

ELearn++¹: یک الگوریتم جدید برای یادگیری افزایشی در شبکه های

چندلایه پرسپترون

علی صادقی نایینی^{*}، حمید بیگی[†]

چکیده

هدف الگوریتم های یادگیری افزایشی در شبکه های چند لایه پرسپترون، حفظ نتایج فازهای آموزشی قبلی و بهبود عملکرد شبکه با آموزش انحصاری آن بر روی نمونه های جدید است. الگوریتمهایی که تاکنون برای یادگیری افزایشی در یک شبکه چند لایه پرسپترون پیشنهاد شده اند قابلیت تعریف کلاسهای جدید را در اختیار نمی گذارند. در این مقاله الگوریتمی برای یادگیری افزایشی در شبکه های چند لایه پرسپترون مورد بررسی قرار می گیرد که این نقطه ضعف را رفع کرده است. این الگوریتم از ترکیب افزایشی تعدادی شبکه یادگیر ضعیف که هریک مربوط به تعدادی از نمونه های آموزشی است که در طول زمان به سیستم ارائه شده اند، یک شبکه یادگیر قوی می سازند و قابلیت پذیرش کلاسهای جدید معرفی شده توسط نمونه های تازه وارد را نیز دارا می باشند. در اینجا تاثیر پارامترهای مختلف بر میزان کارایی الگوریتم مورد بررسی قرار گرفته، و بوسیله نتایج بدست آمده، عملکرد الگوریتم مورد تحلیل قرار خواهد گرفت. به علاوه سه راه برای بهبود عملکرد چنین الگوریتمی پیشنهاد شده است. نتایج حاصل از پیاده سازیهای صورت گرفته، حاکی از موفقیت این روشها در بهبود عملکرد الگوریتم نسبت به نسخه اصلی می باشد تاجایی که استفاده همزمان از این روشها منجر به کاهش پنجاه درصدی خطا نسبت به نسخه اصلی الگوریتم گردیده است.

کلمات کلیدی

شبکه عصبی، چند لایه پرسپترون، یادگیری، افزایشی، همجوشی، رای گیری اکثریت، توزیع احتمال، ترکیب فرضیات.

ELearn++: A novel incremental learning algorithm for MLP neural networks

Ali Sadeghi Naini
Amirkabir University of Technology
a_sadeghi@aut.ac.ir

Hamid Beigy
Sharif University of Technology
beigy@ce.sharif.edu

Abstract

The incremental learning algorithms in multilayer perceptron save the result of previous learning phases and enhance the performance of the network by training it on new samples exclusively. So far, introduced incremental learning algorithms for multilayer perceptron has not been capable of training the network on new samples which reflect a new class. In this article, an incremental learning algorithm is purposed which resolve this problem utilizing incremental combination of a set of weak learners' hypothesizes. First an original algorithm will be described. Next the influence of different parameters on efficiency of the algorithm is studied and exploring the results, the performance of the algorithm is analyzed. In addition, three approaches for the algorithm performance enhancement are proposed. The related results report the accomplished algorithm performance enhancement using these approaches in comparison with the original version of the algorithm. Applying these approaches simultaneously causes total error decreasing up to 50 percent comparing to the original algorithm result.

Keywords

Neural Network, Multilayer Perceptron, Learning, Incremental, Fusion, Weighted Majority, Hypothesis.

^{*} دانشکده مهندسی کامپیوتر، دانشگاه صنعتی امیرکبیر، تهران، ایران. a_sadeghi@aut.ac.ir

[†] دانشکده مهندسی کامپیوتر، دانشگاه صنعتی شریف، تهران، ایران. beigy@ce.sharif.edu

۱- مقدمه

فاز یادگیری در شبکه‌های چند لایه پرسپترونی پر هزینه‌ترین مرحله در بکارگیری این شبکه‌ها به حساب می‌آید. در این فاز با ارائه مجموعه‌ای از داده‌ها آموزشی به شبکه، وزنهای مختلف شبکه تنظیم، و سپس در فاز آزمایش از آنها برای تخمین پاسخ شبکه به داده‌های ناآشنا استفاده خواهد شد. برای این منظور معمولاً از الگوریتم انتشار خطا به عقب استفاده می‌شود که الگوریتمی کاملاً برون خط است، بدین معنا که برای استفاده از این الگوریتم لازم است تمامی داده‌های آموزشی از ابتدا در اختیار الگوریتم قرار گیرد. اگر در کاربردی داده‌های آموزشی به طور قطعی و کامل از ابتدا مشخص نباشند و به مرور زمان بر تعداد آنها افزوده گردد، پس از هر بار به روز رسانی داده‌های آموزشی لازم است فاز یادگیری از ابتدا و با مجموعه جدید آموزشی شامل داده‌های آموزشی قدیمی و جدید پشت سر گذاشته شود. در این حالت نتایج فاز آموزش قبلی کاملاً فراموش شده و از دست می‌رود، در نتیجه ما مجبور به پرداخت مکرر هزینه آموزش برای داده‌های آموزشی قدیمی هستیم. حال آنکه ممکن است در کاربردی داده‌های آموزشی در یک دنباله زمانی و تک تک به سیستم ارائه شوند. درمورد پروسه‌های غیر ایستا^۲ این مشکل بیشتر نمایان میشود. در برخی موارد مشکل حادتر می‌شود و آن هنگامی است که داده‌های آموزشی قدیمی دیگر در دسترس نباشند.

الگوریتم‌های یادگیری افزایشی^۲ برای حل این مشکل معرفی شده‌اند. هدف این الگوریتم‌ها حفظ نتایج فازهای آموزشی قبلی و بهبود عملکرد شبکه با آموزش آن تنها بر روی نمونه‌های جدید است. در واقع این الگوریتمها بدون آنکه دوباره بر روی نمونه‌های قدیمی، که ممکن است دیگر در دسترس نباشند، شبکه را آموزش دهند با رسیدن نمونه‌های جدید آن را با این نمونه‌ها نیز وفق داده و به روز می‌رسانند.

