

بسم الله الرحمن الرحيم



دانشکده مهندسی کامپیوتر و فناوری اطلاعات

گروه هوش مصنوعی

سیستم مکانیزه طبقه‌بندی اخبار در بستر وب

دانشجو: ملیکا یعقوبی

استاد راهنما:

جناب آقای دکتر حمید حسن پور

استاد مشاور:

جناب آقای دکتر مرتضی زاهدی

پایان نامه ارشد جهت اخذ درجه کارشناسی ارشد

زمستان 1391

دانشگاه صنعتی شاهرود

دانشکده مهندسی کامپیوتر و فناوری اطلاعات
گروه هوش مصنوعی

پایان نامه کارشناسی ارشد خانم ملیکا یعقوبی
تحت عنوان: سیستم مکانیزه طبقه‌بندی اخبار در بستر وب

در تاریخ توسط کمیته تخصصی زیر جهت اخذ مدرک کارشناسی ارشد
مورد ارزیابی و با درجه مورد پذیرش قرار گرفت.

امضاء	اساتید مشاور	امضاء	اساتید راهنما
	دکتر مرتضی زاهدی		دکتر حمید حسن پور

امضاء	نماینده تحصیلات تکمیلی	امضاء	اساتید داور
	مهندس علی بازقندی		دکتر علی اکبر پویان
			دکتر علیرضا الفی

تقدیم به پدر و مادرم،

که از نگاهشان صلابت

از رفتارشان محبت

و از صبرشان ایستادگی را آموختم.

تشکر و قدردانی

سپاسگزار کسانی ام که سرآغاز تولدم بودند،

از یکی زاده شدم و از دیگری جاودانه.

پدر و مادری که تار مویی از آنها به پای من سیاه نماند،

استادی که سپیدی علم و اندیشه را بر سیاهی جهل در زندگیم نگاشت.

با سپاس فراوان از استاد راهنمای بزرگوایم، جناب آقای دکتر حسن پور، به پاس رهنمودها، کلام و حضورشان، که گذران سخت‌ترین روزها را برایم ممکن ساخت. و قدردانی از استاد مشاور ارجمند، جناب آقای دکتر زاهدی که با نظرات سازنده و اندرزهای حکیمانه‌شان، مرا راهنمایی نمودند.

همچنین بر خود لازم می‌دانم از اعضای خانواده و تمامی دوستانی که در اتمام این پروژه به من یاری رساندند، کمال سپاس را به جای آورم.

تعهد نامه

اینجانب **ملیکا یعقوبی** دانشجوی دوره کارشناسی ارشد رشته **هوش مصنوعی** دانشکده کامپیوتر و فناوری اطلاعات دانشگاه صنعتی شاهرود نویسنده پایان نامه سیستم مکانیزه طبقه‌بندی اخبار در بستر وب تحت راهنمایی جناب آقای **دکتر حمید حسن‌پور** 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 می باشد 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 بوده اند در مقالات مستخرج از پایان نامه رعایت می گردد.
- 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 .

تاریخ: 1391/11/15

امضای دانشجو: **ملیکا یعقوبی**

مالکیت نتایج و حق نشر

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

چکیده

رشد روزافزون و گسترش چشمگیر تعداد وبسایت‌ها و حجم بالای داده‌های موجود در اینترنت، یکی از چالش‌هایی است که دهکده جهانی وب با وجود مزایای بیشمارش برای انسان به همراه آورده است. از طرف دیگر، امروزه نیاز و تمایل برای کسب دقیق و به موقع اطلاعات پیرامون حوادث جاری، به منزله برنامه‌ریزی جهت گذران زندگی، چه در سطح شخصی و چه در سطح سازمانی و سیاسی از اهمیت بسیار بالایی برخوردار است.

