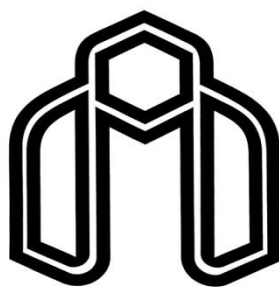


سُبْحَانَكَ يَا عَزِيزُ



دانشگاه صنعتی شاهرود

دانشکده مهندسی کامپیوتر

پایان نامه کارشناسی ارشد مهندسی هوش مصنوعی

کاوش مقیاس پذیر نظرات کاربران با استفاده از محاسبات خوشه‌ای

نگارنده

مجتبی اسداللهی

استاد راهنما

دکتر هدی مشایخی

شهریور ۹۷

شماره: ۵۸۱/ع.ز.
تاریخ: ۹۷/۷/۱

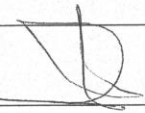
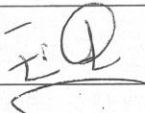
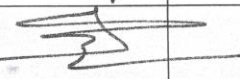
باسمه تعالی

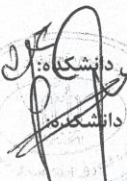


فرم شماره (۳) صورتجلسه نهایی دفاع از پایان نامه دوره کارشناسی ارشد

با نام و یاد خداوند متعال، ارزیابی جلسه دفاع از پایان نامه کارشناسی ارشد آقای **مجتبی اسداللهی** با شماره دانشجویی ۹۴۰۱۹۴۴ رشته مهندسی کامپیوتر گرایش هوش مصنوعی تحت عنوان **کاوش مقیاس پذیر نظرات کاربران با استفاده از محاسبات خوشه‌ای** که در تاریخ ۱۳۹۷/۶/۱۹ با حضور هیأت محترم داوران در دانشگاه صنعتی شاهرود برگزار گردید به شرح ذیل اعلام می گردد:

☐ مردود	☒ قبول (با درجه: خیبوس)
☐ عملی	☒ نظری: نوع تحقیق:

عضو هیأت داوران	نام و نام خانوادگی	مرتبه علمی	امضاء
۱- استاد راهنمای اول	دکتر هدی مشایخی	استادیار	
۲- استاد راهنمای دوم			
۳- استاد مشاور			
۴- نماینده تحصیلات تکمیلی	دکتر منصور فاتح	استادیار	
۵- استاد ممتحن اول	دکتر فاطمه جعفری نژاد	استادیار	
۶- استاد ممتحن دوم	دکتر بهروز مینایی	دانشیار	

نام و نام خانوادگی رئیس دانشگاه: 
تاریخ و امضاء و مهر دانشگاه: 

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).

پاسکزاری

باسپاس از پروردگار،

از استاد کرامیم سرکار خانم دکتر هدی مشایخی بسیار پاسکزارم چرا که بدون راهنماییهای ایشان تائین این پایان نامه بسیار مشکل می نمود. همچنین از تمامی اساتید کرامی که تائین مقطع دسی مرار، نمود فرمودند مراتب پاسکزاری را دارم.

و در انتها با سپاس از دو وجود مقدس:

آنان که ناتوان شدند تا ما به توانایی برسیم...

آنان که موهایشان سپید شد تا ما رو سفید شویم...

و عاشقانه سوختند تا که ما بخش وجود ما و روشنگر راهمان باشند...

پدرانمان

مادرانمان

تعهد نامه

اینجانب **مجتبی اسداللهی** دانشجوی دوره کارشناسی ارشد رشته **هوش مصنوعی و رباتیک** دانشکده کامپیوتر و فناوری اطلاعات دانشگاه صنعتی شاهرود نویسنده پایان نامه **کاوش مقیاس پذیر نظرات کاربران با استفاده از محاسبات خوشه‌ای تحت راهنمایی خانم دکتر هدی مشایخی** متعهد می‌شوم.

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهشهای محققان دیگر به مرجع مورد استفاده استناد شده است.
- مطالب مندرج در پایان نامه تاکنون توسط خود یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ جا ارائه نشده است.
- کلیه حقوق معنوی این اثر متعلق به دانشگاه صنعتی شاهرود است و مقالات مستخرج با نام «دانشگاه صنعتی شاهرود» و یا «Shahrood University of Technology» به چاپ خواهد رسید.
- حقوق معنوی تمام افرادی که در به دست آمدن نتایج اصلی پایان نامه تأثیرگذار بوده‌اند در مقالات مستخرج از پایان نامه رعایت می‌گردد.
- در کلیه مراحل انجام این پایان نامه، در مواردی که از موجود زنده (یا بافتهای آن‌ها) استفاده شده است ضوابط و اصول اخلاقی رعایت شده است.
- در کلیه مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته یا استفاده شده است اصل رازداری، ضوابط و اصول اخلاق انسانی رعایت شده است.

امضای دانشجو

تاریخ

مالکیت نتایج و حق نشر

- کلیه حقوق معنوی این اثر و محصولات آن (مقالات مستخرج، کتاب، برنامه‌های رایانه‌ای، نرم‌افزارها و تجهیزات ساخته شده است) متعلق به دانشگاه صنعتی شاهرود است. این مطلب باید به نحو مقتضی در تولیدات علمی مربوطه ذکر شود.
- استفاده از اطلاعات و نتایج موجود در پایان نامه بدون ذکر مرجع مجاز نمی‌باشد.

چکیده:

در دنیای امروز، گسترش رسانه‌ها و شبکه‌های اجتماعی و همچنین افزایش تارنماهای تحلیلی تجاری، خبری، باعث استقبال و حضور گسترده مردم در این مجموعه‌ها شده است. این امر موجب انتشار و اشتراک داده‌ها و نظرات بسیار زیادی می‌شود. با توجه به این حجم انبوه از داده، با تحلیل این داده‌ها، می‌توان اطلاعات بسیار ارزشمندی استخراج نمود. اما در کنار این مزیت، پردازش این حجم از اطلاعات خود یک معضل را پدید می‌آورد. سرعت پردازش روش‌های گذشته و سنتی در مقابل رشد سریع جریان داده‌های ورودی کند می‌باشد. برای حل این مشکل، می‌توان از روش‌های مقیاس‌پذیر و همچنین ابزارهای نوین در حوزه پردازش کلان داده، برای تسریع روند پردازش استفاده نمود. پژوهش‌هایی در این حوزه انجام شده است اما همچنان این حوزه با توجه به پتانسیل بالای خود، نیازمند پژوهش‌های بیشتری می‌باشد. در پژوهش‌های گذشته علاوه بر کندی سرعت پردازش مشکلاتی از جمله پایین بودن دقت، استفاده از ساختار پیچیده روش پیشنهادی و وابستگی به دامنه موضوع وجود دارد.

اسپارک چارچوبی نسبتاً نوین برای پردازش کلان داده‌ها می‌باشد که پردازش را در حافظه انجام می‌دهد و این موضوع باعث افزایش سرعت پردازش الگوریتم می‌شود. در این پژوهش از رده‌بند ماشین بردار پشتیبان (SVM) و ابزار اسپارک به صورت همزمان جهت رده‌بندی نظرات کاربران استفاده شده است. روش پیشنهادی در کنار حفظ دقت مطلوب، وابسته به دامنه خاصی نمی‌باشد و در پیاده‌سازی نیز از پیچیدگی به‌دور است. برای بهره‌مندی از توانایی اسپارک در پردازش توزیع‌شده از رده‌بند SVM موازی به صورت آبشاری، آبشاری بهبودیافته و گروهی و رأی‌گیری بین چند SVM بهره گرفته شده است. برای مقایسه کارایی روش‌های فوق، SVM استاندارد نیز پیاده‌سازی شده است. برای ارزیابی از مجموعه داده IMDB استفاده شده است. این مجموعه داده شامل نظرات مربوط به فیلم است و به زبان انگلیسی می‌باشد. در نتایج بدست آمده برای رده‌بندی، بهترین دقت حاصل شده ۸۵٫۹۴۶٪ است. همچنین زمان سپری شده برای اجرای الگوریتم SVM استاندارد در محیط اسپارک در مقابل پژوهش دیگری که زمان

را گزارش داده است، بیش از ۸ مرتبه کاهش یافته است. زمان برای SVM های موازی در مقابل SVM استاندارد نیز بهبود قابل توجهی داشته است. در بهترین حالت زمان آموزش روش پیشنهادی نسبت به زمان آموزش SVM استاندارد حدود ۶۰ مرتبه کاهش یافته است.

کلمات کلیدی: نظرکاوی، ماشین بردار پشتیبان، پردازش موازی، اسپارک

فهرست مطالب

فصل اول - مقدمه..... ۱

- ۱-۱ مقدمه..... ۲
- ۲-۱ بیان مسئله و تعاریف..... ۲
- ۳-۱ هدف پایان نامه..... ۴
- ۴-۱ روش پیشنهادی..... ۵
- ۵-۱ روش ارزیابی..... ۵
- ۶-۱ ساختار پایان نامه..... ۶

فصل دوم - ادبیات موضوع..... ۷

- ۱-۲ مراحل اولیه نظرکاوی..... ۸
- ۱-۱-۲ جمع آوری داده..... ۸
- ۲-۱-۲ پیش پردازش متن..... ۹
- ۲-۲ روش های نظرکاوی..... ۱۰
- ۱-۲-۲ روش های مبتنی بر یادگیری ماشین..... ۱۰
- ۲-۲-۲ روش های مبتنی بر واژه نامه..... ۱۴
- ۳-۲-۲ روش های ترکیبی..... ۱۶
- ۳-۲ آپاچی اسپارک..... ۱۶
- ۴-۲ ماشین بردار پشتیبان..... ۱۹
- ۱-۴-۲ مقدمه..... ۱۹
- ۲-۴-۲ ماشین بردار پشتیبان..... ۲۰
- ۳-۴-۲ بهینه سازی متوالی کمینه..... ۲۲
- ۴-۴-۲ هسته در ماشین بردار پشتیبان..... ۲۷
- ۵-۲ مشکلات مرتبط با نظرکاوی و چالش ها..... ۲۹
- ۱-۵-۲ مشکلات مربوط به منبع داده..... ۲۹

۲-۵-۲ مشکلات مربوط به پردازش زبان طبیعی ۳۰

فصل سوم - کارهای پیشین..... ۳۳

۱-۳ روش‌های نظر کاوی ۳۴

۱-۱-۳ روش‌های مبتنی بر یادگیری ماشین..... ۳۴

۲-۱-۳ روش‌های مبتنی بر فرهنگ لغت..... ۳۶

۲-۳ مقیاس‌پذیری ۳۶

۳-۳ ماشین بردار پشتیبان موازی..... ۳۸

۱-۳-۳ ماشین بردار پشتیبان آبخاری..... ۴۲

۲-۳-۳ ماشین بردار پشتیبان آبخاری بهبود یافته..... ۴۳

۳-۳-۳ ماشین بردار پشتیبان گروهی..... ۴۳

۴-۳ تجمیع پژوهش‌های پیشین..... ۴۴

فصل چهارم - روش پیشنهادی..... ۴۹

۱-۴ مقدمه..... ۵۰

۲-۴ روش پیشنهادی..... ۵۰

۱-۲-۴ پیش‌پردازش..... ۵۰

۲-۲-۴ آموزش مدل رده‌بند..... ۵۱

۳-۲-۴ رده‌بندی..... ۵۲

۳-۴ آموزش مدل رده‌بندی..... ۵۲

۱-۳-۴ ماشین بردار پشتیبان مبتنی بر رأی‌گیری..... ۵۲

۲-۳-۴ ماشین بردار پشتیبان آبخاری..... ۵۳

۳-۳-۴ ماشین بردار پشتیبان آبخاری بهبود یافته..... ۵۵

۴-۳-۴ ماشین بردار پشتیبان آبخاری با پالایش بردار پشتیبان..... ۵۵

۵-۳-۴ ماشین بردار پشتیبان گروهی..... ۵۵

۶-۳-۴ مرتبه زمانی..... ۵۶

۴-۴	سایر روش‌های انجام شده	۵۷
-----	------------------------	----

فصل پنجم - نتایج و ارزیابی ۵۹

۱-۵	مقدمه	۶۰
۲-۵	مجموعه داده	۶۰
۳-۵	مدل پیاده‌سازی	۶۰
۴-۵	معیارهای ارزیابی	۶۱
۵-۵	نتایج	۶۲
۱-۵-۵	ماشین بردار پشتیبان مبتنی بر رأی‌گیری	۶۲
۲-۵-۵	مقایسه سایر مدل‌ها	۶۵
۳-۵-۵	ارزیابی پارامترهای مدل‌ها	۶۸
۴-۵-۵	مقایسه مدل‌ها	۷۰
۵-۵-۵	پیاده‌سازی موازی در سایر پژوهش‌ها بر روی مجموعه‌داده IMDB	۷۲
۶-۵	نتایج سایر پژوهش‌ها	۷۳
۷-۵	جمع‌بندی	۷۴

فصل ششم - نتیجه‌گیری ۷۷

۱-۶	جمع‌بندی و نتیجه‌گیری	۷۸
۲-۶	کارهای آینده	۷۸
منابع		۷۹

فهرست جدول‌ها

- جدول ۱-۳ : بررسی و تجمیع کارهای مرتبط پیشین ۴۵
- جدول ۱-۴: شرح پارامترهای موجود در روش پیشنهادی ۶۱
- جدول ۲-۴: نتایج مربوط به SVM مبتنی بر رأی‌گیری و استاندارد ۶۲
- جدول ۳-۴: نتایج مربوط به SVM مبتنی بر رأی‌گیری با مقادیر بهینه و متفاوت σ و C ۶۳
- جدول ۴-۴: مقایسه معیار دقت (ACC) مدل‌ها ۶۵
- جدول ۵-۴: مقایسه معیار دقت و زمان SVM آبشاری برای $C = 4, C = 6$ ۶۸
- جدول ۶-۴: مقایسه معیار تشخیص و حساسیت مدل‌های مختلف مبتنی بر SVM ۷۰
- جدول ۷-۴: نتیجه SVM استاندارد پیاده‌سازی شده و مقایسه با سایر مراجع آزمایش شده بر روی IMDB ۷۲
- جدول ۸-۴: نتایج SVM استاندارد روش پیاده‌سازی شده ۷۳
- جدول ۹-۴: نتایج تعدادی از سایر پژوهش‌های انجام شده بر روی IMDB ۷۵

فهرست شکل‌ها

- شکل ۱-۲: روش‌های موجود در تحلیل احساسات و نظرکاوی ۱۱
- شکل ۲-۲: اجزای Apache Spark [۱۵] ۱۸
- شکل ۳-۲: پارامتر C در SVM ۲۵
- شکل ۳-۲: الگوریتم SMO ساده شده [۱۹] ۲۸
- شکل ۵-۲: استفاده از هسته برای جداسازی داده‌ها ۲۹
- شکل ۱-۳: ساختار SVM موازی ارائه شده در مراجع [۳۵, ۳۶] ۴۰
- شکل ۲-۳: ساختار SVM موازی استفاده شده در مرجع [۳۷] ۴۱
- شکل ۳-۳: ساختار SVM موازی ارائه شده در مرجع [۳۸] ۴۲
- شکل ۱-۳: معماری کلی روش پیشنهادی ۵۱
- شکل ۲-۳: معماری SVM مبتنی بر رأی‌گیری ۵۳
- شکل ۳-۳: معماری SVM آبشاری ۵۴
- شکل ۴-۳: معماری SVM آبشاری بهبودیافته ۵۴
- شکل ۵-۳: معماری SVM گروهی ۵۶
- شکل ۱-۴: نتایج معیار دقت روش SVM مبتنی بر رأی‌گیری ۶۴
- شکل ۲-۴: نتایج زمان روش SVM مبتنی بر رأی‌گیری ۶۴
- شکل ۳-۴: مقایسه معیار تشخیص (SPC or TNR) مدل‌ها ۶۶
- شکل ۴-۴: مقایسه معیار حساسیت (TPR) مدل‌ها ۶۶
- شکل ۵-۴: مقایسه معیار زمان آموزش مدل‌ها ۶۷
- شکل ۶-۴: زمان آموزش SVM آبشاری و آبشاری بهبودیافته با مقادیر مختلف حد آستانه برای ضرایب لاگرانژ ۶۹
- شکل ۷-۴: دقت SVM آبشاری و آبشاری بهبودیافته با مقادیر مختلف حد آستانه برای ضرایب لاگرانژ

۶۹.....

شکل ۴-۸: مقایسه معیار دقت در مدل‌های مختلف مبتنی بر SVM ۷۱

شکل ۴-۹: مقایسه معیار زمان آموزش در مدل‌های مختلف مبتنی بر SVM ۷۱

فهرست واژگان و اختصارات

Accuracy	ACC	دقت
Conventional Neural Network	CNN	شبکه عصبی کانولوشن
International Movie Database	IMDB	پایگاه داده بین المللی فیلم
Karash-Kuhn-Tucker	KKT	کاراش-کوهن-تاگر
Logistic Regression	LR	رگرسیون منطقی
Maximum Entropy	ME	بیشینه آنتروپی
Naïve Bayes	NB	بیز ساده
Natural Language Proccesing	NLP	پردازش زبان طبیعی
Opinion Mining	OM	نظر کاوی
Part of Speech	POS	اقسام جمله
Quadratic Programming	QP	برنامه ریزی مرتبه دوم
Radial Basic Function	RBF	تابع پایه شعاعی
Resilient Distributed Dataset	RDD	مجموعه داده توزیع شده انعطاف پذیر
Sensitivity (True Positive Rate)	TPR	نرخ پیش بینی مثبت صحیح
Sentiment Analysis	SA	تحلیل احساسات
Sequential Minimal Optimization	SMO	بهینه سازی متوالی کمینه
Stochastic Gradient Descent	SGD	گرادیان نزولی تصادفی
Support Vector	SV	بردار پشتیبان
Support Vector Machine	SVM	ماشین بردار پشتیبان
Specificity (True Negative Rate)	TNR	نرخ پیش بینی منفی صحیح

فصل اول

مقدمه

۱-۱ مقدمه

امروزه با گسترش اینترنت، روزانه اطلاعات زیادی در این فضا ایجاد می‌شود که پردازش این حجم انبوه از اطلاعات نیازمند ابزارها و روش‌هایی می‌باشد که در پردازش، سریع و در عملکرد، دقیق هستند. با توسعه طراحی و معماری وب ۲/۰ تعامل کاربران با فضای اینترنت تسهیل شد و باعث ظهور شبکه‌های اجتماعی^۱، وب‌نوشت‌های کوچک^۲، انجمن‌ها^۳ و سایت‌های تحلیلی مختلفی گردید. روزانه نظرات و دیدگاه‌های بسیاری در موضوعات مختلف، توسط کاربران در این محیط‌ها ارسال می‌شود. این نظرات می‌تواند حاوی اطلاعات ارزشمندی برای دولت‌ها، سازمان‌ها، شرکت‌ها و مردم باشد. بعنوان نمونه از گذشته تاکنون، افراد جهت خرید یک محصول یا سرویس، از نظرات دیگران برای تصمیم‌گیری استفاده می‌نمایند. لذا با بررسی نظرات دیگران می‌تواند تصمیم بهتری در این مورد اخذ نمایند. سازمان‌ها و شرکت‌ها نیز با بررسی همین نظرات می‌توانند نقاط قوت و ضعف محصولات و خدمات خود را شناخته و با سلیقه مردم آشنا شوند و در نقشه توسعه خود از آن بهره ببرند. دولت‌ها نیز می‌توانند در زمینه‌های مختلف با بررسی نظرات مردم تفکر کلی جامعه را بدون استفاده از نظرسنجی رصد کنند و این موضوع در تصمیم‌گیری‌ها حائز اهمیت خواهد بود. همین موضوع باعث ایجاد شاخه‌ای جدید به نام نظرکاوی^۴ در متن‌کاوی^۵ شده است.

۱-۲ بیان مسئله و تعاریف

روزانه حجم عظیمی از نظرات کاربران در شبکه‌های اجتماعی مختلف، انجمن‌ها و وب‌نوشت‌های کوچک منتشر می‌شود و به همین دلیل بررسی انسانی آن عملاً غیر ممکن است. با توجه به پتانسیل و اهمیت این موضوع، نظرکاوی به صورت خودکار مورد نیاز می‌باشد. تاکنون پژوهش‌هایی در این زمینه انجام شده

¹ Social Networks

² Microblogs

³ Forums

⁴ Opinion Mining

⁵ Text Mining

است، اما همچنان ارائه روش‌هایی با سرعت و کارایی بیشتر مورد نیاز می‌باشد.

در ادامه نیز تعاریف نظر، نظرکاوی و همچنین سطوح نظرکاوی آورده شده است.

نظر: بر اساس فرهنگ لغت آکسفورد «دیدگاه یا قضاوت شکل گرفته در مورد یک چیز، که لزوماً براساس واقعیت یا دانش نمی‌باشد»^۱. همچنین بینگ لیو این تعریف را ارائه داده است: «نظر، احساسات نظردهنده در مورد یک وجود^۲ (محصول یا شیء) یا ویژگی‌های یک وجود می‌باشد»^۳ [۱].

جمله ذهنی:^۳ این نوع جملات شامل نظر و احساس نظردهنده در مورد وجود و یا ویژگی‌های آن می‌باشد. مثال: «باتری گوشی‌های برند سونی دارای طول عمر بیشتری نسبت به سایر برندها می‌باشد.»

جمله عینی:^۴ این نوع جملات همراه با احساس نمی‌باشد و بیشتر شامل یک حقیقت یا جمله خبری می‌باشند. مثال: «گوشی جدید برند سونی امروز معرفی شد.»

نظرکاوی: نظرکاوی (و یا تجزیه و تحلیل احساس^۵) زیرشاخه‌ای از متن‌کاوی و بخشی از حوزه پردازش زبان طبیعی^۶ می‌باشد که به تشخیص و استخراج اطلاعات ذهنی^۷ از نظر می‌پردازد و قطبیت^۸ نظر شامل مثبت و منفی، یا حالت سوم یعنی طبیعی (نرمال) را تعیین می‌نماید.

نظرکاوی در سه سطح مختلف انجام می‌شود که در ادامه به آن اشاره شده است:

سطح سند:^۹ در این سطح به تجزیه و تحلیل کل سند و متن پرداخته می‌شود و قطبیت آن شامل مثبت، منفی و یا حالت طبیعی تعیین می‌شود. چالش مهم این قسمت تعیین جملات و قسمت‌های ذهنی متن می‌باشد. در این سطح معمولاً قطبیت کل سند در مورد یک ویژگی بیان می‌شود و این یکی

^۱ <https://en.oxforddictionaries.com/definition/opinion>

^۲ Entity

^۳ Subjectivity

^۴ Objectivity

^۵ Sentiment Analysis

^۶ Natural Language Processing (NLP)

^۷ Subjective Information

^۸ Polarity

^۹ Document level

از مشکلات آن می‌باشد.

سطح جمله^۱: در این سطح از پردازش، متن به تک تک جملات تقسیم‌بندی شده و قطبیت هر جمله به طور جداگانه تعیین می‌شود. در این سطح، در مرحله پیش پردازش جملات عینی یا ذهنی^۲ تفکیک شده و پس از آن قطبیت جملات ذهنی تعیین می‌گردد. این سطح نسبت به سطح سند به طور جزئی‌تر به تعیین قطبیت می‌پردازد. شناسایی قطبیت در این سطح از دو روش تعیین می‌گردد: رویکرد نحوی دستوری^۳ و رویکرد معنایی^۴. در رویکرد نحوی دستوری با استفاده از ساختار دستوری جمله مانند اسم و صفت و ... قطبیت جمله تعیین می‌گردد، اما در رویکرد معنایی با استفاده از تعداد تکرار کلمات با قطبیت مثبت و منفی در جمله، قطبیت کلی جمله مشخص می‌گردد.

سطح جنبه یا ویژگی^۵: در این سطح نیز ویژگی‌های وجود استخراج و همچنین کلمات و یا عبارات نظری^۶ مربوط به هر ویژگی استخراج می‌شود و در نهایت قطبیت نظر نسبت به هر ویژگی تعیین می‌شود.

۱-۳ هدف پایان‌نامه

با توجه به حجم عظیم نظرات منتشر شده در فضای اینترنت، پردازش آن نیازمند ابزار و روشی با سرعت پردازش و دقت عملکرد بالا می‌باشد. روش‌های ارائه شده پیشین نیز سعی در بهبود سرعت و دقت نظرکاوی داشته‌اند اما همچنان از معضل زمان زیاد پردازش و کاهش دقت رنج می‌برند. در برخی از پژوهش‌ها به دلیل استفاده از روش‌های پیچیده و دارای محاسبات زیاد، به دقت بالایی دست یافته‌اند اما همین موضوع سبب زمان‌بر شدن روش پیشنهادی شده است. در برخی دیگر از پژوهش‌ها سرعت الگوریتم پیشنهادی بالاست و زمان کمی برای اجرا مصرف می‌نماید اما در زمینه دقت عملکرد، قابل

¹ Sentence Level

² Objective or Subjective Sentences

³ Grammatical Syntactic Approach

⁴ Semantic Approach

⁵ Aspect or Feature Level

⁶ Opinion Words or Phrases

قبول نمی‌باشند و در این مورد ضعف دارند. لذا با این تفاسیر، فقدان روشی با دقت مطلوب در کنار هزینه زمانی معقول و کم، محسوس است. با وجود این محدودیت‌ها، با بهره‌مندی هوشمندانه از امکانات در دسترس می‌توان بر این مشکلات تا حدودی غلبه کرد.

استفاده مناسب از ابزارهای پردازش کلان‌داده و ارائه روش‌های مقیاس‌پذیر^۱ می‌تواند در هزینه زمانی روش پیشنهادی موثر باشد و باعث بهبود و کاهش زمان اجرای الگوریتم شود. همچنین استفاده از الگوریتم‌های مناسب حوزه موردنظر که دارای عملکرد دقیق می‌باشند و سازگار با پردازش موازی هستند نیز می‌تواند دقت روش پیشنهادی را تضمین نماید.