یک سیستم مبتنی بر یادگیری افزایشی سیستمی است که هنگام ورود نمونه یا نمونه‌های آموزشی جدید فرضیات و اطلاعات قبلی خود را بدون استفاده دوباره از نمونه‌های آموزشی قبلی به هنگام در می‌آورد. در واقع چنین سیستمی هنگام ارائه نمونه‌های جدید نتایج حاصل از فازهای قبلی آموزشی خود روی نمونه‌های قدیمی را از دست نمی‌دهد بلکه آنها را با توجه به نمونه‌های جدید بهبود می‌بخشد. به بیان دیگر چنین سیستمی تابع Y را بر اساس مجموعه نمونه آموزشی $X1$ می‌آموزد، سپس تابع بهبود یافته Z را بر اساس تابع Y و مجموعه نمونه آموزشی جدید $X2$ می‌آموزد و همینطور این روند را ادامه می‌دهد [۱].

بر این اساس یک الگوریتم یادگیری افزایشی ایده آل برای شبکه‌های عصبی چند لایه پرسپترونی باید شش ویژگی زیر را دارا باشد [۲، ۵]:

۱. [افشردگی^۴]، از واحد‌ها و اتصالات افزایشی بر اساس یک رابطه خطی یا غیر خطی با تعداد نمونه‌های آموزشی استفاده نکند، در غیر اینصورت شبکه کاری غیر از ذخیره کامل تمام نمونه‌های آموزشی و پیدا کردن تابعی که از تمامی نمونه‌ها عبور می‌کند انجام نمی‌دهد.
 ۲. [رجوع نکردن به داده‌های آموزشی قدیمی]، با ورود نمونه‌های آموزشی جدید، شبکه فاز یادگیری افزایشی را فقط روی این نمونه‌ها، و نه روی مجموعه‌ای شامل این نمونه‌ها و نمونه‌های قدیمی، اجرا نماید.
 ۳. [شکل پذیری^۵]، باید بتوان نمونه‌های آموزشی جدید را در هر زمانی به شبکه ارائه نمود و شبکه باید بتواند نتایج فازهای قبلی آموزش را با آموزش روی این نمونه‌های جدید گسترش دهد. نتایج جدید باید شامل ویژگی‌های نمونه‌های آموزشی قدیمی و جدید باشد، به ویژه اگر نمونه‌ها در یک دنباله زمانی و به مرور به شبکه ارائه شوند، عملکرد شبکه بر روی فضای مساله باید به تدریج بهبود یابد.
 ۴. [پایداری^۶]، اگر الگویی که قبلاً توسط شبکه دیده شده است به آن ارائه شود، پاسخ شبکه با یادگیری افزایشی در شرایط عدم وجود تناقض و اختلال^۷ (وجود تناقض و اختلال) باید (تقریباً) مطابق پاسخ قبلی و صحیح شبکه به آن الگو باشد.
 ۵. [مقاومت در برابر اختلال]، از آنجا که در یادگیری افزایشی تحلیل تمامی نمونه‌های آموزشی از قبل میسر نیست، وجود دو قابلیت ضروری است: اول آنکه شبکه با محلی کردن تاثیر نمونه‌های آموزشی دارای اختلال و تناقض در برابر آنها مقاومت کند، و دوم نتایج اشتباه و غیر قابل قبول ناشی از نمونه‌های متناقض تحت تاثیر آموزش بیشتر بر روی نمونه‌های صحیح قرار گرفته و محو شوند.
 ۶. [قدرت تعمیم^۸]، شبکه آموزش دیده باید قابلیت کافی برای تعمیم داشته باشد تا بتواند به نمونه‌های ناآشنا پاسخ مناسب دهد.
- آنچه مسلم است تامین این شش ویژگی در یک شبکه عصبی به طور همزمان و کامل کار ساده‌ای نیست و در برخی از موارد میسر نمی‌باشد. بنابراین با توجه به هر کاربرد سعی می‌شود در شبکه بین آنها توازن ایجاد شود. بر این اساس الگوریتمهای مختلفی برای یادگیری افزایشی در شبکه‌های چند لایه پرسپترونی پیشنهاد شده است. به عنوان نمونه الگوریتم ارائه شده در [۱] بدون نیاز به مجموعه آموزشی قبلی، ابتدا سعی می‌کند بدون تغییر در معماری شبکه و افزایش نوروها و اتصالات آن، نمونه‌های جدید را به صورت افزایشی به شبکه آموزش دهد، و تنها در صورتی که شبکه بدلیل کم بودن تعداد نوروها و اتصالات قادر به یادگیری افزایشی این نمونه‌ها نبود، نوروها و اتصالات جدید را به شبکه می‌افزاید. به بیان دیگر این

ادامه این مقاله بدین صورت زیر سازمان یافته است: ابتدا در بخش ۲ نسخه اصلی الگوریتم یادگیری افزایشی مورد بحث در این مقاله (Learn+) بررسی می‌شود، سپس در بخش ۳ روشهای پیشنهادی این مقاله برای بهبود کارایی این الگوریتم ارائه و الگوریتم پیشنهادی (ELearn++) بر این اساس معرفی خواهد شد. در بخشهای ۴ و ۵ پیاده سازی صورت گرفته و نتایج حاصل از آن بررسی شده و در مورد این نتایج بحث و نتیجه گیری بعمل خواهد آمد.