سرعت بالای گسترش اخبار در اینترنت و عدم امکان کنترل آنها پس از انتشار، افزایش چشمگیر منابع خبری و حجم بالای اخبار منتشر شده در موضوعات و فیلدهای مختلف، امکان پیگیری مداوم اخبار به صورت دستی را غیر ممکن نموده است. این حجم بالا نه تنها در بدست آوردن اطلاعات مورد نیاز به کاربران کمک نمی‌کند، بلکه باعث سردرگمی و ابهام بیشتر آنها نیز می‌گردد، تا آنجا که نیاز به سیستمی برای مدیریت و در نهایت متمایز ساختن اسناد خبری از یکدیگر و دسته‌بندی آنها در گروه‌های متشابه در این بستر به چشم می‌خورد.

در این پروژه برای رفع این معضل، از تلفیق روش‌های وب‌کاوی و دسته‌بندی داده‌های متنی برای مدیریت بهینه اسناد خبری بهره گرفته شده است. قسمتی از متن اسناد خبری در ابتدا به کمک یک خدمت‌گزار وب از سایت‌های خبری مورد تایید و مشخص برداشته شده و پس از آموزش سیستم، پیش‌پردازش‌های مورد نیاز بر روی اسناد قابل تست صورت گرفته، هر سند خبری به دسته‌های مربوطه و انتخابی کاربر ارجاع داده شده و نتیجه به صورت لیستی از اخبار دسته‌بندی شده نمایش داده می‌شود. از دو عامل پویایی اسناد خبری و ناهمگونی‌های موجود در زبان فارسی می‌توان به عنوان اساسی‌ترین چالش‌های موجود در روند کار نام برد.

تا کنون بیشترین تمرکز در مبحث دسته‌بندی اسناد بر روی استفاده از معیارهای شباهت متفاوت و مقایسه عملکرد آنها بر روی ویژگی‌های انتخابی بوده است. همچنین در اکثر موارد از فرکانس تکرار لغات در متن و ارتباط آنها با مجموعه اسناد تحت آزمایش و در چند مورد نیز از تعداد اسناد درون

گروهی به عنوان ویژگی انحصاری متن استفاده شده است. در حالی که در این پروژه، بیشترین تاکید بر روی آماده‌سازی و نرمال‌سازی داده‌های قابل پردازش، ادغام دیکشنری‌های کمکی به منزله افزایش اهمیت کلمات کلیدی در گروه‌ها و در نهایت توجه به فرکانس تکرار لغات هر سند -به صورت مستقل از دیگر اسناد- در گروه‌های مختلف صورت گرفته شده است. همچنین به منزله انتساب سند حد آستانه‌ای برای تعیین حداقل میزان شباهت در نظر گرفته شده ، که تعلق سند به بیش از یک گروه را ممکن می‌سازد. نتیجه بدست آمده حاکی از موفقیت روش پیشنهادی بر روی داده‌های خبری موجود در وب می‌باشد.

کلمات کلیدی: دسته بندی، ریزاسناد خبری ، متن کاوی، وب کاوی

فهرست مطالب

1	فصل اول: مقدمه
2	1-1 پیشگفتار و ضرورت پروژه
3	2-1 تعریف پروژه
6	3-1 ساختار پروژه
7	فصل دوم: مبانی نظری و پایه
8	1-2 مقدمه
8	2-2 دنیای وب
9	1-2-2 وب کاوی
13	2-2-2 روش های دسترسی به اطلاعات وب
24	3-2-2 کدهای وضعیت HTTP و خطاهای سرور
26	3-2 دسته بندی اسناد
26	1-3-2 خوشه بندی
34	2-3-2 دسته بندی
36	3-3-2 دسته بندی اسناد متنی
37	1-3-3-2 آماده سازی و پیش پردازش داده ها
39	2-3-3-2 وزن دهی به کلمات و ایجاد ماتریس
44	3-3-3-2 کاهش داده و کاهش بعد
44	4-3-3-2 معیارهای شباهت
46	4-2 دسته بندی اسناد وب
49	1-4-2 مراحل خوشه بندی نتایج وب
54	2-4-2 معرفی موتورهای دسته بند
59	3-4-2 سیستم دسته بندی اسناد در مقابل موتورهای دسته بند وب