در این پژوهش روشی مقیاس‌پذیر ارائه شده است و از ابزار پردازش کلان‌داده اسپارک استفاده شده است. این امر باعث کاهش زمان آموزش و بهبود در سرعت تشخیص قطبیت نظر می‌شود. همچنین در این پژوهش با بهره‌گیری از الگوریتم و رده‌بند SVM دقت بدست آمده مطلوب و قابل قبول می‌باشد.

۴-۱ روش پیشنهادی

در این پژوهش، روشی مقیاس‌پذیر برای رده‌بندی نظرات کاربران ارائه شده است. به این منظور دو روش SVM مبتنی بر رأی‌گیری و SVM با پالایش بردارهای پشتیبان پیشنهاد شده و در محیط اسپارک پیاده‌سازی شده است. همچنین برخی از روش‌های موجود مانند SVM آبشاری، SVM آبشاری بهبودیافته و SVM گروهی که در محیطی خارج از اسپارک پیاده‌سازی شده‌اند و بعضاً برای کاربردهایی خارج از حوزه نظرکاوی پیشنهاد شده بودند، در محیط مذکور پیاده‌سازی و برای رده‌بندی نظرات کاربران مورد استفاده قرار گرفت.

۵-۱ روش ارزیابی

برای ارزیابی نتایج مدل‌ها با یکدیگر و همچنین با سایر پژوهش‌ها از چهار معیار استفاده شده است.

^۱ Scalable

اولین معیار، زمان آموزش و آزمون مدل رده‌بند می‌باشد که بر حسب ثانیه است. معیار دیگر نیز دقت^۱، تشخیص^۲ و حساسیت^۳ می‌باشد که در فصل چهارم به طور کامل معرفی شده‌اند.

۱-۶ ساختار پایان‌نامه

این پایان‌نامه در شش فصل تدوین گردیده است. در ادامه، فصل دوم شامل ادبیات موضوع می‌باشد و فصل سوم مربوط به پژوهش‌های گذشته می‌باشد. در فصل چهارم روش ارائه شده در این پایان‌نامه توضیح داده شده است. نتایج و ارزیابی روش موردنظر در فصل پنجم آورده شده است. فصل ششم نیز شامل پیشنهادات و نتیجه‌گیری می‌باشد.

¹ Accuracy

² Specificity

³ Sensivity

فصل دوم

ادبیات موضوع

۱-۲ مراحل اولیه نظر کاوی

روند نظر کاوی شامل مراحل مختلفی می شود که از مراحل مقدماتی آن می تواند به مرحله جمع آوری داده و پیش پردازش متن اشاره کرد. در ادامه به توضیح بیشتر این مراحل می پردازیم.

۱-۱-۲ جمع آوری داده^۱

اولین مرحله و قدم برای نظر کاوی ایجاد مجموعه داده برای این عمل می باشد. دو روش برای این عمل وجود دارد:

- استفاده از رابط برنامه نویسی کاربردی^۲ تارنما^۳هایی نظیر آمازون^۴، توییتر^۵، فیسبوک^۶ برای جمع آوری داده، که بعنوان مثال Twitter4J API^۷ برای توییتر می باشد و از دیگر API های شناخته شده برای توییتر می توان [۴-۲] را نام برد.

- ایجاد و استفاده از خزنده های وب^۸ که در وبسایت ها به جمع آوری نظرات و داده می پردازد. الستون و ناجورک یک بررسی دقیق و قوی روی خزنده های وب انجام داده اند [۵].

هر کدام از روش های فوق دارای مزایا و معایبی می باشند، برای مثال رابط های برنامه نویسی کاربردی دارای محدودیت هایی از طرف میزبان در هنگام جمع آوری داده می باشند و در مقابل خزنده ها نیز دارای مشکلاتی در پیاده سازی می باشند. همچنین تعدادی از پژوهش ها از برخی مجموعه داده های معرفی شده و استاندارد استفاده می کنند که این موضوع باعث ارزیابی بهتر روش ارائه شده نسبت به سایر پژوهش های پیشین می شود. کایجی گئو^۹ و همکاران [۶] به طور خلاصه به مقایسه این دو روش پرداخته اند.

¹ Data Acquisition

² Application Programming Interface

³ Website

⁴ Amazon

⁵ Twitter

⁶ Facebook

⁷ <http://twitter4j.org/en/index.html>

⁸ Web Crawlers

⁹ Kaijie Guo

۲-۱-۲ پیش پردازش متن

مرحله دیگر پیش پردازش متن می باشد که با استفاده از عمل های معمول پردازش زبان طبیعی و استفاده از تجزیه و تحلیل واژه ها انجام می شود. مهمترین تکنیک های معمول در ادامه اشاره شده است:

تقسیم بندی کلمات^۱: به عمل جداسازی و تفکیک جملات یک سند به لیستی از کلمات مجزا گفته می شود. این تکنیک برای زبان هایی که کلمات با استفاده از فاصله جدا می شوند، به راحتی قابل پیاده سازی است اما برای زبان هایی مانند چینی و ژاپنی که کلمات با فاصله از هم جدا نمی شوند با مشکلاتی روبرو است.

ریشه یابی^۲: به فرآیند حذف پسوند و پیشوندهای کلمات و به طور کلی تبدیل آن ها به شکل ریشه و اصلی آن، ریشه یابی گفته می شود. بعنوان مثال کلماتی نظیر شخص، اشخاص، شخصیت، شخصی، همگی پس از ریشه یابی به کلمه شخص تبدیل می شوند. یکی از معروفترین ریشه یاب ها برای زبان انگلیسی Porter's stemmer می باشد [۷].

حذف ایست واژه ها^۳: این تکنیک شامل حذف ایست واژه ها می باشد که جزء ساختار زبانی می باشند اما کمکی به تحلیل نظر نمی کنند. و، یا، با، تا، از، به، است و تعدادی از این نوع کلمات می باشند.

تقسیم بندی جمله^۴: این تکنیک برای تفکیک جملات از یکدیگر استفاده می شود. این روش مشکلات مربوط به خود را دارد. یکی از راه های جداسازی جملات از یکدیگر استفاده از نقطه انتهای جمله می باشد، اما اعداد اعشاری و همچنین کلمات مخفف شده باعث ایجاد مشکل در این حالت می شوند.

برچسب گذاری اقسام جمله^۵: این روش شامل برچسب گذاری نقش هر کلمه در جمله مانند اسم،

^۱ Word Segmentation

^۲ Stemming

^۳ Stopword Removal

^۴ Sentence Segmentation

^۵ Part-of-Speech (POS) Tagging

فعل، صفت، قید، حرف اضافه و... می‌باشد. از این روش در پردازش‌های مراحل بعدی استفاده می‌شود. برای مثال می‌توان از آن به عنوان ویژگی ورودی در روش‌های مبتنی بر یادگیری ماشین^۱ استفاده کرد. پیش‌پردازش‌های دیگری نیز برای حالت‌های خاص انجام می‌گیرد. برای مثال هنگامی که از خزنده‌های وب برای جمع‌آوری داده استفاده می‌شود، مراحل برای حذف لینک‌ها و همچنین موارد غیرممتنی (مانند عکس‌ها و...) انجام می‌شود و یا برای مطالب موجود در توییتر نیز پیش‌پردازش برای حذف هشتگ^۲‌ها، اشاره^۳‌ها، شکلک^۴‌ها، کلمات مختصرشده و خنده‌های نوشته شده می‌باشد.

۲-۲ روش‌های نظر کاوی

روش‌های نظر کاوی دارای دو قسمت اصلی می‌باشند: روش‌های مبتنی بر یادگیری ماشین و روش‌های مبتنی بر واژه نامه^۵ که در ادامه به آن‌ها پرداخته شده است. در شکل ۱-۲ شمایی کلی از روش‌های موجود نشان داده شده است.

۱-۲-۲ روش‌های مبتنی بر یادگیری ماشین

این دسته از روش‌ها، در بخش روش‌های با ناظر قرار می‌گیرند و در بعضی از منابع با این عنوان نیز مطرح می‌گردند. همان‌طور که انسان‌ها از تجربیات گذشته خود یاد می‌گیرند، رایانه‌ها نیز می‌توانند از داده‌ها آموزش ببینند که این داده‌ها مربوط به کاربردهای گذشته می‌باشد. در این روش نیاز به دو دسته از داده‌ها می‌باشد، مجموعه داده‌های آموزش که دارای برچسب کلاس می‌باشند و از آن برای ایجاد و آموزش رده‌بند^۶ استفاده می‌شود و مجموعه داده‌های آموزش که با توجه به رده‌بند آموزش دیده، رده‌بندی روی آن‌ها انجام می‌شود. نکته‌ای که در این روش‌ها حائز اهمیت است، ویژگی‌هایی می‌باشد که از داده‌های مجموعه آموزش استخراج می‌شود تا برای یادگیری سیستم از آن‌ها استفاده شود.

¹ Machine Learning Based Methods

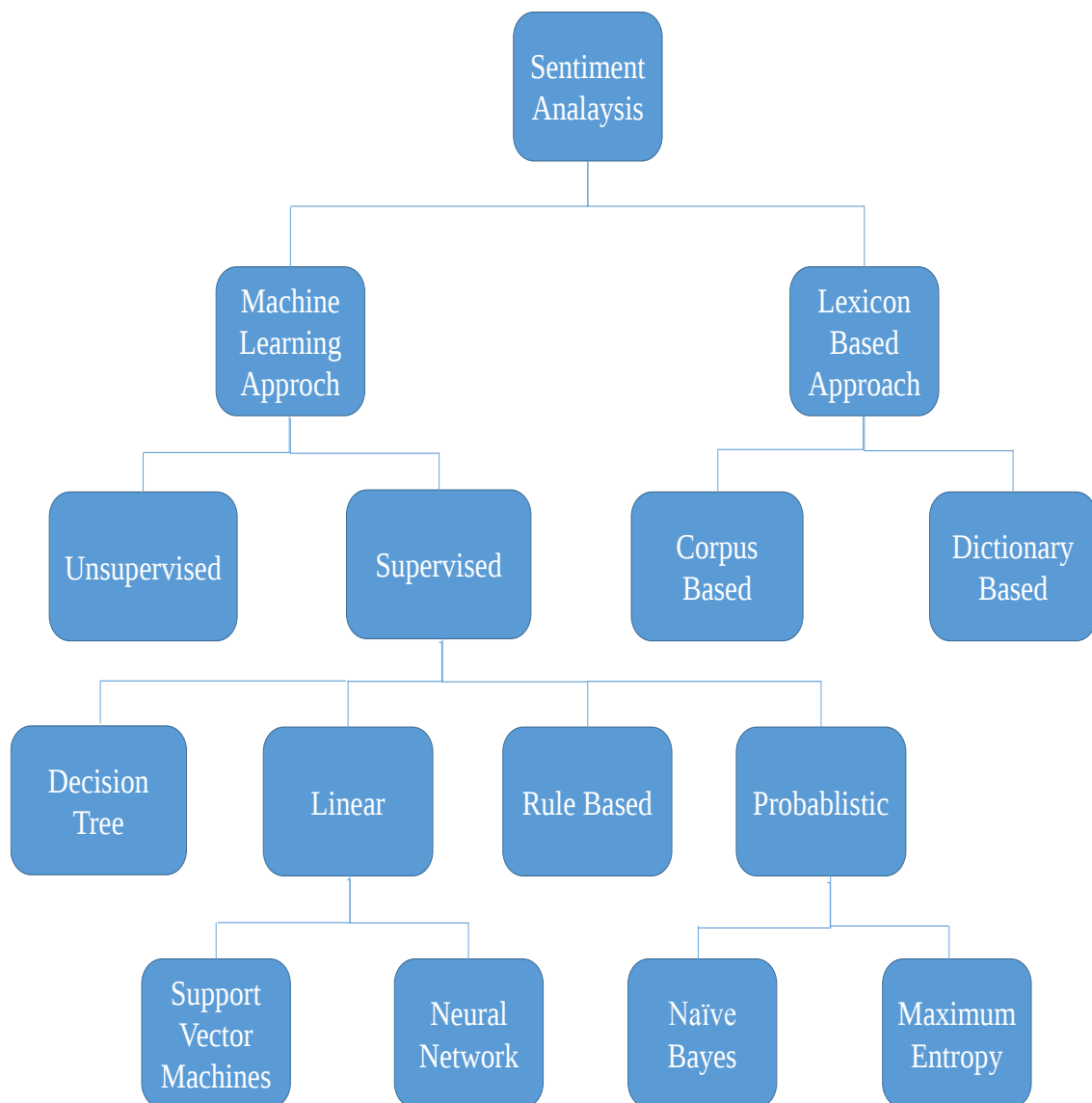
² Hashtag

³ Mention

⁴ Emoji

⁵ Lexicon Based Methods

⁶ Classifier



شکل ۱-۲: روش‌های موجود در تحلیل احساسات و نظرکاوی

در ادامه به تعدادی از مهمترین ویژگی‌ها مورد استفاده، اشاره می‌شود. پس از آن نیز تعدادی از روش‌های مبتنی بر یادگیری ماشین معرفی می‌شود.

وجود و تعداد تکرار کلمات^۱: این ویژگی به کلمات موجود در متن و یا کلمات n-gram متن توجه و به بررسی حضور و تعداد تکرار آن می‌پردازد. این ویژگی بسیار زیاد در زمینه نظرکاوی استفاده می‌شود و کاملاً موثر واقع شده است. در پژوهشی، بو پانگ^۲ و همکاران [۸] از n-gram برای بررسی نظرات

^۱ Terms presence and frequency

^۲ Bo Pang

مربوط به فیلم استفاده شده است و باعث افزایش بازده روش شده است. دیو و همکارانش نیز [۹] از bi-gram و tri-gram برای بررسی نظرات یک محصول استفاده نموده‌اند.

اطلاعات مربوط اقسام جمله^۱: این ویژگی نیز در تحلیل احساسات و نظرکاوی بسیار مهم می‌باشد. با استفاده از برچسب‌گذاری، به شناسایی نقش کلمات در جمله می‌پردازیم. کلماتی مانند صفت‌ها و قیدها تاثیر بسیاری زیادی در تجزیه و تحلیل احساسات صاحب‌نظر دارند.

نفی^۲: کلمات نفی (مانند نا، پیشوند ن و ...) نیز در نظرکاوی و تحلیل احساسات بسیار مهم است. زیرا این نوع از کلمات قطبیت جمله را معکوس می‌نمایند.

عبارات و کلمات نظری^۳: کلمات و عبارات نظری، کلماتی هستند که قطبیت جمله و احساس کاربر را بیان می‌کنند. اکثرا این کلمات شامل صفت‌ها و قیدها می‌شوند اما بعضی از کلمات در نقش اسم مانند آشغال و یا به صورت فعل مانند دوست داشتن و متنفر بودن در متن استفاده می‌شوند.

روش‌های یادگیری ماشین شامل روش‌های زیادی می‌باشد که از معروف‌ترین و کاربردی‌ترین آن‌ها بخصوص در نظرکاوی می‌توان به بیز ساده، ماشین بردار پشتیبان، بیشینه آنتروپی اشاره کرد. در ادامه به معرفی جزئی‌تر این برخی از این روش‌ها می‌پردازیم.

۲-۱-۲-۲ رده‌بندی کننده بیز ساده^۴:

رده‌بند بیز ساده، معمول‌ترین و ساده‌ترین روش مورد استفاده در رده‌بندی متن می‌باشد. این روش مبتنی بر محاسبه احتمال پیشین یک کلاس می‌باشد و بر مبنای توزیع کلمات در متن می‌باشد. این تکنیک به نقش و مکان کلمات در جمله توجهی نمی‌کند. در این روش به طور ساده و خلاصه، ویژگی-های مربوط به یک کلاس را در نظر گرفته و پس از آن با بررسی احتمال تعلق یک متن به همراه

¹ Part of speech information

² Negation

³ Opinion Words and Phrases

⁴ Naive Bayes Classifier (NB)

ویژگی‌هایش در هر کلاس، رده‌بندی را انجام می‌دهد.

۲-۱-۲-۲ ماشین بردار پشتیبان^۱:

این روش یکی از موثرترین و کارآترین روش‌ها برای رده‌بندی متن بوده و می‌باشد. ایده اصلی این روش، تعیین یک جداساز خطی برای تفکیک هرچه بهتر فضای جستجو می‌باشد. به طور خلاصه در این روش، هدف ایجاد یک ابرصفحه (بردار) می‌باشد که دارای بیشینه اطمینان برای تفکیک یک کلاس از سایر کلاس‌ها می‌باشد. داده‌های متنی برای این روش بسیار ایده‌آل می‌باشند و این موضوع به دلیل ماهیت پراکندگی (ویژگی‌های) متن می‌باشد که جداسازی را تقریباً تسهیل می‌نماید. در فصل چهارم، به توضیح بیشتر این روش که مدنظر ماست خواهیم پرداخت.

۳-۱-۲-۲ بیشینه آنتروپی^۲:

این روش نیز از روش‌های کارآمد در زمینه پردازش زبان طبیعی می‌باشد و در اکثر موارد بهتر از روش بیز ساده عمل می‌نماید. در این تکنیک با استفاده از کدگذاری، مجموعه ویژگی‌ها به یک بردار تبدیل می‌شود. سپس از این بردار کدگذاری شده برای تعیین وزن هر ویژگی استفاده می‌شود و در نهایت با توجه به وزن و اهمیت ویژگی‌ها برای کلاس‌ها، رده‌بندی صورت می‌گیرد.

۴-۱-۲-۲ درخت تصمیم^۳:

درخت تصمیم، یک تقسیم‌بندی سلسله مراتبی را از مجموعه داده‌های آموزش با توجه به محدودیت‌ها و شرایط خاصی انجام می‌دهد. این شرایط می‌تواند وجود یا عدم وجود یک یا چند کلمه باشد. این تجزیه و تقسیم به صورت بازگشتی انجام می‌شود تا زمانی که گره‌های برگ به حداقل تعداد مشخصی داده (رکورد) برسد.

^۱ Support Vector Machine (SVM)

^۲ Maximum Entropy (ME)

^۳ Decision Tree

۲-۲-۵ رده‌بندی کننده‌های مبتنی بر قاعده^۱:

در این روش فضای داده با تعدادی قوانین مدل می‌شود. در سمت چپ قوانین، ویژگی‌ها قرار دارند و در سمت راست آن‌ها، برچسب کلاس قرار دارد. قوانین با استفاده از معیارهای زیادی تولید می‌شود که مرحله ساخت و آموزش قوانین به همین معیارها بستگی دارد. دو نوع از مهمترین این معیارها، پشتیبان^۲ و اطمینان^۳ می‌باشد [۱۰].

۲-۲-۶ شبکه عصبی مصنوعی^۴:

در این روش با استفاده از شبکه عصبی مصنوعی، رده‌بندی انجام می‌شود. شبکه عصبی مصنوعی از تعدادی نرون^۵ تشکیل شده که هر نرون دارای یک وزن مشخص می‌باشد. این شبکه و نرون‌ها با استفاده از مجموعه داده آموزش، آموزش می‌بینند و وزن نرون‌ها تعیین می‌شود. پس از آن با اعمال داده‌های آموزش، شبکه اقدام به رده‌بندی داده‌ها می‌نماید.

۲-۲-۲ روش‌های مبتنی بر واژه‌نامه

در بسیاری از موارد ایجاد یک مجموعه داده برچسب‌گذاری شده برای آموزش مشکل می‌باشد و در مقابل ایجاد یک مجموعه داده بدون برچسب نسبتاً آسان‌تر می‌باشد. بنابراین برای غلبه بر این مشکل روش‌های مبتنی بر واژه‌نامه ایجاد و مورد استفاده قرار گرفتند. در این روش‌ها با استفاده از کلمات و عبارات نظری موجود در متن به تعیین قطبیت هر جمله و متن پرداخته می‌شود که نیازی به داشتن مجموعه داده برچسب‌گذاری شده پیشین نمی‌باشد. مزیت این روش‌ها همان‌طور که اشاره شد عدم نیاز به مجموعه آموزشی می‌باشد و همچنین یکی از مهمترین معایب آن مشکل در ایجاد واژه‌نامه می‌باشد که یا باید به صورت انسانی باشد که بسیار وقت‌گیر و هزینه‌بر می‌باشد و یا باید به صورت اتوماتیک و با استفاده از روش‌های مختلف انجام شود که آن‌ها نیز مشکلات مربوط به خود را دارا می‌باشند. روش‌های

¹ Rule Based Classifier

² Support

³ Confidence

⁴ Artificial Neural Network (ANN)

⁵ Neuron

مبتنی بر واژه نامه به دو دسته‌ی روش‌های مبتنی بر فرهنگ لغت^۱ و روش‌های مبتنی بر پیکره^۲ تقسیم می‌شوند.

۲-۲-۲-۱ روش‌های مبتنی بر فرهنگ لغت:

در این روش از یک فرهنگ لغت از پیش تعیین شده که شامل عبارات و کلمات نظری می‌باشد استفاده می‌شود که قطبیت هر کلمه (مثبت یا منفی) و شدت این قطبیت مشخص است. برای تعیین قطبیت یک جمله با این روش، ابتدا کلمات و عبارات نظری جمله استخراج و با توجه قطبیت آن‌ها در فرهنگ-لغت، قطبیت کل جمله تعیین می‌شود. همچنین در این روش از مترادف و متضاد کلمات نیز برای تعیین احساس استفاده می‌شود. فرهنگ لغت ابتدا با تعداد محدودی کلمه ایجاد می‌شود و در ادامه با توجه به مترادف و متضاد کلمات گسترش می‌یابد. مترادف و متضاد کلمات نیز با استفاده از وردنت^۳ [۱۱] قابل دستیابی می‌باشد. از معایب این روش می‌توان به ناتوانی آن در تعیین قطبیت کلمات در زمینه‌های خاص اشاره کرد و از مزایای آن عدم نیاز به دانش قبلی (مجموعه داده آموزش) با جایگزین کردن فرهنگ لغت اشاره نمود.

۲-۲-۲-۲ روش‌های مبتنی بر پیکره:

در این روش یک پیکره عظیم از متون و کلمات مربوط به زمینه‌ای خاص ایجاد می‌شود و پس از آن الگوها و کلمات نظری متن و جمله مورد بررسی، در آن جستجو و با توجه به آن، قطبیت جمله تعیین می‌گردد. این جستجو و تحلیل احساس می‌تواند برحسب روش‌های آماری و یا استفاده از ساختار معنایی و نحوی الگو باشد. برای نمونه و تاثیر ساختار نحوی، هنگامی که دو کلمه (مثلاً صفت) با حرف اضافه «و» در یک الگو قرار می‌گیرند، در اکثر موارد دارای قطبیت یکسان می‌باشند و در هنگامی که از حرف اضافه «اما» استفاده می‌شود، قطبیت جمله تغییر می‌نماید. از جمله نقاط قوت این روش، تعیین قطبیت

^۱ Dictionary Based Approach

^۲ Corpus Based Approach

^۳ WordNet

متون و جملات در زمینه‌های خاص می‌باشد که به طور صحیح این عمل را انجام می‌دهد و از جمله معایب آن نیز می‌توان به مشکل بودن ایجاد یک تنه عظیم متن در زمینه‌ای خاص اشاره کرد.

۳-۲-۲ روش‌های ترکیبی^۱

در این روش از ترکیبی از هر دو روش مبتنی بر یادگیری ماشین و روش مبتنی بر فرهنگ‌گفت استفاده می‌شود. در [۱۲-۱۴] از این نوع روش استفاده شده است.

۳-۲ آپاچی اسپارک

اسپارک یک چارچوب پردازش داده‌های بزرگ به صورت توزیع شده و متن‌باز است که اولین بار در سال ۲۰۰۹ به عنوان یک پروژه در آزمایشگاه ای ام پی^۲ دانشگاه کالیفرنیا، برکلی^۳ طراحی شد و بعدها به بنیاد نرم‌افزاری آپاچی^۴ هدیه گردید. در سال ۲۰۱۳ به یک پروژه انحصاری از بنیاد نرم‌افزاری آپاچی تبدیل شد و در اوایل سال ۲۰۱۴ میلادی به یکی از پروژه‌های برتر بنیاد ارتقا یافت.

یکی از ویژگی‌های اصلی و بزرگترین برتری اسپارک حفظ مجموعه داده در حین اجرای پردازش‌ها در حافظه می‌باشد، همین مورد بازدهی اسپارک را در مقابل سایر چارچوب‌ها مانند آپاچی هدوپ^۵ افزایش می‌دهد. عبارتی دیگر اسپارک به گونه‌ای طراحی و بهینه شده است که عملیات پردازش را در حافظه انجام دهد و این مزیت نسبت به برنامه نگاشت-کاهش که داده‌ها را بر روی دیسک سخت^۶ نوشته و از روی دیسک سخت نیز برای پردازش فراخوانی می‌کند، موجب سرعت فوق‌العاده بالاتری شده است. در آپاچی هدوپ مدل نگاشت-کاهش بر روی دیسک سخت انجام می‌گیرد. در نتیجه، همان‌طور که اشاره شد مقداری از زمان صرف خواندن و نوشتن اطلاعات بر روی دیسک سخت و انتقال آن به حافظه یا

¹ Hybrid Methods

² AMPLab

³ University of California, Berkeley

⁴ Apache Software Foundation

⁵ Apache Hadoop

⁶ Hard Disk

بالعکس می‌شود.

اسپارک دارای اجزا و مولفه‌های مختلفی است که در شکل ۲-۲ مشخص می‌باشد. از دیگر مزایای اسپارک انسجام اجزای آن است. این مزیت، توانایی نوشتن برنامه‌هایی را می‌دهد که در آن از مدل‌های مختلف پردازش به طور یکپارچه استفاده می‌شود. برای مثال در اسپارک می‌توان برنامه‌ای نوشت که رده‌بندی داده‌ها از طریق یادگیری ماشین صورت می‌پذیرد و داده‌های ورودی نیز به صورت همزمان از منابع جریانی^۱ دریافت می‌شوند. در آپاچی هدوپ برای هر کدام از این مولفه‌ها موتوری مجزا تعریف شده است که یکپارچه‌سازی و برنامه‌نویسی را مشکل‌تر می‌نماید. در ادامه به توضیح مختصر این اجزا پرداخته می‌شود.