۲- یادگیری افزایشی مبتنی ترکیب فرضیات شبکه های یادگیر ضعیف

قبل از پرداختن به الگوریتم یادگیری افزایشی، الگوریتمی را معرفی می‌کنیم که فارغ از یادگیری افزایشی، برای ارتقا عملکرد یک شبکه ضعیف بوسیله تولید چندین فرضیه^۱ مختلف و سپس ترکیب آنها، [۳-۶]، توسعه یافته است. سپس نحوه استفاده از این الگوریتم در یادگیری افزایشی، و اصلاحات مورد نیاز در آن برای این منظور را بررسی خواهیم نمود.

۲-۱- الگوریتمی برای ایجاد یک شبکه قوی مبتنی بر ترکیب فرضیات تعدادی شبکه ضعیف

این الگوریتم برای ارتقا عملکرد یک شبکه یادگیر ضعیف، که عمل طبقه بندی را انجام می‌دهد، با تولید چندین فرضیه طبقه بندی ضعیف و سپس ترکیب آنها بوسیله رای گیری وزن دار روی کلاسهای پیشبینی شده به عنوان خروجی توسط هر فرضیه جداگانه، عمل می‌کند. این فرضیات با آموزش مجدد شبکه روی توزیع های مختلف انتخاب شده از پایگاه داده نمونه های آموزشی بدست می‌آید.

ورودی های این الگوریتم عبارتند از: دنباله ای از نمونه های برچسب خورده (داده های آموزشی، S)، یک الگوریتم یادگیری ضعیف، WeakLearn، و یک عدد صحیح، T، که مشخص کننده تعداد فرضیات (تعداد تکرار) تولید شده بوسیله، WeakLearn، می‌باشد. این الگوریتم، به صورت تکراری و با انتساب یک وزن به هر نمونه، توزیع S را بهنگام می‌کند. انتساب وزن در مورد خروجی یادگیر ضعیف، که روی زیر مجموعه انتخاب شده توسط این توزیع آموزش دیده، نیز صورت می‌گیرد. بهنگام سازی توزیع به گونه ای انجام می‌شود که تمرکز بیشتر روی نمونه های مشکل تر باشد.

در تکرار tام، این الگوریتم یک زیر مجموعه از داده های آموزشی را که مطابق توزیع D_t از پایگاه داده آموزشی اصلی، $S = [(x_1, y_1), \dots, (x_m, y_m)]$ ، انتخاب شده، به WeakLearn ارائه می‌کند. پس از آن، WeakLearn، یک فرضیه (طبقه بندی کننده) را مانند، $h_t: X \rightarrow Y$ ، محاسبه می‌کند که سری از مجموعه آموزشی را با توجه به D_t ، بدرستی طبقه بندی می‌نماید. بنابراین هدف WeakLearn، یافتن فرضیه ای است که خطای آموزشی را کمینه کند:

الگوریتم برای یادگیری افزایشی نمونه های تازه وارد دو نوع تغییر را در شبکه اعمال میکند: ۱- تغییر وزنها که اگر موثر تشخیص داده شود اعمال می‌شود؛ ۲- افزایش نوروں ها و اتصالات به شبکه که در صورت تاثیر نداشتن تغییر وزنها اعمال می‌شود. برای هر دوی این تغییرات باید محدودیت هایی در نظر گرفت تا شبکه و نتایج یادگیری افزایشی آن قابل اعتماد و پایدار باشند. در [۲] الگوریتمی برای یادگیری افزایشی یک شبکه چند لایه پرسپترون ارائه شده است که از لحاظ نحوه عملکرد شبیه الگوریتم پیشنهادی در [۱] می‌باشد، یعنی در مرحله اول سعی می‌کند با تغییر وزنها و بدون بزرگ کردن شبکه نمونه های جدید را به صورت افزایشی به شبکه بیاموزد و تنها در صورتی نوروں ها و اتصالات جدید را به شبکه می‌افزاید که مرحله اول آموزش ناموفق باشد. تفاوت این دسته از الگوریتم ها با الگوریتم قبلی در این است که این الگوریتم دارای روابطی برای یادگیری افزایشی است که به نمونه های قبلی آموزشی نیاز دارد، راه حلی که این الگوریتم به کار می‌برد استفاده از شبه نمونه هایی است که از پاسخ شبکه نیمه آموزش دیده در هر زمان، بدست می‌آید.

با آنکه این الگوریتمها و الگوریتمهای مشابهی که برای یادگیری افزایشی در یک شبکه چند لایه پرسپترون پیشنهاد شده اند هر یک در حوزه کاربرد خود عملکرد نسبتا مناسبی داشته اند، اما همگی از یک نقطه ضعف اساسی رنج می‌برند: هیچ یک از این الگوریتمها قابلیت تعریف کلاس جدید را در اختیار نمی‌گذارند. به بیان دیگر با آنکه این الگوریتمها قابلیت ارائه تدریجی نمونه های آموزشی و در نتیجه یادگیری افزایشی را دارا می‌باشند، فرض می‌کنند که تمامی کلاسهای مساله طبقه بندی از ابتدا معلوم بوده و در اختیار آنها قرار می‌گیرد. در نتیجه با این الگوریتمها تنها می‌توان دانش شبکه را درمورد تعداد محدود و از ابتدا معینی از کلاسهها به تدریج افزایش داد و اگر در کاربردی نیاز باشد تا شبکه قابلیت افزایش دانش خود را در مورد تعداد و توزیع کلاسهها دارا باشد، این الگوریتمها از عهده ارائه راه حل ناتوان می‌مانند. در این مقاله الگوریتمی برای یادگیری افزایشی در شبکه های چند لایه پرسپترون مورد بررسی قرار می‌گیرد که این نقطه ضعف را رفع کرده است. این الگوریتم که در [۳-۶]، Learn++ نامیده شده است، از ترکیب افزایشی تعدادی شبکه یادگیر ضعیف که هریک مربوط به تعدادی از نمونه های آموزشی است که در طول زمان به سیستم ارائه شده اند، یک شبکه یادگیر قوی می‌سازند و قابلیت پذیرش کلاسههای جدید معرفی شده توسط نمونه های تازه وارد را نیز دارا می‌باشند. در اینجا تاثیر پارامترهای مختلف بر میزان کارایی الگوریتم مورد بررسی قرار گرفته، و بوسیله نتایج بدست آمده، عملکرد الگوریتم مورد تحلیل قرار خواهد گرفت. به علاوه سه راه برای بهبود عملکرد این الگوریتم پیشنهاد شده است. نتایج حاصل از پیاده سازیهای صورت گرفته، حاکی از موفقیت این روشها در بهبود عملکرد الگوریتم نسبت به نسخه اصلی می‌باشد.

