قوانین پیچیده‌تر خواهد شد. بنابراین برای تولید دسته‌های کمتر و قوانین ساده‌تر، مطلوب این است که گسسته‌سازی با تعداد کمی مقدار گسسته انجام شود. الگوریتم با یک ϵ بزرگ آغاز می‌شود، در صورتی که مقادیر گسسته قادر نباشند کارایی شبکه را حفظ کنند، الگوریتم دوباره برای ϵ کوچک‌تر اجرا می‌شود. این عمل تا جایی ادامه می‌یابد تا شبکه با حداقل تعداد مقادیر گسسته، کارایی مناسبی داشته باشد. خوشه‌بندی با انتخاب یک ϵ ، برای هر نرون لایه میانی به شرح زیر است:

۱. اگر α_1 خروجی اولین نمونه برای نرون لایه میانی باشد، آنگاه $H_1 = \alpha_1$ مرکز اولین خوشه، $Count_1 = 1$ و $Sum_1 = \alpha_1$ ، به ترتیب تعداد و مجموع نمونه‌های عضو این خوشه و $D = 1$ تعداد خوشه‌های تشکیل شده برای این نرون خواهد بود.

۲. برای تمام نمونه‌های $i = 2, 3, \dots, k$ در مجموعه آموزش، اگر یک j وجود داشته باشد بطوریکه $\min_{j=1, \dots, D} |i - H_j| \leq \epsilon$ ، در این صورت $Count_{j+} = 1$ و $Sum_{j+} = \alpha_i$ و در غیر این صورت یک خوشه جدید تشکیل می‌شود.

۳. در هر خوشه، مقدار H با میانگین تمام مقادیر عضو آن خوشه، جایگزین می‌شود.

زمانی که مقادیر گسسته برای تمام نرون‌های لایه میانی بدست آمدند، دقت شبکه با استفاده از مقادیر گسسته دوباره مورد بررسی قرار می‌گیرد. اگر دقت شبکه به پایین‌تر از حد تعیین شده کاهش یافت، الگوریتم دوباره با مقدار کوچک‌تر ϵ اجرا می‌شود. برای ϵ به اندازه کافی کوچک، همواره می‌توان به دقت شبکه با مقادیر پیوسته دست یافت، هر چند که تعداد مقادیر گسسته افزایش می‌یابد.

۶.۳ نتیجه‌گیری

در این فصل مراحل استخراج قانون با استفاده از روش FSRE تشریح شد. چالش‌های موجود در این روش برای بهبود قوانین استخراج شده مطرح و راهکارهایی ارائه شد. هرس‌سازی با اتلاف و گسسته‌سازی نامناسب در این روش موجب می‌شود تا قوانین کیفیت مناسبی را نداشته باشند. در پایان راهکارهای ارائه شده در سه مرحله توضیح داده شد. در مرحله اول وزن‌ها در حین آموزش در جهت تولید ساختار مناسب برای بهبود هرس‌سازی، اصلاح می‌شوند. در مرحله دوم این اصلاح به صورت محلی و برای هر نرون به صورت جداگانه انجام می‌شود. در پایان برای بهبود قوانین تولید شده، یک روش خوشه‌بندی برای گسسته‌سازی خروجی نرون‌های لایه میانی ارائه شد.

فصل ۴

بررسی و تحلیل نتایج

۱.۴ مقدمه

در این فصل به بررسی و آزمایش روش ارائه شده برای استخراج قانون از شبکه عصبی برای برخی از مجموعه داده‌های استاندارد موجود می‌پردازیم. مسائل مورد بررسی در این فصل عمدتاً ساده می‌باشند و جهت مقایسه نتایج روش پیشنهادی نسبت به سایر روش‌های موجود، به خصوص روش FSRE می‌باشد. نتایج بدست آمده برای هر مجموعه داده به صورت مجزا و با ذکر جزئیات در بخش‌های بعد ارائه شده است. در پایان فصل نیز این نتایج با نتایج حاصل از دیگر روش‌های استخراج قانون مورد مقایسه قرار می‌گیرد.

ماتریس درهم‌ریختگی^۲ ابزار مناسبی برای مسائل دسته‌بندی است که قادر است به سادگی میزان کارایی

²Confusion matrix

شبکه و خطای موجود در دسته‌بندی نمونه‌های مختلف را بررسی کند. این ماتریس، یک ماتریس مربعی است که تعداد سطر و ستون آن به تعداد کلاس‌های موجود در مسئله مورد نظر است. مثالی از این ماتریس در جدول ۱۰۴ نمایش داده شده است. در این مثال اگر فرض کنیم الگوریتمی برای کلاس بندی بین گربه‌ها، سگ‌ها و خرگوش‌ها طراحی کرده‌ایم. فرض کنیم در این مثال ۸ گربه، ۶ سگ و ۱۳ خرگوش داریم. در سطر مربوط به گربه‌ها، ۵ مورد به عنوان گربه و ۳ مورد به عنوان سگ دسته بندی شده‌اند. در صورتی که در سطر مربوط به خرگوش‌ها، تنها چند مورد اشتباه وجود دارد. به سادگی مشاهده می‌شود که عملکرد الگوریتم در تمییز دسته‌های خرگوش‌ها نسبت به گربه‌ها بسیار بهتر است. مشخص است که اعداد روی قطر اصلی ماتریس نمایش تعداد کلاس بندی‌های درست هستند. لذا در صورتی که تمام اعداد غیر روی قطر اصلی صفر باشند، الگوریتم دارای دقت حداکثر است.

جدول ۱۰۴: نمونه‌ای از جدول در هم‌ریختگی

کلاس واقعی				
خرگوش	سگ	گربه		
۰	۳	۵	گربه	کلاس پیش‌بینی شده
۱	۳	۲	سگ	
۱۱	۲	۰	خرگوش	

در این ماتریس، مجموع مقادیر یک سطر خاص نشان دهنده تعداد نمونه‌های واقعی در کلاس مورد نظر از مجموعه داده می‌باشد. درحالی‌که مجموع مقادیر هر ستون، بیانگر تعداد نمونه‌های تخمین زده شده توسط شناسایی کننده برای یک کلاس خاص می‌باشد.

در این بخش مجموعه داده‌های Iris flower، Glasses و Wisconsin Breast Cancer برای ارزیابی میزان کارایی روش پیشنهادی از مجموعه پایگاه‌های داده UCI [۸] تهیه و مورد استفاده قرار گرفتند.

در تمام مسائل مورد بررسی تعداد نرون‌های لایه ورودی به تعداد ویژگی‌های موجود در مجموعه داده و تعداد نرون‌های لایه خروجی به تعداد کلاس‌های موجود در مسئله هستند. همچنین همواره مقادیر کلاس‌های

خروجی به صورت ۱ از N گسسته‌سازی می‌شود که سبب می‌شود برای هر کلاس تنها یکی از مقادیر برابر با ۱ و سایر مقادیر صفر باشد. همچنین توابع فعال‌ساز برای نرون‌های لایه میانی تانژانت هیپربولیک ($\tanh$) و برای نرون‌های لایه خروجی سیگموید (sigmoid) در نظر گرفته شده است. مقدار ضریب یادگیری برای آموزش با استفاده از گرادیان نزولی 0.001 قرار داده شده است، همچنین شرط خاتمه یادگیری با استفاده از داده‌های اعتبارسنجی و معیار Cross-Entropy تنظیم می‌شود.

الگوریتم‌های ژنتیک استفاده شده در این بخش از تابع‌های برازندگی تعریف شده در روابط ۲.۳ و ۳.۳ برای تولید پاسخ بهینه استفاده می‌کنند. جمعیت اولیه شامل ۲۰۰ کروموزوم در بازه بین -0.25 تا 0.25 به صورت تصادفی تولید می‌شود. پارامترهای مربوط به الگوریتم‌های ژنتیک مطابق جدول ۲.۴ تعریف شده‌اند.

جدول ۲.۴: پارامترهای الگوریتم ژنتیک

پارامتر	مقدار
جمعیت اولیه	۲۰۰
بازه تولید جمعیت اولیه	$[-0.25, 0.25]$
عملگر انتخاب	رقابتی
احتمال ترکیب	0.8
احتمال جهش	0.5
انحراف معیار توزیع گوسی عملگر جهش	0.4
تعداد کروموزوم‌های نخبه	10
حداکثر تعداد نسل	500
حداکثر زمان اجرا	700 ثانیه

۲.۴ مجموعه داده Iris

مجموعه داده Iris یکی از پرکاربردترین مجموعه داده‌هایی است که در بررسی نحوه کارایی سیستم‌های شناسایی گر استفاده می‌شود. این مجموعه شامل نمونه‌های مختلف گل زنبق می‌باشد که در سه دسته متفاوت تقسیم‌بندی می‌شوند. مشخصات این مجموعه در جدول ۳.۴ ارائه شده است.