هسته اسپارک^۲: مسئولیت هسته اسپارک انجام عملیات‌های اولیه اسپارک همانند مدیریت حافظه، زمان‌بندی وظایف، مقابله و رفع خطا و همچنین تعامل سیستم ذخیره‌سازی با دیگر اجزای اسپارک می‌باشد. مفهوم اصلی برنامه‌نویسی اسپارک یعنی مجموعه داده توزیع شده انعطاف‌پذیر^۳ در هسته تعریف می‌شوند. RDDها مجموعه داده‌هایی هستند که توانایی اجرای توزیع شده بر روی چندین گره^۴ محاسباتی متعدد را دارند و بدین ترتیب می‌توان آن‌ها را به صورت موازی پردازش نمود.

مولفه Spark SQL: این مولفه یک بخش بر روی هسته اسپارک می‌باشد و پشتیبانی از داده‌های ساخت‌یافته^۵ و پرس و جو روی آن‌ها را فراهم می‌آورد.

مولفه Spark Streaming: مولفه پردازش داده‌های جریانی اسپارک یکی از اجزای اسپارک است که

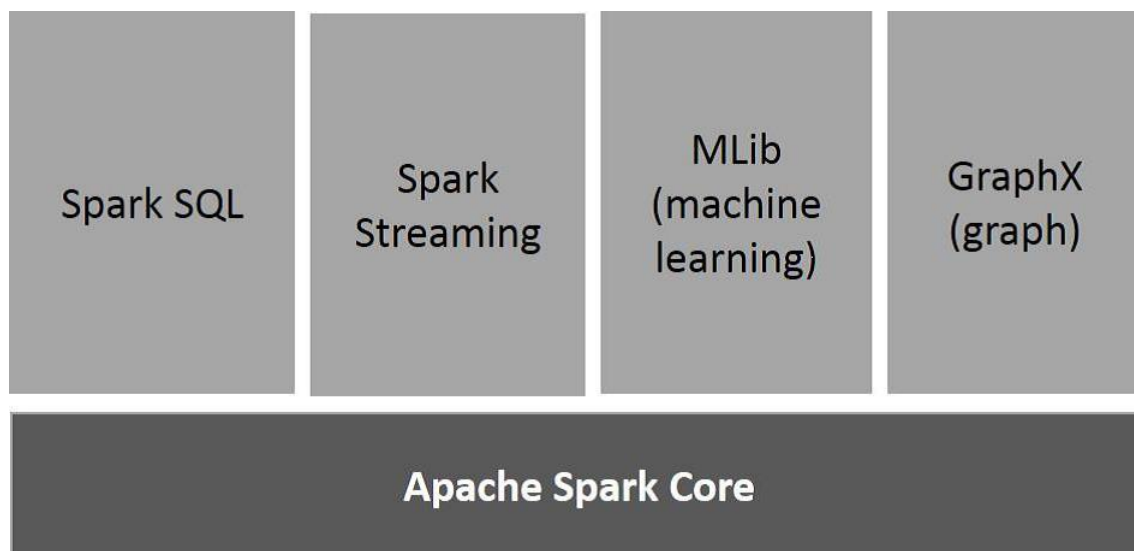
^۱ Streaming

^۲ Apache Spark Core

^۳ Resilient Distributed Dataset (RDD)

^۴ Node

^۵ Structured-Data



شکل ۲-۲: اجزای Apache Spark [۱۵]

از قابلیت زمان‌بندی سریع هسته‌های اسپارک برای فراهم کردن آنالیزهای جریانی^۱ استفاده می‌کند و پردازش جریان داده‌ها را فراهم می‌آورد. یکی دیگر از دلایل محبوبیت اسپارک همین مولفه می‌باشد که امکان پردازش همزمان جریان داده همراه با ورود آن را در اختیار قرار می‌دهد.

مولفه MLlib: این مولفه همان‌طور که از نام آن مشخص است، کتابخانه یادگیری ماشین، برای اسپارک است و توانایی پردازش الگوریتم‌های مختلف داده کاوی و یادگیری ماشین را بر روی بستر اسپارک فراهم می‌آورد. MLlib انواع مختلفی از الگوریتم‌های یادگیری ماشین از جمله رده‌بندی، رگرسیون^۲، خوشه‌بندی^۳ را ارائه می‌دهد و همچنین از قابلیت‌های مثل ارزیابی مدل و ورود داده‌ها پشتیبانی می‌کند.

مولفه GraphX: GraphX، یک مولفه و کتابخانه دیگر از اسپارک است. این کتابخانه، توانایی پردازش و تحلیل داده‌های گراف (نظیر گراف دوستی شبکه‌های اجتماعی) و یا داده‌هایی که به صورت گراف نمایش داده می‌شوند را به صورت استاندارد و موازی دارد. GraphX همچنین عملگرهای مختلفی را برای تغییر گراف‌ها (نظیر subgraph و mapVertices) و کتابخانه‌ای از الگوریتم‌های گراف خاص (نظیر

^۱ Streaming Analytics

^۲ Regression

^۳ Clustering

PageRank و شمارش مثلث‌های گراف) فراهم آورده است.

اسپارک همچنین از چندین زبان برنامه‌نویسی مانند جاوا^۱، اسکالا^۲، پایتون^۳، آر^۴ پشتیبانی می‌کند و SQL نیز از زبان‌هایی است که می‌توانید به وسیله آن، با داده‌های موجود در اسپارک کار کنید.

۴-۲ ماشین بردار پشتیبان

۴-۲-۱ مقدمه

ماشین بردار پشتیبان، یکی از روش‌های یادگیری با نظارت^۵ می‌باشد که در حال حاضر به صورت کارآمد و گسترده برای پاسخ به مسائل رده‌بندی مورد استفاده قرار می‌گیرد. الگوریتم اصلی ماشین بردار پشتیبان توسط ولادیمیر واپنیک^۶ و الکسی چرووننکسیس^۷ در سال ۱۹۶۳ اختراع شد. در سال ۱۹۹۲ ولادیمیر واپنیک و همکاران یک روش برای ایجاد رده‌بندی‌های غیرخطی بوسیله استفاده از ترفند هسته^۸ ارائه دادند [۱۶]. روش استاندارد رایج نیز توسط واپنیک و کورینا کورتس^۹ در سال ۱۹۹۳ مطرح و در ۱۹۹۵ به چاپ رسید^{۱۰} [۱۷]. مبنای کاری رده‌بندی کننده SVM پیدا کردن خط (ابرفضا) با حداکثر حاشیه برای رده‌بندی خطی داده‌ها است. به بیان دیگر در تقسیم خطی داده‌ها سعی می‌کند خطی را انتخاب کند که حاشیه اطمینان بیشتری داشته باشد. می‌توان پیدا کردن خط بهینه برای داده‌ها را به یک مسئله بهینه‌سازی درجه دوم محدب^{۱۱} تبدیل کرد و به وسیله روش‌های QP^{۱۲} که روش‌های شناخته شده‌ای در حل مسائل محدودیت‌دار هستند، حل نمود.

^۱ Java

^۲ Scala

^۳ Python

^۴ R

^۵ Supervised Learning

^۶ Vladimir N. Vapnik

^۷ Alexey Yakovlevich Chervonenkis

^۸ Kernel Trick

^۹ Corinna Cortes

^{۱۰} https://en.wikipedia.org/wiki/Support_vector_machine#History

^{۱۱} Convex Quadratic Optimization Problem

^{۱۲} Quadratic Programming

۲-۴-۲ ماشین بردار پشتیبان

ابتدا برای توضیح بهتر به ماشین بردار پشتیبان خطی پرداخته می‌شود. فرض کنید تعداد N نقطه داده $\{(x_1, y_1), (x_2, y_2), \dots, (x_N, y_N)\}$ در مجموعه داده آموزش قرار دارد که در آن $x_i \in R^d$ و $y_i \in \{\pm 1\}$ می‌باشد. می‌خواهیم یک رده‌بند خطی آموزش دهیم:

$$f(x) = \text{sgn}(w \cdot x - b) \quad (۱-۲)$$

همچنین می‌خواهیم خطی داشته باشیم که بیشترین حاشیه جداسازی با توجه به دو دسته را داشته باشد. به بیانی دیگر خط $H: y = w \cdot x - b = 0$ که دو خط $H_1: y = w \cdot x - b = +1$ و $H_2: y = w \cdot x - b = -1$ موازی و با فاصله برابر نسبت به آن قرار دارند را تعیین نماییم با این شرط که هیچ داده ورودی بین H_1 و H_2 قرار نگیرد و همچنین فاصله بین H_1 و H_2 بیشینه گردد.

برای انتخاب حاشیه دو حالت در نظر گرفته می‌شود. برای داده‌هایی که براحتی قابل جداسازی می‌باشند، حالتی به نام سخت حاشیه^۱ استفاده می‌شود. در این حالت می‌توان همیشه برای هر خط جداساز H و خطوط H_1 و H_2 ، با نرمال‌سازی^۲ ضرایب بردار w خطوط H_1 و H_2 را به ترتیب به صورت $y = w \cdot x - b = +1$ و $y = w \cdot x - b = -1$ تعیین کرد. ما می‌خواهیم فاصله بین H_1 و H_2 بیشینه شوند. تعدادی داده نمونه مثبت بر روی خط H_1 و تعدادی داده نمونه منفی بر روی خط H_2 قرار دارند. به این نمونه‌ها بردارهای پشتیبان می‌گویند، زیرا تنها این نقاط در تعریف خط جداساز تاثیر گذار می‌باشند و سایر نمونه‌ها را می‌توان نادیده گرفت.

به یاد داریم که در فضای دو بعدی، فاصله یک نقطه (x_0, y_0) از خط $Ax + By + C = 0$ برابر با

$$\frac{|Ax_0 + By_0 + C|}{\sqrt{A^2 + B^2}}$$

می‌باشد. به طور مشابه فاصله یک نقطه روی خط H_1 تا خط $H: y = w \cdot x - b = 0$

^۱ Hard-Margin

^۲ Normalize

برابر با $\frac{|w \cdot x - b|}{\|w\|} = \frac{1}{\|w\|}$ می‌باشد و فاصله بین H_1 و H_2 برابر با $\frac{2}{\|w\|}$ می‌باشد. بنابراین برای بیشینه کردن فاصله، باید $\|w\| = w^T w$ را کمینه نمود با این شرط که داده‌ای بین H_1 و H_2 وجود نداشته باشد:

$$y_i = +1 \text{ برای داده‌های مثبت } w \cdot x - b \geq +1$$

$$y_i = -1 \text{ برای داده‌های منفی } w \cdot x - b \leq -1$$

می‌توان این دو شرط را با یکدیگر ترکیب نمود: $y_i (w \cdot x_i - b) \geq 1$

بنابراین می‌توانیم مسئله‌مان را به صورت معادله‌ی بهینه‌سازی مقید زیر بیان نماییم:

$$y_i (w \cdot x_i - b) \geq 1 \quad \text{s.t.} \quad \min_{w, b} \frac{1}{2} w^T w \quad (2-2)$$

و این یک مسئله بهینه‌سازی درجه دوم محدب است.

برای حل این مسئله ساده‌تر آن است که مسئله به صورت دوگان^۱ مطرح و حل شود. ضرایب لاگرانژ یک استراتژی برای پیدا کردن کمینه یا بیشینه یک تابع با توجه به محدودیت‌ها است. برای بدست آوردن شکل دوگان مسئله ضرایب لاگرانژ مثبت $\alpha_i \geq 0$ در قیود ضرب شده و از تابع هدف کسر می‌شود و در نهایت به شکل مسئله اولیه^۲ زیر می‌رسیم:

$$L(w, b, \alpha) \equiv \frac{1}{2} w^T w - \sum_{i=1}^N \alpha_i y_i (w \cdot x_i - b) - \sum_{i=1}^N \alpha_i \quad (2-3)$$

شرایط KKT^۳ نقش مهمی را در مسائل بهینه‌سازی مقید، بازی می‌کنند. این شرایط، شرایط لازم و کافی برای داشتن پاسخ بهینه برای معادلات مقید را بیان می‌دارند و باید مشتق نسبت به متغیرها صفر باشد. با اعمال شرایط KKT روی معادله L و همچنین از L نسبت به w و b مشتق گرفته و برابر صفر

¹ Dual

² Primal Problem

³ Karush-Kuhn-Tucker

قرار دهیم به روابط زیر می‌رسیم:

$$\frac{\partial L}{\partial w} = 0 \Rightarrow w = \sum_{i=1}^N \alpha_i y_i x_i \quad (۴-۲)$$

$$\frac{\partial L}{\partial b} = 0 \Rightarrow \sum_{i=1}^N \alpha_i y_i = 0$$

و با جایگذاری این روابط در رابطه $L(w, b, \alpha)$ ، رابطه زیر را داریم:

$$L_D \equiv \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j y_i y_j (x_i \cdot x_j) \quad (۵-۲)$$

که معادله دوگان^۱ نامیده می‌شود. L و L_D هر دو از شرایط یکسانی بدست آمده‌اند و لذا با محاسبه کمینه L و یا بیشینه L_D با شرط $\alpha_i \geq 0$ می‌توان به جواب رسید.

برای هر نمونه آموزشی یک ضریب لاگرانژ $\alpha_i \geq 0$ وجود دارد. به نمونه‌هایی که برای آنها $\alpha_i > 0$ می‌باشد، بردارهای پشتیبان گفته می‌شود و همان نقاطی هستند که بر روی خط جداساز قرار می‌گیرند.

پس از حل معادله فوق و بدست آوردن مقادیر α_i ، می‌توان $w = \sum_{i=1}^N \alpha_i y_i x_i$ را محاسبه و پس از آن می‌توان داده x جدید را رده‌بندی نمود:

$$\begin{aligned} f(x) &= \text{sgn}(w \cdot x + b) \\ &= \text{sgn}\left(\left(\sum_{i=1}^N \alpha_i y_i x_i\right) \cdot x + b\right) \\ &= \text{sgn}\left(\sum_{i=1}^N \alpha_i y_i (x_i \cdot x) + b\right) \end{aligned} \quad (۶-۲)$$

۳-۴-۲ بهینه‌سازی متوالی کمینه

بهینه‌سازی متوالی کمینه (SMO^۲)، یک الگوریتم برای حل مسائل QP منتج شده در ضمن فرآیند

^۱ Dual Problem

^۲ Sequential Minimal Optimization

آموزش ماشین بردار پشتیبان می‌باشد. این روش در سال ۱۹۹۸ توسط جان پلات^۱ در مرکز تحقیقات مایکروسافت^۲ ابداع گردید [۱۸]. SMO برای تعلیم ماشین بردار پشتیبان به طور گسترده‌ای مورد استفاده قرار گرفت و سبب موج استقبال وسیعی توسط انجمن SVM گردید. روش‌های موجود تا قبل از ارائه SMO برای حل SVM بسیار پیچیده و هزینه‌بر بودند.

SMO الگوریتم ساده‌ایست که به سرعت مسئله QP در SVM را بدون استفاده از راه‌حل عددی بهینه‌سازی درجه دوم، حل می‌کند. SMO اقدام به تجزیه مسئله‌ی اصلی به تعدادی زیرمسئله می‌نماید و به این ترتیب همگرایی روش را نیز تضمین می‌نماید. در هر مرحله SMO دو ضریب لاگرانژ را برای بهینه کردن انتخاب می‌نماید و مقدار بهینه را برای این ضرائب پیدا نموده و SVM را برای یافتن مقدار بهینه به‌روز می‌نماید. مزیت روش SMO در این است که راه‌حل مسئله برای دو ضریب لاگرانژ می‌تواند به صورت تحلیلی حاصل شود. بنابراین در سراسر راه‌حل از بهینه‌سازی QP به صورت عددی پرهیز می‌گردد.

میزان فضای مورد نیاز برای SMO نسبت به اندازه‌ی مجموعه‌ی آموزش خطی می‌باشد که این خود به SMO اجازه‌ی مدیریت مجموعه‌های آموزشی بزرگ را می‌دهد. با توجه به اینکه در این روش از محاسبات ماتریسی اجتناب شده است. لذا SMO بین مرتبه‌های خطی و دوم نسبت به اندازه‌ی مجموعه‌ی آموزشی قرار می‌گیرد، در حالی که در SVM استاندارد در میان مرتبه‌های دوم و سوم قرار می‌گیرند. در دنیای واقعی SMO می‌تواند هزار بار از الگوریتم قطعه‌بندی استاندارد برای SVM سریع‌تر باشد [۱۸].

الگوریتم ساده شده SMO:

الگوریتم SMO دو ضریب α ، یعنی α_i و α_j را در هر مرحله تکرار انتخاب می‌نماید و ارزش عینی هر دو ضریب را به طور مشترک بهینه می‌کند. پس از آن پارامتر b بر اساس α های جدید تنظیم و محاسبه

¹ John Platt

² Microsoft Research

می‌شود. در ادامه این مراحل بیشتر توضیح داده می‌شود [۱۹].

انتخاب دو ضریب α : بخش عمده‌ای از الگوریتم SMO روش اکتشافی^۱ برای انتخاب ضرایب α_i و α_j می‌باشد تا آنها را بهینه نماید. برای مجموعه داده‌های بزرگ، این نکته برای سرعت انجام الگوریتم حیاتی می‌باشد. زیرا $N(N-1)$ تعداد انتخاب‌های ممکن برای ضرایب α_i و α_j می‌باشد و برای برخی از این انتخاب‌ها، بهبود بسیار کمتر نسبت به سایر ضرایب می‌باشد.

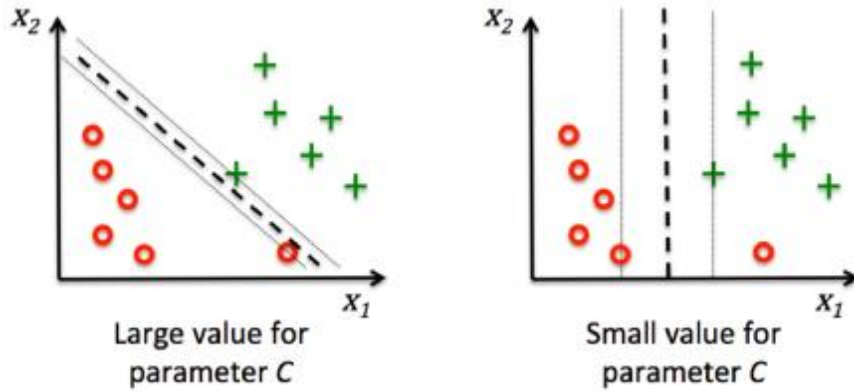
برای این کار از دو انتخاب اکتشافی مجزا برای ضرایب استفاده می‌شود. در الگوریتم ارائه شده [۱۸] از دو حلقه‌ی متداخل استفاده شده است. در حلقه‌ی خارجی در اولین اجرا، تمام مجموعه‌ی داده‌های آموزشی را برای یافتن نمونه‌ای که شرایط KKT را نقض نماید، پویش می‌نماید. اگر نمونه‌ای شرایط KKT را نقض نماید، مناسب برای بهینه‌سازی خواهد بود. بعد از اولین تکرار حلقه‌ی خارجی تکرار را بر روی نمونه‌هایی از ضرایب لاگرانژ که بین ۰ و C قرار دارند ($0 < \alpha_i < C$)، ادامه می‌دهد.

- پارامتر C برای اجتناب از داده‌های آموزشی اشتباه رده‌بندی شده در بهینه‌سازی SVM می‌باشد. برای مقادیر بزرگ C، حاشیه کم برای ابرصفحه^۲ انتخاب می‌شود تا رده‌بندی داده‌ها درست انجام شود. برای مقادیر کوچک C نیز حاشیه زیاد برای ابرصفحه در نظر گرفته می‌شود. در این حالت داده‌های بیشتری اشتباه رده‌بندی می‌شود. شکل ۲-۳ تاثیر مقادیر مختلف C نمایش می‌دهد. سپس مجدداً، هر نمونه با شرایط KKT مقایسه شده و نمونه‌ای که آنرا نقض نماید، مناسب برای بهینه‌سازی خواهد بود. این بهینه‌سازی در حلقه‌ی خارجی تا زمانی اجرا می‌شود که تمام این نمونه‌های مذکور شرایط KKT را با خطای ε برآورده سازند. در ادامه حلقه خارجی بازگشت نموده و تکرار را روی تمام مجموعه‌ی آموزشی انجام می‌دهد. حلقه خارجی به طور مداوم تا زمانی که تمامی نمونه‌های آزمایشی شرایط KKT را با خطای ε برآورده سازند، بین این دو انتخاب این جابجا می‌شود. بعد از انتخاب اولین ضریب توسط حلقه‌ی

^۱ Heurist

^۲ Hyperplane

خارجی، دومین ضریب توسط حلقه داخلی و با تکرار بر روی ضرائب باقیمانده، برای انجام بهینه‌سازی انتخاب می‌شود.



شکل ۳-۲: پارامتر C در SVM

محاسبه مقادیر دو ضریب: برای ضرایب لاگرانژ انتخاب بهینه کردن α_i و α_j ، ابتدا قیود روی مقادیر این پارامترها محاسبه می‌کنیم و سپس مسئله بیشینه‌سازی مقید مربوط به آن را حل می‌کنیم.

ابتدا مرزهای L و H را بدست می‌آوریم به طوری که $L < \alpha_i < H$ و به شرطی قید $0 < \alpha_i < C$ برآورده شود. این مرزها توسط روابط زیر بدست می‌آید:

$$\begin{aligned} \text{if } y_i \neq y_j, \quad L &= \max(0, \alpha_j - \alpha_i), \quad H = \min(C, C + \alpha_j - \alpha_i) \\ \text{if } y_i = y_j, \quad L &= \max(0, \alpha_j + \alpha_i - C), \quad H = \min(C, \alpha_j + \alpha_i) \end{aligned} \quad (7-2)$$

حال می‌خواهیم مقدار α_j را بدست آوریم تا تابع هدف را بیشینه نماییم. اگر این مقدار از مرزهای L و H عبور کند، ما مقدار α_j در این محدوده نگه می‌داریم. مقدار بهینه α_j با استفاده از رابطه زیر بدست می‌آید:

$$\alpha_j := \alpha_j - \frac{y_j (E_i - E_j)}{\eta} \quad (8-2)$$

که در آن E_k و η از روابط زیر بدست می‌آید:

$$E_k = f(x_k) - y_k \quad (9-2)$$

$$\eta = 2(x_i \cdot x_j) - (x_i \cdot x_i) - (x_j \cdot x_j)$$

E_k را می‌توان به عنوان خطایی بین خروجی SVM برای نمونه‌ی k ام و برچسب درست آن دانست که $f(x_k)$ در آن رابطه زیر می‌باشد:

$$f(x) = \sum_{i=1}^N \alpha_i y_i (x_i \cdot x) + b \quad (10-2)$$

همچنین هنگام محاسبه پارامتر η می‌توان در صورت استفاده از هسته، از آن بجای ضرب داخلی به صورت دلخواه استفاده نمود.

پس از آن ما مقدار α_j را با استفاده از رابطه زیر برش می‌دهیم تا در بازه $[L, H]$ قرار گیرد:

$$\alpha_j := \begin{cases} H & \text{if } \alpha_j > H \\ \alpha_j & \text{if } L < \alpha_j < H \\ L & \text{if } \alpha_j < L \end{cases} \quad (11-2)$$

در نهایت پس از اینکه مقدار α_j را بدست آوردیم، مقدار α_i را با توجه به α_i و با استفاده از رابطه زیر بدست می‌آوریم:

$$\alpha_i := \alpha_i + y_i y_j (\alpha_j^{(old)} - \alpha_j) \quad (12-2)$$

$\alpha_j^{(old)}$ مقدار قبل از بروزرسانی در این مرحله می‌باشد. همچنین اگر $\eta = 0$ حالتی است که ما از این جفت α نمی‌توانیم بهبودی داشته باشیم و به سراغ یک جفت دیگر می‌رویم.

محاسبه مقدار b : پس از بهینه‌سازی مقادیر α_i و α_j مقدار آستانه b را انتخاب می‌نماییم به قسمی که شرایط KKT را برای i ام و j امین نمونه برآورده نماید. برای بدست آوردن آستانه b جدید، اگر مقدار α_i پس از بهینه‌سازی در محدوده معین شده قرار گیرد (برای مثال $0 < \alpha_i < C$) بنابراین مقدار آستانه b_1 از رابطه زیر بدست می‌آید:

$$b_1 = b - E_i - y_i (\alpha_i - \alpha_i^{(old)})(x_i \cdot x_i) - y_j (\alpha_j - \alpha_j^{(old)})(x_i \cdot x_j) \quad (۱۳-۲)$$

به طور مشابه، b_2 از رابطه زیر بدست می‌آید اگر $0 < \alpha_j < C$ قرار گیرد:

$$b_2 = b - E_j - y_i (\alpha_i - \alpha_i^{(old)})(x_i \cdot x_j) - y_j (\alpha_j - \alpha_j^{(old)})(x_j \cdot x_j) \quad (۱۴-۲)$$

اگر هر دو ضریب در محدوده $0 < \alpha_i < C$ و $0 < \alpha_j < C$ قرار گیرند، مقادیر آستانه فوق معتبر می-

باشند و آن‌ها با هم برابر خواهند شد. اگر مقادیر جدید دو ضریب α_i و α_j روی مرزها قرار بگیرد

(برای مثال $\alpha_i = 0$ or $\alpha_i = C$ and $\alpha_j = 0$ or $\alpha_j = C$) تمام مقادیر آستانه بین دو مقدار b_1 و b_2 ،

شرایط KKT را برآورده می‌نماید. بنابراین می‌توان مقدار b جدید را بر طبق رابطه زیر بدست آورد:

$$b := \begin{cases} b_1 & \text{if } 0 < \alpha_i < C \\ b_2 & \text{if } 0 < \alpha_j < C \\ (b_1 + b_2) / 2 & \text{otherwise} \end{cases} \quad (۱۵-۲)$$

در شکل ۳-۲ نیز شبه‌کد این روش آمده است.