ای از داده‌های آموزشی، S ، که از توزیع $D_1 \subseteq D$ آمده‌اند به یک شبکه یادگیر، مجموعه‌ای از فرضیات تولید شده است. حال فرض کنید مجموعه جدیدی از داده‌ها که از توزیع $D_2 \subseteq D$ آمده‌اند و برای شبکه ناآشنا هستند، به شبکه داده شود، اگر مجموعه داده قبلی که از توزیع D_1 آمده‌اند، نمایش مناسبی از کل فضای نمونه‌ای D باشند، شبکه روی نمونه‌های جدید نیز به خوبی کار می‌کند، در غیر اینصورت، می‌توان نتیجه گرفت مجموعه داده قبلی نمایش مناسبی از کل فضای نمونه‌ای D نیست و $D_1 \not\subseteq D_2$. به بیان دیگر، شبکه روی نمونه‌های آن قسمت از فضای نمونه‌ای که نمونه‌های جدید از آن آمده‌اند، آموزش ندیده است. بنابراین می‌توان نمونه‌های جدید را به عنوان نمونه‌های سخت D در نظر گرفت و شبکه را مجبور کرد نمونه‌های تازه وارد از D_2 را با تولید فرضیات جدید برای این مجموعه جدید داده‌ها و سپس ترکیب تمام فرضیات تولید شده (برای D_1 و D_2) یاد بگیرد و به این ترتیب نمونه‌های ناآشنا را نیز بدرستی طبقه بندی کند. فرض اساسی آن است که توزیع D ، توزیع D_2 را نیز به خوبی D_1 شامل می‌شود.

از آنجا که هدف الگوریتم ارائه شده در بخش ۲-۱ افزایش دقت طبقه بندی بود، آن الگوریتم روی توزیع‌های مختلف یک مجموعه آموزشی واحد کار می‌کرد و به این دلیل خطا بر روی نمونه‌های همان مجموعه آموزشی که بدرستی طبقه بندی نشده بودند، محاسبه می‌شد. به علاوه بهنگام سازی توزیع بر اساس فرضیه منفرد h_t و خطای نسبی آن، β_t صورت می‌گرفت. اما Learn++ نسخه متفاوتی از این الگوریتم را برای هر مجموعه داده جدید به کار می‌برد. در هر تکرار، t ، روی مجموعه داده D_k زیر مجموعه‌های آموزشی (TRt) و آزمایشی (TEt) به صورت تصادفی و با توجه به D_t از مجموعه داده انتخاب می‌شوند. زیر مجموعه آموزشی در اختیار شبکه چند لایه پرسپترون نسبتاً کوچک (یادگیر ضعیف) قرار می‌گیرد و فرضیه h_t توسط آن تولید می‌شود. خطا روی اجتماع زیر مجموعه‌های آموزشی و آزمایشی محاسبه می‌شود. اگر خطا بزرگتر از ۰.۵ بود، h_t دور ریخته شده و زیر مجموعه‌های جدید آموزشی و آزمایشی تولید می‌شوند. سپس اکثریت وزن دار همه فرضیات تولید شده، H_t محاسبه می‌شود. این اکثریت وزن دار و خطای نسبی مربوط به آن برای بهنگام سازی توزیع، D_t ، به کار برده می‌شوند. برای بدست آوردن خروجی شبکه به یک نمونه، از اکثریت وزن دار روی اکثریت وزن دار مربوط به هر مجموعه داده D_k استفاده می‌شود.

۲-۳- نقاط ضعف الگوریتم Learn++ و راههای بهبود آن

در نسخه اصلی الگوریتم Learn++ از رابطه $\log(\frac{1}{\beta_t})$ برای محاسبه امتیاز هر کلاس استفاده شده است. بدین ترتیب وزن بیشتری به فرضیه با خطای کمتر داده خواهد شد، اما اگر خطای یک فرضیه صفر باشد، که بسیار مطلوب است، این رابطه خطای تقسیم بر صفر را در

$$\mathcal{E}_t = \sum_{i: h_t(x_i) \neq y_i} D_t(i) \quad (1)$$

این الگوریتم برای عملکرد مناسب نیاز دارد که کارایی هر فرضیه (طبقه بند ضعیف) حداقل ۵۰٪ باشد، یعنی $\mathcal{E}_t < \frac{1}{2}$. توزیع اولیه، D_1 ، به صورت یکنواخت بر روی S انتخاب می‌شود، یعنی برای تمامی i ها $D_1(i) = \frac{1}{m}$. این مقدار دهی اولیه، منجر به شانس مساوی قرار گرفتن در زیر مجموعه آموزشی انتخاب شده، برای تمامی نمونه‌های S می‌گردد. رابطه بهنگام سازی این توزیع به شکل زیر است:

$$D_{t+1}(i) = \frac{D_t(i)}{Z_{t+1}} \times \begin{cases} \beta_t & \text{if } h_t(x_i) = y_i \\ 1 & \text{Otherwise} \end{cases} \quad (2)$$

که $Z_t = \sum_i D_t(i)$ یک ثابت نرمال سازی است تا مطمئن

شویم که D_{t+1} یک توزیع معتبر خواهد بود و $\beta_t = \frac{\mathcal{E}_t}{(1 - \mathcal{E}_t)}$ بدین

ترتیب، نمونه‌های ساده‌تر که توسط h_t بدرستی طبقه بندی شده‌اند، از احتمال کمتر و نمونه‌های مشکل‌تر که نادرست طبقه بندی شده‌اند از احتمال بیشتری برای انتخاب در زیر مجموعه داده‌های آموزشی بعدی برخوردار خواهند شد.

