

ارائه روشی برای پیش‌بینی خطای محتوا در محاسبات ابری با روش نزدیکترین همسایه

رضا ایمانی

چکیده

رایانش ابری یک تکنولوژی جدید برای مدیریت و ارائه خدمات بر روی اینترنت است. رایانش ابری، گروهی بزرگ از کامپیوترهای متصل به هم است. این کامپیوترها می‌توانند رایانه‌های شخصی یا سرورهای شبکه باشند. همچنین آن‌ها می‌توانند عمومی یا خصوصی باشند. پیش‌بینی خطای محتوا یکی از چالش‌های اساسی در محاسبات ابری می‌باشد. در این پژوهش سعی بر این است، با استفاده از الگوریتم نزدیکترین همسایه دقت پیش‌بینی خطای محتوا در محاسبات ابری افزایش داده شود. نتایج شبیه‌سازی در نرم‌افزار متلب نشان می‌دهد که روش برای مجموعه داده دارای کارایی بسیار بالایی می‌باشد. نتایج را با عملکرد سه شبکه عصبی پرسپترون، بر اساس عملکرد شعاعی و نزدیکترین همسایه مقایسه شده است و در نهایت، بهترین ترکیب در بین این سیستم‌ها، سیستم نزدیکترین همسایه شناسایی خطا را با دقت نزدیک ۷۴ درصد به دست داد.

کلید واژه : رایانش ابری، نزدیکترین همسایه، پیش‌بینی خطای محتوا، تابع عملکرد شعاعی

مقدمه

به موازات دسترسی پردازش دستی
قدرتمند، سیار و ابزارهای
حسگر، مسیر فعل و انفعال،

با رشد محبوبیت اینترنت و وب،

اداره کردن زندگی، مدیریت تجارت و دسترسی یا دریافت خدمات را عوض کرده است. کاهش هزینه‌های محاسبات و ارتباطات توجه را از شخصی بودن به محاسبات مرکزداده^۱ محوری جلب می‌کند. همچنین محاسبات توزیع‌شده و موازی حدود چندین سال است که وجود دارند که شکل‌های جدید آن رایانش ابری^۲، و چند هسته^۳ تغییرات جامعی را در صنعت به همراه داشته است. این گرایش‌ها کانون توجه صنعت را از توسعه نرم افزارها برای کامپیوترهای شخصی به مراکز داده ابری که میلیون‌ها کاربر را جهت استفاده از نرم افزار به طور همزمان، سوق می‌دهد [۱].

محاسبات ابری مدلی است برای دسترسی به یک شبکه یکپارچه و مناسب و برحسب تقاضا جهت به اشتراک‌گذاری حجم زیادی از منابع محاسباتی قابل تنظیم (مانند شبکه، سرورها، ذخیره‌گاه‌ها، برنامه‌ها و سرویس‌ها) که می‌تواند با کمترین نیاز به مدیریت و تنها به وسیله یک فراهم‌کننده سرویس ابر به سرعت قانونمند شده و انتشار یابد [۲]. رایانش ابری یک سبک ابتکاری و مهیج برای برنامه‌نویسی و استفاده از رایانه‌ها می‌باشد. رایانش ابری فرصت‌های فوق‌العاده‌ای را برای توسعه دهندگان نرم‌افزار ایجاد می‌کند [۳]. این معماری جدید از پهنای باند به صورت بهینه استفاده می‌کند و هزینه‌ی ارتباطی برای دسترسی به اطلاعات

پرطرفدار را کاهش می‌دهد. این شبکه همچنین ازدحام در شبکه را کاهش داده و پیش‌بینی خطا، تحویل داده و تحمل قطعی بهتری را فراهم می‌کند [۴].

با پیشرفت محاسبات ابری، پیش‌بینی خطای داده‌ها به امری مهم در محاسبات ابری تبدیل شده است. بطوریکه پیش‌بینی خطاهای ابری مهمترین مانع سرعت و گسترش نرم‌افزارهای محاسبات ابری می‌باشد به همین دلیل پیش‌بینی خطا یکی از فاکتورهای کلیدی برای موفقیت ابر می‌باشد [۵]. رشد فناوری اطلاعات و ارتباطات، تغییرات پرشتاب محیط کسب و کار و همچنین گسترش روزافزون سرویس‌های ذخیره‌سازی ابری، شرکت‌ها و کاربران این عرصه را تشویق نموده تا نگهداری و مدیریت داده‌های خود را به ارائه دهندگان خدمات فضای ذخیره‌سازی ابری واگذار نمایند. اما نگرانی از پیش‌بینی خطا و حفظ حریم شخصی کاربران، همواره به عنوان چالشی مهم و بحث برانگیز در ابر مطرح است [۶].