۴-۴-۲ هسته در ماشین بردار پشتیبان

ماشین بردار پشتیبان با تشکیل خطی جداساز بین دو دسته‌ی داده‌های آموزش، اقدام به رده‌بندی داده‌های آزمون می‌نماید. این مدل پایه را ماشین بردار پشتیبان خطی^۱ می‌گویند. در برخی از مسائل داده‌ها به راحتی قابل تفکیک از یکدیگر نمی‌باشند و در این حالت یک ماشین بردار پشتیبان خطی موثر باشد. در این حالت اگر داده‌ها را به فضایی با ابعاد دیگر انتقال دهیم می‌توان راه‌حلی برای جداسازی داده‌ها بدست آورد. برای این عمل می‌توان از توابع مختلفی از جمله چندجمله‌ای^۲، گوسین^۳ یا تابع پایه شعاعی^۴ (RBF) استفاده نمود که آن را هسته می‌نامند. شکل ۵-۲ نمونه‌ای از کاربرد هسته را نشان می‌دهد.

^۱ Linear SVM

^۲ Polynomial

^۳ Gaussian

^۴ Radial Basic Function

Algorithm: Simplified SMO**Input:**

C: regularization parameter

tol: numerical tolerance

max_passes: max # of times to iterate over α 's without changing $(x_1, y_1), \dots, (x_N, y_N)$: training data**Output:** $\alpha \in \mathbb{R}^N$: Lagrange multipliers for solution $b \in \mathbb{R}$: threshold for solution**Pseudo-Code:**Initialize $\alpha_i = 0, \forall i, b = 0$

Initialize passes = 0

while (passes < max_passes)

num_changed_alphas = 0

for i = 1,...,m

 Calculate $E_i = f(x_i) - y_i$ using (9-2) and (10-2) if $((y_i E_i < -tol \ \&\& \ \alpha_i < C) \parallel (y_i E_i > tol \ \&\& \ \alpha_i > 0))$ Select j $\neq i$ randomly Calculate $E_j = f(x_j) - y_j$ using (9-2) and (10-2) Save old α 's: $\alpha_i^{(old)} = \alpha_i, \alpha_j^{(old)} = \alpha_j$

Compute L and H by (7-2)

if (L == H)

continue to next i

 Compute η by (9-2) If ($\eta \geq 0$)

continue to next i

 Compute and clip new value for α_i using (8-2) and (11-2) If $(|\alpha_j - \alpha_j^{(old)}| < 10^{-5})$

continue to next i

 Determine value for α_i using (12-2) Compute b_1 and b_2 using (13-2) and (14-2) respectively

Compute b by (15-2)

num_changed_alphas := num_changed_alphas + 1

end if

end for

if (num_changed_alphas == 0)

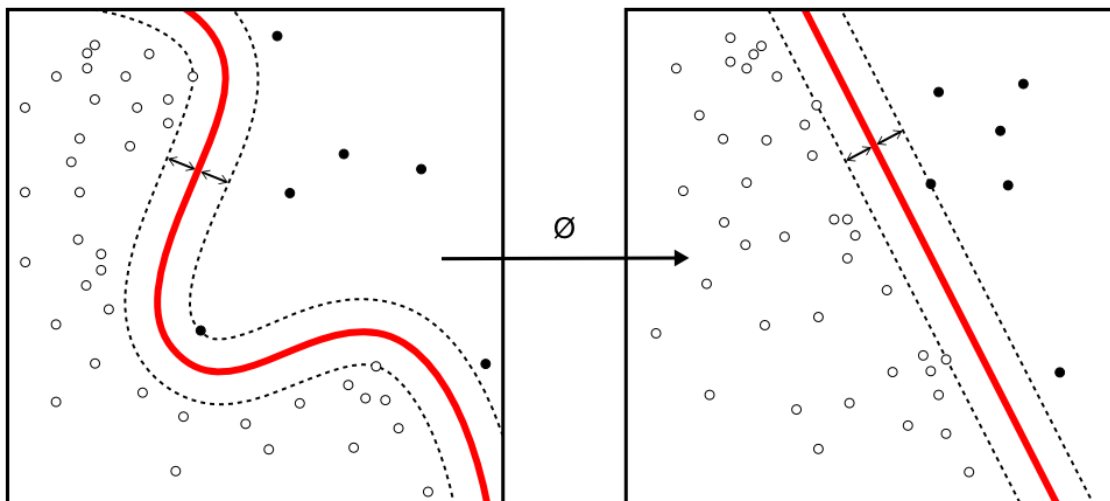
passes := passes + 1

else

passes := 0

end while

شکل ۲-۳: الگوریتم SMO ساده شده [۱۹]



شکل ۵-۲: استفاده از هسته برای جداسازی داده‌ها

۵-۲ مشکلات مرتبط با نظرکاوی و چالش‌ها

در حوزه نظرکاوی و تجزیه و تحلیل احساسات، تعدادی مشکل وجود دارد که به آن اشاره می‌شود و می‌توان برای کارهای آینده بر روی این موضوعات تحقیق نمود. اکثراً این مشکلات مربوط به زمینه پردازش زبان طبیعی می‌باشند.

۵-۲-۱ مشکلات مربوط به منبع داده

در زمینه نظرکاوی فقدان مجموعه داده‌ای معیار برای ارزیابی نتایج محسوس می‌باشد. همچنین برای بعضی از زمینه‌های نظرکاوی و برخی زبان‌ها مانند زبان‌های فارسی و عربی و ... مشکل تهیه مجموعه داده وجود دارد. میکالای^۱ [۲۰] تعداد کمی از معروف‌ترین مجموعه داده‌ها را برای نظرکاوی معرفی کرده است که تا حدودی می‌توان از آن‌ها بعنوان معیار استفاده نمود (به دلیل استفاده فراوان). همچنین وبسایت‌های^۲ IMDB، آمازون و توییتر از مهم‌ترین منابع برای نظرکاوی می‌باشند.

^۱ Mikalai Tsytarau

^۲ International Movie Database

۲-۵-۲ مشکلات مربوط به پردازش زبان طبیعی

مشکلاتی در پردازش زبان طبیعی وجود دارد که در این بخش تعدادی از این مشکلات را مرور می‌نماییم.

نفی: یکی از بخش‌های مهم تشخیص احساسات توجه و اداره کردن کلمات نفی می‌باشد. برای مثال نظر «رستوران خوبی بود.» دارای قطبیت مثبت می‌باشد اما نظر «رستوران خوبی نبود.» دارای قطبیت منفی می‌باشد، علی‌رغم استفاده از صفتی مانند «خوب» که بار مثبت دارد. همچنین به این نکته نیز باید توجه شود که همیشه با دیدن کلمات نفی نباید بار معنایی جمله را معکوس نمود.

برای مثال دو نظر «جای تعجب ندارد که رستوران خوبی است.» و «نه تنها فضای رستوران جالب بود، بلکه سرویس‌دهی نیز عالی بود.» دارای کلمات نفی می‌باشند اما قطبیت نظرات کاملاً مثبت می‌باشد و همین موضوع باعث ایجاد مشکل در این بخش می‌شود.

وابستگی به دامنه: برخی از کلمات با توجه به دامنه و جمله‌ای که در آن قرار می‌گیرند دارای قطبیت و بار معنایی متفاوتی می‌باشند، بنابراین توجه به این موضوع نیز حائز اهمیت می‌باشد. برای مثال «طول عمر باتری گوشی طولانی است.» کلمه طولانی دارای بار مثبت می‌باشد اما در «زمان بارگذاری سیستم عامل گوشی، طولانی است.» کلمه طولانی، دارای بار منفی می‌باشد.

استفاده از ضمیر: در برخی از نظرات از ضمیر استفاده می‌شود و یا اشاره به موجودیتی دارد که مشخص نیست و اداره کردن این موضوع مشکل می‌باشد. برای مثال، تعیین قطبیت نظر «این فیلم کارگردان نسبت به فیلم قبلی بهتر است.» وابسته به قطبیت فیلم قبلی کارگردان دارد.

استفاده از چند زبان: در برخی از تارنماها، نظرات با زبان‌های متفاوتی بیان می‌شود که طبیعتاً با استفاده از روش‌های متکی به یک زبان نمی‌توان به نظرکاوای آن‌ها پرداخت و یا اینکه در یک نظر از چند زبان استفاده می‌شود. برای مثال «این گوشی موبایل از کابل رابط OTG پشتیبانی می‌کند.»

ابهام فهم جمله برای رایانه: برخی نظرات دارای بار مثبت و منفی به طور همزمان می‌باشند و با این حال برای انسان‌ها به راحتی قابل فهم می‌باشد اما درک و تحلیل آن برای رایانه مشکل می‌باشد. برای مثال نظر «فیلم مانند بمب ترکیب با وجود اینکه بازیگر اصلی آن را لرزاند (و خوب کار نکرد)».

فصل سوم

کارهای پیشین

در این فصل به مرور کارهای پیشین پرداخته می‌شود.

۳-۱ روش‌های نظر کاوی

همانطور که در فصل قبل توضیح داده شد، روش‌های نظر کاوی دارای دو قسمت اصلی روش‌های مبتنی بر یادگیری ماشین و روش‌های مبتنی بر واژه نامه^۱ می‌باشند. در این قسمت پژوهش‌های مرتبط پیشین به تفکیک روش‌های استفاده شده آورده شده است.

۳-۱-۱ روش‌های مبتنی بر یادگیری ماشین

۳-۱-۱-۱ بیز ساده، ماشین بردار پشتیبان و بیشینه آنتروپی

پانگ و همکاران برای سه رده‌بندی کننده فوق مقایسه‌ای را با استفاده ویژگی‌های متفاوت انجام داده‌اند [۸]. در نتایج این پژوهش، هنگامی که مجموعه ویژگی‌ها کوچک باشد، رده‌بند بیز ساده بهتر عمل می‌نماید. همچنین هنگامی که مجموعه ویژگی‌ها بزرگ باشد، ماشین بردار پشتیبان نسبت به دو روش دیگر بهتر عمل می‌کند و همچنین در حالت کلی با افزایش اندازه مجموعه ویژگی، روش بیشینه آنتروپی بهتر از بیز ساده عمل می‌نماید.

۳-۱-۱-۲ درخت تصمیم:

زینیا سینگلا^۲ و همکاران [۲۱] برای رده‌بندی برای تحلیل نظرات کاربران از سه روش رده‌بند SVM، Naïve Bayes(NB) و همچنین درخت تصمیم استفاده کرده است. آنوجا پی جین^۳ و همکارش از دو رده‌بند NB و درخت تصمیم برای تحلیل و رده‌بندی نظرات کاربران استفاده نموده‌اند [۲۲].

۳-۱-۱-۳ رده‌بندی کننده‌های مبتنی بر قاعده:

یولیا بیدولیا^۴ و برونوا^۵ [۲۳] یک رده‌بند مبتنی بر قانون را برای تحلیل احساسات در حوزه کیفیت

^۱ Lexicon Based Methods

^۲ Zeenia Singla

^۳ Anuja P Jain

^۴ Yuliya Bidulya

^۵ Brunova

خدمات بانکی انجام داده است. آیمن سیدیکو^۱ و همکارانش از ترکیب رده‌بند مبتنی بر قانون و رده‌بند NB تحت نظارت ضعیف برای تحلیل نظرات کاربران در تویتر استفاده نموده است [۲۴].

۳-۱-۴ شبکه عصبی مصنوعی:

رودریگو مورائز^۲ و همکاران در پژوهشی به مقایسه تکنیک شبکه عصبی و ماشین بردار پشتیبان پرداخته‌اند و همچنین از رده‌بند NB نیز استفاده نموده‌اند [۲۵]. در این مقایسه که بر روی مجموعه داده‌هایی از جمله نظرات مربوط به فیلم، دوربین و کتاب انجام شده است، شبکه عصبی نتایج بهتری را نسبت به ماشین بردار پشتیبان بدست آورده است. همچنین هر کدام از تکنیک‌های فوق دارای محدودیت‌هایی می‌باشند. برای مثال هزینه زمانی و محاسباتی مرحله آموزش برای شبکه عصبی مصنوعی و همچنین هزینه محاسباتی برای ماشین بردار پشتیبان از محدودیت‌ها و معایب این روش‌ها می‌باشند.

در پژوهش دیگری، هائو لینگ^۳ و ژائو هوئی^۴ [۲۶] از شبکه عصبی کانولوشن استفاده کرده‌اند و سعی در بهبود آن داشته‌اند. همچنین در این پژوهش از سایر روش‌های یادگیری از جمله SVM، NB، MaxEntropy و CNN استاندارد استفاده شده است.

تانجیم الهق^۵ و همکاران [۲۷] به تحلیل نظرات کاربران بر روی مجموعه داده آمازون پرداخته‌اند. در این پژوهش از چندین روش مبتنی بر یادگیری ماشین استفاده شده است. SVM خطی، Mutlinimoial، NB، SGD^۶، Random Forest، Logistic Regression و درخت تصمیم روش‌های استفاده در این پژوهش می‌باشد. در ادامه به بررسی سایر پژوهش‌هایی که بر روی مجموعه داده آمازون اجرا شده، پرداخته شده است.

¹ Umme Aymun Siddiqua

² Rodrigo Moraes

³ Hao Lidong

⁴ Zhao Hui

⁵ Tanjim Ul Haque

⁶ Stochastic Gradient Descent

۳-۱-۲ روش‌های مبتنی بر فرهنگ لغت

سحر سوهان‌گیر و همکاران در پژوهشی از روش‌های مبتنی بر فرهنگ لغت برای تشخیص احساسات استفاده نموده‌اند [۲۸]. در پژوهش انجام شده از SentiWordNet و TextBlob برای رده‌بندی استفاده کرده است. همچنین در این پژوهش از سه روش مبتنی بر یادگیری ماشین از جمله Linear SVM، NB و Logistic Regression برای مقایسه استفاده شده است.

۳-۲ مقیاس‌پذیری^۱

«مقیاس‌پذیری توانایی یک سامانه^۲، شبکه^۳ و یا فرآیند^۴ برای رسیدگی (پاسخ‌گویی) به مقداری کار می‌باشد که در حال افزایش است و یا آمادگی آن برای بزرگ شدن (توسعه و افزایش عملکرد) در مقابل رشد کار است [۲۹].»

امروزه با توجه به گسترش شبکه‌های اجتماعی، حجم زیادی داده در فضای مجازی تولید می‌شود. این داده‌ها در بسیاری از زمینه‌ها می‌توانند مفید باشند و اطلاعات ارزشمندی را شامل شوند. اما پردازش این حجم عظیم از داده، بخصوص در مواردی که نیازمند به پردازش بلادرنگ^۵ می‌باشد، روش‌های نوینی را نسبت به روش‌های گذشته می‌طلبد. پژوهش‌گران راه‌کارها و روش‌های گوناگونی را برای غلبه بر این مشکل ارائه داده‌اند. برای نمونه پژوهش‌گران در حوزه نظیرکاوی به ارائه چارچوب‌های مقیاس‌پذیر^۶ پرداخته‌اند که درون خود از روش‌های مختلفی از جمله روش‌های مبتنی بر یادگیری ماشین و یا روش‌های مبتنی بر پردازش زبان طبیعی (NLP) بهره می‌برند. در ادامه به برخی از این پژوهش‌ها اشاره و تا حدودی به کلیت روش ارائه شده می‌پردازیم.

¹ Scalability

² System

³ Network

⁴ Process

⁵ Real-Time

⁶ Scalable

ماریا کاراناسو^۱ و همکاران [۳۰] روشی مقیاس‌پذیر برای غلبه بر مشکل پردازش عظیم داده‌ها در زمینه نظرکاوی ارائه داده‌اند. روشی که در این مرجع بیان شده دارای دو قسمت کلی می‌باشد. در قسمت اول که مرحله برون‌خطی^۲ نامیده می‌شود، رده‌بند با استفاده از مجموعه داده آموزشی برچسب‌گذاری شده، آموزش داده می‌شود. در مرحله دوم که مرحله درون‌خطی^۳ نامیده می‌شود، داده‌ها به صورت بلادرنگ دریافت می‌شوند و پیش‌پردازش و انتخاب ویژگی بر رویشان اعمال می‌شود و در نهایت از همان رده‌بند آموزش دیده برای رده‌بندی داده‌های بلادرنگ استفاده می‌شود. در روش پیشنهادی از یادگیری گروهی استفاده شده است. استفاده از رأی اکثریت^۴ برای بهبود عملکرد رده‌بندها استفاده شده است. همچنین در مرحله دوم از بازخورد^۵ برای ارزیابی استفاده می‌شود و در صورت عدم دقت کافی مرحله اول تکرار می‌شود. این روش بر روی چارچوب استورم^۶ پیاده‌سازی شده است.

آسو حمزه‌ای و همکاران، با ارائه روشی به نام SSSA^۷ سعی در بهبود مشکل پردازش عظیم داده‌ها در زمینه نظرکاوی داشته‌اند [۳۱]. در روش پیشنهادی پس از مراحل پیش‌پردازش، به هر کلمه امتیازی با استفاده از ابزارهای وردنت و SentiWordNet تعلق می‌گیرد و در نهایت با استفاده از مجموع امتیازات بدست آمده، قطبیت جمله یا نظر تعیین می‌گردد. در روش ارائه شده سعی شده کندی روش‌های مبتنی بر معنا^۸، با استفاده از ویژگی‌های NLP و روش‌های مقیاس‌پذیر بهبود یابد.

باستوری باتاچارجی^۹ و لیندا پتزولد^{۱۰}، روشی مقیاس‌پذیر روی مجموعه نظرات مشتریان رستوران‌ها مطرح کرده‌اند [۳۲]. در روش بیان شده با ارائه تعریفی جدید به نام فراویژگی^{۱۱} سعی در ایجاد ویژگی جدیدی نموده‌اند تا باعث کاهش بردار ویژگی و افزایش سرعت و دقت شوند. فراویژگی‌ها با توجه به هر

^۱ Maria Karanasou

^۲ Offline

^۳ Online

^۴ Maxvote

^۵ Feedback

^۶ Storm Framwork

^۷ Semantic Scoring Sentiment Analysis

^۸ Semantic-based

^۹ Kasturi Bhattacharjee

^{۱۰} Linda Petzold

^{۱۱} Meta-Feature

مجموعه داده متفاوت هست و وابسته به نوع مجموعه نظر می‌باشد و در روش مورد نظر ابتدا جنبه^۱ها و توصیف‌گر^۲های نظرات استخراج و رده‌بندی می‌شوند و سپس فراویژگی‌ها با ترکیب جنبه‌ها و توصیف-گرها ایجاد می‌شود و در نهایت با توجه به فراویژگی‌های موجود در نظرات، قطبیت نظر تعیین می‌گردد. برای مثال نظر «غذاهای اینجا عالی‌ست.» فراویژگی (غذا، عالی) وجود دارد که غذا خود شامل یک دسته و عالی نیز خود دارای یک زیرمجموعه هم معنی می‌باشند.

کوالین^۳ و همکاران [۳۳] یک معماری مقیاس‌پذیر برای تحلیل بلادرنگ داده‌های وب‌نوشت‌های کوچک^۴ ارائه کرده‌اند. این معماری بر روی زیرساخت جریانی IBM InfoSphere^۵ پیاده‌سازی شده است و توانایی مقابله با حجم زیاد نظر را داراست. این روش به تحلیل نظرات کاربران توئیتر در طی برگزاری بازی نهایی^۶ تیم فوتبال برزیل در جام کنفدراسیون‌ها ۲۰۱۳^۷ به صورت درون‌خطی پرداخته است. در روش ارائه شده رده‌بند با استفاده از نظرات کاربران در بازی‌های پیشین آموزش دیده است و از آن برای رده‌بندی نظرات کاربران در بازی نهایی استفاده می‌شود.

۳-۳ ماشین بردار پشتیبان موازی^۸

یکی از راهکارهای اخیر که برای مقابله با هزینه پردازش داده‌های حجیم و کاهش آن مورد توجه قرار گرفته استفاده از چندین گره برای پردازش سریع‌تر و موازی است. در این راه‌کار از برنامه‌نویسی نگاشت-کاهش^۹ استفاده می‌شود. مراحل ماند پیش‌پردازش، استخراج و انتخاب ویژگی به راحتی به صورت موازی قابل پیاده‌سازی می‌باشد اما آموزش رده‌بند و استفاده از آن به صورت موازی یکی از دغدغه‌های این راه‌کار می‌باشد.

¹ Aspect

² Descriptor

³ P. R. Cavalin

⁴ Microblogging data

⁵ IBM InfoSphere Streams platform

⁶ Final

⁷ 2013 FIFA Confederations Cup

⁸ Parallel SVM

⁹ MapReduce

از طرفی رده‌بند SVM یکی از رده‌بندهای پرکاربرد و دارای کارایی بالا می‌باشد و با توجه به اینکه رده‌بند SVM در زمینه نظرکاوی نتایج و دقت بالاتری را داشته است، در پژوهش‌ها بیشتر مورد توجه قرار گرفته است. در بسیاری از پژوهش‌های انجام شده، پژوهش‌گران برای افزایش سرعت و مقیاس-پذیری روش پیشنهادی، سعی در بهبود عملکرد رده‌بند SVM با استفاده از برنامه‌نویسی نگاشت-کاهش دارند. در ادامه به معرفی تعدادی از این پژوهش‌ها می‌پردازیم.

بو یان^۱ و همکارانش به رده‌بندی احساسات نظرات کاربران بر روی چارچوب اسپارک^۲ پرداخته‌اند [۳۴]. آنها ابتدا نظرات را پیش‌پردازش نموده و با استفاده از TF-IDF بردار ویژگی را ایجاد کرده و بعنوان ورودی رده‌بند در نظر گرفته‌اند. در ادامه نیز به آموزش رده‌بند SVM موازی همراه با هسته RBF با استفاده از برنامه‌نویسی نگاشت-کاهش نموده‌اند. در این مرجع موازی‌سازی در قسمت ساخت هسته RBF انجام شده و سایر قسمت‌ها همانند SVM استاندارد می‌باشد. در نهایت روش پیشنهادی را بر روی چند مجموعه داده متفاوت و با استفاده از چند گره پردازشی پیاده‌سازی کرده‌اند.

پوجاسومن تریپسی^۳ و همکاران به تجزیه و تحلیل ریسک^۴ با استفاده از SVM موازی بر روی چارچوب هادوپ^۵ پرداخته‌اند [۳۵]. در روش پیشنهادی مجموعه داده به چند قسمت تقسیم شده و برای پردازش-های اولیه در اختیار گره‌ها قرار می‌گیرد. در ادامه هر قسمت از زیرمجموعه داده اولیه برای آموزش رده‌بند SVM استفاده و بردارهای پشتیبان آموزش دیده^۶، توسط نگارنده‌ها^۷ به کاهنده‌ها^۸ فرستاده می‌شود. در ادامه تمامی بردارهای پشتیبان به کاهنده‌ها ارسال می‌شود و توسط آنها ترکیب و بردارهای پشتیبان نهایی و رده‌بند نهایی آماده می‌شود. در پژوهش دیگری [۳۶] نیز ابتدا توضیحی کوتاه در مورد SVM آورده شده، در ادامه اشاره‌ای به ساختار SVM موازی، مروری بر برنامه‌نویسی نگاشت-کاهش و نگاهی

¹ Bo Yan

² Spark

³ Pujasuman Tripathy

⁴ Risk analysis

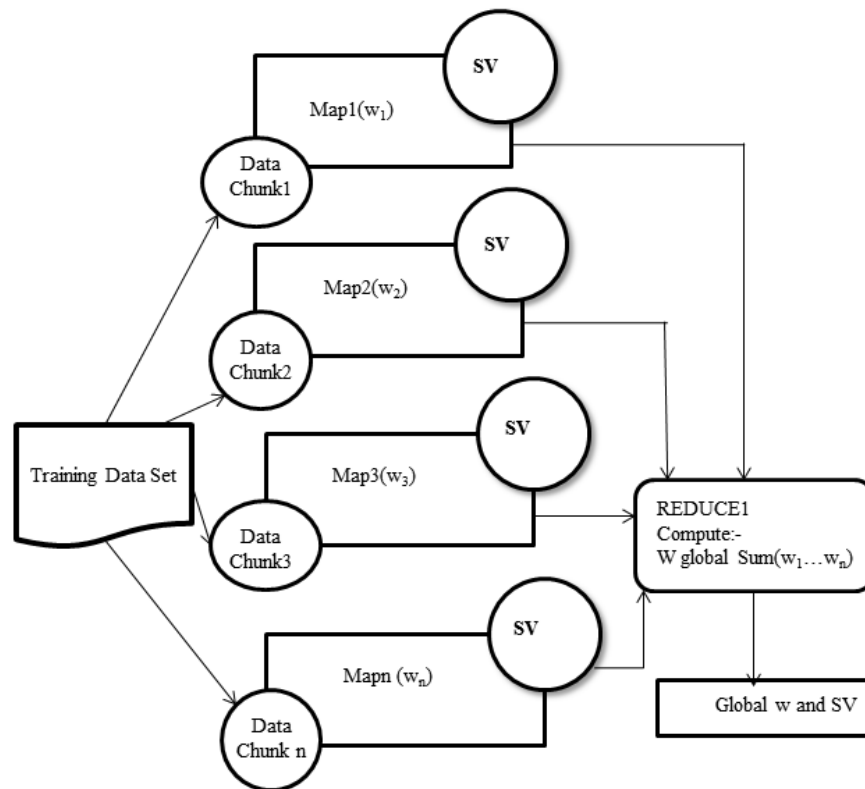
⁵ Hadoop

⁶ Trained Support Vectors

⁷ Mappers

⁸ Reducers

به چارچوب هادوپ داشته است. در نهایت به ارائه شبه‌کد SVM موازی مبتنی بر نگاشت-کاهش مشابه مرجع قبل داشته است. در شکل ۳-۱ ساختار روش ارائه شده در مراجع [۳۵, ۳۶] نشان داده شده است.