پس از خاتمه T-امین تکرار، این الگوریتم، فرضیات ضعیف، $h_1, \dots, h_t$ ، را با محاسبه اکثریت وزن دار این فرضیات ضعیف ترکیب کرده، و یک فرضیه نهایی h_{final} را بصورت زیر بدست می‌آورد:

$$h_{final} = \operatorname{argmax}_{y \in Y} \sum_{t: h_t(x) = y} \log\left(\frac{1}{\beta_t}\right) \quad (3)$$

برای یک نمونه مانند x ، h_{final} برچسب y را که، حاصل جمع وزنهای فرضیات ضعیفی را که این برچسب را پیشبینی کرده‌اند بیشینه کند، به عنوان خروجی تحویل می‌دهد. وزن فرضیه h_t ، به صورت $\log(\frac{1}{\beta_t})$ ، تعریف شده تا وزن بیشتری به فرضیه با خطای کمتر داده شود.

نکته‌ای که باید به آن دقت شود این است که حداقل دقت هر شبکه یادگیر ضعیف باید ۵۰٪ باشد، اما استفاده از شبکه‌های یادگیر خیلی قوی به جای شبکه‌های یادگیر ضعیف در این الگوریتم توصیه نمی‌شود، زیرا این شبکه‌ها منجر به برازش بیش از حد روی داده‌ها می‌شوند.

۲-۲- Leran++ : یک الگوریتم یادگیری افزایشی

در الگوریتم بخش ۲-۱ نیازی به دانستن توزیع اصلی، D ، که داده‌های آموزشی، S ، از آن برداشته شده‌اند نیست. فرض کنید با ارائه مجموعه

۳. فرضیه خروجی شبکه، $h_t: X \rightarrow Y$ را دریافت کرده و خطا،

$$\varepsilon_t = \sum_{i: h_t(x_i) \neq y_i} D_t(i)$$

را برقرار نبود، T را برابر $t-1$ قرار بده، h_t را دور بریز و به مرحله

۱ برگرد. در غیر این صورت خطای نسبی را محاسبه کن: $\beta_t = \frac{\varepsilon_t}{(1 - \varepsilon_t)}$

۴. اکثریت وزن دار را و خطای عمومی را با توجه به همه فرضیات بدست آمده محاسبه کن:

$$H_t = \arg \max_{y \in Y} \sum_{i: h_t(x_i) = y} (1 - \beta_t), \quad E_t = \sum_{i: h_t(x_i) \neq y_i} D_t(i)$$

$$B_t = \frac{E_t}{(1 - E_t)}, \quad D_t = \frac{E_t}{(1 - E_t)}$$

۵. اکثریت وزن دار را برای محاسبه کن و توزیع D_t را بهنگام کن:

$$D_{t+1}(i) = \frac{D_t(i)}{Z_{t+1}} \times \begin{cases} B_t & \text{if } H_t(x_i) = y_i \\ 1 & \text{Otherwise} \end{cases}, \quad Z_t = \sum_i D_t(i)$$

خروجی: اکثریت وزن دار را برای محاسبه خروجی نهایی (فرضیه نهایی) محاسبه کن:

$$h_{final}(x) = \arg \max_{y \in Y} \sum_k \left(\frac{C}{C_Y} * \sum_{i: h_t(x_i) = y} (1 - \beta_t) \right)$$

که C تعداد کل مراحل، و C_Y تعداد مراحل است که کلاس Y در آن وجود داشته است. $\frac{C}{C_Y}$ همان ضریب تناسب می باشد

الگوریتم ۱- الگوریتم ELearn++

در هر تکرار، t ، روی مجموعه داده D_k زیر مجموعه های آموزشی (TRt) و آزمایشی (TEt) به صورت تصادفی و با توجه به D_t از مجموعه داده انتخاب می شوند. زیر مجموعه آموزشی در اختیار مجموعه ای از شبکه های چند لایه پرسپترون نسبتا کوچک و با یک خروجی (یادگیر ضعیف) قرار می گیرد و فرضیه h_t توسط آن تولید می شود. خطا روی اجتماع زیر مجموعه های آموزشی و آزمایشی محاسبه می شود. اگر خطا بزرگتر از ۰.۵ بود، h_t دور ریخته شده و زیر مجموعه های جدید آموزشی و آزمایشی تولید می شوند. سپس اکثریت وزن دار همه فرضیات تولید شده، H_t محاسبه می شود. این اکثریت وزن دار و خطای نسبی مربوط به آن برای بهنگام سازی توزیع، D_t ، به کار برده می شوند. برای بدست آوردن پاسخ شبکه به یک نمونه جدید، از اکثریت وزن دار روی اکثریت وزن دار هنجار سازی شده مربوط به هر مجموعه داده D_k استفاده می شود.

۴- پیاده سازی و نتایج

مجموعه داده های بکار رفته جهت آموزش و آزمایش سیستم یادگیر در این پیاده سازی، همان مجموعه داده های بکار رفته در الگوریتم اصلی است تا امکان مقایسه نتایج فراهم آید. این مجموعه داده از پنج کلاس با ناحیه دایره ای و در یک فضای دو بعدی تشکیل شده است (شکل ۱). داخلی ترین دایره ناحیه کلاس ۱ و خارجی ترین دایره ناحیه کلاس ۵ است. ۶ مجموعه داده $S1 \sim S6$ از این مجموعه داده انتخاب شده که $S1, S2$ شامل نمونه هایی از کلاس ۱، ۳ و ۵ هستند. نمونه های $S3, S4$ علاوه بر کلاسهای قبلی کلاس ۴ و نمونه های

سیستم به دنبال خواهد داشت. در این حالت فرضیه مطلوب از رای گیری حذف خواهد شد. پیشنهاد ما برای رفع این مشکل استفاده از رابطه ساده $(1 - \beta)$ برای محاسبه امتیاز هر کلاس است که منجر به چنین خطایی نخواهد شد.