به طور کلی، پردازش ابری با چندین خطر پیش‌بینی خطای مجزا، اما مرتبط با هم مواجه است. علاوه بر این که ممکن است اطلاعات ذخیره شده، توسط مهاجمان به سرقت رفته یا در اثر خرابی سیستم از بین برود، ممکن است تأمین‌کننده خدمات ابری نیز از این اطلاعات سوء استفاده کند یا در اثر فشارهای دولت آن را افشا کند واضح است که این خطرات پیش‌بینی خطا پیامدهای وخیمی را در پی دارند [۷]. تحقیقاتی نیز در این زمینه انجام شده است. در [۸] برای محافظت از محیط محاسبات

^۱ Data Center

^۲ Cloud Computing

^۳ Multicore

ابری یک قالب برای استفاده از سیستم تشخیص نفوذ^۴ طراحی گردید که در آن بتوان حملات از نوع DDoS را به صورت متمرکز مدیریت نمود. آن‌ها بر این عقیده بودند که اگر به ازای هر ماشین مجازی یک سیستم تشخیص نفوذ استفاده شود می‌توان مدیریت حملات را به خوبی انجام داد و با تبادل اطلاعات بین آن‌ها به حملات پاسخ دهند. در [۷] به طراحی یک معماری مدیریتی سیستم تشخیص نفوذ توسعه یافته پرداخته شده است که در آن سیستم تشخیص نفوذ در کل سیستم توزیع گردیده است که همگی زیر نظر یک مرکز کنترل و واحد مدیریت می باشند. در این تحقیق سنسورهای به عنوان سیستم تشخیص نفوذ در سیستم فراهم کننده خدمات و کاربران کار گذاشته می شود که به محض بروز هر نوع ناهنجاری به مرکز کنترل و مدیریت گزارش داده می شوند و مرکز در مورد آن تصمیم گیری می کند. روش دیگری که برای محافظت از محیط ابری طراحی شد در [۹] می باشد. در این مقاله حملات با توجه به میزان خسارت و همچنین سطح حملات آن‌ها در محیط ابری دسته بندی شدند بر این اساس برای هر حمله یک ریسک خطر تهیه و به آن‌ها داده شد و بر اساس آن روشی برای مقابله با آن در نظر گرفته شده است.

در [۱۰] نیز همانند [۹] از روش تعیین میزان ریسک برای منابع استفاده شده است و در سیستم از یک سیستم تشخیص نفوذ مرکزی استفاده گردیده است. مقاله [۱۱] به طراحی یک سیستم امنیتی برای تبادل امن اطلاعات با استفاده

از احراز هویت و رمزنگاری پرداخته است. در سیستم پیشنهاد شده دو ابر (دوقلو) در نظر گرفته شده است. یک ابر معتبر^۵ و یک ابر ارتباطی^۶. یکی از روش های سنجش موفقیت سیستم ابر روش K نزدیکترین همسایه است. روش K نزدیکترین همسایه^۷ گروه شامل K رکورد از مجموعه رکوردهای آموزشی که نزدیکترین رکوردها به رکورد آزمایشی باشند را انتخاب کرده و بر اساس برتری رده یا برجسب مربوط به آن‌ها در مورد دسته رکورد آزمایشی مزبور تصمیم گیری می نماید. به عبارت ساده تر این روش رده ای را انتخاب می کند که در همسایگی انتخاب شده بیشترین تعداد رکورد منتسب به آن دسته باشند. بنابراین رده ای که از همه رده ها بیشتر در بین K نزدیکترین همسایه مشاهده شود، به عنوان رده رکورد جدید در نظر گرفته می شود [۱۲]. با توجه به مطالب فوق پیش بینی خطای محتوای داده و حفظ حریم خصوصی که از چالش های پیش بینی خطا و مسئله مهم در محیط محاسبات ابری هستند لذا در این پژوهش سعی داریم، با استفاده از الگوریتم نزدیکترین همسایه خطای محتوا در محاسبات ابری را پیش بینی کنیم.

روش پژوهش

الگوریتم نزدیکترین همسایه

یکی از الگوریتم های معروف دسته بندی، الگوریتم نزدیک همسایه

^۵ Trust

^۶ Commodity

^۷ – K Nearest Neighbor

^۴ IDS

(۱)

$$d(X, Y) = \sqrt{\sum_{i=1}^m (x_i - y_i)^2}$$

که در آن m تعداد ویژگی‌های نمونه‌ها و x_i و y_i مقدار ویژگی i ام دو نمونه است.