61	ارزیابی عملکرد موتورهای دسته بندی	4-4-2
63	دسته بندی اسناد خبری در وب	5-2
64	تفاوت اسناد وب و اسناد خبری	1-5-2
65	چالش های موجود در وب	2-5-2
65	بررسی دیدگاه های متفاوت	3-5-2
68	فصل سوم: مروری بر تحقیقات	
69	مقدمه	1-3
78	فصل چهارم: روش تحقیق	
79	مقدمه	1-4
80	اخبار موجود در وب	2-4
82	پویش صفحات و جستجو لینک های خبری	3-4
84	مدیریت خطاهای احتمالی به کمک کدهای وضعیت HTTP	1-3-4
86	انتخاب قطعه خبری	2-3-4
88	فاز آموزش	4-4
88	جمع آوری اسناد آموزشی	1-4-4
90	تطبیق و تعیین گروه ها	2-4-4
91	آماده سازی اسناد آموزشی	3-4-4
95	ایجاد پایگاه لغات	4-4-4
95	ادغام پایگاه های لغت کمکی	1-4-4-4
96	تعیین و حذف کلمات پر کاربرد	2-4-4-4
97	تعلق کلمات به گروه ها	3-4-4-4
98	فاز تست	5-4
98	تهیه قطعات	1-5-4

99	آماده‌سازی اسناد آزمایشی	2-5-4
99	وزندهی به کلمات	3-5-4
100	انتساب اسناد و تعیین مقدار آستانه	4-5-4
104	بازسازی مجدد پایگاه لغات	5-5-4
104	نمای کلی برنامه	6-4
105	ارزیابی	7-4
109	فصل پنجم: نتیجه‌گیری و پیشنهادات	
110	مقدمه	1-5
111	تحلیل نتایج	2-5
114	اقدامات آینده	3-5
116	مراجع	

فهرست اشکال

23	تغییر ساختار ارائه اسناد وب از حالت لیست رتبه‌بندی شده به دسته‌های مجزا	شکل (1-2)
40	ماتریس سند-لغت، نمایانگر وزن هر لغت در سند متناظر	شکل (2-2)
50	مراحل خوشه بندی نتایج وب	شکل (3-2)
52	مقایسه میانگین زمان اجرا و دقت برای تعداد دسته و درخواست مشابه در کل متن و قطعات وب	شکل (4-2)
53	میانگین دقت الگوریتم های دسته بندی برای 10 مجموعه سند	شکل (5-2)
54	مقایسه روش های STC , Lingo	شکل (6-2)
58	تصویری از خروجی موتور دسته بند carrot2 در مورد درخواست Iran	شکل (7-2)
66	ایجاد فضایی برای تشخیص گروه خبر توسط موتور دسته بند نتایج	شکل (8-2)
66	دسته بندی موضوعات خبری توسط موتور دسته بند	شکل (9-2)
67	ایجاد جریان خبری توسط موتور دسته بند نتایج	شکل (10-2)
105	نمای کلی برنامه	شکل (1-4)
112	نمونه ای از تناقض در گروه بندی پیشنهادی خبرگزاری‌ها	شکل (1-5)