با توجه به پراکندگی مقادیر ویژگی‌های مختلف، این داده‌ها برای اعمال به شبکه‌عصبی نیاز به آماده‌سازی

جدول ۳.۴: اطلاعات کلی در مورد مجموعه داده Iris Flower

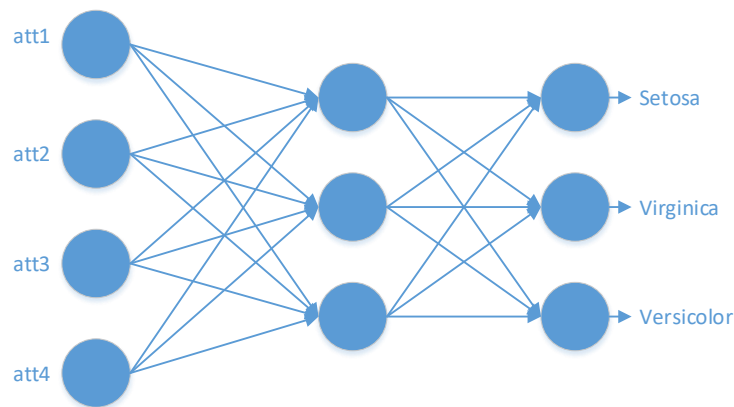
نام مجموعه داده	Flower Iris
تعداد کلاس‌های خروجی	۳
تعداد ویژگی‌های ورودی	۴
تعداد نمونه‌ها	۱۵۰
تعداد نمونه در هر کلاس	Versicolor=۵۰ Virginica=۵۰ Setosa=۵۰

دارند. عملیات نرمال‌سازی با استفاده از رابطه ۱.۳ بیان شده در بخش ۱.۲.۳ برای یکنواخت کردن پراکندگی تمام ویژگی‌ها انجام می‌شود. همچنین خروجی شبکه به صورت ۱ از N گسسته‌سازی می‌شود تا هر نرون خروجی بیانگر یک کلاس باشد.

داده‌ها برای اعمال به شبکه به سه دسته مجموعه آموزش، اعتبار سنجی و ارزیابی تقسیم می‌شوند. برای این مجموعه داده، ۵۰٪ داده‌ها برای آموزش شبکه در مجموعه آموزش، ۲۵٪ داده‌ها برای اعتبارسنجی در مجموعه داده اعتبار سنجی و ۲۵٪ باقی مانده داده‌ها در مجموعه ارزیابی برای بررسی کارایی شبکه برای نمونه‌های دیده نشده قرار می‌گیرند.

با توجه به تعداد ویژگی‌های ورودی و کلاس‌های خروجی که در جدول ۳.۴ بیان شده است، ۴ نرون در لایه ورودی و ۳ نرون در لایه خروجی قرار می‌گیرد. تعداد نرون‌های لایه میانی با توجه به کارایی شبکه تعیین می‌شود، بدین صورت که در ابتدا کارایی شبکه با ۱ نرون در لایه میانی محاسبه می‌شود، در صورتی که کارایی شبکه مطلوب نباشد یک نرون به لایه میانی اضافه می‌شود. این عملیات تا رسیدن به کارایی مطلوب، ادامه پیدا می‌کند. در این مسئله با ۳ نرون لایه میانی به کارایی مطلوبی دست یافتیم. ساختار اولیه شبکه در شکل ۱.۴ نمایش داده شده است. این شبکه قادر است با دقت ۹۶٪ عملیات دسته‌بندی را انجام دهد.

ابتدا آموزش شبکه با استفاده از این ساختار به روش گرادیان نزولی و الگوریتم پس انتشار خطا انجام می‌شود. روند آموزش شبکه و کارایی آن در شکل ۲.۴ نمایش داده شده است. همانطور که مشاهده می‌شود این شبکه قادر است نمونه‌های هر کلاس را به خوبی شناسایی کند.



شکل ۱.۴: ساختار اولیه شبکه برای آموزش برای مجموعه داده Iris

Training Confusion Matrix

Output Class \ Target Class	1	2	3	
1	21 28.4%	2 2.7%	0 0.0%	91.3% 8.7%
2	0 0.0%	23 31.1%	0 0.0%	100% 0.0%
3	0 0.0%	0 0.0%	28 37.8%	100% 0.0%
	100% 0.0%	92.0% 8.0%	100% 0.0%	97.3% 2.7%

Validation Confusion Matrix

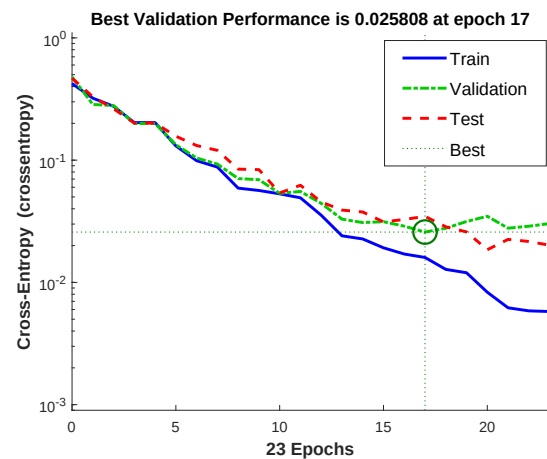
Output Class \ Target Class	1	2	3	
1	14 36.8%	0 0.0%	0 0.0%	100% 0.0%
2	1 2.6%	15 39.5%	0 0.0%	93.8% 6.3%
3	0 0.0%	0 0.0%	8 21.1%	100% 0.0%
	93.3% 6.7%	100% 0.0%	100% 0.0%	97.4% 2.6%

Test Confusion Matrix

Output Class \ Target Class	1	2	3	
1	13 34.2%	2 5.3%	0 0.0%	86.7% 13.3%
2	1 2.6%	8 21.1%	0 0.0%	88.9% 11.1%
3	0 0.0%	0 0.0%	14 36.8%	100% 0.0%
	92.9% 7.1%	80.0% 20.0%	100% 0.0%	92.1% 7.9%

All Confusion Matrix

Output Class \ Target Class	1	2	3	
1	48 32.0%	4 2.7%	0 0.0%	92.3% 7.7%
2	2 1.3%	46 30.7%	0 0.0%	95.8% 4.2%
3	0 0.0%	0 0.0%	50 33.3%	100% 0.0%
	96.0% 4.0%	92.0% 8.0%	100% 0.0%	96.0% 4.0%



(ب) ماتریس درهم‌ریختگی بدست آمده برای مجموعه داده Iris

(آ) روند آموزش شبکه برای مجموعه داده Iris

شکل ۲.۴: روند آموزش و کارایی شبکه برای مجموعه داده Iris

در مرحله بعد پس از آموزش شبکه، با استفاده از الگوریتم‌های ارائه شده در بخش‌های ۱.۵.۳ و ۲.۵.۳

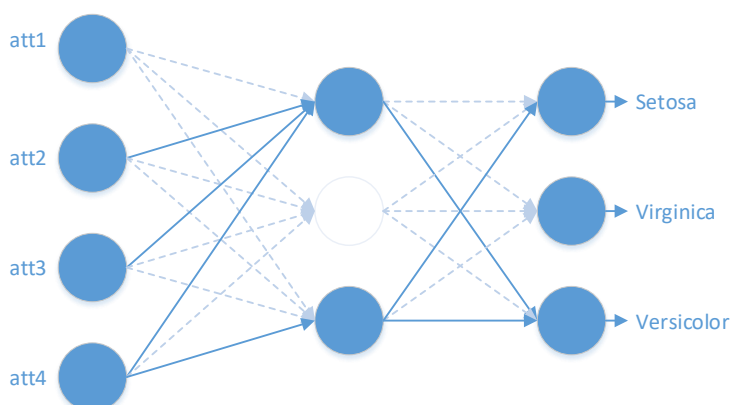
وزن‌های شبکه را برای هرس‌سازی آماده می‌کنیم. پس از اصلاح وزن‌ها شبکه حاصل را با استفاده از الگوریتم

معرفی شده در ۳.۲.۳ هرس‌سازی می‌کنیم تا به یک ساختار بهینه هرس شده برای استخراج قانون دست یابیم.

نتایج هرس‌سازی در شکل ۳.۴ نمایش داده شده است. همانطور که مشاهده می‌شود تعداد کمی از اتصالات در

شبکه هرس شده باقی مانده است اما این شبکه همچنان قادر است تا نمونه‌های ورودی را با دقت ۹۶/۶۶٪

دسته‌بندی نماید.



شکل ۳.۴: ساختار نهایی هرس شده شبکه برای مجموعه داده Iris

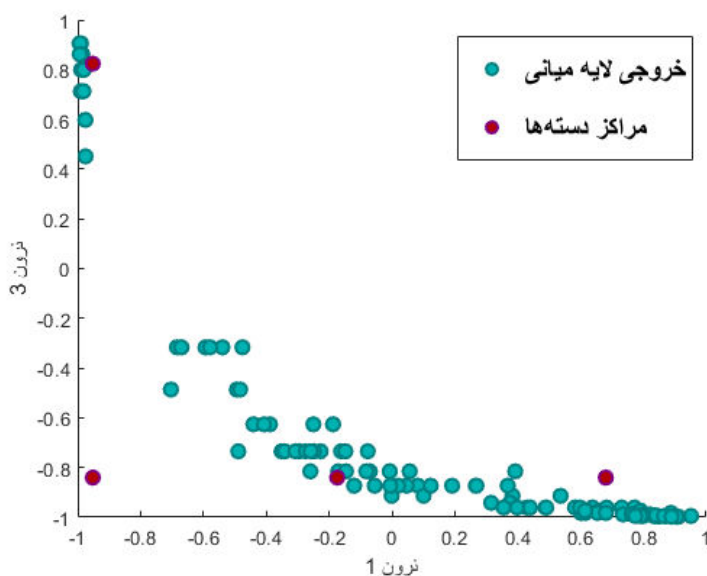
در نهایت استخراج قانون با استفاده از الگوریتم FSRE بیان شده در بخش ۴.۲.۳ انجام خواهد شد.