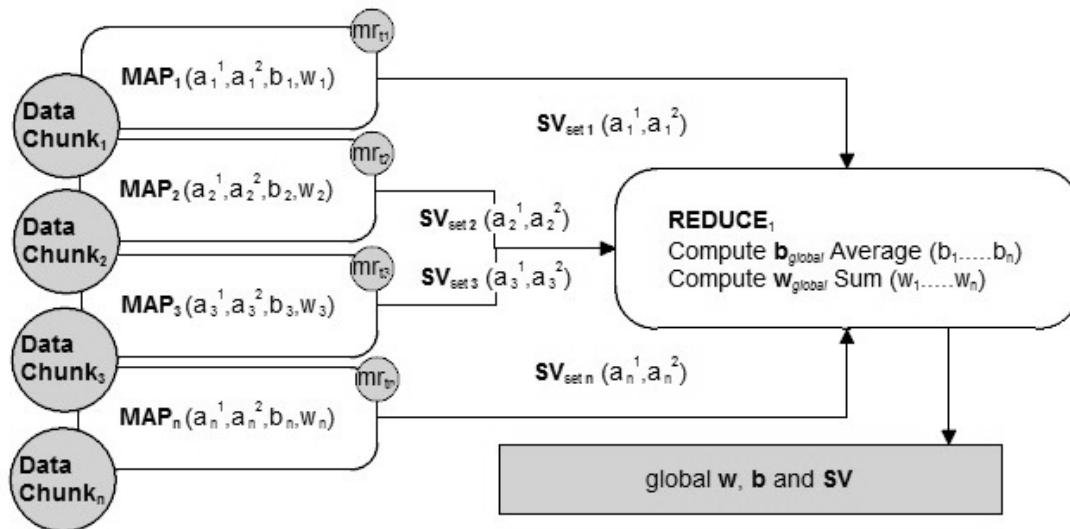
شکل ۳-۱: ساختار SVM موازی ارائه شده در مراجع [۳۵, ۳۶]

گادوین کاروانا^۱ و همکارانش [۳۷] روشی برای فیلتر کردن هرزنامه^۲ با استفاده از SVM موازی مبتنی بر مدل نگاشت-کاهش ارائه کرده‌اند که پیاده‌سازی آن بر روی هادوپ انجام شده است. پس از تقسیم داده‌ها به چند بخش، هر بخش به یک نگارنده ارسال می‌شود و با استفاده از الگوریتم SMO، بردارهای پشتیبان (مقادیر b و w) برای هر بخش محاسبه می‌شود و این بردارهای پشتیبان برای بدست آوردن بردار پشتیبان کلی به کاهنده(ها) فرستاده می‌شود. در کاهنده میانگین مقادیر b بعنوان b کلی و

^۱ Godwin Caruana

^۲ Spam

مجموع مقادیر W بعنوان W کلی در نظر گرفته شده است. پس از بدست آوردن مدل رده‌بند از آن برای تشخیص هرزنانه استفاده می‌گردد. ساختار و معماری روش ارائه شده در مرجع [۳۷] در شکل ۲-۳ نمایش داده شده است.



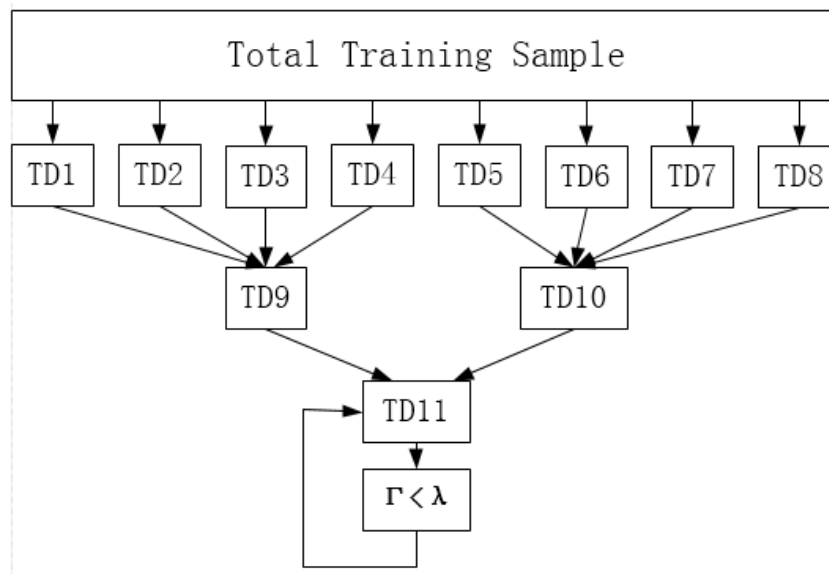
شکل ۲-۳: ساختار SVM موازی استفاده شده در مرجع [۳۷]

وی نی^۱ و همکارانش [۳۸] سعی در بهینه سازی رده‌بند SVM موازی با استفاده از چند تابع هسته داشته‌اند. ابتدا اشاره ای به SVM موازی آبشاری^۲ و SVM گروهی^۳ داشته‌اند و در روش خود، ترکیبی از این دو مدل استفاده کرده‌اند. همچنین برای تابع هسته مدل پیشنهادی خود از ترکیب تابع هسته چندجمله‌ای و RBF بهره برده‌اند. در SVM استفاده شده در لایه آخر از شرط توقف λ (تعداد بردارهای پشتیبان مرحله کنونی تعداد بردارهای پشتیبان مرحله قبلی) برای اطمینان از همگرایی الگوریتم بهره گرفته شده است. ساختار SVM پیشنهادی مرجع [۳۸] را در شکل ۳-۳ می‌توانید مشاهده نمایید.

^۱ Wei Nie

^۲ Cascade-PSVM

^۳ Grouped-SVM



شکل ۳-۳: ساختار SVM موازی ارائه شده در مرجع [۳۸]

پس از معرفی پژوهش‌هایی که برای موازی‌سازی SVM از نگاشت-کاهش استفاده نموده‌اند، حال به معرفی تعداد دیگری از پژوهش‌ها می‌پردازیم. در این پژوهش‌ها سعی شده است تا ساختار و معماری جدیدی برای رده‌بند SVM معرفی گردد.

۳-۳-۱ ماشین بردار پشتیبان آبشاری^۱

در این روش مجموعه داده آموزشی به چند (توانی از دو) بخش تقسیم می‌شود و بر روی هر قسمت یک SVM مجزا اعمال می‌شود. خروجی SVMها (بردارهای پشتیبان) دو به دو با هم ترکیب شده و بعنوان ورودی SVM مرحله بعد در نظر گرفته می‌شوند. این روش از مرجع [۳۹] اقتباس شده است که برای رده‌بندی تصاویر و خارج از محیط اسپارک پیاده‌سازی شده است. در این مرجع از نظر ریاضی نیز همگرایی روش اثبات شده است. برای فرموله کردن این روش ابتدا پارامترهای این روش معرفی می‌شود. n تعداد بخش‌های مجموعه داده، z سطح z ام در SVM آبشاری، m مدل آموزش دیده برای هر بخش مجموعه داده، sv مجموعه‌ی بردارهای پشتیبان و L نیز سطح پایانی در SVM آبشاری می‌باشد.

^۱ Cascade SVM

روش ماشین بردار پشتیبان آبخاری به صورت زیر فرموله می‌شود:

$$\begin{aligned}
 & m_1^j, m_2^j, \dots, m_i^j, \dots, m_n^j && \text{مدل های آموزش دیده سطح } j \\
 \Rightarrow & sv_1^j, sv_2^j, \dots, sv_i^j, \dots, sv_n^j && \text{بردارهای پشتیبان خروجی مدل های آموزش دیده سطح } j \\
 \{sv_1^j, sv_2^j\}, \dots, \{sv_{2*i-1}^j, sv_{2*i}^j\}, \dots, \{sv_{n-1}^j, sv_n^j\} \rightarrow & && \text{بردارهای پشتیبان سطح } j, \text{ ورودی مدل های سطح } j+1 \\
 & m_1^{j+1}, m_2^{j+1}, \dots, m_i^{j+1}, \dots, m_{n/2}^{j+1} && \text{مدل های آموزش دیده سطح } j+1 \\
 \Rightarrow & sv_1^{j+1}, sv_2^{j+1}, \dots, sv_i^{j+1}, \dots, sv_{n/2}^{j+1} && \text{بردارهای پشتیبان خروجی مدل های آموزش دیده سطح } j+1 \\
 & \dots && \dots \\
 & M = m_1^{n+1} && \text{در نهایت مدل آموزش دیده نهایی}
 \end{aligned}$$

۳-۳-۲ ماشین بردار پشتیبان آبخاری بهبودیافته^۱

در این روش مجموعه داده آموزشی به چند بخش تقسیم می‌شود و بر روی هر قسمت یک SVM مجزا اعمال می‌شود. تفاوت روش SVM آبخاری بهبودیافته با روش SVM آبخاری در این است که خروجی SVM های مرحله اول، دو به دو ترکیب نمی‌شود بلکه خروجی تمام SVM های مرحله اول باهم ترکیب شده و به تنها SVM مرحله بعدی (آخر) فرستاده می‌شود. این روش اقتباسی نیز در مرجع [۴۰] ارائه شده است و روش پیشنهادی در مرجع [۳۹] را بهبود داده و از این روش برای تشخیص نفوذ^۲ استفاده کرده است. فرموله کردن این روش نیز به صورت زیر انجام می‌گیرد:

$$\begin{aligned}
 & m_1^1, m_2^1, \dots, m_i^1, \dots, m_n^1 && \text{مدل های آموزش دیده سطح یک} \\
 \Rightarrow & sv_1^1, sv_2^1, \dots, sv_i^1, \dots, sv_n^1 && \text{بردارهای پشتیبان خروجی مدل های آموزش دیده سطح یک} \\
 & sv_1^1, sv_2^1, \dots, sv_i^1, \dots, sv_n^1 \rightarrow m_1^2 && \text{بردارهای پشتیبان سطح یک، ورودی مدل سطح دو} \\
 & M = m_1^2 && \text{در نهایت مدل آموزش دیده سطح دو، مدل نهایی}
 \end{aligned}$$

۳-۳-۳ ماشین بردار پشتیبان گروهی^۳

روش SVM گروهی تا حدودی مشابه با روش SVM آبخاری بهبود یافته می‌باشد، با این تفاوت که

^۱ Improved Cascade SVM

^۲ Intrusion Detection

^۳ Grouped-SVM

SVM در مرحله بعدی وجود ندارد و خروجی SVM های مرحله اول با یکدیگر ترکیب می شوند. خروجی مرحله آموزش برای هر SVM بردار w (ضرایب لاگرانژ) و ثابت b می باشد. w کلی ترکیب w های مرحله قبل و b کلی، میانگین b های مرحله قبل می باشد. این روش در [۴۱] معرفی گردیده است. در مرجع [۴۲] از این روش با استفاده از MapReduce برای تشخیص و رده بندی نوع فایل (فرمت فایل) استفاده کرده بود و در مرجع [۴۳] نیز از این روش برای تفسیر تصویر (Image Annotation) استفاده کرده بود.

فرموله بندی روش SVM گروهی به صورت زیر می باشد:

$m_1^1, m_2^1, \dots, m_i^1, \dots, m_n^1$	مدل های آموزش دیده سطح یک
$P \text{ } sv_1^1, sv_2^1, \dots, sv_i^1, \dots, sv_n^1$	بردارهای پشتیبان خروجی مدل های آموزش دیده سطح یک
$M = \bigcup_{i=1}^n sv_i^1$	در نهایت مدل نهایی، اجتماع بردارهای پشتیبان سطح یک

در رابطه بالا n تعداد بخش های مجموعه داده، m مدل آموزش دیده برای هر بخش مجموعه داده، sv مجموعه ی بردارهای پشتیبان می باشد.

۳-۴ تجمیع پژوهش های پیشین

در این قسمت به بررسی و رده بندی کارهای مرتبط و پیشین می پردازیم. در جدول ۳-۱ نتایج این قسمت آورده شده است. در تمامی پژوهش ها از روش های مبتنی بر یادگیری ماشین استفاده شده است. معیارهایی همچون مقیاس پذیری روش، استفاده از محاسبات خوشه ای و همچنین کاربرد روش پیشنهادی در تحلیل احساسات در بررسی مراجع و پژوهش های پیشین مورد توجه قرار گرفته است. همچنین خلاصه ای از روش پیشنهادی نیز آورده شده است.

جدول ۱-۳: بررسی و تجمیع کارهای مرتبط پیشین

مقیاس پذیری	محاسبات	تحلیل احساسات	روش مورد استفاده
خوشه‌ای			
مرجع [۲۵]	خیر	خیر	بله
مرجع [۱۲]	خیر	خیر	بله
[۱۴]			ترکیب روش های مبتنی یادگیری ماشین و واژه نامه
مرجع [۳۰]	بله	بله	بله
			بهره‌گیری از دو فاز برون و دورن-خطی برای غلبه بر پردازش داده‌های کلان و پیاده‌سازی در استورم ^۱
مرجع [۳۱]	بله	خیر	بله
			ارائه روشی به نام SSSA و استفاده از SVM و SentiWN
مرجع [۳۲]	بله	خیر	بله
			ارائه تعریف جدیدی به نام فراویژگی و امتیازدهی به کلمات و جملات
مرجع [۳۳]	بله	بله	بله
			استفاده از رده‌بند آموزش دیده برای تحلیل نظرات کاربران در طی مسابقه فوتبال، پیاده‌سازی در IBM InfoSphere
مرجع [۳۴]	بله	بله	بله
			موازی‌سازی فرآیند تشکیل هسته در مرحله آموزش SVM، پیاده‌سازی در Spark
مرجع [۳۵]	بله	بله	خیر
[۳۷]			موازی‌سازی SVM مبتنی بر MapReduce و پیاده‌سازی در محیط Hadoop

^۱ Apache Storm

مرجع [۳۸]	بله	بله	خیر	استفاده از ترکیب چند هسته برای آموزش SVM و استفاده از SVM آبشاری و گروهی، پیاده‌سازی در Hadoop
مرجع [۳۹]	خیر	خیر	خیر	معرفی و استفاده از SVM آبشاری
مرجع [۴۰]	خیر	خیر	خیر	بهبود روش SVM آبشاری
مرجع [۴۱]	خیر	خیر	خیر	معرفی و استفاده از SVM گروهی
مرجع [۴۲]	بله	بله	خیر	استفاده از SVM گروهی برای رده‌بندی نوع داده
مرجع [۴۳]	بله	بله	خیر	استفاده از SVM گروهی برای تفسیر تصویر
مرجع [۴۴]	بله	بله	بله	موازی‌سازی SVM و پیاده‌سازی در محیط Cloudera
مرجع [۴۵]	بله	بله	بله	ارائه الگوریتمی به نام STING و پیاده‌سازی موازی آن در Cloudera
مرجع [۴۶]	بله	بله	بله	استفاده از گرامر وابسته به مجموعه‌داده برای استخراج ویژگی و استفاده از SVM، NB و Logistic Regression، پیاده‌سازی در Spark
مرجع [۴۷]	بله	خیر	بله	استفاده از شبکه عصبی CNN و ترکیب با روش‌های دیگر
مرجع [۴۸]	بله	بله	بله	موازی‌سازی الگوریتم Fuzzy C-Mean، پیاده‌سازی در محیط

Cloudera				
ارائه روشی برای پیش‌پردازش و مقایسه آن با استفاده از SVM، ANN و NB	بله	خیر	خیر	مرجع [۴۹]
استفاده از شبکه عصبی بازگشتی و همچنین ترکیب مجموعه داده‌ها برای مرحله آموزش	بله	خیر	خیر	مرجع [۵۰]
استفاده از یادگیری عمیق و شبکه عصبی کانولوشن	بله	خیر	خیر	مرجع [۵۱]
استفاده از Unigram و Bigram برای استخراج ویژگی، همچنین استفاده از رده‌بند های SVM و NB و ترکیب هر دو	بله	خیر	خیر	مرجع [۵۲]

فصل چهارم

روش پیشنهادی

۱-۴ مقدمه

حجم زیاد اطلاعات در دنیای امروز برای پردازش و تحلیل، نیازمند روش‌های سریع و کارآمد می‌باشد. نظرات کاربرانی که در فضای اینترنت منتشر می‌شود نیز از این قاعده مستثنی نیست و رده‌بندی آنها محتاج روش‌های پرسرعت و در عین حال دقیق می‌باشد.

هدف این پژوهش نیز ارائه روشی جهت بهبود سرعت رده‌بندی نظرات کاربران در کنار حفظ دقت بالا می‌باشد. در این پژوهش چندین روش مبتنی بر SVM پیشنهاد شده است. این روش‌ها بر روی چارچوب اسپارک پیاده‌سازی شده اند. در ادامه به توضیح بیشتر این روش‌ها می‌پردازیم.

۲-۴ روش پیشنهادی

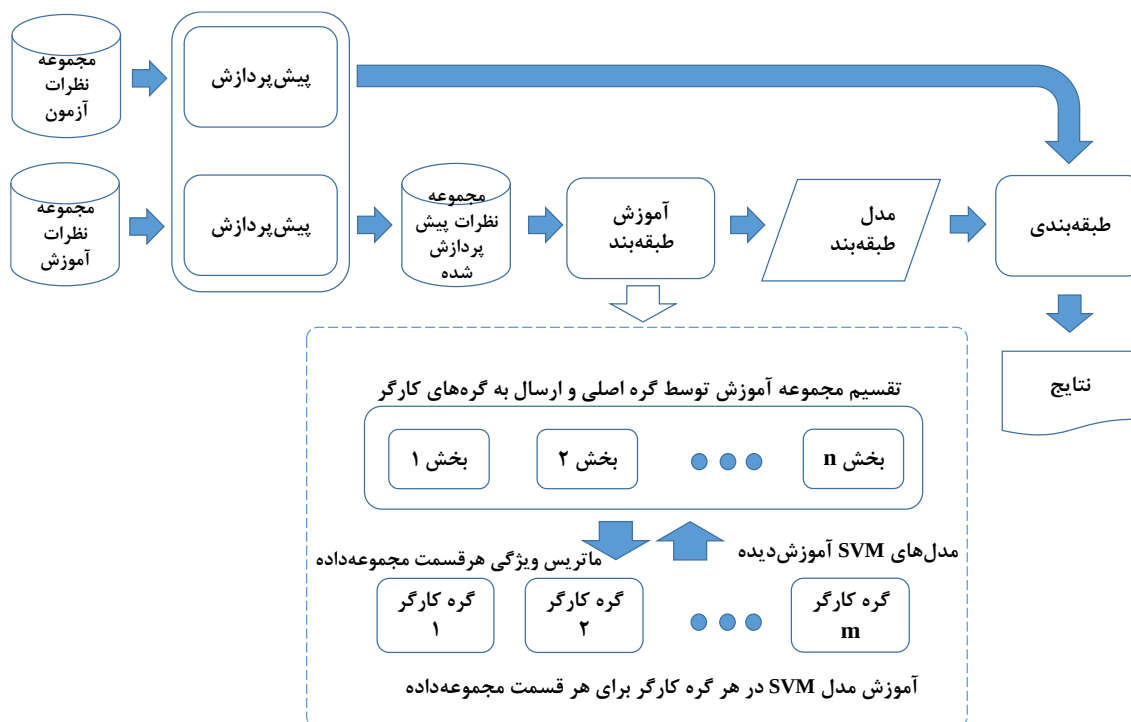
در این پژوهش روشی برای رده‌بندی نظرات کاربران به دو دسته مثبت و منفی بر روی مجموعه داده IMDB ارائه شده است. پس از انجام مراحل پیش‌پردازش به آموزش مدل رده‌بند پرداخته می‌شود و در ادامه از مدل آموزش دیده برای رده‌بندی استفاده می‌شود. پیاده‌سازی روش پیشنهادی در محیط اسپارک انجام شده است. در شکل ۱-۴ معماری و ساختار روش پیاده‌سازی شده آورده شده است.

۱-۲-۴ پیش‌پردازش

در این پژوهش مراحل پیش‌پردازشی زیادی انجام نشده است که این موضوع به سادگی پیاده‌سازی و سرعت آن کمک می‌نماید. ابتدا ایست‌واژه‌ها^۱ از نظرات حذف می‌شوند. پس از آن برداری شامل تمام کلمات مجموعه داده آموزش تشکیل می‌شود. کلماتی که تعداد تکرار آنها در تمام مجموعه نظرات کمتر از ۷۰ مرتبه می‌باشد نیز در نظر گرفته نمی‌شوند. در نهایت با توجه به کلمات موجود در هر نظر و تعداد تکرار آن، بردار ویژگی برای هر نظر، تشکیل می‌شود. تقریباً این نوع پیش‌پردازش به دلیل عدم استفاده از نقش کلمات در جمله (و در صورتی که کلمات پرتکرار نادیده گرفته شوند) می‌تواند برای زبان‌های

^۱ Stopword

نیز استفاده شود. این قسمت در نرم افزار متلب^۱ پیاده سازی شده است.



شکل ۴-۱: معماری کلی روش پیشنهادی

۴-۲-۲ آموزش مدل رده بند

در این مرحله با استفاده از ماتریس ویژگی به آموزش مدل رده بند پرداخته می شود. در این پژوهش، از رده بند SVM برای رده بندی دسته نظرات کاربران استفاده کرده ایم. این رده بند در حوزه متن کاوی و نظر کاوی بسیار موثر و کارا عمل می نماید. بنابراین در این پژوهش از رده بند SVM استفاده گردد. برای آموزش رده بند از روش های SVM استاندارد، آبخاری، آبخاری بهبود یافته، گروهی و SVM مبتنی بر رأی گیری برای آموزش مدل به صورت مستقل استفاده شده است. در ادامه به طور مفصل تر به معرفی ویژگی های استفاده شده در هر یک از روش ها می پردازیم.

برای آموزش مدل رده بند SVM از روش آموزشی SMO استفاده کرده ایم و به دلیل زیاد بودن مولفه های بردار ویژگی، از هسته RBF استفاده نموده ایم. در مجموعه داده های بزرگ و در صورت استفاده از هسته،

^۱ Matlab

قبل از اجرای SMO نیاز به تشکیل ماتریس هسته^۱ می‌باشد که زمان زیادی از مرحله آموزش را نیازمند است. این ماتریس متقارن می‌باشد و می‌توان به راحتی با نیمی از زمان مورد نیاز، آن را ایجاد کرد. در مقالات مطالعه شده، مقاله‌ای که به این نکته اشاره نماید، مشاهده نشد.

۳-۲-۴ رده‌بندی

پس از مرحله آموزش مدل رده‌بند، مرحله آخر یعنی رده‌بندی قرار دارد. در این حالت نظر برای رده‌بندی و تعیین قطبیت به مدل آموزش دیده فرستاده می‌شود و خروجی پیش‌بینی بدست می‌آید. این قسمت برای تمام روش‌های SVM موازی به جز روش SVM مبتنی بر رأی‌گیری، توسط یک مدل آموزش دیده انجام می‌شود. در روش SVM مبتنی بر رأی‌گیری، چند مدل آموزش دیده وجود دارد و نظر توسط هر مدل پیش‌بینی و رده‌بندی شده و در نهایت رأی‌گیری بین پیش‌بینی مدل‌ها انجام می‌شود.

۳-۴ آموزش مدل رده‌بندی

برای رده‌بندی و آموزش مدل رده‌بند از روش SVM مبتنی بر رأی‌گیری و روش‌های SVM موازی آبشاری، آبشاری بهبودیافته، گروهی استفاده شده است که در ادامه به توضیح بیشتر این روش‌ها می‌پردازیم.

