
بهینه سازی در الگوریتم K-MEANS مبتنی بر کرنل

حامد ملایی سربیزن

www.ketab.ir



سرشناسه
عنوان و نام پدیدآور
شخصیات نشر
مشخصات ظاهری
شابک
وضعیت فهرست نویسی
یادداشت
موضوع
موضوع
موضوع
موضوع
رده بندی کنگره
رده بندی دیویی
شماره کتابشناسی ملی
وضعیت رکورد
فبا

مولف: حامد ملای سريون

عنوان: بهینه سازی در الگوریتم K-MEANS مبتنی بر کرنل / حامد ملایی، سرینژاد

طراح جلد: فہیمہ پور حسنخانی

نوبت چاپ: اول بهار ۱۴۰۰

شایک: ۷-۹-۹۷۸۹۷-۶۲۲-۹۷۸

شماره: ۱۰۰۰ نسخه

قیمت: ۴۵۰۰۰ تومان

زرنند، خیابان فردوس، نشن فردوس ۲۰، طبقه فوقانی، دارالترجمه الفبا

تلفن: ۳۳۴۳۹۹۵۶ - ۰۹۱۳۵۱۸۸۷۱۶ - ۰۹۲۲۱۱۸۳۴۹۵: شماره



تمام حقوق این اثر متعلق به نشر نوید دانش دارالفنون می باشد. هر گونه استفاده از عناصر صوری و محتوایی آن کلا و

جزئیات، به هر زبانی و به هر شکلی بدون اجازه ناشر ممنوع می باشد.

فهرست مطالب

۶.....	مقدمه
۱۲.....	الگوریتم کرنل
۱۸.....	خوشه بندی K - means
۱۹.....	روال الگوریتم K - Means
۲۶.....	الگوریتم خوشه بندی K - Means
۳۴.....	شیوه بهبود یافته K - Means
۳۵.....	بهبود روش k-means با به کارگیری کرنل
۳۶.....	استفاده از کرنل در k-means
۴۳.....	بهینه سازی الگوریتم k-means با آموزش کرنل
۴۴.....	ساختار پیشنهادی ۱
۴۸.....	ساختار پیشنهادی ۲
۵۱.....	ارزیابی الگوریتم های پیشنهادی
۵۲.....	معرفی داده های استفاده شده در شیوه سازی ها
۵۳.....	معیار ارزیابی کیفیت خوشه بندی
۵۵.....	نتایج ارزیابی الگوریتم (۲)
۶۵.....	بیمبستگی محاسباتی الگوریتم (۱)
۶۶.....	نتایج ارزیابی الگوریتم (۲) (DML-MKL)
۷۰.....	پیاده سازی الگوریتم

مقدمه

$K - means$ یک روش پایه، برای بسیاری از روش های خوشه بندی است. با تحقیقات انجام شده توسط کواترو همکارانش در سال ۲۰۱۲ این روش می تواند در خوشه بندی داده ها در داده کاوی مورد استفاده قرار گیرد و عملکرد مطلوب داشته باشد. الگوریتم $K - means$ را به منظور خوشه بندی به کار خواهیم گرفت. داده های نزدیک یک گروه را باهم می سازند که این گروه یک خوشه در شبکه در نظر گرفته خواهد شد. جاذبه مرکزی خوشه به عنوان مرکز خوشه در نظر گرفته می شود. در موقعیت اولیه اجرا، هر دو داده نزدیک به هم یک گروه یا خوشه را در زمان اولیه می سازند. گره های دیگر به صورت تدریجی به خوشه ها ملحق می شوند و این امر تا زمانی که تعداد خوشه ها به یک حد مشخصی برسد ادامه پیدا می کند. همه داده ها در ابتدای اجرای انرژی یکسانی هستند. داده ای که نزدیکترین فاصله را با داده های دیگر در خوشه دارد به عنوان مرکز خوشه در آن خوشه انتخاب می شود [۱].

خوشه بندی در واقع تعریف یک مجموعه داده بر اساس تشابهات بین اعضای آن در راستای رسیدن به تشخیص یک مجموعه از دسته بندی ها می باشد. هدف اصلی خوشه بندی،

بزرگنمایی همجنس سازی در هر خوشه و عدم تجانس بین خوشه های مختلف می باشد که وابسته به نوع الگوریتم بکار رفته است. پس درواقع، موجودیت هایی که به یک دسته تعلق دارند دارای اشتراک های بیشتر با یکدیگر می باشند تا موجودیت هایی که متعلق به خوشه های مختلف هستند.

با مسئله سنجش شباهت معمولاً به صورت غیرمستقیم برخورد می شود، یعنی مقیاس های مسافت برای تعیین میزان تفاوت بین اشیاء مورد استفاده قرار می گیرد، در چنین روشی اشیاء بسیار مشابه دارای کمترین مقدار تفاوت می باشند. برای اجرای تفکیک، از معیارهای متفاوت بسیاری می توان استفاده کرد. هر معیار ویژگی های خود را داشته و با منفعتهای ایرادهای خود همراه است.