برای اینکار ابتدا مقادیر خروجی نرون‌های لایه میانی با استفاده از الگوریتم خوشه‌بندی ارائه شده در ۳.۵.۳

گسسته‌سازی می‌شوند. همانطور که مشاهده می‌شود، تنها ۲ نرون در شبکه هرس شده باقی مانده است. نتایج

حاصل از گسسته‌سازی بر روی این مجموعه داده در شکل ۴.۴ نمایش داده شده است. شبکه با استفاده از این

مقادیر گسسته قادر است تا نمونه‌های موجود را با دقت $97/33\%$ دسته‌بندی نماید.



شکل ۴.۴: نتایج حاصل از گسسته‌سازی نرون‌های لایه میانی برای مجموعه داده Iris

با توجه به اتصالات باقی مانده در شبکه هرس شده، ویژگی اول به طور کامل غیر موثر تشخیص داده شده است و از شبکه حذف شده است. با توجه به اینکه اولین کلاس خروجی فقط از سومین نرون لایه میانی تاثیر می پذیرد و این نرون تنها از ویژگی چهارم ورودی می گیرد، در نتیجه قوانین تولید شده برای این کلاس فقط به این ویژگی وابسته خواهد بود. ورودی های اولین کلاس خروجی در ساختار هرس شده به طور کامل حذف شده اند، در نتیجه در این ساختار نمی توان برای این کلاس قانونی استخراج کرد و این کلاس در قانون پیش فرض قرار خواهد گرفت. قوانین مربوط به آخرین کلاس نیز به همین ترتیب استخراج خواهد شد. قوانین استخراج شده برای این مجموعه داده در قالب رابطه ۱.۴ بیان شده اند که قادرند نمونه های این مجموعه داده را با دقت مطلوب $97/33\%$ دسته بندی نمایند. ماتریس درهم ریختگی برای بررسی کارایی قوانین استخراج شده در شکل ۵.۴ ارائه شده است.

```

if ( $att_4 \leq 0/6$ ) then Setosa

else if ( $2 \leq att_2 \leq 3/4$ ) and ( $3 \leq att_3 \leq 5/1$ )
and ( $1 \leq att_4 \leq 1/6$ ) then Versicolour

else Virginica

```

(۱.۴)

۳.۴ مجموعه داده Wisconsin Breast Cancer

مجموعه داده سرطان سینه یا Wisconsin Breast Cancer یکی از مجموعه داده های دسته بندی موجود در زمینه پزشکی و تشخیص بیماری می باشد. در این مجموعه داده دو نوع سرطان سینه خوش خیم^۱ و بدخیم^۲

^۱ Benign

^۲ Malignant

Confusion Matrix				
Output Class	1	2	3	
	48 32.0%	2 1.3%	0 0.0%	96.0% 4.0%
	2 1.3%	48 32.0%	0 0.0%	96.0% 4.0%
	0 0.0%	0 0.0%	50 33.3%	100% 0.0%
	96.0% 4.0%	96.0% 4.0%	100% 0.0%	97.3% 2.7%
	1	2	3	
Target Class				

شکل ۵.۴: ماتریس درهم ریختگی برای قوانین استخراج شده مجموعه داده Iris

وجود دارد که با استفاده از ۹ ویژگی موجود در این مجموعه داده، عملیات دسته‌بندی داده‌ها انجام می‌شود. جدول ۴.۴ اطلاعات کلی را در رابطه با این مجموعه داده نشان می‌دهد. همانطور که مشاهده می‌شود، تعداد نمونه‌های دو دسته موجود اختلاف زیادی دارند، بطوریکه تعداد نمونه‌های خوش‌خیم حدود ۶۵/۵٪ از این مجموعه را تشکیل داده و مابقی نمونه‌ها نیز متعلق به حالت بدخیم هستند.

جدول ۴.۴: اطلاعات کلی در مورد مجموعه داده Wisconsin Breast Cancer

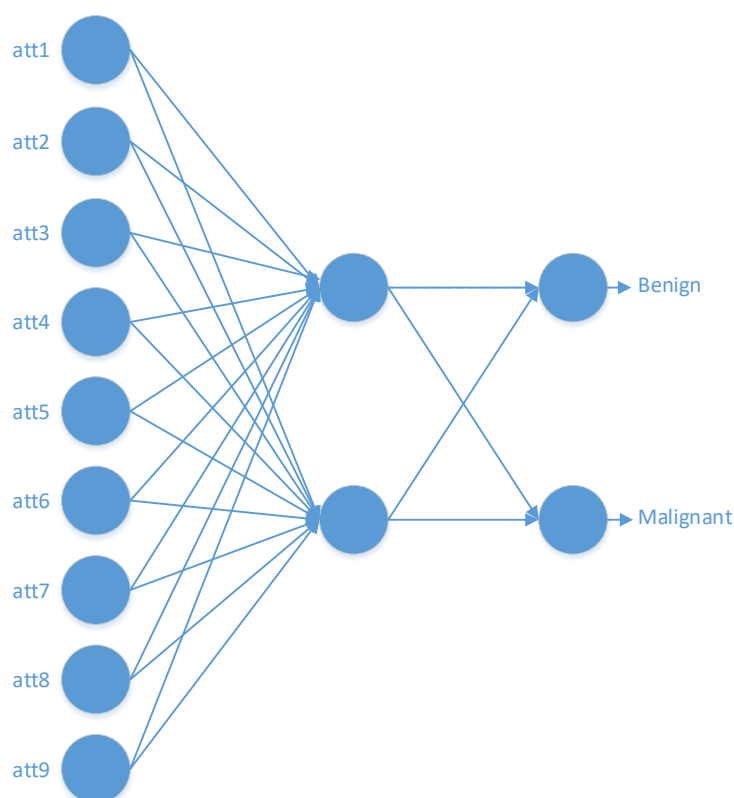
Cancer Breast Wisconsin		نام مجموعه داده
۲		تعداد کلاس‌های خروجی
۹		تعداد ویژگی‌های ورودی
۶۹۹		تعداد نمونه‌ها
Malignant=۲۴۱	Benign=۴۵۸	تعداد نمونه در هر کلاس

این مجموعه داده شامل ویژگی‌ها با مقادیر گسسته بوده که نیاز است تا قبل از اعمال به شبکه عصبی، نرمال‌سازی شوند. برای این کار از رابطه‌ی ۱.۳ استفاده می‌گردد. در نهایت این مجموعه داده به سه مجموعه

آموزش، اعتبارسنجی و ارزیابی با نسبت ۵۰-۲۵-۲۵ تقسیم می‌شود.

ساختار اولیه بدست آمده برای آموزش شبکه با ۹ نرون در لایه ورودی، ۲ نرون در لایه میانی و ۲ نرون در

لایه خروجی است. شکل ۶.۴ ساختار اولیه شبکه را نشان می‌دهد.



شکل ۶.۴: ساختار اولیه شبکه برای آموزش برای مجموعه داده Wisconsin Breast Cancer

ساختار اولیه حاصل را با استفاده از روش گرادیان نزولی و الگوریتم پس انتشار خطا بر روی نمونه‌های

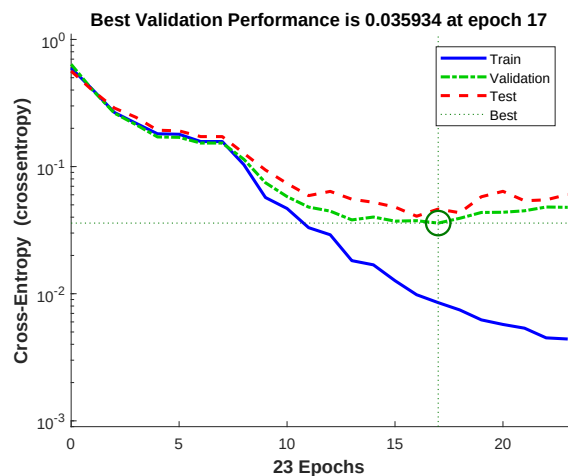
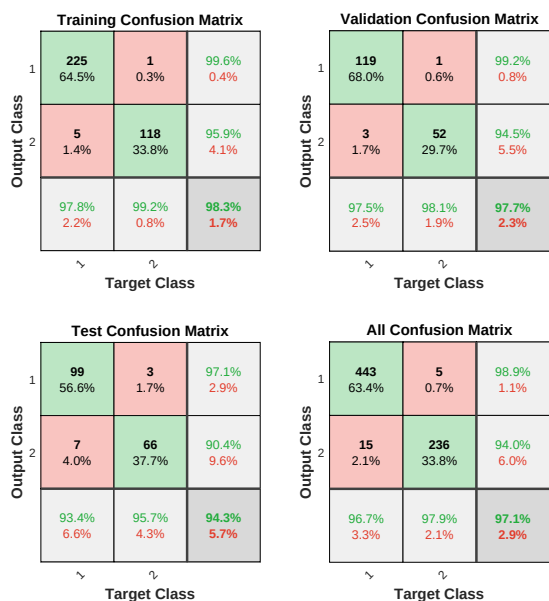
آموزشی، آموزش می‌دهیم. نتایج حاصل از آموزش و روند یادگیری شبکه در شکل ۷.۴ ارائه شده است.