۱-۳-۴ ماشین بردار پشتیبان مبتنی بر رأی‌گیری

در این روش پیشنهادی پژوهش حاضر، مجموعه داده آموزش را به چند بخش (تعداد فرد) تقسیم نموده و از هر بخش برای آموزش یک مدل SVM استفاده می‌نماییم. پس از مرحله آموزش، از تمام SVM‌ها برای رده‌بندی مجموعه داده آزمون استفاده می‌کنیم. برای این عمل، کل مجموعه داده آزمون توسط هر یک از مدل‌های SVM رده‌بندی می‌شود. خروجی رده‌بندی هر SVM برای هر نظر، دو دسته مثبت و منفی می‌باشد که به ترتیب با ۱ و ۱- مشخص می‌شوند. در نهایت با تجمیع خروجی SVM‌ها برای هر

^۱ Kernel Matrix

داده‌ی مجموعه آزمون، نتیجه نهایی رده‌بندی بدست می‌آید که رأی‌گیری بین SVM‌ها می‌باشد. مدل پیشنهادی به صورت زیر فرموله می‌شود:

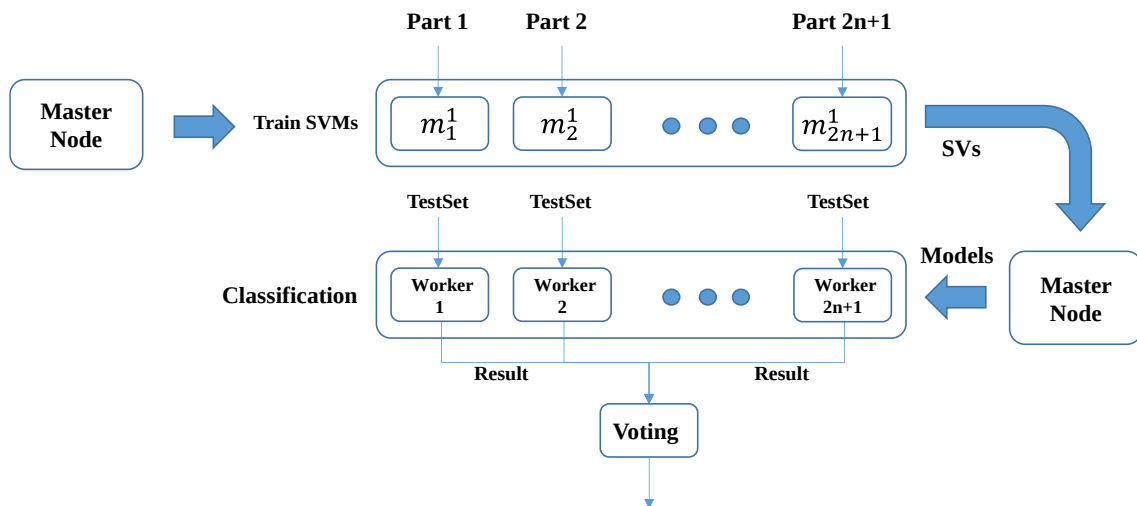
$$m_1^1, m_2^1, \dots, m_i^1, \dots, m_n^1$$

مدل‌های آموزش‌دیده سطح یک

$$\text{classify}\{M, o\} = \underset{m \in M}{\operatorname{argmax}}_i \sum \{\text{classify}(m, o) = i\}$$

در نهایت رده‌بندی نظر بر مبنای رأی‌گیری مدل‌های آموزش‌دیده سطح یک

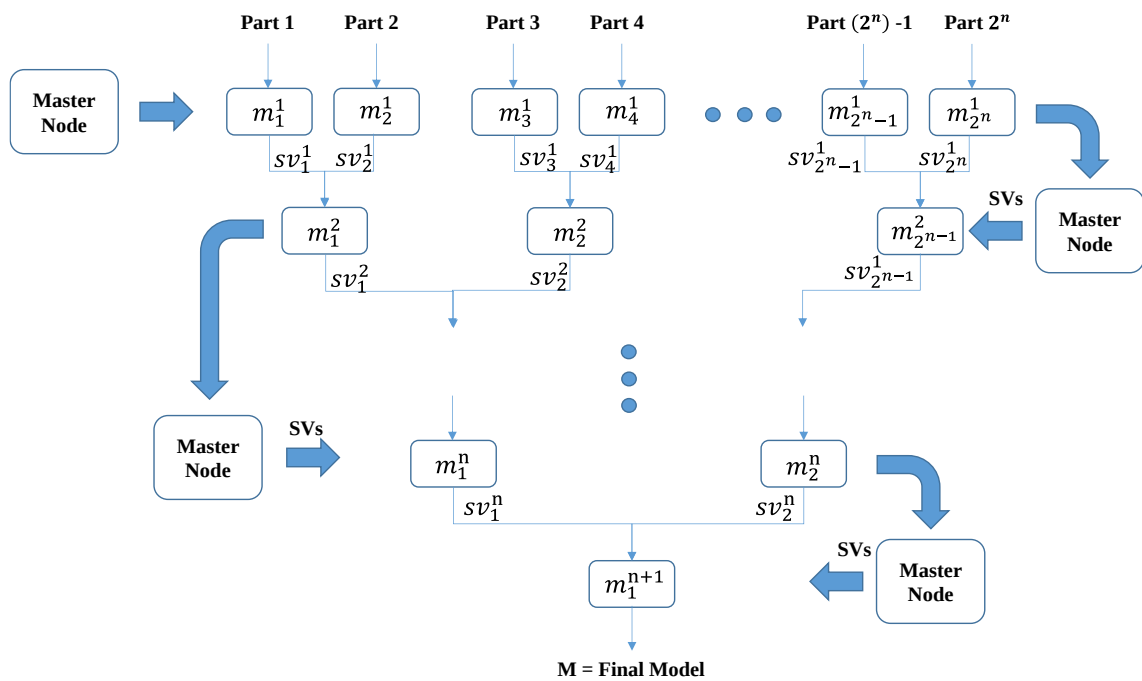
که در آن n تعداد بخش‌های مجموعه‌داده، o نظر ورودی برای رده‌بندی، M مدل آموزش‌دیده و m هر کدام از SVM‌های آموزش‌دیده می‌باشند. معماری و ساختار روش SVM مبتنی بر رأی‌گیری در شکل ۲-۴ بر پایه معماری اسپارک، نمایش داده شده است.



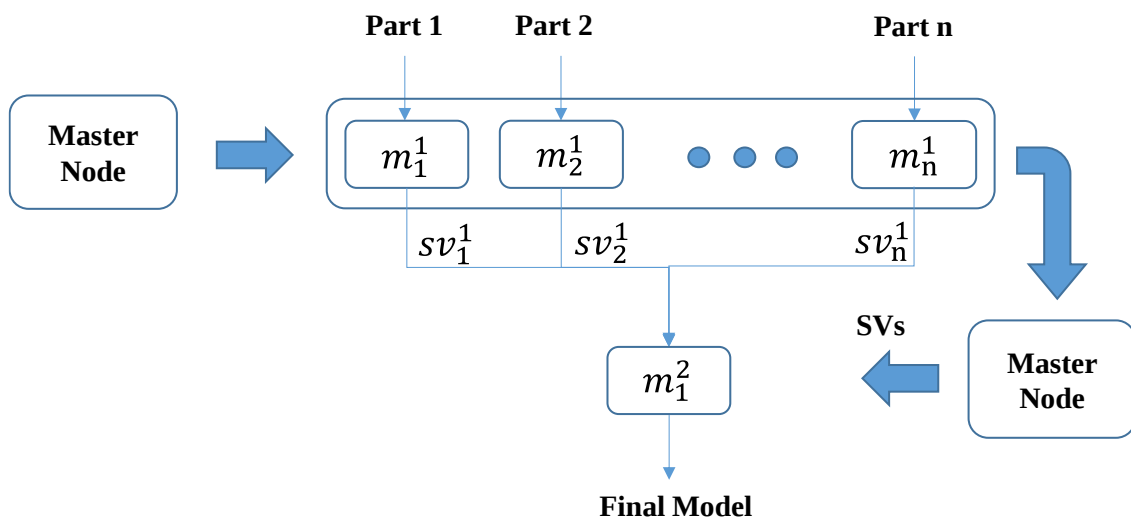
شکل ۲-۴: معماری SVM مبتنی بر رأی‌گیری

۲-۳-۴ ماشین بردار پشتیبان آبخاری

همانطور که در فصل قبل اشاره شد این روش توسط هانس پیتر گرفت و همکاران [۳۹] معرفی شده است. در پژوهش حاضر از این روش برای استفاده در حوزه نظرکاوی استفاده شده است. پیاده‌سازی این روش نیز در محیط اسپارک انجام شده است. در شکل ۳-۴ معماری و ساختار این روش مبتنی بر اسپارک، نشان داده شده است.



شکل ۳-۴: معماری SVM آبشاری



شکل ۴-۴: معماری SVM آبشاری بهبود یافته

۳-۳-۴ ماشین بردار پشتیبان آبخاری بهبودیافته

این روش توسط دوهانگلی و همکاران در سال ۲۰۰۹ معرفی شده است [۴۰]. دوهانگلی و همکاران سعی در بهبود روش SVM آبخاری داشته‌اند. در پژوهش حاضر، این روش را در محیط اسپارک پیاده‌سازی و همچنین برای کاربرد در حوزه نظرکاوی استفاده شده است. در شکل ۴-۴ ساختار روش ماشین بردار پشتیبان آبخاری بهبود یافته با توجه به معماری اسپارک را نشان می‌دهد.

۴-۳-۴ ماشین بردار پشتیبان آبخاری با پالایش بردار پشتیبان

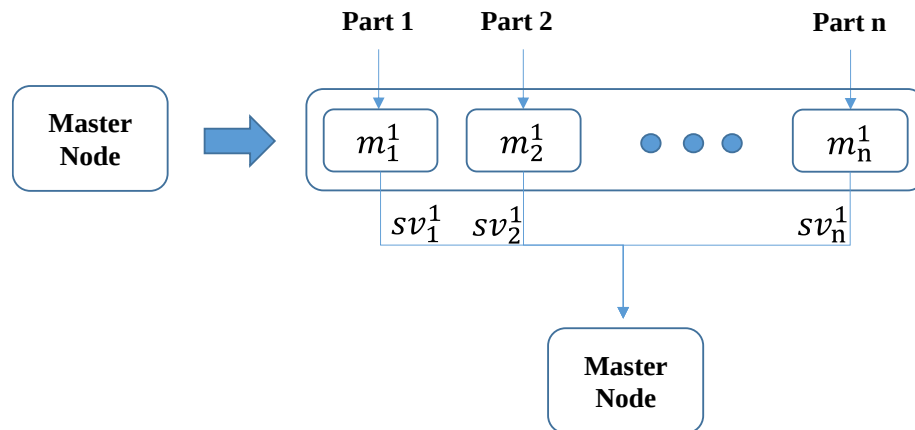
این روش در پژوهش حاضر معرفی و استفاده شده است. ایده این روش سعی در کاهش داده‌های ورودی SVMهای مرحله دوم به بعد می‌باشد. در SVM آبخاری و آبخاری بهبود یافته، در هر مرحله بردارهای پشتیبان جدا شده و به عنوان ورودی SVM مرحله بعدی در نظر گرفته می‌شوند. بردارهای پشتیبان، نمونه‌هایی می‌باشند که برای آن‌ها ضرایب لاگرانژ خروجی از روش آموزشی SMO، بزرگتر از صفر باشند ($0 < \alpha_i \leq C$). با توجه به این موضوع می‌توانیم با تعیین حد آستانه برای انتخاب ضرایب لاگرانژ، بخشی از ضرایب لاگرانژی که بزرگتر از حد آستانه باشند را به عنوان ورودی مرحله بعد در نظر بگیریم. با این عمل تعداد داده‌های ورودی SVM مرحله بعدی کاهش می‌یابد و در نتیجه باعث افزایش سرعت آموزش می‌شود. از طرفی دیگر ضرایب لاگرانژ با مقادیر بزرگتر دارای اهمیت بیشتری بوده و با توجه به این موضوع از دقت نهایی نیز مقدار زیادی کاسته نخواهد شد. این روش به صورت زیر فرموله می‌شود:

$$\begin{aligned}
 & m_1^1, m_2^1, \dots, m_i^1, \dots, m_n^1 && \text{مدل‌های آموزش دیده سطح یک} \\
 & \Rightarrow sv_1^1, sv_2^1, \dots, sv_i^1, \dots, sv_n^1 && \text{بردارهای پشتیبان خروجی مدل‌های آموزش دیده سطح یک} \\
 & \bigcup_{i=1}^n \left\{ sv_i^1 \mid \alpha_i > \frac{C}{2} \right\} && \text{بردارهای پشتیبان سطح یک با ضرایب لاگرانژ بزرگتر از } C/2, \text{ بعنوان} \\
 & \rightarrow m_1^2 && \text{ورودی مدل سطح دو} \\
 & M = m_1^2 && \text{در نهایت مدل آموزش دیده سطح دو، مدل نهایی}
 \end{aligned}$$

۵-۳-۴ ماشین بردار پشتیبان گروهی

همانطور که در فصل سوم اشاره شد، این روش توسط دانگ و همکاران در سال ۲۰۰۳ معرفی گردیده

است [۴۱]. در پژوهش حاضر، از این روش نیز برای استفاده در حوزه نظردکاوای بهره گرفته شده است. پیاده‌سازی روش ذکر شده، در محیط اسپارک انجام شده است. شکل ۴-۵ مربوط به ساختار و معماری این روش مبتنی بر چارچوب اسپارک می‌باشد.



شکل ۴-۵: معماری SVM گروهی

۴-۳-۶ مرتبه زمانی

مرتبه زمانی رده‌بند SVM غیرخطی (استفاده از هسته) با توجه به کم یا زیاد بودن مقدار پارامتر C ، بین دو مقدار $O(N^2)$ و $O(N^3)$ می‌باشد که در آن N تعداد داده‌های ورودی می‌باشد [۵۳]. با این شرایط مرتبه زمانی برای SVM آبشاری (برای حالتی که مجموعه داده به چهار بخش تقسیم شده باشد) در کمترین حالت برابر با:

$$O(4 * (\frac{N}{4})^2 + (|sv_1^1| + |sv_2^1|)^2 + (|sv_3^1| + |sv_4^1|)^2 + (|sv_1^2| + |sv_2^2|)^2) \quad (۴-۱)$$

و در بیشترین حالت برابر با:

$$O(4 * (\frac{N}{4})^3 + (|sv_1^1| + |sv_2^1|)^3 + (|sv_3^1| + |sv_4^1|)^3 + (|sv_1^2| + |sv_2^2|)^3) \quad (۴-۲)$$

می‌باشد. مرتبه زمانی برای SVM آبشاری بهبودیافته در کمترین حالت برابر با:

$$O(n * (\frac{N}{n})^2 + (|sv_1^1| + |sv_2^1| + \dots + |sv_n^1|)^2) \quad (۴-۳)$$

و در بیشترین حالت برابر با رابطه زیر می باشد:

$$O(n * (\frac{N}{n})^3 + (|sv_1^1| + |sv_2^1| + \dots + |sv_n^1|)^3) \quad (4-4)$$

مرتبه زمانی برای SVM گروهی نیز در کمترین حالت $O(n * (\frac{N}{n})^2)$ و بیشترین حالت $O(n * (\frac{N}{n})^3)$

می باشد. مرتبه زمانی برای SVM مبتنی بر رأی گیری نیز مشابه حالت گروهی می باشد.

۴-۴ سایر روش های انجام شده

در این پژوهش دو روش دیگر نیز مورد بررسی قرار گرفت اما به دلیل نامطلوب بودن نتایج بدست آمده از بررسی بیشتر و ذکر نتایج خودداری شد.

- روش اول: ابتدا مجموعه داده به چند بخش تقسیم می شود، سپس آموزش SVM توسط بخشی از مجموعه داده (برای مثال بخش اول) انجام شده و در ادامه مابقی مجموعه داده (برای مثال بخش دوم) با مدل SVM آموزش دیده آزمایش می شود، در نهایت بردارهای پشتیبان و همچنین داده های اشتباه پیش بینی شده برای آموزش SVM مرحله بعد ارسال می شود.
- روش دوم: ابتدا مجموعه داده به چند بخش تقسیم می شود، سپس آموزش SVM توسط بخشی از مجموعه داده (برای مثال بخش اول) انجام می شود. سپس فاصله (Cosine و RBF) مابقی داده ها با بردارهای پشتیبان بدست آمده در مرحله اول تعیین می گردد، در نهایت بردارهای پشتیبان و داده های دارای کمترین فاصله (بیشترین شباهت) برای آموزش SVM مرحله بعد فرستاده می شود.

فصل پنجم

نتایج و ارزیابی

۵-۱ مقدمه

در این فصل نتایج حاصل از اجرای الگوریتم‌های SVM موازی در محیط اسپارک آورده شده و با نتایج سایر مقالات مقایسه و ارزیابی می‌شود. در این پژوهش دقت و زمان آموزش برای بررسی مورد استفاده قرار گرفته است. در ادامه نیز نتایج مربوط به هر مدل از SVM به صورت مجزا آمده است.

۵-۲ مجموعه داده

در این پژوهش ما از مجموعه‌داده IMDB برای بررسی و ارزیابی نتایج استفاده کرده‌ایم. این مجموعه‌داده شامل ۲۵۰۰۰ نظر برای آموزش و ۲۵۰۰۰ نظر برای آزمون می‌باشد. نظرات در مورد فیلم و به زبان انگلیسی می‌باشند. در هر قسمت مجموعه‌داده آموزش و آزمون ۱۲۵۰۰ نظر مثبت و ۱۲۵۰۰ نظر منفی وجود دارد. این مجموعه‌داده در پژوهش‌های دیگری نیز مورد استفاده قرار گرفته است و به همین دلیل برای مقایسه با سایر مقالات نیز مناسب می‌باشد.

۵-۳ مدل پیاده‌سازی

برای پیاده‌سازی این پژوهش از چارچوب اسپارک نسخه ۲,۲,۱ استفاده شده است. مشخصات سیستم رایانه به این شرح می‌باشد:

CPU: Intel(R) Xeon(R) CPU E5-2660 @ 2.20GHz

Cashe: 3MB

RAM: 32GB DDR3

برای پیش‌پردازش نیز از نسخه متلب ۲۰۱۴ استفاده شده است. پارامترهای مربوط به اجرای روش پیشنهادی به شرح جدول ۵-۱ می‌باشد.

جدول ۵-۱: شرح پارامترهای موجود در روش پیشنهادی

پارامتر	مقدار
n	1 – 25
C	1 – 48 (6)
σ	1 – 44 (32)

۴-۵ معیارهای ارزیابی

برای ارزیابی روش‌های پیاده‌سازی شده از معیارهای زیر استفاده شده است. اولین مورد معیار زمان آموزش و زمان آزمون بر حسب ثانیه می‌باشد. معیارهای دیگری که برای ارزیابی عملکرد تشخیص قطبیت نظر استفاده شده است در ادامه آورده شده است:

معیار دقت (ACC):

$$\text{Accuracy} = \frac{\text{True Positive Predictions} + \text{True Negative Predictions}}{\text{Total Predictions}} \quad (۱-۵)$$

معیار تشخیص (SPC) یا نرخ پیش‌بینی‌های صحیح منفی^۱ (TNR):

$$\text{Specificity} = \frac{\text{Number of True Negatives(TN)}}{\text{Number of True Negatives(TN)} + \text{Number of False Positives(FP)}} \quad (۲-۵)$$

معیار حساسیت یا نرخ پیش‌بینی‌های صحیح مثبت^۲ (TPR) یا Recall:

$$\text{Sensitivity} = \frac{\text{Number of True Positives(TP)}}{\text{Number of True Positives(TP)} + \text{Number of False Negatives(FN)}} \quad (۳-۵)$$

^۱ True Negative Rate

^۲ True Pasitive Rate

۵-۵ نتایج

۵-۵-۱ ماشین بردار پشتیبان مبتنی بر رأی‌گیری

در این بخش نتایج مربوط به اجرای روش SVM مبتنی بر رأی‌گیری آورده و به بررسی آن پرداخته شده است. در جدول ۵-۲ نتایج مربوط به این قسمت آورده شده است. در جدول ۵-۳ مقادیر پارامترهای σ و C برای هر سطر متفاوت می‌باشد و مقادیری که بهترین دقت را بدست آورده‌اند، گزارش شده است. بیشینه زمان برای مرحله آموزش، بیشترین زمان سپری شده برای آموزش بین SVM‌های موجود می‌باشد.

جدول ۵-۲: نتایج مربوط به SVM مبتنی بر رأی‌گیری و استاندارد

IMDB مجموعه داده	دقت (ACC)	حساسیت (TPR)	تخصص (TNR)	بیشینه زمان آموزش	زمان آزمون	زمان کل
رأی‌گیری SVM – ۳ بخش	84.872%	0.8386	0.8541	340	4378	5381
رأی‌گیری SVM – ۵ بخش	84.040%	0.8311	0.8502	119	4396	4984
رأی‌گیری SVM – ۷ بخش	83.560%	0.8222	0.8500	63	4381	4797
رأی‌گیری SVM – ۹ بخش	82.964%	0.8158	0.8447	38	4383	4709
رأی‌گیری SVM – ۲۵ بخش	79.928%	0.7830	0.8175	6	4370	4496
SVM استاندارد	85.946%	0.8543	0.8658	3052	4385	7437

همان‌طور که مشخص می‌باشد زمان‌های آموزش نسبت به SVM استاندارد تا حدود ۹۸,۷٪ کاهش داشته است و دقت نیز کاهش داشته است اما تا حدودی با توجه به کاهش زمان، مطلوب است. این زمان گزارش شده برای حالتی است که مجموعه داده به نه قسمت تقسیم شده است. در حالتی که مجموعه داده به ۲۵ قسمت تقسیم شده است زمان آموزش در حدود ۹۹,۸٪ درصد کاهش داشته است

که رقم قابل توجهی می باشد اما با این وجود، دقت ارائه شده دقت مطلوبی نیست و کاهش داشته است. در صورت استفاده از گره های کارگر^۱ برای پردازش، زمان آموزش به میزان گزارش شده، کاهش می یابد. زمان مرحله آزمون نیز با توجه به تعداد گره های کارگر کاهش می یابد.

جدول ۳-۵: نتایج مربوط به SVM مبتنی بر رأی گیری با مقادیر بهینه و متفاوت σ و C

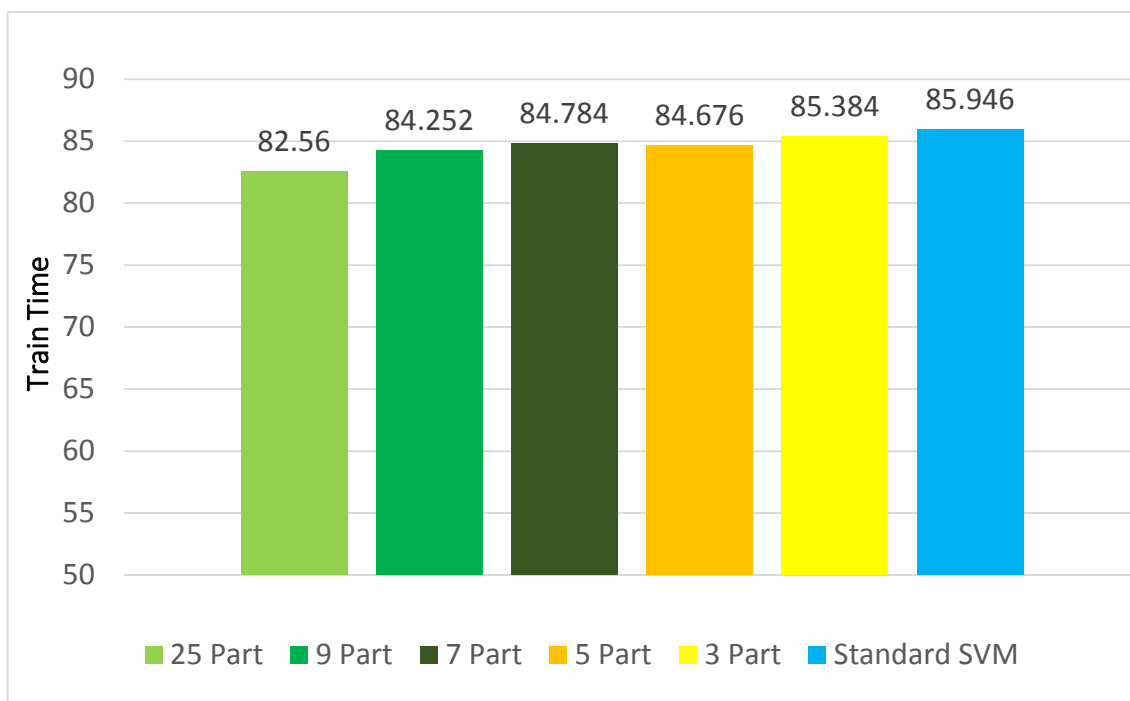
زمان کل	زمان آزمون	بیشینه زمان آموزش	تشخیص (TNR)	حساسیت (TPR)	دقت (ACC)	مجموعه داده IMDB
5663	4382	433	0.8626	0.8454	85.384%	رأی گیری SVM - ۳ بخش
5087	4396	147	0.8557	0.8381	84.676%	رأی گیری SVM - ۵ بخش
5018	4391	92	0.8613	0.8353	84.784%	رأی گیری SVM - ۷ بخش
4823	4373	51	0.8563	0.8297	84.252%	رأی گیری SVM - ۹ بخش
4579	4388	9	0.8412	0.8113	82.560%	رأی گیری SVM - ۲۵ بخش

در نتایج بالا نیز مشاهده می شود که زمان آموزش نسبت به نتایج جدول ۲-۵ کمی افزایش یافته است اما نسبت به SVM استاندارد کاهش قابل توجهی داشته است. این کاهش زمان در کنار این مورد که دقت نیز تا حدود زیادی حفظ شده است، ارزش بیشتری پیدا می کند.

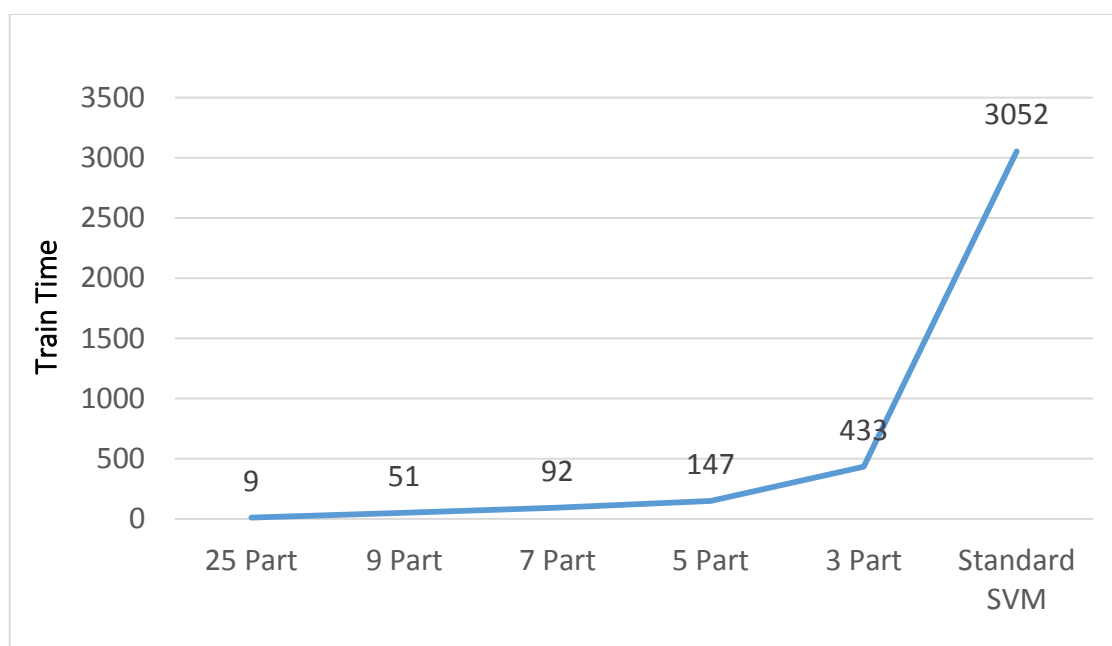
نمودار مربوط به نتایج روش SVM مبتنی بر رأی گیری در شکل ۵-۱ و شکل ۲-۵ نشان داده شده است.