پیشنهاد دیگری که به نظر می رسد موجب عملکرد سیستم گردد استفاده از WeakLearner هایی با یک خروجی می باشد. در این حالت برای هر کلاس مساله یک شبکه وجود خواهد داشت که تعیین می کند یک داده به این کلاس تعلق دراد یا خیر. در نتیجه فارغ از تعداد کلاسهای مساله که ممکن است به تدریج بر تعداد آنها افزوده گردد، هر شبکه مامور شناسایی اعضای تنها یک کلاس خواهد بود.

روش دیگری که در اینجا برای بهبود عملکرد الگوریتم به کار گرفته شده است، هنجار سازی امتیازات کلاسها در رای گیری اکثریت وزندار می باشد. در الگوریتم Learn++ هنگامیکه کلاسهای جدیدی به سیستم معرفی می شوند، از آن مرحله به بعد WeakLearner ها دارای یک خروجی برای هر کلاس جدید خواهند بود. حال آنکه کلاسهای قدیمی در WeakLearner های مراحل قبلی نیز دارای خروجی بوده اند و لذا از شانس بیشتری برای پیروزی در رای گیری اکثریت برخوردار خواهند بود. در مرحله هنجار سازی امتیازات کلاسها، امتیاز کلاسهایی که جدیداً تعریف شده اند در ضریب تناسبی ضرب می شود تا با کلاسهای قدیمی که دارای امتیاز تجمعی بیشتری هستند در یک سطح قرار گیرند. انتظار می رود این هنجار سازی موجب کاهش خطای طبقه بندی گردد.

۳- ELearn++ نسخه بهبود یافته الگوریتم Learn++

با توجه به مطالب بیان شده در قسمت قبل، نسخه بهبود یافته الگوریتم Learn++ که در این مقاله ELearn++ نامیده شده است به صورت زیر است:

ورودی: برای هر مجموعه داده جدید گرفته شده از توزیع D_k ، $k=1,2,\dots,K$

- دنباله ای از m نمونه: $S = [(x_1, y_1), \dots, (x_m, y_m)]$

- به ازاء هر کلاس موجود در مجموعه داده جدید یک شبکه چند لایه پرسپترون نسبتا ضعیف (کوچک) با یک خروجی (WeakLearn)، در هر مرحله هر داده ورودی به کلاسی نسبت داده می شود که شبکه مربوط به آن کلاس در ازای ارائه آن داده ورودی بزرگترین خروجی را تولید نماید.

- عدد صحیح T_k که نمایانگر تعداد تکرار است.

برای هر $k=1,2,\dots,K$ مراحل زیر را انجام بده:

$$D_1(i) = \frac{1}{m}, \forall i$$

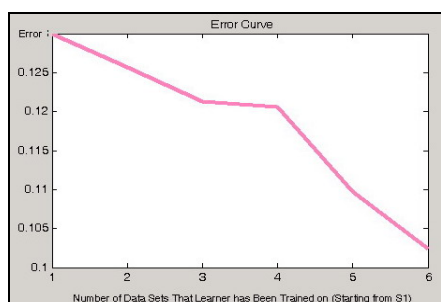
برای هر $t=1,2,\dots,T_k$ مراحل زیر را انجام بده:

۱. به طور تصادفی و با توجه به D_t ، زیر مجموعه های TRt و TEt (آموزش و آزمایش) را از S انتخاب کن.

۲. TRt را به WeakLearn ارائه کن.

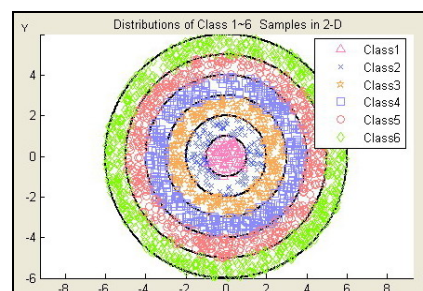
در اینجا نیز استفاده از تعداد WeakLearnerهای بسیار زیاد و مجموعه داده‌های آموزشی و آزمایشی بسیار بزرگ برای آموزش هریک، با آنکه موجب ارتقا چشمگیری در عملکرد سیستم یادگیری نمی‌شوند، فاز یادگیری را بسیار طولانی می‌کنند. مقادیر مناسب برای این پارامترها، مقادیری هستند که تعداد داده‌های یک مجموعه داده جدید را بپوشانند و فرصت تمرکز بیشتر بر روی نمونه‌های مشکل‌تر را نیز داشته باشند. انتخاب مقدار ۱۰ برای تعداد WeakLearnerها مناسب به نظر می‌رسد. اندازه مجموعه آموزشی مناسب نیز مقداری است که تعداد اعضای آن ضرب در تعداد WeakLearnerها حدود ۲ تا ۳ برابر تعداد اعضای مجموعه داده جدید باشد. اندازه مناسب برای مجموعه آزمایشی نیز حدود ۰/۳ تا نصف اندازه مجموعه داده آموزشی می‌باشد. با توجه به این موارد، می‌توان گفت قدرت یادگیری (اندازه) مناسب برای یک WeakLaner، باید با توجه به اندازه مجموعه آموزشی تعیین گردد. به گونه‌ای که هر WeakLearner نه آنقدر ضعیف باشد که به ندرت به آستانه خطای پذیرش دست یابد، و نه آنقدر قوی باشد که بهبود فاز یادگیری را طولانی نماید.