فهرست جداول

12	جدول (1-2) نحوه تشخیص عملیات مناسب برای حل مشکلات موتورهای جستجو
13	جدول (2-2) ابزارها و روش‌های وب کاوی
28	جدول (3-2) خلاصه ای از روش‌های مورد استفاده در خوشه‌بندی اسناد
41	جدول (4-2) روش‌های وزندهی بر اساس فرکانس تکرار لغت در سند
42	جدول (5-2) روش‌های وزندهی بر اساس معکوس فرکانس تکرار لغت در اسناد و ترکیب آن با روش محلی
43	جدول (6-2) روش‌های وزندهی بر اساس فرکانس تکرار لغت در دسته
45	جدول (7-2) روش‌های اندازه‌گیری معیار شباهت بر مبنای فاصله
46	جدول (8-2) روش‌های اندازه‌گیری معیار شباهت بر مبنای نزدیکی
55	جدول (9-2) خلاصه‌ای از سیستم‌های پژوهشی دسته‌بندی نتایج جستجو
60	جدول (10-2) دسته‌بندی نتایج جستجو در مقابل دسته‌بندی‌های عمومی و سنتی متون
82	جدول (1-4) نرخ به روز رسانی اخبار در سایت‌های پایه
84	جدول (2-4) نمودار خبری مربوط به سه سازمان خبرگزاری پایه
88	جدول (3-4) چکیده ساختار صفحه منبع کد سایت‌های پایه و سیستم نهایی
89	جدول (4-4) تعداد اسناد آموزشی به تفکیک گروه‌ها
90	جدول (5-4) نمودار گروه بندی سازمان‌های خبرگزاری و نمودار خبری نهایی
92	جدول (6-4) یکسان‌سازی حروف مشابه
98	جدول (7-4) تعداد اسناد آزمایشی به تفکیک گروه‌ها
102	جدول (8-4) تشابه/عدم تشابه گروه‌های خبرگزاری، انتسابی و انتخابی
103	جدول (9-4) نمونه ریز متون تست و دسته‌های نسبت داده شده به آنها
105	جدول (10-4) پارامترهای مربوط به معیارهای ارزیابی
107	جدول (11-4) نتایج حاصل از ارزیابی سیستم نسبت به گروه انتخابی و خبرگزاری

فصل اول: مقدمه

1-1 پیشگفتار و ضرورت پروژه

تمایل به کسب اطلاع از اوضاع جامعه و پیگیری روند امور جاری به ویژه در محیط زندگی، از گذشته تاکنون یکی از اساسی‌ترین علاقه‌مندی‌های انسان‌ها را به خود معطوف داشته است. بر همین مبنا صنعت خبرگزاری به یکی از پرقدمت‌ترین صنایع در هر جامعه‌ای تبدیل شده است. امر خبرگزاری بسته به دوره زندگی و امکانات موجود در هر زمان به گونه‌ای متفاوت صورت پذیرفته است. در گذشته به دلیل گستره پراکندگی کم جمعیت، انتقال اخبار به صورت دهان به دهان از فردی به فرد دیگر انجام می‌شد، پس از مدتی، افراد مشخصی وظیفه این امر را به عهده گرفته و با اعلان همگانی در مراکز شهرها اهم اخبار را به گوش سایرین می‌رساندند. بعد از ایجاد صنعت چاپ و گسترش جغرافیایی شهرها، اعلان‌های کاغذی جای نیروهای انسانی را گرفته و اخبار مهم به صورت مکتوب در اختیار شهروندان قرار داده می‌شد و عملاً به جای یک فرد، یک گروه مشترک تحت عنوان یک سازمان خبرگزاری این وظیفه را به عهده گرفتند. وظیفه سازمان خبرگزاری تهیه و تدوین اخبار و قرار دادن آنها در دسترس عموم مردم است. این اخبار نه به معنای اعلامیه و یا اعلان بلکه تنها به صورت گزارشی از سیر وقایع اتفاقیه در زمینه‌های مختلف در اقصی نقاط جهان است که عموماً به صورت روزنامه در اختیار مردم قرار داده می‌شود. در عصر جدید با ظهور و گسترش کاربردی تکنولوژی وب، امر اطلاع رسانی در این بستر و به صورت برخط، بلافاصله پس از اطلاع از وقوع حوادث صورت می‌پذیرد. هرچند این روش مرزهای جغرافیایی و زمانی را شکسته و اطلاع‌رسانی را در کمترین زمان و با کمترین هزینه میسر می‌سازد، اما معضلات بسیاری را نیز به همراه آورده است.