همان طور که در این نمودارهای مشاهده می گردد زمان آموزش رده بند SVM مبتنی بر رأی گیری، نسبت به SVM استاندارد کاهش چشم گیری دارد؛ این مورد به دلیل تقسیم مجموعه داده به بخش های کوچک تر می باشد که همین عامل به دو دلیل باعث افزایش سرعت اجرای روش می شود. همان طور که

^۱ Worker Nodes



شکل ۵-۱: نتایج معیار دقت روش SVM مبتنی بر رأی‌گیری



شکل ۵-۲: نتایج زمان روش SVM مبتنی بر رأی‌گیری

در فصل سوم اشاره شد، مرتبه زمانی رده‌بند SVM وابسته به اندازه مجموعه داده آموزش است و کاهش

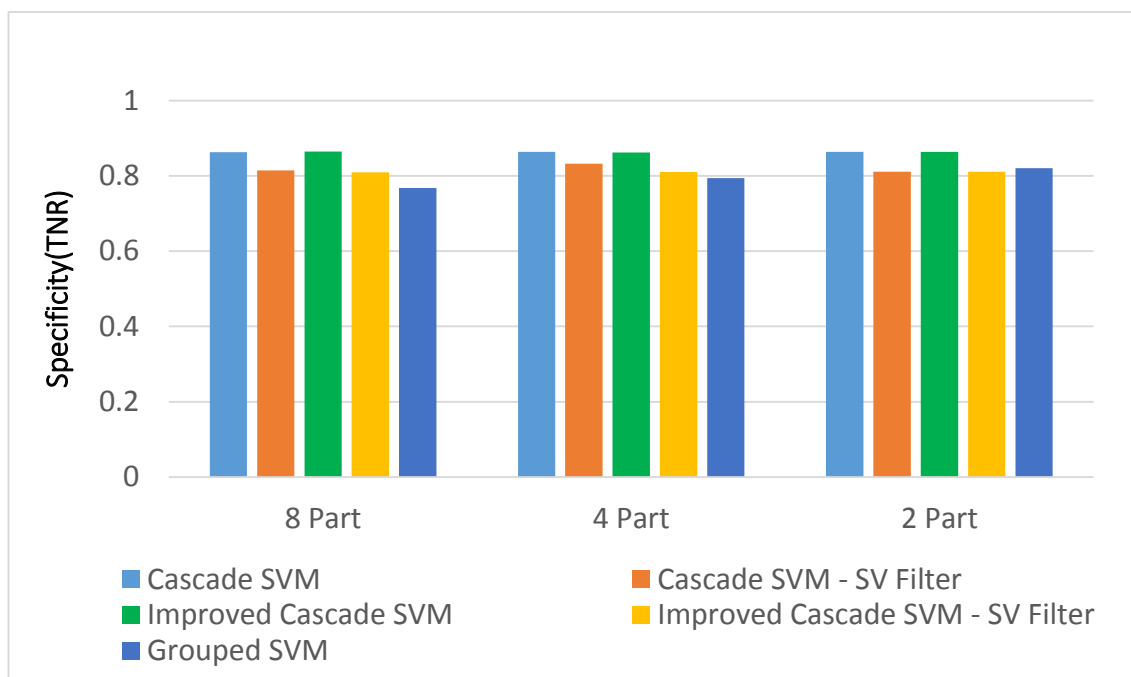
مجموعه داده آموزش باعث افزایش سرعت مرحله آموزش رده بند می گردد. همچنین با تقسیم مجموعه داده به چند قسمت و استفاده از هر قسمت برای آموزش رده بند موجب می شود تا بتوان از پردازش موازی بهره برد. مزیت استفاده از الگوریتم SVM استاندارد نسبت به SVM مبتنی بر رأی گیری و سایر روش های مبتنی بر SVM که در آن مجموعه داده تقسیم می شود، استفاده از کل مجموعه داده برای آموزش رده بند می باشد. این موضوع باعث افزایش دقت رده بند می شود. در روش SVM مبتنی بر رأی گیری به دلیل این که از رأی گیری بین SVM ها استفاده می شود، می توان گفت مشابه استفاده از تمام مجموعه داده برای آموزش رده بند است و به همین دلیل دقت در روش SVM مبتنی بر رأی گیری نیز تا حد مطلوبی حفظ شده است.

۵-۲ مقایسه سایر مدل ها

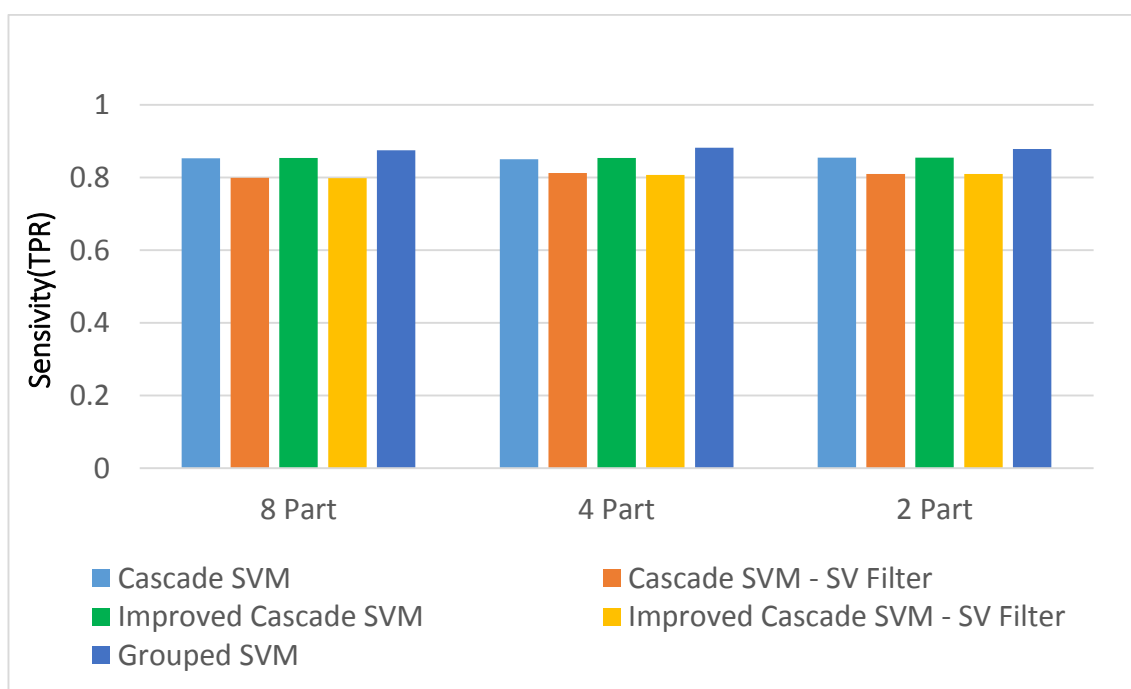
در این قسمت نتایج سایر روش های SVM آبخاری، آبخاری بهبود یافته، آبخاری با پالایش بردارهای پشتیبان و گروهی آورده شده و به مقایسه آن ها می پردازیم. برای ارزیابی، مجموعه داده به دو، چهار و هشت قسمت تقسیم شده است. در جدول ۴-۵ معیار دقت برای این روش ها گزارش شده است.

جدول ۴-۵: مقایسه معیار دقت (ACC) مدل ها

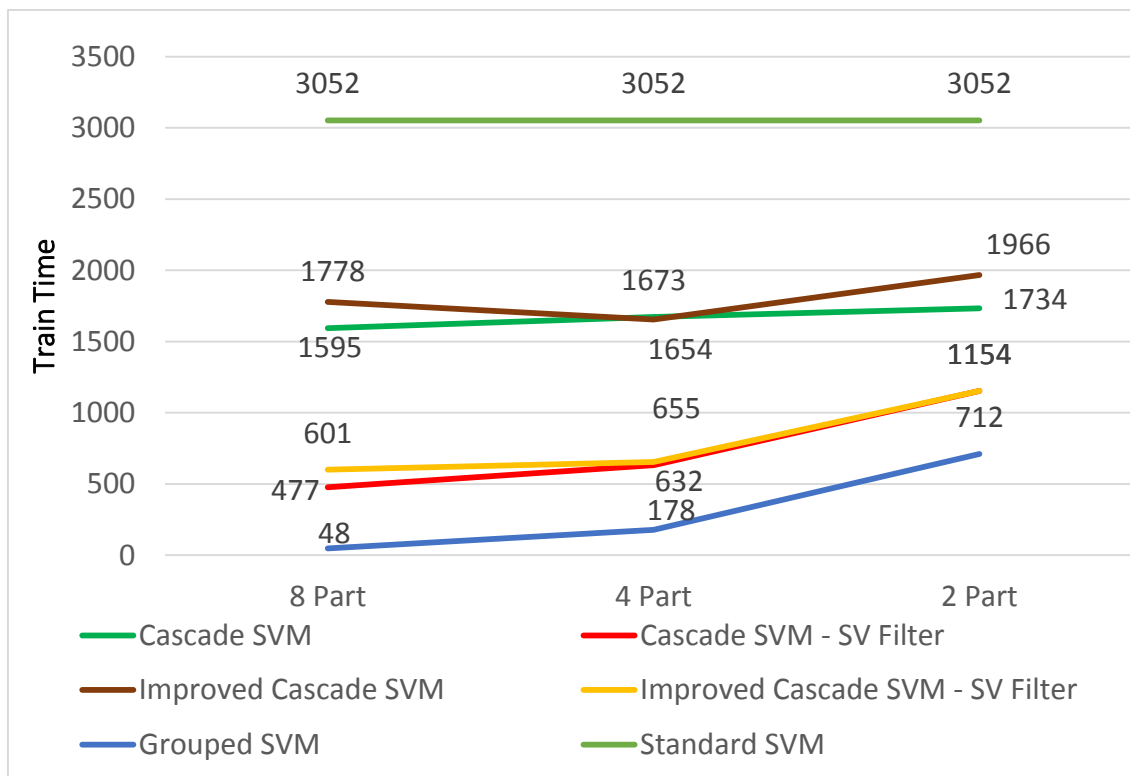
مدل SVM – دقت (ACC)	بخش ۲	بخش ۴	بخش ۸
آبخاری	85.972%	85.844%	85.816%
آبخاری بهبود یافته	85.972%	85.924%	85.916%
آبخاری با پالایش بردارهای پشتیبان	81.08%	82.248%	80.696%
آبخاری بهبود یافته با پالایش بردارهای پشتیبان	81.08%	80.9%	80.396%
گروهی	84.712%	83.276%	81.222%



شکل ۵-۳: مقایسه معیار تشخیص (SPC or TNR) مدل‌ها



شکل ۵-۴: مقایسه معیار حساسیت (TPR) مدل‌ها



شکل ۵-۵: مقایسه معیار زمان آموزش مدل‌ها

شکل ۵-۳ و شکل ۵-۴ به ترتیب مربوط به نتایج معیار تشخیص و حساسیت می‌باشند. همچنین در شکل ۵-۵ نیز زمان آموزش مدل رده‌بند روش‌های مدنظر آمده است. همانطور که مشخص است زمان آموزش برای SVM آبخاری با پالایش بردارهای پشتیبان و آبخاری بهبودیافته با پالایش بردارهای پشتیبان، نسبتاً نزدیک است. این مورد برای SVM آبخاری و آبخاری بهبود یافته نیز صدق می‌نماید. زمان آموزش برای SVM گروهی نسبتاً به سایرین کمتر می‌باشد زیرا نسبت به سایر روش‌ها تعداد زیررده‌بندهای کمتری مورد استفاده قرار می‌گیرد. پس از آن، روش‌های SVM آبخاری با پالایش بردارهای پشتیبان و آبخاری بهبودیافته با پالایش بردارهای پشتیبان زمان کمتری را نسبت به سایرین سپری کرده‌اند. با توجه به نتایج مشاهده می‌گردد روش‌هایی که زمان آموزش کمتری داشته‌اند، دارای دقت کمتری نیز هستند و بالعکس روش‌هایی که دقت بیشتری داشته‌اند، زمان مرحله آموزش رده‌بند آن‌ها بیشتر بوده است.

۵-۳ ارزیابی پارامترهای مدل‌ها

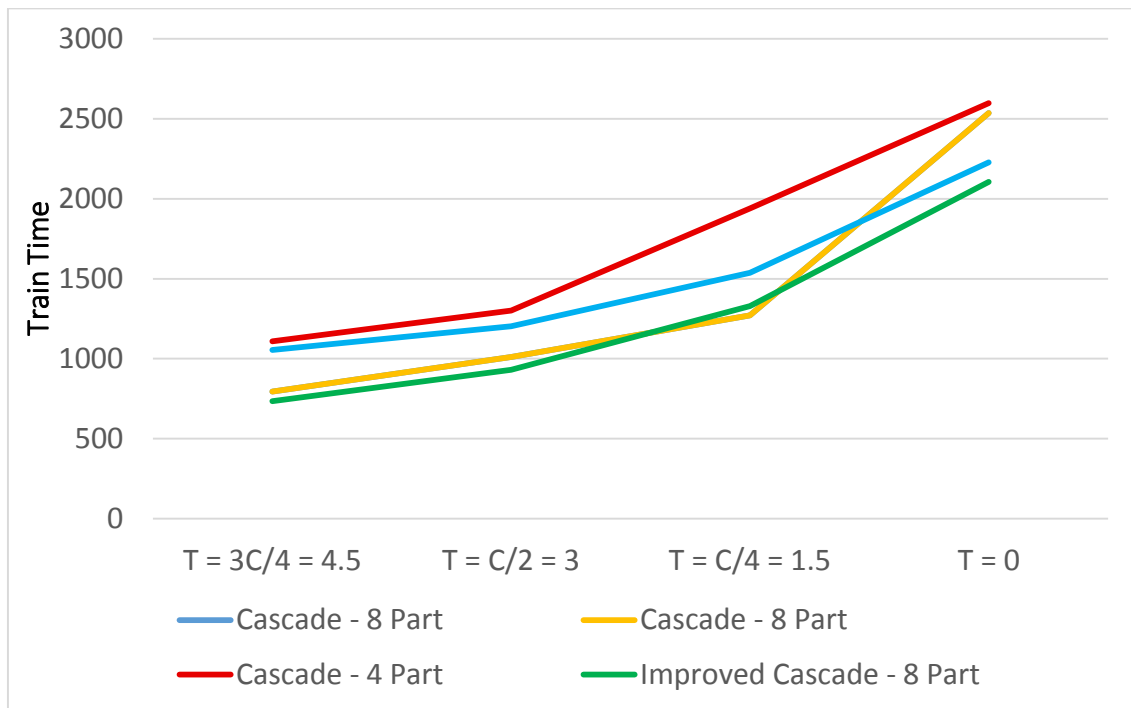
در ارزیابی مدل‌ها با روش‌های مختلف پارامترهای C و σ موثر می‌باشند. برای روش SVM استاندارد، مقادیر $C = 6$ و $\sigma = 32$ دارای بیشترین دقت بودند. برای انجام مقایسه، در سایر روش‌ها نیز نتایج با در نظر گرفتن این مقادیر بدست آورده شده است اما در مقادیر متفاوتی از C و σ ، دقت خروجی بهبود می‌یابد. برای نمونه در روش SVM مبتنی بر رأی‌گیری نتایج با تغییر C و σ ، بهبود یافت. نتایج نیز در بخش مربوط به همین روش گزارش شده است. در روش SVM آبشاری نیز در $C = 4$ ، دقت کمی بهبود داشته است که در جدول ۵-۵ آورده شده است.

جدول ۵-۵: مقایسه معیار دقت و زمان SVM آبشاری برای $C = 4, C = 6$

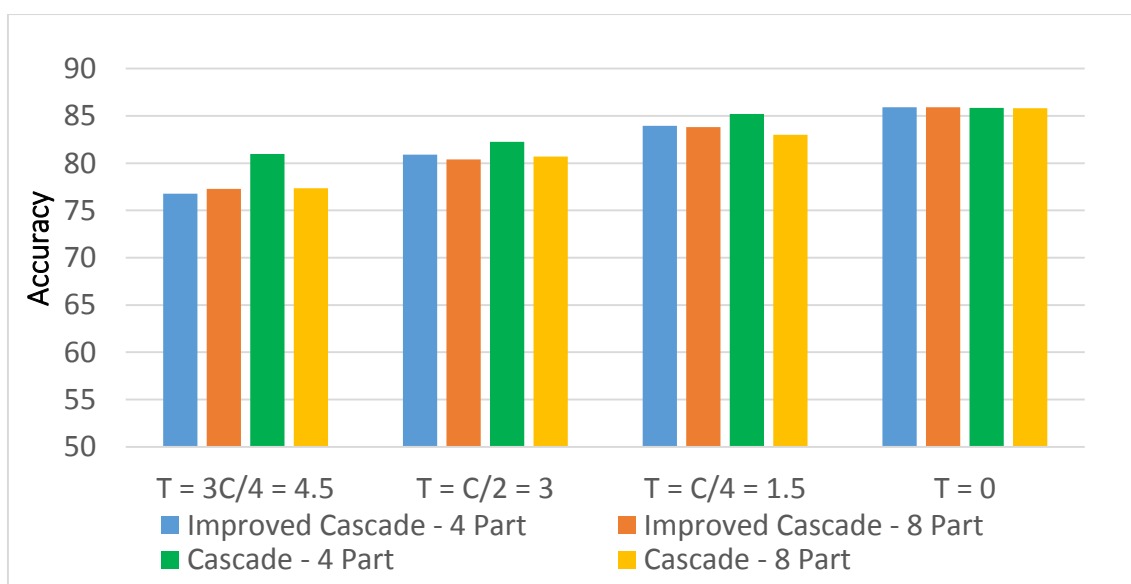
C	زمان کل	زمان آزمون	متوسط زمان آموزش	تعداد SV	دقت (ACC)	مجموعه داده IMDB
6	6027	4432	1595	12159	85.816%	SVM آبشاری - ۸ بخش
6	6092	4419	1673	12269	85.844%	SVM آبشاری - ۴ بخش
6	7542	5576	1966	12443	85.972%	SVM آبشاری - ۲ بخش
4	5868	4409	1459	12697	85.832%	SVM آبشاری - ۸ بخش
4	5934	4422	1512	12768	85.904%	SVM آبشاری - ۴ بخش
4	7251	5460	1791	12889	85.848%	SVM آبشاری - ۲ بخش

در ادامه بررسی پارامترها و تاثیر آن‌ها در روند یادگیری مدل‌ها، به تاثیر مقدار C در روش SVM آبشاری با پالایش بردارهای پشتیبان می‌پردازیم. همان‌طور که در معرفی شد، در روش SVM آبشاری با پالایش بردارهای پشتیبان، تمام بردارهای پشتیبان با ضرایب لاگرانژ بزرگتر از صفر به مرحله بعد فرستاده نمی‌شود. ما با تعیین حد آستانه، سعی در پالایش بردارهای پشتیبان داشته و بردارهای پشتیبان موثرتر که مقادیر ضرایب لاگرانژ بیشتر را دارا می‌باشند، برای مرحله بعد انتخاب می‌نماییم. بنابراین با کاهش

مقدار بردارهای پشتیبان ورودی SVM مرحله بعد، زمان آموزش مدل رده‌بند کاهش یافته و دقت نیز تا حد مطلوبی حفظ خواهد شد. حال ما توجه به نمودارهای شکل ۵-۶ و شکل ۵-۷ این حد آستانه را برابر $T = C/2$ در نظر می‌گیریم. با در نظر گرفتن این مقدار تا حدودی توازنی بین حفظ دقت و کاهش زمان آموزش مدل رده‌بند برقرار می‌شود.



شکل ۵-۶: زمان آموزش SVM آبشاری و آبشاری بهبودیافته با مقادیر مختلف حد آستانه برای ضرایب لاگرانژ



شکل ۵-۷: دقت SVM آبشاری و آبشاری بهبودیافته با مقادیر مختلف حد آستانه برای ضرایب لاگرانژ

۴-۵-۵ مقایسه مدل‌ها

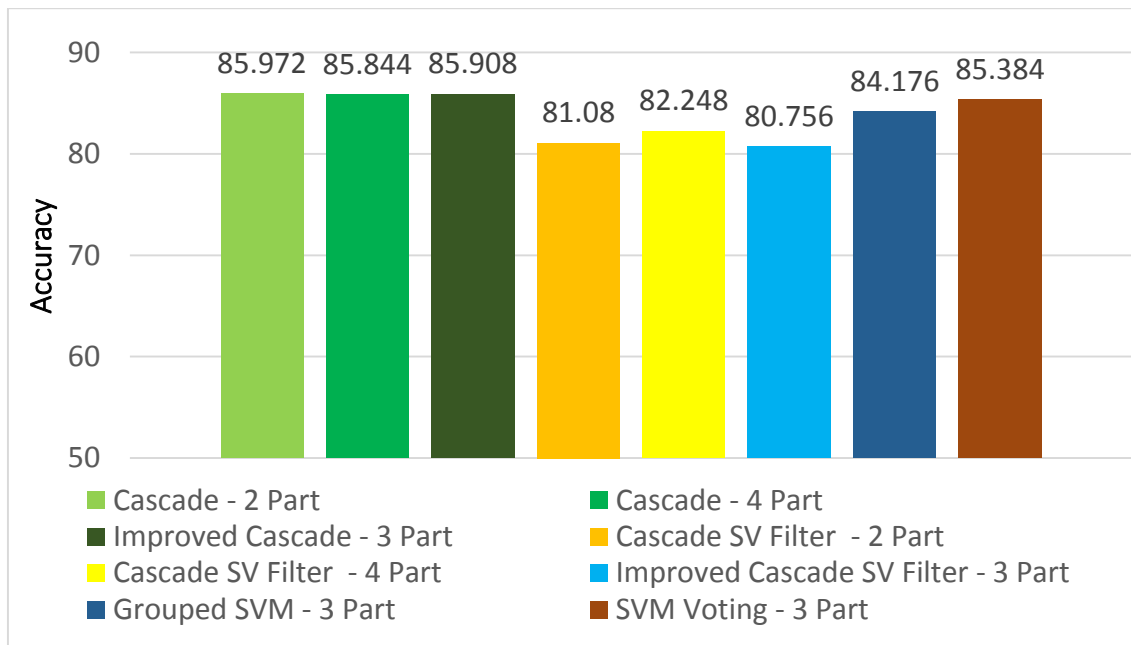
در این قسمت به مقایسه تمام روش‌های پیاده‌سازی شده در این پژوهش می‌پردازیم. برای تمام حالت مجموعه داده به سه بخش تقسیم شده است. برای روش آبخاری و آبخاری با پالایش بردارهای پشتیبان که نمی‌توان مجموعه داده را به سه بخش تقسیم کرد، نتایج مربوط به دو حالت دو و چهار بخشی مجموعه داده آموزش آورده شده است. در جدول ۶-۵ دو معیار تشخیص و حساسیت آورده شده است.

جدول ۶-۵: مقایسه معیار تشخیص و حساسیت مدل‌های مختلف مبتنی بر SVM

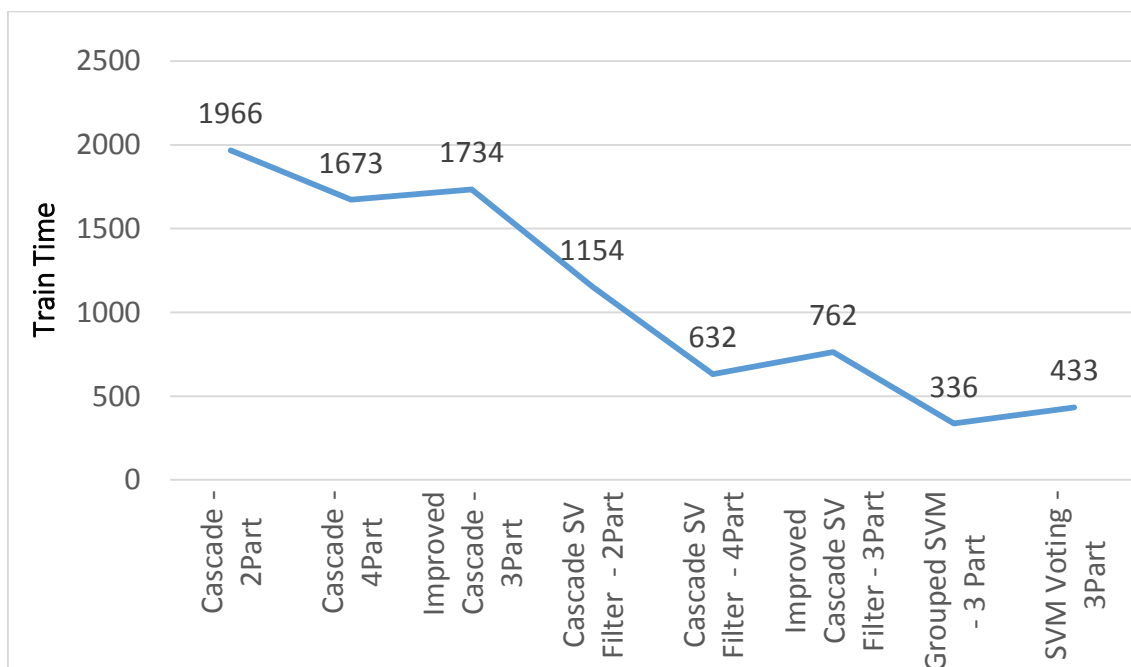
مدل SVM	تشخیص (TNR)	حساسیت (TPR)
آبخاری - ۲ بخش	0.8644	0.8550
آبخاری - ۴ بخش	0.8638	0.8507
آبخاری بهبودیافته - ۳ بخش	0.8647	0.8563
آبخاری با پالایش بردارهای پشتیبان - ۲ بخش	0.8113	0.8102
آبخاری با پالایش بردارهای پشتیبان - ۴ بخش	0.8324	0.8131
آبخاری بهبودیافته با پالایش بردارهای پشتیبان - ۳ بخش	0.8028	0.8124
گروهی - ۳ بخش	0.8035	0.8909
SVM مبتنی بر رأی‌گیری - ۳ بخش	0.8626	0.8454

با توجه به نتایج جدول فوق، بیشترین مقدار معیار حساسیت متعلق به روش SVM گروهی می‌باشد اما مقدار معیار تشخیص این روش نیز جزو کمترین مقادیر می‌باشد. بیشترین مقدار معیار تشخیص نیز به روش SVM آبخاری بهبودیافته می‌باشد که دارای اختلاف کمی با روش SVM آبخاری و SVM مبتنی بر رأی‌گیری می‌باشد. بالا بودن هر دو معیار تشخیص و حساسیت در کنار هم منجر به ایجاد دقت بالا می‌شود.