بکارگیری سطوح متفاوتی از خطا برای پذیرش یک WeakLearner نشان داد استفاده از سطح آستانه خطا برای پذیرش یک WeakLearner با مقادیر کمتر از ۰/۵ بدلیل آنکه موجب رد کردن بیشتر WeakLearnerها خواهد شد، فاز یادگیری را طولانی مینماید، هرچند موجب بهبود چشمگیری در عملکرد نهایی سیستم نخواهد شد. در آزمایشی دیگر یکبار از رابطه $\log(\frac{1}{\beta_t})$ و بار دیگر از رابطه $(1 - \beta)$ برای محاسبه امتیاز هر کلاس استفاده شد. نتایج این آزمایش حاکی از برتری عملکرد الگوریتمی داشت که از رابطه دوم استفاده می‌کرد. بهترین نتایجی بدست آمده از نسخه اصلی الگوریتم مربوط به پارامترهایی با مقادیر زیر بودند: تعداد نوروهای لایه مخفی به ازاء هر نورو خروجی = ۵، تعداد Epochs = ۵۰، تعداد WeakLearner برای هر مجموعه داده جدید = ۱۰، تعداد داده آموزشی = ۱۰۰، تعداد داده آزمایشی = ۵۰، سطح پذیرش: ۰/۵، استفاده از رابطه $(1 - \beta)$ برای محاسبه امتیاز هر کلاس. شکل ۲ نمودار خطای نسخه اصلی الگوریتم را با مقادیر پارامترهایی که منجر به بهترین کارایی شدند نشان می‌دهد.



شکل ۲. نمودار خطای الگوریتم Learn++ بر روی مجموعه داده آزمایشی مستقل

S5, S6 کلاس ۲ را نیز شامل می‌شوند. یک مجموعه داده به نام Test نیز از تمامی کلاسها تولید شده است. این مجموعه‌ها به ترتیب از S1 به الگوریتم داده شده‌اند و عملکرد الگوریتم پس از ارائه هر مجموعه داده جدید روی مجموعه داده Test سنجیده شده است. در نمودارهای خطا که در زیر مشاهده می‌کنید، محور عمودی میزان خطا و محور افقی تعداد مجموعه داده‌هایی را نشان می‌دهد که سیستم تاکنون بر روی آنها آموزش دیده است. وقتی کار الگوریتم روی هر یک از مجموعه‌های داده به پایان رسید، پس از آن و در مراحل بعد، این مجموعه داده مجدداً به الگوریتم ارائه نشده است تا از اجرای یک یادگیری افزایشی مطمئن باشیم.



شکل ۳- توزیع داده‌های بکار رفته جهت بررسی عملکرد نسخه پیاده‌سازی شده الگوریتم Learn++

در ابتدای آزمایشات، تاثیر پارامترهای مختلف الگوریتم بر روی کارایی نسخه اصلی آن مورد بررسی قرار گرفت تا بدین طریق پارامترهای بهینه تعیین گردد. برای این منظور پارامترهای زیر هر یک با مقادیر متفاوت، که در پرانتز آمده‌اند، در الگوریتم بکار رفتند و خطای طبقه‌بندی الگوریتم در هر مرحله بدست آمده و ثبت شد: تعداد نوروهای لایه مخفی هر WeakLearner به ازاء هر نورو خروجی (۱، ۲، ۵، ۱۰)، تعداد Epoch برای هر WeakLearner (۵۰)، تعداد داده آموزشی (۴۰، ۱۰۰)، تعداد داده آزمایشی (۲۰، ۵۰)، سطح آستانه خطا برای پذیرش هر WeakLearner (۰/۳، ۰/۵).

نتایج این آزمایشها نشان می‌دهد استفاده از WeakLearnerهایی با قدرت یادگیری بیشتر (تعداد نوروهای لایه مخفی و تعداد Epochهای بیشتر) تا حدودی موجب بهبود عملکرد کلی سیستم می‌گردد. هرچند استفاده از WeakLearnerهای بسیار بزرگ با آنکه فاز یادگیری را بسیار طولانی می‌کنند، موجب کاهش قابل توجهی در خطای نهایی سیستم نخواهد شد. از طرف دیگر استفاده از WeakLearnerهای بسیار کوچک نیز به علت آنکه به ندرت به دستیابی به حد آستانه خطای قابل قبول برای پذیرش WeakLearner می‌شوند، فاز یادگیری را طولانی می‌نماید. به علاوه استفاده از تعداد WeakLearnerهای بیشتر و مجموعه داده‌های آموزشی و آزمایشی بزرگتر برای آموزش هریک، در یادگیری یک مجموعه داده جدید نیز موجب بهبود عملکرد کلی سیستم می‌گردد.

۵- بحث و نتیجه گیری

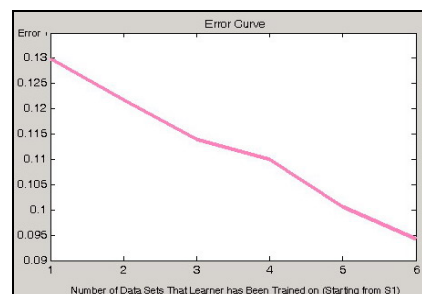
شکل نزولی نمودارهای خطا نشان دهنده موفقیت یادگیری افزایشی سیستم می باشد. در آزمایشهای صورت گرفته، تنها در مواردی که سیستم بر روی مجموعه آموزشی و در ابتدای فاز یادگیری به خطای صفر رسیده است، نمودار خطا بر روی این مجموعه آموزشی حالت صعودی (و البته با شیب کم) دارد، که این مساله نیز بدلیل وجود رای گیری اکثریت در سیستم طبیعی به نظر می رسد. مقایسه شکلهای ۲ و ۳ نشان می دهد استفاده از WeakLearnerهایی با یک خروجی موجب کاهش خطای نهایی سیستم یادگیری شده است. این کاهش خطا در حدود یک درصد بوده است. همچنین مقایسه شکلهای ۲ و ۴ نشان می دهد استفاده از هنجارسازی امتیازات برای رای گیری اکثریت وزندار موجب کاهش چشمگیری (حدود ۴ درصد) در خطای نهایی سیستم شده است.