با توجه به ماهیت پویای اطلاعات در صنعت خبرگزاری و به‌روزرسانی ثانیه‌ای اخبار در بستر وب در ساعات مختلف شبانه روز، نیاز به مراجعه مکرر به وب‌سایت‌های خبری برای مطلع شدن از جدیدترین اخبار وجود دارد. اولین راهکار استفاده از افرادی خبره به منزله دسته‌بندی اسناد و تیتراهای خبری سازمان‌های خبرگزاری مختلف و بازپخش آنها در قالب گروه‌های مجزا است. از جمله مشکلات دسته‌بندی دستی مدارک توسط انسان، می‌توان به نیاز برای داشتن نیروی کار زیاد و متخصص در هر

زمینه، خطر تاثیرگذاری تصورات شخصی، عدم تشابه گروه‌بندی توسط دو انسان و در نتیجه زمانگیری و خطاپذیری بالا اشاره کرد.

مشکلاتی از قبیل ساعات کاری و نیروی انسانی محدود، تفاوت زمانی در نقاط مختلف جغرافیایی، تاثیر غیر مستقیم اخبار هر منطقه بر تحولات دیگر کشورها، سرعت بالای گسترش اخبار در اینترنت و عدم امکان کنترل آن پس از انتشار، افزایش چشمگیر منابع خبری با نمودارهای خبری متفاوت و حجم بالای اخبار منتشر شده در آنها در موضوعات و فیلدهای مختلف باعث سردرگمی و ابهام بیشتر در درک اخبار گردیده، تا آنجا که نیاز به یک سیستم برای متمایز ساختن داده‌ها از یکدیگر و دسته‌بندی آنها در گروه‌های متشابه در بستر وب به چشم می‌خورد.

نمایش و ارائه اخبار دسته‌بندی شده علاوه بر کاهش زمان اطلاع از اخبار، سهولت مطالعه و پیگیری عناوین مورد تمایل را نیز در برخواهد داشت.

2-1- تعریف پروژه

سرویس‌های خبری به صورت کلی به سه دسته خبرگزاری رسمی، پورتال اطلاع‌رسانی و سایت‌های خبری-تحلیلی تقسیم می‌شوند. هر چند تنها مورد اول به تولید و پخش خبر می‌پردازد، اما تعداد بالای این سازمان‌ها، روزانه بیش از صدها خبر مختلف را انتشار می‌دهد، تا جایی که می‌توان این چنین بیان نمود که تضاد عمیقی میان حجم اخبار و رشد سریع منابع قابل جستجو و توانایی در کنترل و زمانبندی ورود این اطلاعات به شبکه و توان، زمان و ساعات کاری نیروی انسانی وجود دارد و برای مقابله با آن، طراحی و پیاده‌سازی موتور جستجوگر و دسته‌بند اسناد خبری در وب ضروری است. در این پروژه با استفاده از مفهوم وب‌کاوی و به کمک یک خدمت‌گزار وب به استخراج و کشف خودکار اسناد خبری از وب سایت متعلق به خبرگزاری‌های رسمی می‌پردازیم. بدین منظور پس از دریافت اخبار، جهت ارائه اخبار در دسته‌های مورد نظر از تکنیک‌های دسته‌بندی استفاده نموده و هر کدام از اسناد را پس از انجام پیش پردازش‌های مورد نیاز اولیه و نهایی، داده کاوی مرتبط و استخراج ویژگی‌های مورد نظر، به یکی از شاخه‌های تعریف شده در نمودار خبری نهایی نسبت می‌دهیم.

هدف نهایی از انجام پروژه طراحی و پیاده‌سازی سیستمی برای دسته‌بندی خودکار اسناد خبری به دست آمده از سایت‌های اینترنتی جهت ارائه هرچه سریعتر و بهتر اطلاعات بوده و به عنوان یک سیستم میانی به تسهیل کار افراد می‌پردازد. هر چند در ادامه می‌توان از این روش به منزله شخصی سازی اسناد خبری برای کاربران صفحات خبرگزاری‌ها نیز استفاده نمود.