این شبکه قادر است نمونه‌های موجود را با دقت مناسبی دسته‌بندی نماید. برای استخراج قانون از این

شبکه ابتدا به اصلاح وزن‌های شبکه در دو مرحله سراسری و محلی می‌پردازیم. سپس با استفاده از الگوریتم

هرس‌سازی ارائه شده، اتصالات بی‌اهمیت را از شبکه حذف می‌کنیم. ساختار هرس شده شبکه در شکل ۸.۴

نمایش داده شده است.

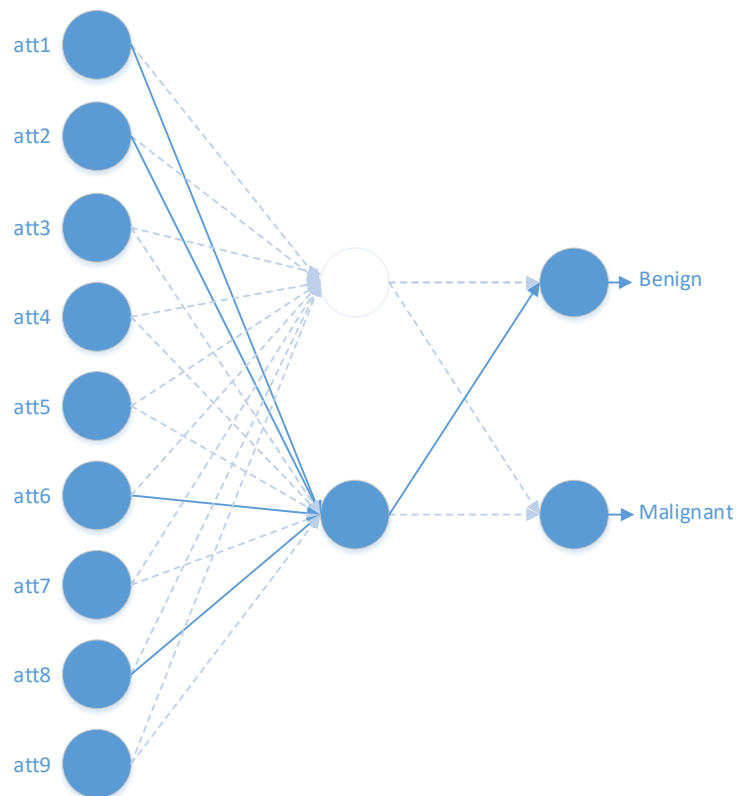


(آ) روند آموزش شبکه برای مجموعه داده Wisconsin Breast Cancer (ب) ماتریس درهم‌ریختگی بدست آمده برای مجموعه داده Wisconsin Breast Cancer

شکل ۷.۴: روند آموزش و کارایی شبکه برای مجموعه داده Wisconsin Breast Cancer

دقت دسته‌بندی با استفاده از شبکه هرس شده برای تمام نمونه‌ها $96/85\%$ است. برای استخراج قانون ابتدا خروجی نرون‌های لایه میانی با استفاده از الگوریتم خوشه‌بندی معرفی شده گسسته‌سازی می‌شوند، سپس قوانین نهایی با استفاده از این مقادیر گسسته شده و با بهره‌گیری از الگوریتم ارائه شده در ۴.۲.۳ استخراج می‌شوند. شبکه با استفاده از مقادیر گسسته در نرون‌های لایه میانی نمونه‌های موجود را با دقت $96/65\%$ دسته‌بندی می‌کند.

همانطور که مشخص است در ساختار هرس شده شبکه فقط کلاس خوش‌خیم دارای ورودی از لایه مخفی می‌باشد و از ۴ ویژگی ورودی تاثیر می‌پذیرد. در نتیجه قوانین مربوط به این کلاس استخراج و کلاس بدخیم به صورت قانون پیش فرض در نظر گرفته می‌شود. قوانین استخراج شده برای این مجموعه داده در رابطه ۲.۴



شکل ۸.۴: ساختار نهایی هرس شده شبکه برای مجموعه داده Wisconsin Breast Cancer

بیان شده‌اند.

if $(att_1 \leq 0.6)$ **and** $(att_2 \leq 0.4)$ **and** $(att_6 \leq 0.6)$
and $(att_8 \leq 0.8)$ **then** *Benign*
else *Malignant*

(۲.۴)

این قوانین قادرند نمونه‌های موجود را با دقت ۹۷/۱۳٪ دسته‌بندی نمایند. نتایج حاصل از دسته‌بندی

این مجموعه داده با استفاده از قوانین استخراج شده در قالب ماتریس درهم‌ریختگی در شکل ۹.۴ قرار گرفته

است.

Confusion Matrix			
Output Class	1	2	
	445 63.7%	7 1.0%	98.5% 1.5%
	13 1.9%	234 33.5%	94.7% 5.3%
Target Class			
1	97.2% 2.8%	97.1% 2.9%	97.1% 2.9%
2			

شکل ۹.۴: ماتریس درهم ریختگی برای قوانین استخراج شده مجموعه داده Wisconsin Breast Cancer

۴.۴ مجموعه داده Glasses

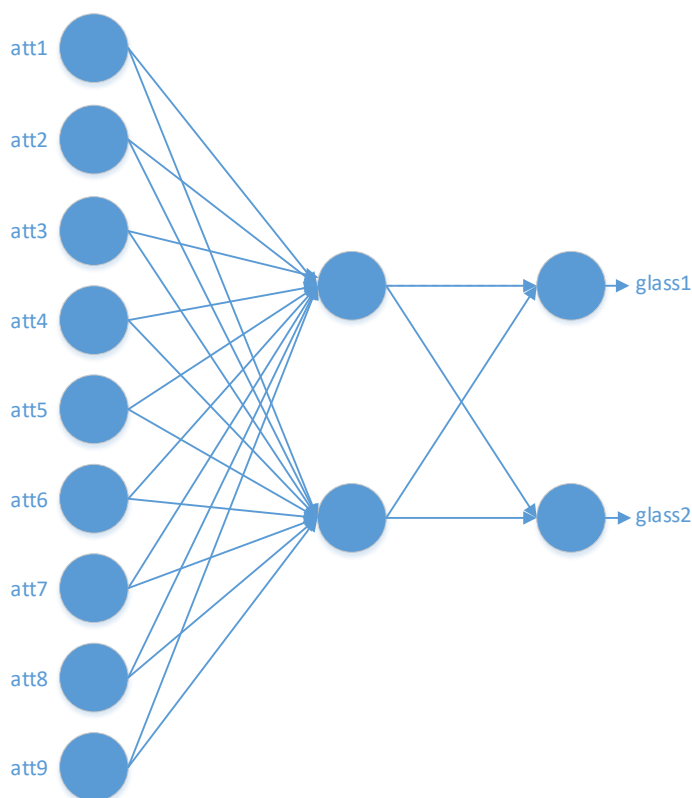
این مجموعه داده یکی دیگر از مجموعه داده‌های دسته‌بندی در UCI بوده که شامل ۲۱۴ نمونه در ۲ کلاس متفاوت است. در این مجموعه داده به منظور دسته‌بندی نمونه‌ها از ۹ ویژگی ورودی استفاده می‌شود. مشخصات کلی این مجموعه داده در جدول ۵.۴ نشان داده شده است. ویژگی‌های موجود در این مجموعه داده دارای مقادیر پیوسته می‌باشد. در فاز پیش پردازش ابتدا این مقادیر نرمال‌سازی شده و سپس نمونه‌های موجود را به سه مجموعه آموزش، اعتبار سنجی و ارزیابی با نسبت ۵۰-۲۵-۲۵ تقسیم شده است. سعی شده است نسبت توزیع نمونه‌ها در این مجموعه‌ها رعایت شود.