واضح است که این معیار زمانی قابل استفاده و کاراست که در فضای ویژگی‌ها، تمام ویژگی‌ها عددی و مقادیر حقیقی بوده و دارای بازه تغییرات مشابه باشند. در مقابل در حالتی که ویژگی‌ها دارای مقادیر صحیح باشند، معیار فاصله منتهی^۸ کارایی بهتر را از خود نشان می‌دهد. این معیار فاصله به صورت زیر تعریف می‌شود:

(۲)

$$d(X, X) = \sum_{i=1}^m |x_i - y_i|$$

که همانند معیار فاصله اقلیدسی، m تعداد ویژگی‌های نمونه‌ها و x_i و y_i مقدار ویژگی i ام دو نمونه است.

معیار فاصله ماهالانوبیس^۹ یکی دیگر از معیارهای فاصله‌ی پرکاربرد و معرف است. در این معیار، علاوه بر مقادیر ویژگی‌ها، چگونگی پخش نمونه‌ها در فضای ویژگی در نظر گرفته می‌شود. این پخش، توسط ماتریس کواریانس^{۱۰} در تابع معیار فاصله تأثیر می‌گذارد. این معیار به صورت زیر تعریف می‌شود:

است؛ با اینکه از معرفی آن چندین دهه می‌گذرد. این روش برای آموزش، مدلی را بر روی نمونه‌ها فرض نمی‌کند. همچنین این روش نیازی به تنظیم پارامتر ورودی توسط فرد متخصص ندارد. در واقع این روش زمانی را برای آموزش صرف نکرده و تا آمدن درخواست برای برچسب نمونه-ی بدون برچسب صبر می‌نماید. بنابراین این روش هزینه زمان برای فاز آموزش ندارد و در مقابل به زمان بیشتری در فاز تست نیاز دارد. همچنین برای تعیین مقدار شباهت بین نمونه‌ها، نیاز به تعیین یک معیار فاصله است. معیارهای فاصله مختلف در فضای ویژگی‌های متفاوت، کارایی‌های متفاوتی از خود نشان می‌دهند [۱۳].

اولین نکته قابل توجه در مورد الگوریتم نزدیکترین همسایه، نحوه به دست آوردن فاصله بین نمونه‌ها است. معیارهای فاصله مختلفی وجود دارد که می‌توان در این روش استفاده کرد. همچنین در بعضی از کارهای گذشته، معیارهای خاص این روش نیز معرفی شده‌اند که با استفاده از آن‌ها می‌توان کارایی بهتری را به دست آورد. در ادامه، چند معیار فاصله پرکاربرد معرفی شده‌اند [۱۴].

فاصله اقلیدسی از گذشته، پرکاربردترین معیار فاصله در زمینه‌های مختلف بوده است. همچنین در دسته‌بند نزدیکترین همسایه سنتی نیز این معیار مورد استفاده شده است. معیار فاصله اقلیدسی بین دو نمونه‌ی Y و X به صورت زیر تعریف می‌شود:

^۸ Manhattan

^۹ Mahalanobis

^{۱۰} Covariance Marix

(۳)

که در آن، $N_{f,a}$ برابر تعداد تکرار مقدار a در ویژگی f است و $N_{f,a,c}$ همین تکرار است در نمونه هایی که دارای برچسب c هستند.

بر اساس تعریف VDM معیار $HVDM$ به صورت زیر تعریف می شود:

(۵)

$$d(X, Y) = \sqrt{\sum_{f=1}^m d_f^2(x_f, y_f)}$$

که در آن:

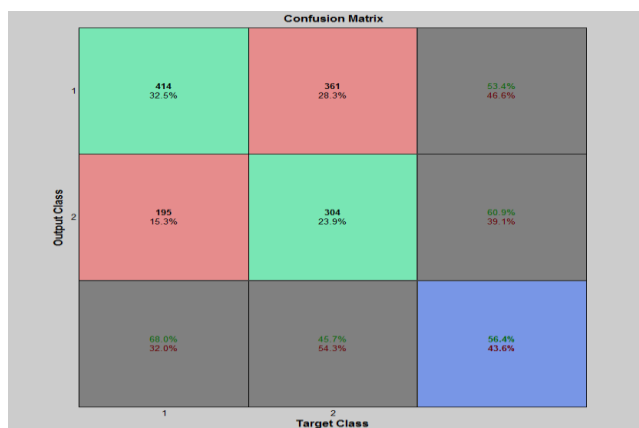
(۶)

$$d_f(x_f, y_f) = \begin{cases} 1 & \text{if } x \text{ or } y \text{ is unknown} \\ vdm_f(x, y) & \text{if } f \text{ is a nominal} \\ \frac{|x - y|}{4\sigma_f} & \text{if } f \text{ is numeric} \end{cases}$$

یافته ها

۱. عملکرد شبکه عصبی پرسپترون

برای این سیستم در این تحقیق از تابع آماده متلب استفاده کرده ایم.



$$d(X, X) = \sqrt{(\bar{x} - \bar{y})^T S^{-1} (\bar{x} - \bar{y})}$$

که در آن $\bar{x}$ و $\bar{y}$ دو بردار ویژگی و S ماتریس کواریانس ویژگی ها است.