با توجه به این موارد، به نظر می رسد تلاش برای بهبود عملکرد الگوریتم یادگیری افزایشی Learn++ موفق بوده است. استفاده از رابطه $(1 - \beta)$ برای محاسبه امتیازات به جای استفاده از رابطه $\log(\frac{1}{\beta_i})$ ، به تنهایی موجب کاهش اندکی در میزان خطا شده است. پس از آن استفاده از WeakLearnerهایی با یک خروجی، کاهش یک درصدی خطای نهایی سیستم یادگیری را موجب گردیده است. هنجار سازی امتیازات برای رای گیری اکثریت وزندار موجب کاهش قابل توجه ۴ درصدی در خطای نهایی سیستم شده است. در انتها، استفاده همزمان از تمامی این موارد (شکل ۵)، خطای نهایی الگوریتم یادگیری افزایشی ELearn++ را نسبت به Learn++ تقریباً نصف کرده است.

مراجع

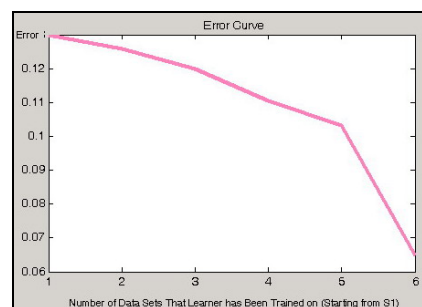
- [1] LiMin Fu, Hui-Huang Hsu, Jose C. Principe, "Incremental Backpropagation Learning Networks", IEEE Trans. On Neural Networks, Vol. 7, No.3, pp. 757-761, May 1996.
- [2] Hown-Wen Chen, Von-Wun Soo, "Design of Adaptive and Incremental Feed-Forward Neural Networks", IEEE International Conference on Neural Networks, Vol. 1, pp. 479-484, 1993.
- [3] R. Polikar, L. Upda, S. Upda, V. Honavar, "LEARN++: An Incremental Learning Algorithm for Multilayer Perceptron Neural Networks.", Proceedings of ICASSP 2000, vol. 6, pp. 3414-3417, 2000.
- [4] R. Polikar, L. Upda, S. Upda, V. Honavar, "LEARN++: An Incremental Learning Algorithm for Supervised Neural Networks", IEEE Trans. On Systems, Man, and Cybernetics, Part c: Applications and Reviews, Vol. 31, No. 4, pp. 497-508, Nov. 2001.
- [5] R. Polikar, J. Byorick, S. Krause, A. Marino, M. Moreton, "Learn++: a Classifier Independent Incremental Learning Algorithm for Supervised Neural Networks", IJCNN 2002, Proc. Of the 2002 Int. Joint Conference on Neural Networks, IEEE, Vol. 2, pp. 1742-1747, May 2002.
- [6] R. Polikar, "Learn++: An Incremental Learning Algorithm Based On Psycho-physiological Model of Learning", Proc. Of the 23rd Annual EMBS Int. Conference, IEEE, Oct. 2001.

پس از این مرحله، و مقداردهی پارامترها با مقداری که منجر به بهترین نتایج بر روی نسخه اصلی الگوریتم شده بود، الگوریتم با استفاده از WeakLearnerهایی با یک خروجی به کار گرفته شد. شکل ۳ نمودار خطای الگوریتم را در این حالت نشان می دهد.

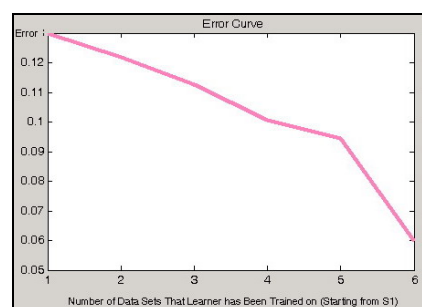


شکل ۳. نمودار خطای سیستم یادگیری افزایشی بر روی مجموعه داده آزمایشی مستقل، استفاده از WeakLearnerهایی با یک خروجی

در آزمایشی دیگر پس از مقدار دهی پارامترها با همان مقداری که منجر به بهترین کارایی در نسخه اصلی الگوریتم شده بود، الگوریتم نسخه اصلی با اضافه کردن مرحله هنجار سازی امتیازات کلاسها اجرا شد. شکل ۴ نمودار خطای الگوریتم را در این حالت نشان می دهد. در نهایت و در آزمایشی دیگر، ایده های استفاده از WeakLearnerهایی با یک خروجی و هنجار سازی امتیازات کلاسها با هم بکار گرفته شد. بهترین نتایج بدست آمده در آزمایشها متعلق به این حالت می باشد که در آن کمترین خطای طبقه بندی بدست آمد. شکل ۵ نمودار خطای الگوریتم را در این حالت نشان می دهد.



شکل ۴. نمودار خطای سیستم یادگیری افزایشی بر روی مجموعه داده آزمایشی مستقل، با استفاده از هنجار سازی امتیازات



شکل ۵. نمودار خطای نهایی الگوریتم ELearn++ بر روی مجموعه داده آزمایشی مستقل: کمترین خطای طبقه بندی در این حالت بدست آمده است

زیرنویس‌ها

¹ Enhanced Learn++

² Non stationary

³ Incremental learning

⁴ Compactness

⁵ Plasticity

⁶ Stability

⁷ Noise

⁸ Generalization

⁹ Hypothesis

¹⁰ Weighted Majority

¹¹ Over Fitting

¹² Instance Space