ساختار اولیه شبکه عصبی همانند قبل با توجه به تعداد ویژگی‌های ورودی و کلاس‌های خروجی برای

جدول ۵.۴: اطلاعات کلی در مورد مجموعه داده Glasses

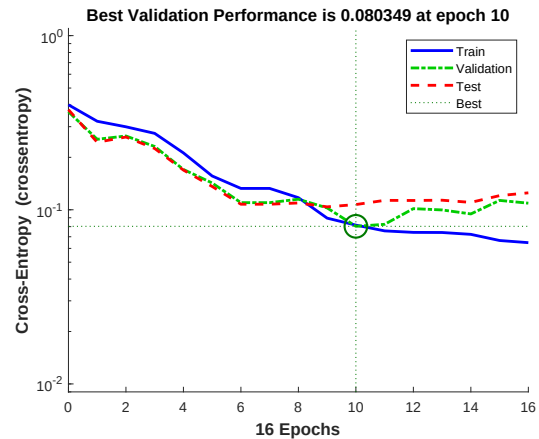
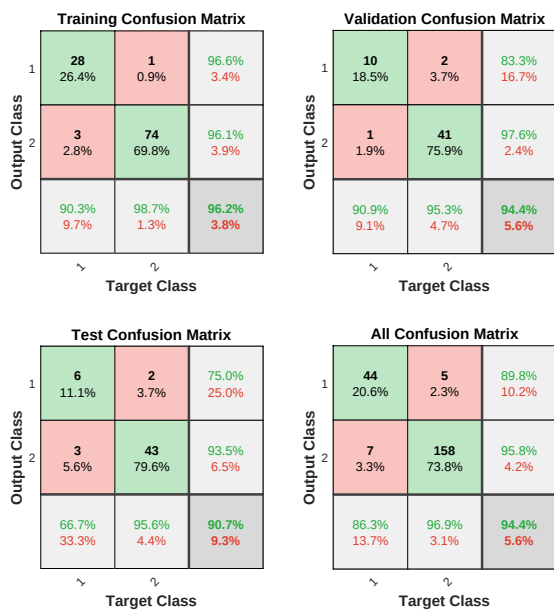
نام مجموعه داده	Glasses
تعداد کلاس‌های خروجی	۲
تعداد ویژگی‌های ورودی	۹
تعداد نمونه‌ها	۲۱۴
تعداد نمونه در هر کلاس	Glass ₁ =۵۱ Glass ₂ =۱۶۳

لایه‌های ورودی و خروجی تعیین می‌شود. همچنین تعداد نرون‌های لایه میانی با توجه به کارایی شبکه با کمترین تعداد نرون در لایه میانی تعیین می‌شود. ساختار اولیه شبکه در شکل ۱۰.۴ نمایش داده شده است. همچنین کارایی شبکه در قالب ماتریس درهم‌ریختگی و روند آموزش شبکه در شکل ۱۱.۴ بیان شده است. همانطور که مشاهده می‌شود شبکه قادر است با دقت مطلوبی داده‌های موجود را دسته‌بندی نماید.



شکل ۱۰.۴: ساختار اولیه شبکه برای آموزش برای مجموعه داده Glasses

در گام بعد وزن‌های آموزش دیده به روش معمول توسط الگوریتم ژنتیک و در دو مرحله اصلاح و برای هرس‌سازی آماده می‌شوند. پس از اعمال الگوریتم هرس‌سازی شبکه حاصل به صورت شکل ۱۲.۴ بدست آمده



(ب) ماتریس درهم‌ریختگی بدست آمده برای مجموعه داده Glasses

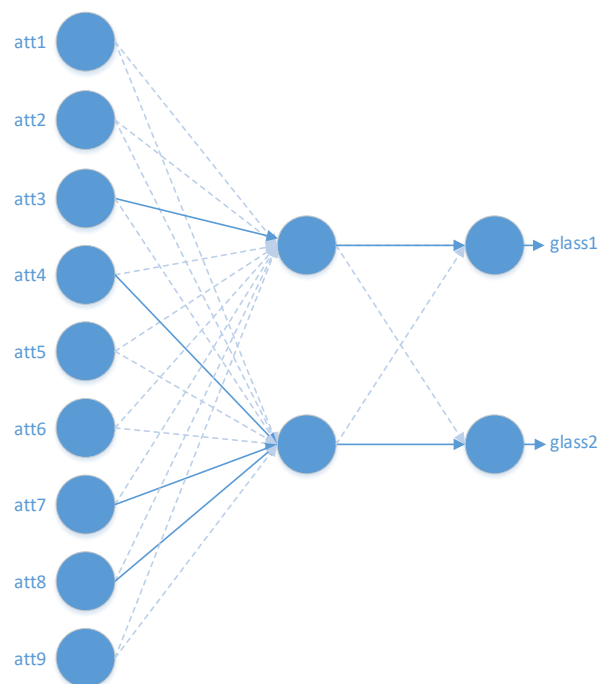
(آ) روند آموزش شبکه برای مجموعه داده Glasses

شکل ۱۱.۴: روند آموزش و کارایی شبکه برای مجموعه داده Glasses

است. دقت شبکه حاصل برابر ۹۴/۳۹٪ برای تمام نمونه‌ها است. همانطور که مشاهده می‌شود ساختار هرس شده همچنان قادر است با دقت بسیار بالایی عملیات دسته‌بندی را انجام دهد. همچنین دقت این شبکه با استفاده از مقادیر گسسته شده ۹۳/۴۵٪ بدست آمده است.

با توجه به ساختار هرس شده بدست آمده تعداد و قوانین به دست آمده برای کلاس اول ساده‌تر است. در نتیجه کلاس دوم را به عنوان قانون پیش فرض در نظر گرفته و قوانین برای کلاس اول استخراج می‌شوند. همانطور که مشاهده می‌شود کلاس اول تنها از ویژگی سوم تاثیر می‌پذیرد. قوانین نهایی استخراج شده در رابطه ۳.۴ آمده است.

$$\begin{aligned} &\text{if } (attr_x \leq 2/7) \text{ then } glass_1 \\ &\text{else } glass_2 \end{aligned} \quad (3.4)$$

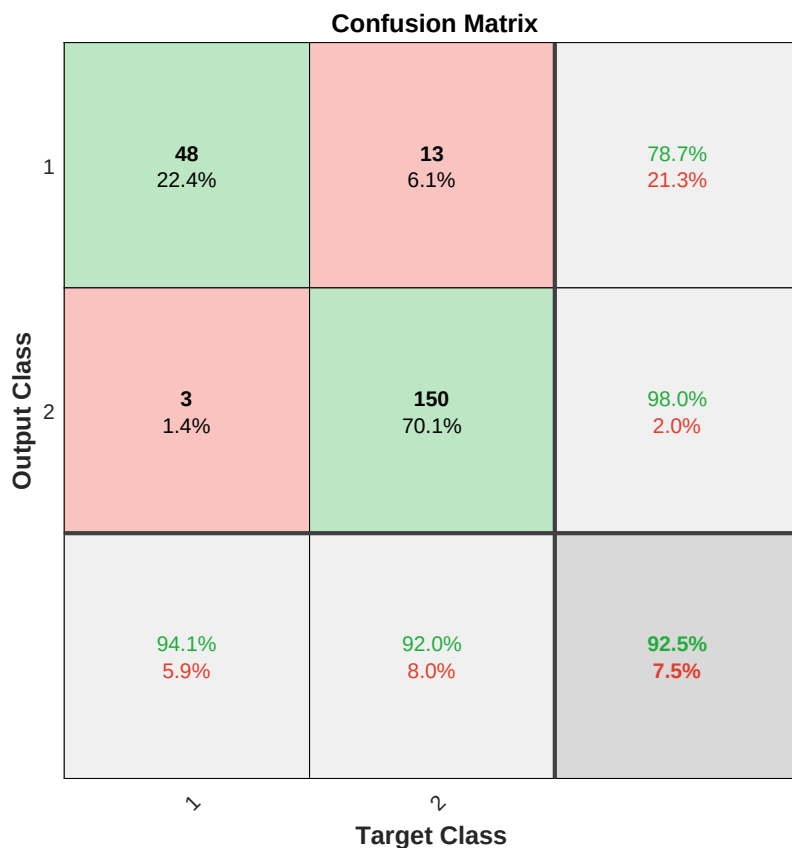


شکل ۱۲.۴: ساختار نهایی هرس شده شبکه برای مجموعه داده Glasses

کارایی این قوانین برای دسته‌بندی نمونه‌ها در قالب ماتریس درهم‌ریختگی در شکل ۱۳.۴ بیان شده است. این قوانین قادرند تا نمونه‌های ورودی را با دقت مطلوب $92/52\%$ دسته‌بندی نمایند.

۵.۴ تحلیل و بررسی نتایج

در آزمایشات انجام شده در این فصل، از مجموعه داده‌های استاندارد جهت بررسی میزان کارایی روش مورد نظر استفاده شده است. این آزمایشات که شامل مجموعه‌های Iris flower، Glasses، Wisconsin Breast و Cancer است که از پایگاه داده UCI جمع‌آوری شده و مورد استفاده قرار گرفته‌اند. همچنین روش‌هایی را که به منظور مقایسه در این بخش مورد بررسی قرار گرفته‌اند شامل روش‌های REANN [۲۵]، RGANN [۲۲]، ESRNN [۲۶]، DIFACONN [۲۹]، Rex [۲۱]، Rex-P [۳۶]، Rex-M [۳۶]، SV-DT [۱۵]، RBFGD [۵۸]، ERBF [۷]، TB-RBF [۲۸]، MMDT [۵۲]، MDTF [۱۸]، GRBF [۷]، FSRE [۵۳] است.