در مسائل دنیای واقعی، مقادیر ویژگی ها ممکن است که همیشه مقدار حقیقی نداشته باشند. در بسیاری از مسائل، بعضی از ویژگی ها دارای مقدار نامی می باشند. در این گونه ویژگی ها، برای دو مقدار، رابطه بزرگتری و کوچکتری و اختلاف دو مقدار مفهومی ندارد و تعریف نمی شود. در این گونه مقادیر، تنها تفاوت و یا شباهت دو مقدار، تفاوت بین دو نمونه را نشان می دهد. برای چنین فضای ویژگی که هم دارای ویژگی های نامی هست و هم مقادیر حقیقی، $Wilson$ و $Martinez$ معیاری به نام معیار فاصله مقادیر ناهمگن^{۱۱} ($HVDM$) را ارائه داده اند. این معیار توسعه یافته معیاری بانام معیار فاصله مقادیر^{۱۲} (VDM) است که قبل از آن ارائه شده بود.

در VDM تفاوت بین دو مقدار a و b برای یک ویژگی f به صورت زیر تعریف می شود:

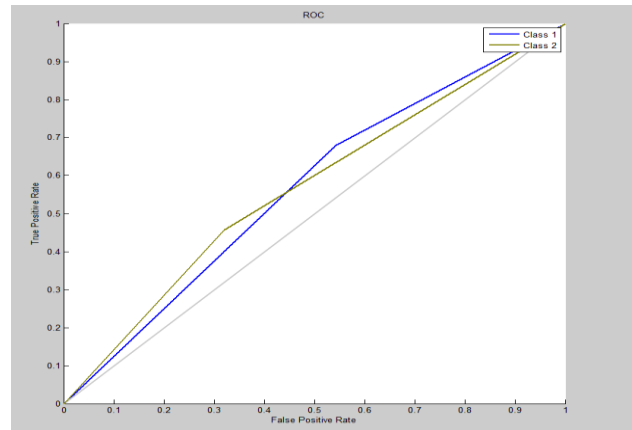
(۴)

$$vdm_f(X, Y) = \sum_{c=1}^C \left(\frac{N_{f,a,c}}{N_{f,a}} - \frac{N_{f,b,c}}{N_{f,b}} \right)^2$$

^{۱۱} Heterogeneous Value Difference Metric (HVDM)

^{۱۲} Value Difference Metric (VDM)

نمودار ۱: کانفیوژن برای شبکه
عصبی پرسپترون

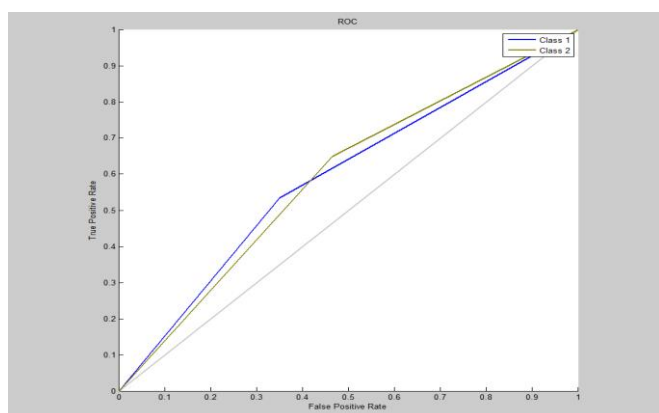


نمودار ۲: مشخصه عامل گیرنده
برای شبکه عصبی پرسپترون

همچنین نتایج بدست آمده از مقایسات کلی در جدول زیر آورده شده است.
جدول ۱: نتایج بدست آمده از ترکیب شبکه عصبی پرسپترون