از میان روشهای مختلف شناسایی الگو میتوان به روش های مبتنی بر کرنل به عنوان یکی از حوزه های بسیار فعال در سالهای اخیر اشاره کرد. در روشهای مبتنی بر کرنل می توان به الگوریتم های طبقه بند از قبیل ماشین بردار پشتیبان¹ (SVM)، الگوریتم های رگرسیون

¹ Support Vector Machine

مانند رگرسیون بردار پشتیبان (SVR)^۱ و الگوریتم های کاهش بعد تعمیم یافته از قبیل آنالیز مؤلفه های ویژه مبتنی بر کرنل (KPCA)^۲ و آنالیز تفکیک کننده خطی تعمیم یافته^۳ اشاره کرد. تعداد زیادی از این روشهای مبتنی بر کرنل به طور گسترده در کاربردهای مختلف از قبیل فناوری اطلاعات زیستی^۴، بینائی ماشین^۵، بازشناسی گفتار^۶، داده کاوی^۷، بازاریابی اطلاعات^۸ و شناسایی الگو مورد استفاده قرار می گیرند و این حاکی از این است که روش های مبتنی بر کرنل ابزارهای قوی برای کاربردهای مختلف می باشند. در الگوریتم های مبتنی بر کرنل، داده ها به فضای ویژگی جدیدی با ابعاد بالاتر نگاشت می شوند و سپس در فضای ویژگی جدید فرایند مورد نظر (به عنوان مثال جستجوی مرزهای خطی بین طبقات دادگان) اجرا می شود. در این الگوریتمها از تابعی به نام تابع کرنل استفاده میشود که میزان شباهت بین دو نمونه در فضای ویژگی جدید را محاسبه می کند. اطلاعات شباهت بین تمام زوج نمونه ها در یک ماتریس به نام ماتریس کرنل جمع آوری می شود که ماتریسی متقارن

^۱ Support Vector Regression

^۲ Kernel Principal Component Analysis

^۳ Generalized Discriminant Analysis

^۴ Bioinformatics

^۵ Computer Vision

^۶ Speech Recognition

^۷ Data Mining

^۸ Information Retrieval

و نیمه معین مثبت است و همگرایی الگوریتم های آموزشی به دلیل همین ویژگی ماتریس کرنل تضمین می شود [۳]

در روشهای مبتنی بر کرنل ذکر شده، تنها از یک تابع کرنل و یا به طور معادل از یک ماتریس کرنل استفاده میشود و بنابراین کیفیت عملکرد چنین الگوریتم هایی به نوع کرنل و یا پارامترهای کرنل استفاده شده وابسته است. یکی از راهکارهای ارائه شده برای تعیین کرنل مناسب و یا پارامترهای آن، استفاده از روشهای مبتنی بر اعتبار سنجی متقابل^{۱۰} است که بسیار زمان بر بوده و از طرفی در این روش ممکن است تعیین مناسب ترین تابع کرنل برای داده های موجود امکان پذیر نباشد. از طرفی ممکن است برخی از فضاهای ویژگی پایه دارای اهمیت بیشتری در مقایسه با دیگر فضاهای ویژگی باشند و به همین دلیل می توان برای فضاهای ویژگی و در نتیجه کرنل های مورد بررسی در ترکیب کرنلی، وزن هایی را اختصاص داد. در روش های یادگیری کرنل مرکب (MKL)^{۱۱}، هدف محاسبه میزان اهمیت هر یک از این فضاهای ویژگی و یا وزن کرنل ها است که در اکثر موارد از یک الگوریتم

^{۱۰} Cross Validation

^{۱۱} Multiple Kernel Learning

تکراری برای محاسبه وزن کرنل ها و پارامترهای سیستم یادگیرنده استفاده می شود. کرنل ها را می توان به طرق مختلف خطی و غیر خطی ترکیب نمود که ترکیب خطی به دلیل سادگی، بیشتر مورد استقبال محققین بوده و روش های زیادی بر مبنای این ترکیب وزن دار ارائه شده است. با نگاهی عمیق تر می توان برخی از روش های ارائه شده برای MKL را به عنوان روش هایی برای یادگیری ماتریس کرنل تفسیر نمود. به عبارتی می توان روشهای MKL را این چنین تفسیر نمود که هدف، یادگیری ماتریسی است که به طور غیر صریح داده ها را به فضای جدیدی نگاشت دهد به گونه ای که داده ها در فضای جدید توزیعی متناسب با کاربرد مورد نظر داشته باشند. به عنوان مثال ممکن است هدف طبقه بندی دادگان در فضای جدید با استفاده از روش های خطی باشد. در این دیدگاه، روش های MKL شباهت بسیار زیادی به روش های یادگیری معیار فاصله دارند. از طرفی، مقادیر کرنل و فاصله، هر دو معیاری برای سنجش شباهت میان نمونه ها بوده و از این رو شباهت زیادی در اهداف مورد بررسی در دو روش یادگیری کرنل و یادگیری معیار فاصله وجود دارد. همانطور که نیاز به یافتن کرنل بهینه در روشهای مبتنی بر کرنل منجر به پیدایش روشهای

MKL گردید، نیاز به تعیین تابع فاصله مناسب برای نمایش بهتر داده‌گان نیز سبب پیدایش و گسترش روش‌های یادگیری معیار فاصله شده است [۴].