شکل ۱۳.۴: ماتریس درهم ریختگی برای قوانین استخراج شده مجموعه داده Glasses

نتایج آزمایشات برای مجموعه داده Iris به صورت جدول ۶.۴ بدست آمده است.

جدول ۶.۴: نتایج مقایسه روش پیشنهادی نسبت به سایر روش‌ها برای مجموعه داده Iris

نام مجموعه داده	ویژگی‌ها	روش پیشنهادی	FSRE	DIFACONN	ESRNN	SV-DT	Rex-P	Rex-M	MDTF
Iris	تعداد قوانین میانگین تعداد شرایط دقت قوانین	۳ ۱/۳۳ ۹۷/۳۳%	۳ ۱ ۹۶/۶۶%	۳/۲۷ ۹۸/۰۷%	۳ ۱ ۹۸/۶۷%	۷ ۳/۵۷ ۸۰%	۳/۸ ۹۸/۶۷%	۸/۸ ۹۷/۶%	۵ ۹۷/۳۳%

همانطور که مشاهده می‌شود روش پیشنهادی توانسته با تولید تعداد کم قوانین عملکرد مناسبی برای

دسته‌بندی نمونه‌های این مجموعه داده داشته باشد. همچنین این روش عملکرد روش FSRE را بهبود بخشیده است.

نتایج آزمایشات برای مجموعه داده Glasses در جدول ۷.۴ قرار داده شده است. با توجه به نتایج بدست

آمده، روش پیشنهادی با حداقل تعداد قوانین بهترین عملکرد را برای دسته‌بندی نمونه‌های این مجموعه داشته

جدول ۷.۴: نتایج مقایسه روش پیشنهادی نسبت به سایر روش‌ها برای مجموعه داده Glasses

نام مجموعه داده	ویژگی‌ها	روش پیشنهادی	FSRE	DIFACONN	RBF-GD	GRBF	ERBF	TB-RBF	MMDT
Glasses	تعداد قوانین	۲	۲	۱۷/۶۳	۶	۱۶	۱۶	۱۶/۴	۱۳/۲۷
	میانگین تعداد شرایط	۰/۵	۳	—	۳/۳۳	—	—	—	۶۵
	دقت قوانین	۹۲/۵۲%	۸۷/۸۵%	۸۴/۱۳%	۸۶/۲۱%	۶۴/۷۶%	۶۸/۵۷%	۶۲/۸%	۸۶/۵۸%

است و به میزان قابل توجهی روش پایه را بهبود بخشیده است.

نتایج آزمایشات برای مجموعه داده Wisconsin Breast Cancer نیز در جدول ۸.۴ قرار دارد.

جدول ۸.۴: نتایج مقایسه روش پیشنهادی نسبت به سایر روش‌ها برای مجموعه داده Wisconsin Breast Cancer

نام مجموعه داده	ویژگی‌ها	روش پیشنهادی	FSRE	REANN	RGANN	Rex	SV-DT	Rex-p	Rex-M
Cancer Breast Wisconsin	تعداد قوانین	۲	۲	۲	۲	۲	۴	۳	۱۵/۶
	میانگین تعداد شرایط	۲	۳	۳	۳	۳	۲/۲۵	—	—
	دقت قوانین	۹۷/۱۳%	۹۶/۵۶%	۹۶/۲۸%	۹۶/۲۸%	۹۶/۲۸%	۹۲/۹۸%	۹۵/۹۷%	۹۳/۴۳%

همانطور که مشاهده می‌شود روش پیشنهادی به خوبی توانسته در این مسئله نیز عملکرد مناسبی داشته

باشد و با تعداد کمتری از قوانین قادر است نمونه‌های این مجموعه داده را به خوبی دسته‌بندی نماید.

۶.۴ نتیجه‌گیری

در این فصل روش پیشنهادی بر روی برخی از پایگاه‌های داده UCI اعمال شد. نتایج بدست آمده نشان می‌دهد

که این روش قادر است با حداقل تعداد قوانین، نمونه‌ها را با دقت مطلوبی دسته‌بندی نماید. همچنین در

ادامه به مقایسه روش پیشنهادی با سایر روش‌های موجود، به خصوص روش FSRE پرداختیم. نتایج بدست

آمده نشان می‌دهد که روش پیشنهادی به طور میانگین به میزان دو درصد روش پایه را بهبود داده است.

فصل ۵

جمع‌بندی و پیشنهادات

۱.۵ مقدمه

در این پایان نامه برخی از مهم‌ترین روش‌های موجود مبتنی بر شبکه‌های عصبی برای استخراج قانون بررسی شده است. روش پایه مورد بررسی در این پایان نامه یک روش جدید چهار مرحله‌ای تحت عنوان FSRE است. این روش توانایی کشف الگوهای بین داده‌ای برای مسائل پیچیده را دارد. همچنین برخلاف برخی روش‌های موجود، انواع داده‌ای گسسته و پیوسته را پشتیبانی می‌کند. چالش‌های اساسی مورد بررسی در این مقاله، کاهش دقت قوانین استخراج شده از شبکه‌عصبی، به دلیل متکی بودن این روش‌ها به عملیات هرس‌سازی شبکه‌عصبی و اهمیت گسسته‌سازی خروجی تابع فعال‌ساز برای نرون‌های لایه میانی است.

۲.۵ نوآوری تحقیق

در این پژوهش، با استفاده از الگوریتم ژنتیک، یک راهکار نوین برای کاهش خسارات ناشی از هرس سازی بر دقت شبکه عصبی ارائه شده است. الگوریتم ژنتیک با استفاده از یک تابع برازندگی مناسب، وزن‌های غیرمتوازی در دو مرحله سراسری و محلی فراهم می‌کند. در مرحله اول به هنگام آموزش شبکه اعمال می‌شود و سعی دارد تا به تولید وزن‌های شبکه جهت‌دهی کند. در مرحله دوم و بعد از آموزش شبکه، این الگوریتم به صورت محلی، به اصلاح وزن‌های موجود می‌پردازد تا در نهایت ساختار بهینه‌ای در اختیار الگوریتم هرس‌سازی قرار بگیرد. به کمک این وزن‌های غیرمتوازن، شبکه هرس شده بهینه‌تر و با وزن‌های موثرتری تولید خواهد شد که سبب می‌شود، قوانین استخراج شده کارایی بهتری در مقایسه با روش‌های موجود ارائه دهد.

اهمیت گسسته‌سازی خروجی تابع فعال‌ساز برای نرون‌های لایه میانی در روش‌های مبتنی بر دیدگاه تجزیه‌ای بیان و چند روش گسسته‌سازی استفاده شده در این حوزه بیان شد. در این پژوهش با استفاده از یک روش گسسته‌سازی با استفاده از خوشه‌بندی، کارایی روش FSRE را بهبود بخشیدیم. این عملیات به صورت دستی و غیر بهینه انجام می‌شد. استفاده از روش خوشه‌بندی معرفی شده سبب شد تا این عملیات به صورت خودکار و بهینه انجام شود.

۳.۵ پیشنهادات

کارهایی که می‌توان برای بهبود عملکرد این روش در آینده انجام داد عبارت‌اند از:

- اصلاح وزن‌ها در این روش با استفاده از یک تابع برازندگی تجربی انجام می‌شود. در صورتی که بتوان با استفاده از روش‌های ریاضی یک تابع برازندگی مناسب برای این روش تعریف کرد، می‌توان وزن‌ها را به

شکل مناسب‌تری برای مسائل پیچیده اصلاح و نتایج قابل قبول‌تری تولید کرد.

- استخراج قانون در این روش با استفاده از یک شبکه‌عصبی چندلایه پرسپترون با سه لایه استاندارد انجام می‌شود. امروزه استفاده از شبکه‌های عمیق برای حل مسائل پیچیده رونق زیادی پیدا کرده است. با تعمیم این روش برای شبکه‌های عمیق می‌توان از این روش برای کشف دانش در مسائل کاربردی و پیچیده‌تری استفاده کرد.

- الگوریتم هرس‌سازی استفاده شده در این روش ساختار ساده و کارایی دارد. اگر بتوان با اعمال تغییراتی این الگوریتم را هوشمندسازی کرد، ساختار هرس شده بهینه‌تری تولید خواهد شد که منجر به تولید قوانین بهتری می‌شود.

- استخراج قانون با استفاده از یک ساختار هرس شده برای تمام کلاس‌های موجود انجام می‌شود. در صورتی که برای هر کلاس خروجی یک شبکه هرس شده جداگانه تولید کنیم، وزن‌های موثر هر کلاس به صورت جداگانه شناسایی خواهند شد. در این صورت قوانین تولید شده برای هر کلاس با معناتر خواهند بود.