در ابتدا، کار خود را با تعداد محدود و نابرابر اسناد مربوط به هر دسته از خبرگزاری‌های پایه آغاز می‌نماییم. در روند انجام کار باید به نکاتی همچون نحوه بدست آوردن اسناد خبری، نرمال و یکدست بودن متن‌های ورودی، روش استخراج و پردازش کلمات کلیدی و سنجش میزان تعلق کلمات به دسته‌ها و به‌روزرسانی پایگاه داده کلمات کلیدی پرداخت. عملاً ما فرآیند استخراج متن را برای بکارگیری اطلاعات درون مدارک انجام می‌دهیم که برای این کار باید با جستجوی متن، الگوهای درون آنرا یافته و روابط میان کلمات با یکدیگر و ارتباط با کلمات در دسته‌های نهایی را بیابیم. از آنجا که ماشین به صورت خودکار این عمل را انجام می‌دهد، در نتیجه پروسه‌ای کامل برای یادگیری آن تعریف نموده‌ایم.

قسمت اول به طراحی موتور جستجوگر هماهنگ با سیاست‌های سازمان می‌پردازد. در این موتور به منزله کنترل صحت اخبار تنها از سایت‌های معرفی شده به سیستم استفاده می‌نماییم. یک خدمتگذار وب به صفحه کدینگ هر کدام از وب‌سایت‌ها رفته و با جستجو لینک‌ها و تشخیص لینک‌های مفید از نظر برنامه، اقدام به خواندن و ثبت محتوای آنها می‌کند. بدون شک یکی از چالش‌های پیش رو در این مرحله ایجاد تمایز میان اسناد خبری و غیرخبری موجود در صفحات است، که این عمل با تعریف مقادیر پیش‌فرض در عامل جستجوگر صورت می‌گیرد. به طور کلی هر خبر شامل سه مولفه اساسی تیتر، متن یا بدنه خبر و مرجع اولیه است، که در این پروژه تنها از بخشی از بدنه متن استفاده نموده و کل محتویات را مورد پردازش قرار نمی‌دهیم. در نهایت اسناد به دست آمده از این مرحله در پایگاه داده مربوطه ذخیره گشته و ویژگی‌های مورد نظر آن‌ها استخراج می‌گردد.

قسمت دوم کار اختصاص به مطالعه و تحلیل اسناد آرشیو خبرگزاری به منزله استخراج کلمات کلیدی مورد نظر در هر زیر شاخه و یافتن الگوی پر تکرار خبری دارد. در این بخش با علم بر شاخه‌های متناظر داده‌ها، کلمات کلیدی مربوط به هر بخش با وزن متناسب برای هر گروه بدست می‌آید. برای افزایش دادن جامعیت کار از نظر کاربران انسانی و دایرکتوری‌های معتبر نیز کمک گرفته و به صورت ترکیبی، کلمات حاصل از آنها را به سیستم اضافه می‌کنیم.

سپس با ادغام دو قسمت قبل فاز آموزش و تهیه داده اصلی مورد آزمایش اتمام یافته، برای تست سیستم به دریافت مجموعه خبری جدیدی رو می‌آوریم. برای سنجش میزان تشابه، کلمات کلیدی متن تست ورودی را با پایگاه‌های داده انحصاری گروه‌ها تطابق داده و گروه با بیشترین میزان شباهت را به عنوان گروه برنده انتخاب می‌کنیم. باید توجه داشت که در برخی موارد، سند خبری به بیش از یک گروه تعلق دارد. در انتها نیز بر اساس تمایل کاربر به مشاهده گروه‌های خبری خاص و یا تمام گروه‌ها، اسناد لیبل گذاری شده نمایش داده می‌شوند.