بر این اساس که تفکیک داده‌ها اصولاً ذات غیر نظارتی دارد، به عنوان یکی از سخت‌ترین و پرچالش‌ترین پرسش‌های مجامع علمی در اتوماسیون شمرده می‌شود. طبیعت غیر نظارتی دلالت بر این امر دارد که ویژگی‌های ساختاری مسئله k -means به وضوح برای ما شناخته شده نیستند، مگر اینکه نمونه‌هایی از دانش پیرامون دامنه به شکل خاص توزیع فضای داده به صورت عدد، حجم، جرم، چگالی، تصویر و جهت خوشه‌ها در دسترس باشد. حجم زیاد اطلاعات و داده‌ها به دلیل استفاده از تکنولوژی روز و استفاده از کامپیوتر در تجارت، دانش و خدمات و پیشرفت تکنیک‌ها و ابزار مورد استفاده در جمع‌آوری و پردازش داده‌ها، گسترش روزافزون وب و اینترنت به عنوان یک رسانه جهانی نیاز آشکاری به تکنولوژی‌ها و ابزارهای جدید و خودکار برای تبدیل هوشمندانه حجم زیاد داده به اطلاعات را ایجاد می‌کنند. بر این اساس خوشه‌بندی k -means باعث از بین رفتن جزئیات داده‌های ورودی خواهد شد اما در نهایت این کار ساده‌سازی امر استفاده از داده در تحلیل را در پی دارد. با این روش، پایگاه‌های داده با حجم زیاد بر اساس تعداد کمتری از خوشه

ها، ارائه می شوند، پس در واقع می توان گفت این تکنیک، داده ها را با خوشه هایشان مدل می کند. این رشد در مقدار داده یک مشکل اساسی در به کار گیری علوم و روشهای تجزیه و تحلیل ایجاد کرده است.

در میان این تکنیک ها، تجزیه و تحلیل خوشه بندی یکی از پر استفاده ترین ها است، و به این صورت است که با استفاده از الگوریتم محبوب $K - Means$ انجام می شود. بسیاری از پژوهش های دانش بنیان نشان میدهد که الگوریتم های خوشه بندی در تحقیقات امروز نقش مؤثری دارند.

الگوریتم کرنل

روشهای مبتنی بر کرنل، الگوریتم های قدرتمندی هستند که از حوزه تئوری علم آمار به مسائل کاربردی یادگیری ماشین وارد شده اند. ایده اصلی در روشهای مبتنی بر کرنل، نگاشت مجموعه داده های ورودی به فضایی برداری است که معمولاً تعداد ابعاد آن خیلی

زیاد است (فضای ویژگی یا فضای هیلبرت^{۱۲}) که پس از آن از الگوریتم های خطی آنالیز الگو برای کشف روابط بین داده های نگاشت یافته استفاده میشود [۴].

یکی از مهم ترین خصوصیات روشهای کرنلی این است که با به کارگیری ترفند کرنل^{۱۳} نیازی به محاسبه صریح نگاشت داده ها به فضای ویژگی نیست. به عبارتی اعمال ترفند کرنل منجر به تعمیم الگوریتم های خطی به نسخه ی غیر خطی آنها می شود. از این رو، نوع کرنلی بسیاری از الگوریتم ها ارائه شده است. اما انتخاب کرنل مناسب و پارامترهای آن (به عنوان مثال درجه چندجمله ای^{۱۴} هر کرنل چند جمله ای، پارامتر واریانس در کرنل گوسی^{۱۵} و...) یکی از مهم ترین مسائل در روشهای کرنلی است چرا که تابع کرنل باید متناسب با داده های ورودی انتخاب شود و نوع کرنل انتخاب شده تأثیر زیادی در کیفیت عملکرد الگوریتم یادگیری مورد بررسی خواهد داشت. بنابراین مهم ترین مسئله در ارتباط با روشهای مبتنی بر کرنل، انتخاب کرنل مناسب است. در سال های اخیر استفاده از ترکیب مناسب تعدادی از توابع کرنل به جای استفاده از یک تابع کرنل پیشنهاد شده است.

¹² Hilbert

¹³ Kernel Trick

¹⁴ Polynomial Kernel

¹⁵ Gaussian Kernel