- [1] Aharkava, L. (2010). Artificial neural networks and self-organization for knowledge extraction.
- [2] Ahn, B. S., S. Cho, and C. Kim (2000). The integrated methodology of rough set theory and artificial neural network for business failure prediction. *Expert systems with applications* 18(2), 65–74.
- [3] Amato, F., A. López, E. M. Peña-Méndez, P. Vaňhara, A. Hampf, and J. Havel (2013). Artificial neural networks in medical diagnosis. *Journal of Applied Biomedicine* 11, 47–58.
- [4] Andrews, R., J. Diederich, and A. B. Tickle (1995). Survey and critique of techniques for extracting rules from trained artificial neural networks. *Knowledge-based systems* 8(6), 373–389.
- [5] Azcarraga, A. and R. Setiono (2018). Neural network rule extraction for gaining insight into the characteristics of poverty. *Neural Computing and Applications* 30(9), 2795–2806.
- [6] Baesens, B., R. Setiono, C. Mues, and J. Vanthienen (2003). Using neural network rule extraction and decision tables for credit-risk evaluation. *Management science* 49(3), 312–329.
- [7] Behloul, F., B. P. Lelieveldt, A. Boudraa, and J. H. Reiber (2002). Optimal design of radial basis function neural networks for fuzzy-rule extraction in high dimensional data. *Pattern recognition* 35(3), 659–675.
- [8] Blake, C. L. (1998). Uci repository of machine learning databases, irvine, university of california. <http://www.ics.uci.edu/~mlearn/MLRepository.html>.
- [9] Bondarenko, A., T. Zmanovska, and A. Borisov (2010). Decompositional rules extraction methods from neural networks. *publication. editionName*, 256–262.
- [10] Cano, A., D. T. Nguyen, S. Ventura, and K. J. Cios (2016). ur-caim: improved caim discretization for unbalanced and balanced data. *Soft Computing* 20(1), 173–188.
- [11] Carpenter, G. A. and A.-H. Tan (1995). Rule extraction: From neural architecture to symbolic representation. *Connection Science* 7(1), 3–27.
- [12] de Fortuny, E. J. and D. Martens (2012). Active learning based rule extraction for regression. In *2012 IEEE 12th International Conference on Data Mining Workshops*, pp. 926–933. IEEE.
- [13] Diederich, J. (1992). Explanation and artificial neural networks. *International Journal of Man-Machine Studies* 37(3), 335–355.
- [14] ElAlami, M. (2012). Destructive algorithm for rule extraction based on a trained neural network. *International Journal of Computer Applications* 42(21), 8–14.
- [15] Farquard, M., J. Sultana, G. Nagalaxmi, and G. Savankumar (2014). Knowledge discovery from data: comparative study. *Training* 1(3), 4.

- [16] Gupta, A., S. Park, and S. M. Lam (1999). Generalized analytic rule extraction for feed-forward neural networks. *IEEE transactions on knowledge and data engineering* 11(6), 985–991.
- [17] Hayashi, Y. (2016). Application of a rule extraction algorithm family based on the re-rx algorithm to financial credit risk assessment from a pareto optimal perspective. *Operations Research Perspectives* 3, 32–42.
- [18] Hong, T.-P. and J.-B. Chen (2000). Processing individual fuzzy attributes for fuzzy rule induction. *Fuzzy sets and systems* 112(1), 127–140.
- [19] Hruschka, E. R. and N. F. Ebecken (2006). Extracting rules from multilayer perceptrons in classification problems: A clustering-based approach. *Neurocomputing* 70(1-3), 384–397.
- [20] Jivani, K., J. Ambasana, and S. Kanani (2014). A survey on rule extraction approaches based techniques for data classification using neural network. *International Journal of Futuristic Trends in Engineering and Technology* 1(1), 4–7.
- [21] Kamruzzaman, S. (2010a). Rex: An efficient rule generator. *arXiv preprint arXiv:1009.4988*.
- [22] Kamruzzaman, S. (2010b). Rgann: An efficient algorithm to extract rules from anns. *arXiv preprint arXiv:1009.4962*.
- [23] Kamruzzaman, S. and A. R. Hasan (2010). Rule extraction using artificial neural networks. *arXiv preprint arXiv:1009.4984*.
- [24] Kamruzzaman, S., M. Islam, et al. (2010a). An algorithm to extract rules from artificial neural networks for medical diagnosis problems. *arXiv preprint arXiv:1009.4566*.
- [25] Kamruzzaman, S., M. Islam, et al. (2010b). Extraction of symbolic rules from artificial neural networks. *arXiv preprint arXiv:1009.4570*.
- [26] Kamruzzaman, S. and A. Sarkar (2011). A new data mining scheme using artificial neural networks. *Sensors* 11(5), 4622–4647.
- [27] Kolman, E. and M. Margaliot (2005). Knowledge extraction from neural networks using the all-permutations fuzzy rule base.
- [28] Kubat, M. (1998). Decision trees can initialize radial-basis function networks. *IEEE Transactions on Neural Networks* 9(5), 813–821.
- [29] Kulluk, S., L. Özbakır, and A. Baykasoğlu (2013). Fuzzy difaconn-miner: A novel approach for fuzzy rule extraction from neural networks. *Expert Systems with Applications* 40(3), 938–946.
- [30] Kurgan, L. A. and K. J. Cios (2003). Fast class-attribute interdependence maximization (caim) discretization algorithm. In *ICMLA*, pp. 30–36.
- [31] Kurgan, L. A. and K. J. Cios (2004). Caim discretization algorithm. *IEEE transactions on Knowledge and Data Engineering* 16(2), 145–153.
- [32] Lakshmi, K. and G. S. Kumar (2014). Association rule extraction from medical transcripts of diabetic patients. In *The Fifth International Conference on the Applications of Digital Information and Web Technologies (ICADIWT 2014)*, pp. 201–206. IEEE.
- [33] Liu, H. and R. Setiono (1995). Chi2: Feature selection and discretization of numeric attributes. In *Tools with artificial intelligence, 1995. proceedings., seventh international conference on*, pp. 388–391. IEEE.

- [34] Lu, H., R. Setiono, and H. Liu (1996). Effective data mining using neural networks. *IEEE transactions on knowledge and data engineering* 8(6), 957–961.
- [35] Malone, J., K. McGarry, S. Wermter, and C. Bowerman (2006). Data mining using rule extraction from kohonen self-organising maps. *Neural Computing & Applications* 15(1), 9–17.
- [36] Markowska-Kaczmar, U. and W. Trelak (2005). Fuzzy logic and evolutionary algorithm—two techniques in rule extraction from neural networks. *Neurocomputing* 63, 359–379.
- [37] McGarry, K. and J. Malone (2004). Analysis of rules discovered by the data mining process. In *Applications and Science in Soft Computing*, pp. 219–224. Springer.
- [38] Molina, J. M., P. Isasi, A. Berlanga, and A. Sanchis (2000). Hydroelectric power plant management relying on neural networks and expert system integration. *Engineering Applications of Artificial Intelligence* 13(3), 357–369.
- [39] Nayak, R. (1999). *GYAN: A methodology for rule extraction from artificial neural networks*. Ph.D. disseration , Queensland University of Technology.
- [40] Núñez, H., C. Angulo, and A. Català (2002). Rule extraction from support vector machines. In *Esann*, pp. 107–112.
- [41] Schmitz, G. P., C. Aldrich, and F. S. Gouws (1999). Ann-dt: an algorithm for extraction of decision trees from artificial neural networks. *IEEE Transactions on Neural Networks* 10(6), 1392–1401.
- [42] Sethi, K. K., D. K. Mishra, and B. Mishra (2012). Kdruleex: A novel approach for enhancing user comprehensibility using rule extraction. In *Intelligent Systems, Modelling and Simulation (ISMS), 2012 Third International Conference on*, pp. 55–60. IEEE.
- [43] Setiono, R. (1996). Extracting rules from pruned neural networks for breast cancer diagnosis. *Artificial Intelligence in Medicine* 8(1), 37–52.
- [44] Setiono, R. (1997). A penalty-function approach for pruning feedforward neural networks. *Neural computation* 9(1), 185–204.
- [45] Setiono, R. (2000). Extracting m-of-n rules from trained neural networks. *IEEE Transactions on Neural Networks* 11(2), 512–519.
- [46] Setiono, R. and A. Gaweda (2000). Neural network pruning for function approximation. In *Neural Networks, 2000. IJCNN 2000, Proceedings of the IEEE-INNS-ENNS International Joint Conference on*, Volume 6, pp. 443–448. IEEE.
- [47] Setiono, R. and W. K. Leow (2000). Fernn: An algorithm for fast extraction of rules from neural networks. *Applied Intelligence* 12(1-2), 15–25.
- [48] Setiono, R., W. K. Leow, and J. M. Zurada (2002). Extraction of rules from artificial neural networks for nonlinear regression. *IEEE transactions on neural networks* 13(3), 564–577.
- [49] Setiono, R. and H. Liu (1995). Understanding neural networks via rule extraction. In *IJCAI*, Volume 1, pp. 480–485.
- [50] Setiono, R. and H. Liu (1997). Neurolinear: From neural networks to oblique decision rules. *Neurocomputing* 17(1), 1–24.
- [51] Setiono, R. and M. Tanaka (2010). Neural network rule extraction and the led display recognition problem. In *2010 22nd International Conference on Tools with Artificial Intelligence*, pp. 53–56. IEEE.