از مزایای این طرح می‌توان به استفاده بهینه از نیروی کار انسانی، به روز رسانی آنلاین اخبار بدون نیاز به حضور فیزیکی در محل، و در نهایت بهبود عملکرد سازمانی و کسب رضایتمندی کاربران اشاره کرد. تاکنون در چندین کشور تلاش برای دسته‌بندی خودکار اسناد موجود و ذخیره شده در پایگاه داده صورت گرفته است که از آن جمله می‌توان به سیستم آزمایشگاهی دسته‌بندی اخبار بر مبنای تیتلر در دو کشور سوئد و کرواسی اشاره نمود. در کشور اسپانیا نیز پروژه ای مشابه تحت استاندارد کنسول بین المللی انتشار اخبار،¹ IPTC در حال تکمیل است. اقداماتی نیز جهت دسته‌بندی اسناد و آرشیوهای خبری در داخل کشور بر روی داده‌های به زبان فارسی صورت گرفته است که اکثر آنها بر روی داده‌های ثابت و غیرپویا آزمایش شده‌اند. باید به این نکته توجه داشت که در این پروژه تنها به موتورهای دسته بند اسناد که به معرفی الگوریتم‌ها و نحوه عملکرد پرداخته‌اند، اشاره نموده و تحلیل و مقایسه خروجی کار با سیستم‌های بسته تجاری مشابه در هیچ یک از بخش‌ها صورت نمی‌پذیرد.

¹ International Press Telecommunication Council

3-1 ساختار پروژه

در این فصل تنها به معرفی و نمایش رویکرد کلی، اهمیت، اهداف و سیر تکاملی پروژه، بدون توجه به ساختارهای فنی و جزئیات کار پرداخته شد. فصل دوم به معرفی مبانی نظری و پایه کار پرداخته و در مواردی به صورت تئوری مقایسه و تحلیل روش‌های مختلف آورده شده است. این فصل به ترتیب مباحث مربوط به موتورهای جستجوگر وب، سیستم‌های دسته‌بندی اسناد و نهایتاً سیستم‌های تلفیقی به نام موتورهای جستجوگر دسته‌بند اسناد در وب را پوشش داده است.

در فصل سوم نگاهی انحصاری به پروژه‌ها و تحقیقات انجام شده در مباحث فصل دوم پرداخته و مروری بر اقدامات پیشین خواهیم داشت. فصل چهارم به شرح روش تحقیق و معرفی مراحل تکمیل کار پرداخته و به صورت تفصیلی جزئیات و معماری بکار گرفته شده در سیستم را نمایش می‌دهد. در انتها نیز ارزیابی سیستم صورت گرفته است.

در فصل پنجم پس از ارایه جمع‌بندی و تحلیل نتایج به دست آمده پروژه، راهکارهایی برای اقدامات آینده جهت ادامه پژوهش پیشنهاد شده است.

فصل دوم: مبانی نظری و پایه

1-2 مقدمه

با افزایش رو به رشد اطلاعات در بستر وب، یافتن سریع اطلاعات مورد نظر درخواستی بسیار دشوارتر از قبل گشته است. یکی از راهکارهای حل این معضل استفاده از تکنیک‌های دسته‌بندی اسناد در گروه‌های مشخص به منزله مدیریت داده‌ها، تسهیل در نمایش و ارائه نتایج به صورت منسجم است. در این بخش نگاهی کلی به دنیای وب و موجودیت‌های آن پرداخته و در ادامه مبحث دسته‌بندی اسناد خبری را مورد مطالعه قرار می‌دهیم. دسته بندی نتایج یکی از قابلیت‌هایی است که به موتورهای جستجوگر امروزی اضافه شده است، در حالی که تا سال‌ها تنها قابلیت موجود در آنها مرتب‌سازی اسناد بوده است. در ادامه با ترکیب دو مورد فوق تعریفی از موتورهای دسته‌بند اسناد را ارائه داده و کار را با محوریت داده‌های خبری و نیازمندی‌های خاص این حوزه پیش خواهیم گرفت.