MEAN		k ^۱	k ^۲	k ^۳	k ^۴
۰,۲۰۴۶۹ ۳	نرخ کلاس بندی اشتباه	۰,۲۵	۰,۱۳۱۵۷۸ ۹	۰,۲۱۰۵۲ ۶	۰,۲۲۶۶۶ ۷
۰,۲۸۷۱۲	میزان هزینه کلاس بندی اشتباه	۰,۳۵۹۱۳۹۷ ۸	۰,۲۱۶۸۴۵ ۹	۰,۳۴۲۵۰ ۹	۰,۲۲۹۹۸ ۴
۰,۱۴۳۵۶	نرمال شده میزان هزینه کلاس بندی اشتباه	۰,۱۷۹۵۶۹۸ ۹	۰,۱۰۸۴۲۲ ۹	۰,۱۷۱۲۵ ۴	۰,۱۱۴۹۹ ۲
۰,۷۸۶۳۴ ۲	حساسیت	۰,۷۳۳۳۳۳۳ ۳	۰,۸۸۸۸۸۸ ۹	۰,۸۲۹۲۶ ۸	۰,۶۹۳۸۷ ۸
۰,۸۱۹۷۰ ۹	نرخ اختصاصی	۰,۷۷۴۱۹۳۵ ۵	۰,۸۳۸۷۰۹ ۷	۰,۷۴۲۸۵ ۷	۰,۹۲۳۰۷ ۷
۰,۷۰۹۱۲ ۵	دقت	۰,۷۱	۰,۶۹	۰,۷۲	۰,۷۱۶۵
۰,۸۶۲۲۵ ۸	صحت	۰,۸۲۵	۰,۸۸۸۸۸۸ ۹	۰,۷۹۰۶۹ ۸	۰,۹۴۴۴۴ ۴
۰,۷۸۶۳۴ ۲	فراخوانی مجدد	۰,۷۳۳۳۳۳۳ ۳	۰,۸۸۸۸۸۸ ۹	۰,۸۲۹۲۶ ۸	۰,۶۹۳۸۷ ۸
۰,۸۱۸۷۲ ۱	معیار ترکیبی F	۰,۷۷۶۴۷۰۵ ۹	۰,۸۸۸۸۸۸ ۹	۰,۸۰۹۵۲ ۴	۰,۸
۰,۴۵۵۰۱	سازگاری	۰,۳۴۶۲۳۶۵	۰,۷۲۷۵۹۸	۰,۶۲۹۲۶	۰,۱۱۶۹۵

۴		۶	۶	۸	۴
۰,۸۰۳۰۲ ۶	AUC	۰,۷۵۳۷۶۳۴ ۴	۰,۸۶۳۷۹۹ ۳	۰,۷۸۶۰۶ ۳	۰,۸۰۸۴۷ ۷
۰,۷۹۳۲۴ ۴	معیار توازن	۰,۷۵۲۹۱۷۳ ۶	۰,۸۶۱۵۰۷ ۷	۰,۷۸۱۷۴ ۴	۰,۷۷۶۸۰ ۹
۴,۰۸۰۳۳ ۱	زمان اجرا	۴,۲۰۲۶۹۷۸ ۶	۴,۳۱۸۴۴۲ ۸	۳,۵۹۱۴۴ ۱	۴,۲۰۸۷۴ ۲

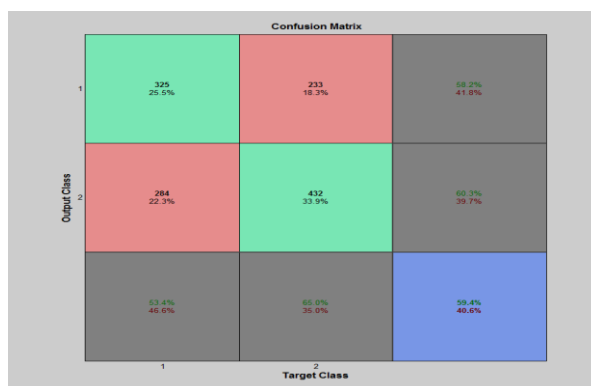


نمودار ۳: مشخصه عامل گیرنده
برای سیستم RBF

همان طور که از نتایج پیدا است در هر ۴ اجرا نتایج ثبات شبکه عصبی پرسپترون را نمایش می دهد و در این اجرای سوم بهترین نتیجه را به همراه داشته است. بنابراین تا این مرحله بهترین دقت بدست آمده ۰,۷۰۹۱ است.

۲. عملکرد سیستم RBF^{۱۳}

برای سیستم RBF در این تحقیق از تابع آماده متلب استفاده کرده ایم. در این تابع شعاع پوشش برابر ۲ قرار داده ایم. نتایج بدست آمده به شرح زیر است.