- [52] Shaabani, B. and H. Sajedi (2014). Mmdt: Multi-objective memetic rule learning from decision tree. *Journal of Computer & Robotics* 7(2), 37–46.
- [53] Shahabi, M., H. Hassanpour, and H. Mashayekhi (2017). Rule extraction for fatty liver detection using neural networks. *Neural Computing and Applications*, 1–11.
- [54] Storn, R. and K. Price (1997). Differential evolution—a simple and efficient heuristic for global optimization over continuous spaces. *Journal of global optimization* 11(4), 341–359.
- [55] Tsai, C.-J., C.-I. Lee, and W.-P. Yang (2008). A discretization algorithm based on class-attribute contingency coefficient. *Information Sciences* 178(3), 714–731.
- [56] Visser, U., R. Nayak, and M. T. Wongy (1998). Rule extraction from trained neural networks & connectionist knowledge representation for the determination of pesticide mixtures. In *Proceedings of the Ninth Australian conference on neural networks*, pp. 138–142.
- [57] Vora, S. V. and M. S. R. M. Tech (2012). Mcaim : Modified caim discretization algorithm for classification.
- [58] Wang, L. and X. Fu (2005). A simple rule extraction method using a compact rbf neural network. In *International Symposium on Neural Networks*, pp. 682–687. Springer.

واژه‌نامه فارسی به انگلیسی

Hyperspace	ابرفضا
Specificity	اختصاصی بودن
Rule Extraction	استخراج قانون
Knowledge induction	استقراء دانش
Reformulation	اصلاح فرموله‌سازی
Genetic algorithm	الگوریتم ژنتیک
Pedagogical	آموزشی
Selection	انتخاب
Boltzmann Selection	انتخاب بولتزمن
Rank Selection	انتخاب ترتیبی
Roulette wheel Selection	انتخاب چرخ رولت
Steady-State Selection	انتخاب حالت پایدار
Tournament Selection	انتخاب رقابتی
Representation	بازنمایی
Bias	بایاس
Critical	بحرانی
Application	برنامه کاربردی
Boolean	بولی
Information gain	بهره اطلاعاتی
Entropy	بی‌نظمی
Knowledge refinement	پالایش دانش
Backpropagation	پس انتشار
Preprocessing	پیش پردازش

Fitness Function.....	تابع برازندگی.....
Radial Basis Function.....	تابع پایه شعاعی.....
Activation Function	تابع فعال ساز
Decompositional.....	تجزیه‌ای.....
Aggregation.....	تجمع.....
Regression	تخمین توابع
Non-linear approximators	تخمین‌گر غیرخطی
CrossOver	ترکیب
Conjunction	ترکیب عاطفی
Disjunction	ترکیب فصلی
Eclectic.....	ترکیبی.....
White Box.....	جعبه سفید.....
Black Box.....	جعبه سیاه.....
Oblique	چند جمله‌ای
Multi-layer perceptron	چند لایه پرسپترون
Mutation	جهش
Sensitivity	حساسیت
Clustering	خوشه‌بندی
Data mining.....	داده‌کاوی.....
Decision tree.....	درخت تصمیم.....
Corresponding grade	درجه تعلق
Classification.....	دسته‌بندی.....
Accuracy	دقت
Regression	رگرسیون
Overlap.....	روی هم افتادگی.....
Subspace.....	زیرفضا.....
Global	سراسری
Decision system.....	سیستم تصمیم‌گیری.....
Artificial neural network.....	شبکه عصبی مصنوعی.....
Classifier	طبقه‌بند
Classification.....	طبقه‌بندی.....
Unbalance.....	غیرمتوازن.....

Fuzzy	فازی
Performance	کارایی
User friendly	کاربر پسند
Dimensionality reduction	کاهش بعد
Chromosome	کروموزوم
Quality	کیفیت
Gradient descent	گرادیان نزولی
Propositional	گزاره‌ای
Discretization	گسسته‌سازی
Support vector machine	ماشین بردار پشتیبان
State machines	ماشین‌های حالت
Penalty	مجازات
Local	محلی
Connectionist modeling	مدل‌سازی پیوندگرا
Trade off	مصالحه
Linguistic concepts	مفاهیم زبانی
Entity	موجودیت
Elitism	نخبه سالاری
Normalization	نرمال‌سازی
Neuron	نرون
Self-Organizing Map	نگاشت خود سازمان یافته
Visual representation	نمایش بصری
Sampling	نمونه‌برداری
Pruning	هرس‌سازی
Seven segment	هفت بخشی
Artificial intelligence	هوش مصنوعی
Machine learning	یادگیری ماشین
Symbolic machine learning	یادگیری ماشین نمادین

واژه‌نامه انگلیسی به فارسی

Accuracy	دقت
Activation Function	تابع فعال‌ساز
Aggregation	تجميع
Application	برنامه کاربردی
Artificial intelligence	هوش مصنوعی
Artificial neural network	شبکه عصبی مصنوعی
Backpropagation	پس انتشار
Bias	بایاس
Black Box	جعبه سیاه
Boltzmann Selection	انتخاب بولتزمن
Boolean	بولی
Chromosome	کروموزوم
Classification	طبقه‌بندی، دسته‌بندی
Classifier	طبقه‌بند
Clustering	خوشه‌بندی
Conjunction	ترکیب عطفی
Critical	بحرانی
CrossOver	ترکیب
Connectionist modeling	مدل‌سازی پیوندگرا
Corresponding grade	درجه تعلق
Data mining	داده‌کاوی
Decision system	سیستم تصمیم‌گیری
Decision tree	درخت تصمیم

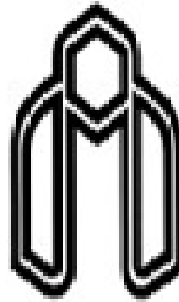
Decompositional	تجزیه‌ای
Dimensionality reduction	کاهش بعد
Discretization	گسسته‌سازی
Disjunction	ترکیب فصلی
Eclectic	ترکیبی
Elitism	نخبه سالاری
Entity	موجودیت
Entropy	بی‌نظمی
Fitness Function	تابع برازندگی
Fuzzy	فازی
Genetic algorithm	الگوریتم ژنتیک
Global	سراسری
Gradient descent	گرادیان نزولی
Hyperspace	ابرفضا
Information gain	بهره اطلاعاتی
Knowledge induction	استقراء دانش
Knowledge refinement	پالایش دانش
Linguistic concepts	مفاهیم زبانی
Local	محلی
Machine learning	یادگیری ماشین
Multi-layer perceptron	چند لایه پرسپترون
Mutation	جهش
Neuron	نرون
Non-linear approximators	تخمین‌گر غیرخطی
Normalization	نرمال‌سازی
Oblique	چند جمله‌ای
Overlap	روی هم افتادگی
Pedagogical	آموزشی
Penalty	مجازات
Performance	کارایی
Preprocessing	پیش‌پردازش
Propositional	گزاره‌ای

Pruning	هرس‌سازی
Radial Basis Function.....	تابع پایه شعاعی
Reformulation	اصلاح فرموله‌سازی
Regression	رگرسیون، تخمین توابع
Rule Extraction.....	استخراج قانون
Rank Selection	انتخاب ترتیبی
Representation	بازنمایی
Roulette wheel Selection.....	انتخاب چرخ رولت
Sampling	نمونه‌برداری
Selection	انتخاب
Self-Organizing Map.....	نگاشت خود سازمان یافته
Sensitivity	حساسیت
Seven segment.....	هفت بخشی
Specificity	اختصاصی بودن
State machines	ماشین‌های حالت
Steady-State Selection	انتخاب حالت پایدار
Subspace.....	زیرفضا
Support vector machine.....	ماشین بردار پشتیبان
Symbolic machine learning.....	سیستم‌های یادگیری ماشین نمادین
Tournament Selection	انتخاب رقابتی
Trade off.....	مصالحه
Unbalance.....	غیرمتوازن
User friendly	کاربر پسند
Visual representation	نمایش بصری
Quality	کیفیت
White Box.....	جعبه سفید
Entity	موجودیت

Aabstract

Nowadays with the growing volume of data, analyzing and discovering the relationships between them, has become very important in data mining science. Rule Extraction is one of the important applications for discovering knowledge and relationships between data. It is done by investigating the internal structure of the network, connections and the output of the hidden layer neurons. Neural networks, especially, have received a lot of attention because of their high accuracy and not requiring prior knowledge about data. Most algorithms in this field of study use decompositional approach and are dependent on pruning phase. However, pruning reduces the efficiency and accuracy of a network which lead to weakening the extracted rules. We proposed a new method to reduce pruning damage using genetic algorithm at two general and local levels. In this method, neural network weights are adjusted in a way, so that the damage caused by the pruning phase will be reduced significantly. The effectiveness of the proposed approach is clearly demonstrated by the experimental results on a set of standard data mining test problems. The accuracy of the rules extracted by our proposed method is improved about 2 percent on average.

Keywords: Rule Extraction, Genetic Algorithm, Artificial Neural Network, Discovering inter-data relationships, Neural Network Pruning, Knowledge discovery



Shahrood University of Technology

Faculty of Computer Engineering

M.Sc. Thesis in Artificial Intelligence Engineering

Improving accuracy of the extracted rules in multilayer neural networks using evolutionary algorithm in training and pruning the network

By: Yasser Irani

Supervisor

Dr. Hamid Hassanpour

February, 2019