در ادامه، نتایج مربوط به معیار دقت در شکل ۸-۵ و نتایج مربوط به زمان آموزش در شکل ۹-۵ آورده شده است.



شکل ۸-۵: مقایسه معیار دقت در مدل‌های مختلف مبتنی بر SVM



شکل ۹-۵: مقایسه معیار زمان آموزش در مدل‌های مختلف مبتنی بر SVM

روش‌های SVM آبخاری، آبخاری بهبودیافته و SVM مبتنی بر رأی‌گیری بالاترین دقت را در رده‌بندی

قطبیت نظرات کاربران دارند. کمترین زمان آموزش نیز متعلق به روش‌های SVM گروهی، SVM مبتنی بر رأی‌گیری و آبشاری با پالایش بردارهای پشتیبان دارند. در این بین روش SVM گروهی کمترین دقت را داراست و روش SVM آبشاری با پالایش بردارهای پشتیبان نیز در نسبت به سایر روش‌ها دقت پایین‌تری دارد. روش SVM مبتنی بر رأی‌گیری در بین این روش‌ها دارای بهترین نتیجه می‌باشد. قرار گرفتن این مدل SVM، در بین روش‌های با دقت بالا و همچنین قرار داشتن در بین روش‌هایی با زمان آموزش کم، دلیلی محکم بر این موضوع می‌باشند.

۵-۵-۵ پیاده‌سازی موازی در سایر پژوهش‌ها بر روی مجموعه‌داده IMDB

در تعدادی از پژوهش‌های دیگر نیز روش‌های موازی و توزیع شده بر روی مجموعه‌داده IMDB انجام شده است. در این قسمت نتایج مربوط به این پژوهش‌ها به همراه نتایج پیاده‌سازی SVM استاندارد در محیط اسپارک توسط ما، آورده شده است. نتایج در جدول ۷-۵ آورده شده است.

جدول ۷-۵: نتیجه SVM استاندارد پیاده‌سازی شده و مقایسه با سایر مراجع آزمایش شده بر روی IMDB

مجموعه‌داده IMDB	دقت	زمان در محیط ترتیبی ^۱	زمان در محیط توزیع شده
[48]	60.2%	150,590	37,659
[45]	61.2%	185,000	46,310
[44]	63.7%	222,559	55,633
[46]	56.84% (63.58%)	-	7.94
[47]	82.58%	-	110
استاندارد SVM	85.946%	≈18,000	-

در پژوهش اول، روشی برای تحلیل نظرات انگلیسی مبتنی بر رده‌بند SVM در محیط موازی Cloudera ارائه کرده است [۴۴]. در این پژوهش برای آموزش از ۹۰ هزار نظر برای آموزش مدل رده‌بند استفاده

^۱ Sequential Environment Time

شده است. پژوهش دوم مربوط به الگوریتم Fuzzy C-Mean به صورت موازی می‌باشد که بر روی IMDB و در Cloudera پیاده‌سازی شده است [۴۸]. مجموعه آموزش در این پژوهش شامل ۶۰ هزار نظر می‌باشد. در پژوهش بعد، الگوریتمی به نام STING معرفی و ارائه شده است و به دو صورت موازی و ترتیبی در محیط Cloudera پیاده‌سازی شده است [۴۵]. مجموعه داده آموزشی در این پژوهش نیز شامل ۹۰ هزار نظر انگلیسی می‌باشد. روش‌های SVM و LR در پژوهش چهارم استفاده شده است که نتایج مربوط به SVM در جدول فوق آورده شده است. در این پژوهش از محیط Spark برای پیاده‌سازی استفاده شده است [۴۶]. در آخرین پژوهشی که نتایج آن گزارش شده است، از چند نوع شبکه عصبی برای تحلیل احساسات بهره گرفته شده است. همچنین برای پردازش نیز از واحد گرافیک استفاده شده است [۴۷].

همان‌طور که مشاهده می‌شود روش پیاده‌سازی شده توسط ما از نظر دقت نسبت به سایر مقالات بهتر عمل کرده است و از نظر زمان آموزش نیز بهبود داشته است. در جدول ۵-۸ نتایج جزئی‌تر پیاده‌سازی SVM استاندارد آورده شده است. زمان اجرای کامل الگوریتم SVM تقریباً ۷۵۰۰ ثانیه می‌باشد و ۱۰۵۰۰ ثانیه نیز زمان پیش‌پردازش می‌باشد.

جدول ۵-۸: نتایج SVM استاندارد روش پیاده‌سازی شده

مجموعه داده IMDB	دقت	تعداد SV	زمان آموزش	زمان آزمون	زمان کل
SVM استاندارد	85.946%	12740	3052	4385	7437

۵-۶ نتایج سایر پژوهش‌ها

در این قسمت نیز نتایج سایر پژوهش‌هایی که بر روی مجموعه داده IMDB آزمایش شده‌اند، در جدول ۵-۹ آورده شده است. در تعدادی از این پژوهش‌ها، دقت‌های خوبی گزارش شده است. اما لازم به ذکر است در این مقالات روش‌های استخراج ویژگی متفاوت و پیچیده‌تری استفاده شده است و همچنین در

مورد زمان اجرای روش پیشنهادی اطلاعاتی گزارش نشده است. در ادامه نتایج این روش‌ها آورده شده است.

در مرجع [۴۹] یک روش پیش‌پردازش ارائه شده است که بر روی پنج دیتاست از جمله IMDB اعمال شده است. و برای ارزیابی از سه رده‌بند SVM، Naïve Bayes و ANN استفاده شده است.

در مرجع [۵۰] از شبکه عصبی بازگشتی برای رده‌بندی متن استفاده شده است و بر روی چهار دیتاست انجام شده است. همچنین در این مرجع از چند دیتاست برای آموزش استفاده می‌شود و مرحله تست برای هر کدام از دیتاست‌ها به صورت جداگانه انجام می‌گیرد.

در مرجع [۵۱] نیز از مدلی مبتنی بر شبکه عصبی کانولوشن و یادگیری عمیق استفاده شده است و در قسمت ارزیابی با چند مقاله دیگر مقایسه شده است.

۷-۵ جمع‌بندی

در این فصل نتایج حاصل از پیاده‌سازی روش‌های پیشنهادی در محیط اسپارک ارائه شد. با بررسی و ارزیابی نتایج پژوهش و سایر پژوهش‌های موجود، زمان مربوط به روش‌های ارائه‌شده از سرعت بیشتری برخوردار است و دقت نیز در حد مطلوب و قابل قبول حفظ شده است.

جدول ۵-۹: نتایج تعدادی از سایر پژوهش‌های انجام شده بر روی IMDB

استناد	سال	دقت	روش مورد استفاده	
1	[49] 2017	84.74	SVM	نتایج ارائه شده در [۴۹]
		86.44	SVM + preprocessing	
		74.91	NB	
		76.94	NB + preprocessing	
		78.64	ANN	
		82.03	ANN + preprocessing	
67	[50] 2016	88.5	Train: IMDB	نتایج ارائه شده در [۵۰]
		89.5	Train: IMDB + SST1	
		89.8	Train: IMDB + SST2	
		89.3	Train: IMDB + SUBJ	
480	[52] 2012	83.55	Multinomial NB-uni [52]	نتایج ارائه شده در [۵۱]
		86.59	Multinomial NB-bi [52]	
		86.95	SVM-unigram [52]	
		89.16	SVM-bigram [52]	
		88.29	NBSVM-unigram [52]	
		91.22	NBSVM-bigram [52]	
72	[54] 2012	87.8	WEEB + BoW [54]	
		89.2	WRBB + BoW(bnc) [54]	
957	[55] 2011	87.8	BoW(bnc) [55]	
		88.3	Full + BoW [55]	
		88.9	Full + Unlabeled +	
2382	[56] 2014	92.5	BoW[55]	
13	[57] 2015	94.4	Paragraph Vector(PV) [56]	
		94.5	PV(LogReg) [57]	
2	[51] 2018	93.2	PV (2-Layer MLP) [57] Paper's Approach[51]	

فصل ششم

نتیجه‌گیری

۱-۶ جمع‌بندی و نتیجه‌گیری

تحلیل نظرات کاربران در دنیای امروز، بسیار کارآمد می‌باشد و در زمینه‌های مختلفی مورد استفاده قرار می‌گیرد. همین موضوع باعث توجه پژوهش‌گران به این شاخه شده است. اما مشکلاتی در این شاخه وجود دارد از جمله حجم عظیم نظرات که نیازمند روش‌های قدرتمند و سریع‌تر می‌باشد. در این بین رده‌بندی‌های موازی، برای کاهش زمان آموزش و افزایش سرعت رده‌بندی ارائه شده‌اند. در این پژوهش از SVM‌های موازی برای تحلیل نظرات کاربران استفاده شده است. پیاده‌سازی نیز در چارچوب اسپارک انجام شده است که مناسب برای داده‌های عظیم می‌باشد. در قسمت نتایج و ارزیابی روش پیشنهادی، زمان‌های بدست آمده نسبت به سایر پژوهش‌ها، بهبود قابل توجهی داشته است و دقت بدست آمده نیز نسبت به تعدادی از پژوهش‌ها کمتر و نسبت به تعدادی بیشتر می‌باشد، اما در کل با در نظر گرفتن کاهش زمان مصرفی، قابل قبول می‌باشد.

۲-۶ کارهای آینده

برای کارهای آینده می‌توان همچنان روی کاهش سرعت در کنار حفظ دقت و یا افزایش آن تحقیق نمود. در پژوهش‌های آینده می‌توان از روش‌های پیش‌پردازش دیگر مانند در نظر گرفتن نقش کلمات در جمله، استفاده نمود تا دقت افزایش یابد و بر روی کاهش سرعت آن‌ها سعی نمود.

در این پژوهش از رده‌بند SVM که بهتر در این زمینه عمل می‌نماید و بیشتر مورد بهره قرار می‌گیرد، استفاده شده است. برای کارهای آتی می‌توان بر روی رده‌بندی‌های دیگر همانند Naïve Bayes یا Maximum Entropy و سایر روش‌های رده‌بندی مبتنی بر یادگیری ماشین کار کرد و سعی در موازی‌سازی آن‌ها نمود. همچنین می‌توان از ترکیب چند رده‌بند برای بهبود عملکرد، بهره برد.

در صورت استفاده از SVM و بهره از تابع هسته در فرایند آموزش، می‌توان از ترکیب توابع هسته مختلف برای بهبود دقت، سود برد.

- [1] B. Liu, "Sentiment analysis and opinion mining," *Synthesis lectures on human language technologies*, vol. 5, pp. 1-167, 2012.
- [2] A. Bakliwal, J. Foster, J. van der Puil, R. O'Brien, L. Tounsi, and M. Hughes, "Sentiment analysis of political tweets: Towards an accurate classifier," 2013.
- [3] L. Jiang, M. Yu, M. Zhou, X. Liu, and T. Zhao, "Target-dependent twitter sentiment classification," in *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies-Volume 1*, 2011, pp. 151-160.
- [4] A. Agarwal, B. Xie, I. Vovsha, O. Rambow, and R. Passonneau, "Sentiment analysis of twitter data," in *Proceedings of the workshop on languages in social media*, 2011, pp. 30-38.
- [5] C. Olston and M. Najork, "Web crawling," *Foundations and Trends® in Information Retrieval*, vol. 4, pp. 175-246, 2010.
- [6] K. Guo, L. Shi, W. Ye, and X. Li, "A survey of internet public opinion mining," in *Progress in Informatics and Computing (PIC), 2014 International Conference on*, 2014, pp. 173-179.
- [7] M. F. Porter, "An algorithm for suffix stripping," *Program*, vol. 14, pp. 130-137, 1980.
- [8] B. Pang, L. Lee, and S. Vaithyanathan, "Thumbs up?: sentiment classification using machine learning techniques," in *Proceedings of the ACL-02 conference on Empirical methods in natural language processing-Volume 10*, 2002, pp. 79-86.
- [9] K. Dave, S. Lawrence, and D. M. Pennock, "Mining the peanut gallery: Opinion extraction and semantic classification of product reviews," in *Proceedings of the 12th international conference on World Wide Web*, 2003, pp. 519-528.
- [10] K. Lai and N. Cerpa, "Support vs. confidence in association rule algorithms," in *Proceedings of the OPTIMA Conference, Curicó*, 2001, pp. 1-14.
- [11] G. A. Miller, R. Beckwith, C. Fellbaum, D. Gross, and K. J. Miller, "Introduction to WordNet: An on-line lexical database," *International journal of lexicography*, vol. 3, pp. 235-244, 1990.
- [12] A. Z. Khan, M. Atique, and V. Thakare, "Combining lexicon-based and learning-based methods for Twitter sentiment analysis," *International Journal of Electronics, Communication and Soft Computing Science & Engineering (IJECSCE)*, p. 89, 2015.
- [13] A. Mudinas, D. Zhang, and M. Levene, "Combining lexicon and learning based approaches for concept-level sentiment analysis," in *Proceedings of the First International Workshop on Issues of Sentiment Discovery and Opinion Mining*, 2012, p. 5.
- [14] J. Fang and B. Chen, "Incorporating Lexicon Knowledge into SVM Learning to Improve Sentiment Classification," in *Proceedings of the Workshop on Sentiment Analysis where AI meets Psychology (SAAIP)*, 2011, pp. 94-100.
- [15] A. Bhattacharya and S. Bhatnagar, "Big data and apache spark a review," *Int. J. Eng. Res. Sci*, vol. 2, p. 206, 2016.
- [16] B. E. Boser, I. M. Guyon, and V. N. Vapnik, "A training algorithm for optimal margin classifiers," in *Proceedings of the fifth annual workshop on Computational learning theory*, 1992, pp. 144-152.
- [17] C. Cortes and V. Vapnik, "Support-vector networks," *Machine learning*, vol. 20,

- pp. 273-297, 1995.
- [18] J. Platt, "Sequential minimal optimization: A fast algorithm for training support vector machines," 1998.
 - [19] "The Simplified SMO Algorithm." Available: <http://cs229.stanford.edu/materials/smo.pdf>, 2009 [Sep. 17, 2018].
 - [20] M. Tsytsarau and T. Palpanas, "Survey on mining subjective data on the web," *Data Mining and Knowledge Discovery*, vol. 24, pp. 478-514, 2012.
 - [21] Z. Singla, S. Randhawa, and S. Jain, "Sentiment analysis of customer product reviews using machine learning," in *Intelligent Computing and Control (I2C2), 2017 International Conference on*, 2017, pp. 1-5.
 - [22] A. P. Jain and P. Dandannavar, "Application of machine learning techniques to sentiment analysis," in *Applied and Theoretical Computing and Communication Technology (iCATccT), 2016 2nd International Conference on*, 2016, pp. 628-632.
 - [23] Y. Bidulya and E. Brunova, "Sentiment analysis for bank service quality: A rule-based classifier," in *Application of Information and Communication Technologies (AICT), 2016 IEEE 10th International Conference on*, 2016, pp. 1-4.
 - [24] U. A. Siddiqua, T. Ahsan, and A. N. Chy, "Combining a rule-based classifier with weakly supervised learning for twitter sentiment analysis," in *Innovations in Science, Engineering and Technology (ICISSET), International Conference on*, 2016, pp. 1-4.
 - [25] R. Moraes, J. F. Valiati, and W. P. G. Neto, "Document-level sentiment classification: An empirical comparison between SVM and ANN," *Expert Systems with Applications*, vol. 40, pp. 621-633, 2013.
 - [26] H. Lidong and Z. Hui, "A new short text sentimental classification method based on multi-mixed convolutional neural network," in *2018 IEEE 3rd International Conference on Cloud Computing and Big Data Analysis (ICCCBDA)*, 2018, pp. 93-99.
 - [27] T. U. Haque, N. N. Saber, and F. M. Shah, "Sentiment analysis on large scale Amazon product reviews," in *Innovative Research and Development (ICIRD), 2018 IEEE International Conference on*, 2018, pp. 1-6.
 - [28] S. Sohagir, N. Petty, and D. Wang, "Financial Sentiment Lexicon Analysis," in *Semantic Computing (ICSC), 2018 IEEE 12th International Conference on*, 2018, pp. 286-289.
 - [29] A. B. Bondi, "Characteristics of scalability and their impact on performance," in *Proceedings of the 2nd international workshop on Software and performance*, 2000, pp. 195-203.
 - [30] M. Karanasou, A. Ampla, C. Doulkeridis, and M. Halkidi, "Scalable and Real-time Sentiment Analysis of Twitter Data," in *Data Mining Workshops (ICDMW), 2016 IEEE 16th International Conference on*, 2016, pp. 944-951.
 - [31] A. Hamzehei, M. Ebrahimi, E. Shafiee, R. K. Wong, and F. Chen, "Scalable Sentiment Analysis for Microblogs Based on Semantic Scoring," in *Services Computing (SCC), 2015 IEEE International Conference on*, 2015, pp. 271-278.
 - [32] K. Bhattacharjee and L. Petzold, "What Drives Consumer Choices? Mining Aspects and Opinions on Large Scale Review Data Using Distributed Representation of Words," in *Data Mining Workshops (ICDMW), 2016 IEEE 16th International Conference on*, 2016, pp. 908-915.
 - [33] P. R. Cavalin, M. A. d. C. Gatti, T. Moraes, F. S. Oliveira, C. S. Pinhanez, A. Rademaker, et al., "A scalable architecture for real-time analysis of

- microblogging data," *IBM Journal of Research and Development*, vol. 59, pp. 16: 1-16: 10, 2015.
- [34] B. Yan, Z. Yang, Y. Ren, X. Tan, and E. Liu, "Microblog Sentiment Classification Using Parallel SVM in Apache Spark," in *Big Data (BigData Congress), 2017 IEEE International Congress on*, 2017, pp. 282-288.
 - [35] P. Tripathy, S. S. Rautaray, and M. Pandey, "Map-reduce based parallel support vector machine for risk analysis," in *I-SMAC (IoT in Social, Mobile, Analytics and Cloud)(I-SMAC), 2017 International Conference on*, 2017, pp. 300-303.
 - [36] A. Priyadarshini, "A map reduce based support vector machine for big data classification," *International Journal of Database Theory and Application*, vol. 8, pp. 77-98, 2015.
 - [37] G. Caruana, M. Li, and M. Qi, "A MapReduce based parallel SVM for large scale spam filtering," in *Fuzzy Systems and Knowledge Discovery (FSKD), 2011 Eighth International Conference on*, 2011, pp. 2659-2662.
 - [38] W. Nie, B. Fan, X. Kong, and Q. Ma, "Optimization of multi kernel parallel support vector machine based on Hadoop," in *Advanced Information Management, Communicates, Electronic and Automation Control Conference (IMCEC), 2016 IEEE*, 2016, pp. 1602-1606.
 - [39] H. P. Graf, E. Cosatto, L. Bottou, I. Dourdanovic, and V. Vapnik, "Parallel support vector machines: The cascade svm," in *Advances in neural information processing systems*, 2005, pp. 521-528.
 - [40] H. Du, S. Teng, X. Fu, W. Zhang, and Y. Pu, "A cooperative intrusion detection system based on improved parallel SVM," in *Pervasive Computing (JCPC), 2009 Joint Conferences on*, 2009, pp. 515-518.
 - [41] D. Kun, L. Yih, and A. Perera, "Parallel SMO for training support vector machines," *SMA5505 Project Final Report*, 2003.
 - [42] C. Jiang, T. Wu, J. Xu, N. Zheng, M. Xu, and T. Yang, "A MapReduce-Based Distributed SVM for Scalable Data Type Classification," in *International Conference on Collaborative Computing: Networking, Applications and Worksharing*, 2016, pp. 115-126.
 - [43] N. K. Alham, M. Li, S. Hammoud, Y. Liu, and M. Ponraj, "A distributed SVM for image annotation," in *Fuzzy Systems and Knowledge Discovery (FSKD), 2010 Seventh International Conference on*, 2010, pp. 2983-2987.
 - [44] V. N. Phu, V. T. N. Chau, and V. T. N. Tran, "SVM for English semantic classification in parallel environment," *International Journal of Speech Technology*, vol. 20, pp. 487-508, 2017.
 - [45] N. D. Dat, V. N. Phu, V. T. N. Tran, V. T. N. Chau, and T. A. Nguyen, "STING algorithm used english sentiment classification in a parallel environment," *International Journal of Pattern Recognition and Artificial Intelligence*, vol. 31, p. 1750021, 2017.
 - [46] M. S. Bhuvan, V. D. Rao, S. Jain, T. Ashwin, and R. M. R. Guddeti, "Semantic sentiment analysis using context specific grammar," in *Computing, Communication & Automation (ICCCA), 2015 International Conference on*, 2015, pp. 28-35.
 - [47] X. Xu, H. Ge, and S. Li, "An improvement on recurrent neural network by combining convolution neural network and a simple initialization of the weights," in *Online Analysis and Computing Science (ICOACS), IEEE International Conference of*, 2016, pp. 150-154.
 - [48] V. N. Phu, N. D. Dat, V. T. N. Tran, V. T. N. Chau, and T. A. Nguyen, "Fuzzy C-

- means for english sentiment classification in a distributed system," *Applied Intelligence*, vol. 46, pp. 717-738, 2017.
- [49] R. Coughlin, J.-C. Coetsier, and R. Jiamthapthaksin, "Integrating Labeled Latent Dirichlet Allocation into sentiment analysis of movie and general domains," in *Knowledge and Smart Technology (KST), 2017 9th International Conference on*, 2017, pp. 18-22.
 - [50] P. Liu, X. Qiu, and X. Huang, "Recurrent neural network for text classification with multi-task learning," *arXiv preprint arXiv:1605.05101*, 2016.
 - [51] A. Hassan and A. Mahmood, "Convolutional Recurrent Deep Learning Model for Sentence Classification," *IEEE Access*, vol. 6, pp. 13949-13957, 2018.
 - [52] S. Wang and C. D. Manning, "Baselines and bigrams: Simple, good sentiment and topic classification," in *Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics: Short Papers-Volume 2*, 2012, pp. 90-94.
 - [53] L. Bottou and C.-J. Lin, "Support vector machine solvers," *Large scale kernel machines*, vol. 3, pp. 301-320, 2007.
 - [54] G. E. Dahl, R. P. Adams, and H. Larochelle, "Training restricted boltzmann machines on word observations," *arXiv preprint arXiv:1202.5695*, 2012.
 - [55] A. L. Maas, R. E. Daly, P. T. Pham, D. Huang, A. Y. Ng, and C. Potts, "Learning word vectors for sentiment analysis," in *Proceedings of the 49th annual meeting of the association for computational linguistics: Human language technologies-volume 1*, 2011, pp. 142-150.
 - [56] Q. Le and T. Mikolov, "Distributed representations of sentences and documents," in *International Conference on Machine Learning*, 2014, pp. 1188-1196.
 - [57] J. Hong and M. Fang, "Sentiment analysis with deeply learned distributed representations of variable length texts," Technical report, Stanford University 2015.

Abstract

The ever increasing growth of media and social networks, alongside the spread of analytical commercial and news webpages, have led to an extensive presence of people all around the world in such environments. This has in turn, resulted in enormous propagation and sharing of data and comments. Valuable information could be extracted by analyzing this massive amount of data. But besides this advantage, processing this huge amount of information encounters several challenges. The processing speed of the conventional methods is too slow against the rapid growth of the input data flow. To solve this problem, scalable methods as well as novel tools in the big data processing area, may be employed to accelerate the processing procedure. Although a lot of researches have been already carried out in this area, there is still room for more improvements. Low processing speed, inadequate accuracy rate, high complexity, and domain dependency are among the issues in similar previous studies.

Spark is a relatively novel framework for processing big data, which carries out processes in main memory to increase the processing speed of the algorithm. In this research, the support vector machine (SVM) and Spark are used simultaneously to classify user opinions. The proposed method, along with maintenance of desirable accuracy, is not dependent on a specific domain and can be simply implemented. To benefit from Spark's ability in distributed processing, some methods of parallel SVMs, including cascade SVM, improved cascade SVM, grouped SVM and voting among multiple SVMs have been utilized. To compare the efficiency of the above-mentioned methods, the standard SVM has also been implemented. In the present study, IMDB dataset is used for the evaluation. This dataset contains various different comments in English on the movies and series. Results for classification indicated the best accuracy to be 85.946%. Moreover, the elapsed time to run the standard SVM algorithm in the Spark environment has been dropped more than 8 times compared to similar research results. The elapsed time for parallel SVMs against the standard SVM has also been remarkably improved. In the best case scenario, the training time of the proposed method is reduced by 60 times compared to the standard SVM training time.

Keywords: opinion mining, Support Vector Machine, parallel processing, Spark



Shahrood University of Technology

Faculty of Computer Engineering

M.Sc. Thesis in Artificial Intelligence Engineering

Scalable opinion mining using cluster computing

By:

Mojtaba Asadollahi

Supervisor:

Dr. Hoda Mashayekhi

Sep 2018