نمودار ۳: کانفیوژن برای سیستم RBF

جدول ۲: نتایج بدست آمده از ترکیب سیستم RBF

MEAN		k ^۱	k ^۲	k ^۳	k ^۴
۰,۴۳۲۵ ۴۴	نرخ کلاس بندی اشتباه	۰,۴۳۴۲۱ ۰۵۳	۰,۳۹۴۷۳ ۶۸	۰,۴۰۷۸ ۹۵	۰,۴۹۳۳ ۳۳
۰,۵۹۵۰ ۸۶	میزان هزینه کلاس بندی اشتباه	۰,۶۲۱۳۱ ۱۴۸	۰,۶۴۶۸۵ ۳۱	۰,۴۵۸۳ ۳۳	۰,۶۵۳۸ ۴۶
۰,۲۹۷۵ ۴۳	نرمال شده میزان هزینه کلاس بندی اشتباه	۰,۳۱۰۶۵ ۵۷۴	۰,۳۲۳۴۲ ۶۶	۰,۲۲۹۱ ۶۷	۰,۳۲۶۹ ۲۳
۰,۵۶۲۱ ۰۱	حساسیت	۰,۵۵۷۳۷ ۷۰۵	۰,۶۱۵۳۸ ۴۶	۰,۵۸۳۳ ۳۳	۰,۴۹۲۳ ۰۸
۰,۶۲۳۸ ۶۴	نرخ اختصاصی	۰,۶	۰,۵۴۵۴۵ ۴۵	۰,۷۵	۰,۶
۰,۵۶۷۴ ۵۶	دقت	۰,۵۶۵۷۸ ۹۴۷	۰,۶۰۵۲۶ ۳۲	۰,۵۹۲۱ ۰۵	۰,۵۰۶۶ ۶۷
۰,۹۰۱۱ ۳	صحت	۰,۸۵	۰,۸۸۸۸۸ ۸۹	۰,۹۷۶۷ ۴۴	۰,۸۸۸۸ ۸۹
۰,۵۶۲۱ ۰۱	فراخوانی مجدد	۰,۵۵۷۳۷ ۷۰۵	۰,۶۱۵۳۸ ۴۶	۰,۵۸۳۳ ۳۳	۰,۴۹۲۳ ۰۸
۰,۶۹۱۱ ۶	معیار ترکیبی F	۰,۶۷۳۲۶ ۷۳۳	۰,۷۲۷۲۷ ۲۷	۰,۷۳۰۴ ۳۵	۰,۶۳۳۶ ۶۳
- ۳,۱۵۶۰ ۸	سازگاری	- ۱,۲۴۲۶۲ ۳	- ۱,۶۵۷۳۴ ۲۷	- ۶,۹۱۶۶ ۷	- ۲,۸۰۷۶ ۹
۰,۵۹۲۹ ۸۲	AUC	۰,۵۷۸۶۸ ۸۵۲	۰,۵۸۰۴۱ ۹۶	۰,۶۶۶۶ ۶۷	۰,۵۴۶۱ ۵۴
۰,۵۸۹۱ ۲۳	معیار توازن	۰,۵۷۸۱۴ ۹۸۶	۰,۵۷۸۹۶ ۵۲	۰,۶۵۶۴ ۰۸	۰,۵۴۲۹ ۷۱
۲۲,۲۱۱ ۹۸	زمان اجرا	۲۷,۸۵۲۸ ۰۶۸	۲۳,۶۳۰۴ ۳۹	۲۴,۰۵۷ ۲۷	۱۳,۳۰۷ ۳۹

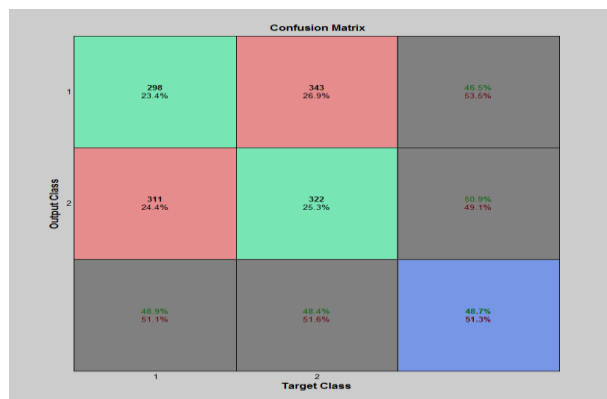
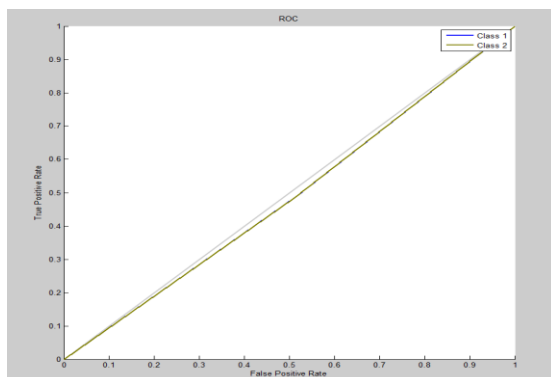
۳. عملکرد سیستم K نزدیکترین همسایه

سیستم K نزدیکترین همسایه در این تحقیق بر مبنای خوشه بندی K نزدیکترین همسایه عمل می کند.

بنابراین تا این مرحله بهترین دقت به دست آمده ۰,۶۰۵ است که از روش قبلی نتیجه بهتری است.

نتایج بدست آمده به شرح زیر است.

نمودار ۴: کانفیوژن برای سیستم K نزدیکترین همسایه



نمودار ۵: مشخصه عامل گیرنده برای سیستم K نزدیکترین همسایه

جدول ۳: نتایج بدست آمده از ترکیب سیستم K نزدیکترین همسایه

MEAN		k ^۱	k ^۲	k ^۳	k ^۴
۰,۲۵۴۱ ۶۷	نرخ کلاس بندی اشتباه	۰,۲۲۳۶۸ ۴۲۱	۰,۱۱۸۴۲ ۱۱	۰,۴۰۷۸ ۹۵	۰,۲۶۶۶ ۶۷
۰,۳۱۴۷ ۳۸	میزان هزینه کلاس بندی اشتباه	۰,۳۲۸۴۰ ۰۲۸	۰,۲۱۲۱۸ ۴۹	۰,۴۵۸۳ ۳۳	۰,۲۶۰۰ ۳۳
۰,۱۵۷۳ ۶۹	نرمال شده میزان هزینه کلاس بندی اشتباه	۰,۱۶۴۲۰ ۰۱۴	۰,۱۰۶۰۹ ۲۴	۰,۲۲۹۱ ۶۷	۰,۱۳۰۰ ۱۷
۰,۷۳۳۲ ۹۸	حساسیت	۰,۷۶۷۴۴ ۱۸۶	۰,۹۲۸۵۷ ۱۴	۰,۵۸۳۳ ۳۳	۰,۶۵۳۸ ۴۶
۰,۸۱۸۶ ۱۳	نرخ اختصاصی	۰,۷۸۷۸۷ ۸۷۹	۰,۸۲۳۵۲ ۹۴	۰,۷۵ ۰	۰,۹۱۳۰ ۴۳
۰,۷۴۵۸ ۳۳	دقت	۰,۷۷۶۳۱ ۵۷۹	۰,۸۸۱۵۷ ۸۹	۰,۵۹۲۱ ۰۵	۰,۷۳۳۳ ۳۳
۰,۹۰۳۲ ۱۴	صحت	۰,۸۲۵ ۰	۰,۸۶۶۶۶ ۶۷	۰,۹۷۶۷ ۴۴	۰,۹۴۴۴ ۴۴
۰,۷۳۳۲ ۹۸	فراخوانی مجدد	۰,۷۶۷۴۴ ۱۸۶	۰,۹۲۸۵۷ ۱۴	۰,۵۸۳۳ ۳۳	۰,۶۵۳۸ ۴۶
۰,۷۹۸۷ ۲۴	معیار ترکیبی F	۰,۷۹۵۱۸ ۰۷۲	۰,۸۹۶۵۵ ۱۷	۰,۷۳۰۴ ۳۵	۰,۷۷۲۷ ۲۷
- ۱,۴۳۵۱	سازگاری	۰,۴۶۴۴۱ ۰	۰,۸۴۰۳۳ ۰	- ۶,۹۱۶۶	- ۰,۱۲۸۷

۷		۱۵۶	۶۱	۷	۶
۰,۷۷۵۹ ۵۶	AUC	۰,۷۷۷۶۶ ۰۳۲	۰,۸۷۶۰۵ ۰۴	۰,۶۶۶۶ ۶۷	۰,۷۸۳۴ ۴۵
۰,۷۶۱۷ ۱۱	معیار توازن	۰,۷۷۷۴۲ ۵۶۳	۰,۸۶۵۳۸ ۲۲	۰,۶۵۶۴ ۰۸	۰,۷۴۷۶ ۲۷
۹,۷۷۳۶ ۹۶	زمان اجرا	۴,۴۴۶۳۷ ۸۵۳	۴,۴۱۳۵۷ ۹۶	۲۶,۰۳۷ ۳۷	۴,۱۹۷۴ ۵۶

۲. H. Brian, T. Brunschwiler, H. Dill, H. Christ, B. Falsafi, M. Fischer, et al., "Cloud computing," Communications of the ACM, vol. ۵۱, pp. ۹-۱۱, ۲۰۰۸.

۳. Magoulès, F., Pan, J., & Teng, F. (۲۰۱۲). Cloud Computing : Data-Intensive Computing and Scheduling. Boca Raton: CRC Press.

۴. A. Soule, K. Salamatian, and N. Taft, "Combining filtering and statistical methods for anomaly detection," in Proceedings of the ۵th ACM SIGCOMM conference on Internet Measurement, pp. ۳۱-۳۱, ۲۰۰۵.

۵. P. Barford, J. Kline, D. Plonka, and A. Ron, "A signal analysis of network traffic anomalies," in Proceedings of the ۲nd ACM SIGCOMM Workshop on Internet measurement, pp. ۷۱-۸۲, ۲۰۰۲.

۶. A. Lakhina, M. Crovella, and C. Diot, "Diagnosing network-wide traffic anomalies ",in ACM SIGCOMM Computer Communication Review, pp. ۲۱۹-۲۳۰, ۲۰۰۴.

۷. S. Roschke, F. Cheng, and C. Meinel, "Intrusion detection in the

همان طور که از نتایج پیدا است هر ۴ اجرا باعث بهبود سیستم K نزدیکترین همسایه نسبت به حالت قبلی شده‌اند و از بین این سه روش، روش K نزدیکترین همسایه بهترین نتیجه را به همراه داشته است.

نتیجه‌گیری

با توجه به موضوع مورد پژوهش، در ابتدا شناسایی خطا در محاسبات ابری و انواع شاخص‌های ارزیابی شرح داده شد و در ادامه به مرور عملکرد سه شبکه عصبی پرسپترون، براساس عملکرد شعاعی و K نزدیکترین همسایه پرداختیم در نهایت بهترین ترکیب در این بین سیستم K نزدیکترین همسایه با دقت نزدیک ۷۴ درصد است. با توجه به نتایج به دست آمده کاملاً مشهود است که در بین سیستم‌های پیش‌بینی مورد استفاده بهترین عملکرد را K نزدیکترین همسایه دارا است.

منابع

۱. Buyya, R., Vecchiola, C., & Selvi, S. T. (۲۰۱۳). Mastering cloud computing : foundations and applications programming.

- Sixth International Conference on, pp[^]. 260-270, 2010.
11. S. Bugiel, S. Nürnberger, A.-R. Sadeghi, and T. Schneider, "Twin clouds: An architecture for secure cloud computing," in Workshop on Cryptography and Security in Clouds (WCSC 2011), 2011.
 12. Manvi, S. S., & Shyam, G. K. (2014). Resource management for Infrastructure as a Service (IaaS) in cloud computing: A survey. Journal of Network and Computer Applications, 41, 444-449.
 13. Muja, M., & Lowe, D. G. (2014). Scalable nearest neighbor algorithms for high dimensional data. IEEE Transactions on Pattern Analysis and Machine Intelligence, 36(11), 2227-2240.
 14. R. O. Duda, et al., Pattern classification: John Wiley & Sons, 2012.
 - cloud," in Dependable, Autonomic and Secure Computing, 2009. DASC'09. Eighth IEEE International Conference on, pp[^]. 729-734, 2009.
 15. C.-C. Lo, C.-C. Huang, and J. Ku, "A cooperative intrusion detection system framework for cloud computing networks," in Parallel Processing Workshops (ICPPW), 2010 39th International Conference on, pp[^]. 280-284, 2010.
 16. J.-H. Lee, M.-W. Park, J.-H. Eom, and T.-M. Chung, "Multi-level Intrusion detection system and log management in cloud computing," in Advanced Communication Technology (ICACT), 2011 13th International Conference on, pp[^]. 552-555, 2011.
 17. C. Mazzariello, R. Bifulco, and R. Canonico, "Integrating a network ids into an open source cloud computing environment," in Information Assurance and Security (IAS), 2010.

Abstract

Cloud computing is a new technology for managing and providing services over the Internet. Cloud computing is a large group of connected computers. These computers can be personal computers or network servers. They can also be public or private. The prediction of content error is one of the fundamental challenges in cloud computing. In this research, we try to increase the accuracy of prediction error of content in cloud computing using the nearest neighbor algorithm. The simulation results in MATLAB software show that the method for the data set has a very high performance. The results are compared with the performance of three perceptron neural networks based on radial performance and nearest neighbor. Finally, the best combination between these systems gave the nearest neighbor the fault detection with close accuracy of 94%.

Keyword: cloud computing, closest neighbor, content error prediction, radial function function