2-2 دنیای وب

دنیای وب منبع عظیمی از اطلاعات است که روز به روز بر حجم آن افزوده می‌شود. در حال حاضر میلیاردها صفحه که اطلاعات فراوانی از موضوعات مختلف را دربردارند، بر روی سرورهای مختلف ذخیره شده‌اند. این در حالی است که تولید سایت‌های جدید و گسترش سایت‌های موجود نیز به صورت چشمگیری به حجم این اطلاعات می‌افزاید. در بررسی‌های گوناگون انجام شده در زمینه‌ی گسترش وب تخمین زده شده است که روزانه بیش از یک میلیون صفحه به وب اضافه می‌شود و بیش از 600 گیگابایت از صفحات در هر ماه تغییر می‌کنند [1] و [2]. نرخ رشد اطلاعات تا به آنجا رسیده که امروزه مشکل دسترسی به اطلاعات جدی‌تر از نبود اطلاعات است و چالش عمده کاربران، دستیابی مستقیم به اطلاعات مورد نظر است. به عبارت دیگر اگر کاربری دنبال موضوعی خاص باشد، کدام صفحه را باید بخواند؟ از میان این تعداد عظیم صفحات موجود، کدامین صفحه نیاز او را برآورده می‌کند؟ در نتیجه باید ابزاری وجود داشته باشد که به کاربران در پیدا کردن اطلاعات کمک کند که ما آن‌ها را با نام موتور جستجوگر می‌شناسیم. هر چه بر محبوبیت وب افزوده می‌گردد، نیاز به بایگانی کردن اطلاعات آن نیز بیشتر می‌شود. از طرف دیگر داده‌ها به سرعت در حال تغییرند و رشد وب از

لحاظ تغییر اطلاعات و به روز شدن آنها در کسری از ثانیه رخ می‌دهد. به دلیل حجم بالا و امکان دسترسی عمومی سازماندهی بسیار بد و ناکاملی بر روی صفحات و اطلاعات مندرج بر روی آنها وجود دارد، در کل می‌توان چنین بیان داشت که بستر وب کاملاً غیر ساخت یافته است. نیازمندی‌ها کاربران وب نیز بسیار متفاوت بوده و هر کاربری ممکن است تنها به بخش کوچکی از وب علاقمند باشد، این امر مشکل یافتن اطلاعات موردنظر را چندین برابر می‌نماید. وجود این مشکلات کارشناسان را بر آن داشته است تا در حد ممکن به کاوش وب و کشف روابط و ساختارهای آن بپردازند. به صورت کلی وب کاوی به سه دسته مختلف زیر تقسیم می‌گردد که در هر کدام به دنبال روابط و خصوصیات مشخصی از وب هستیم.

1-2-2 وب کاوی

وب کاوی از مهمترین کاربردها در کاوش شبکه گسترده جهانی برای کشف و استخراج الگوهای مفید است. وب علاوه بر اینکه دارای مجموعه عظیمی از اطلاعات است، حاوی مجموعه‌ای پویا از پیوندها برای دسترسی به صفحات وب و استفاده از اطلاعات نیز می‌باشد که یک مجموعه غنی برای داده کاوی ایجاد می‌کند. به طور کلی عملیات مورد نظر در وب کاوی شامل کاوش محتوایی وب، کاوش ساختاری وب و کاوش کاربردی وب است. باید توجه داشت که این شاخه‌های مختلف، کاملاً از هم مجزا نبوده و از این عملیات برای حل مشکلات مطرح شده در کاوش شبکه گسترده جهانی استفاده می‌شود. جدول (1-2) مقایسه بین عملیات موتورهای جستجوگر در هر کدام از جنبه‌های وب کاوی را نشان داده است.

کاوش محتوایی وب

کاوش محتوایی وب کشف اطلاعات مفید از اسناد داده و محتوای وب می‌باشد. آنچه با عنوان محتوای وب شناخته می‌شود، می‌تواند شامل انواع مختلفی از داده باشد. محتوای وب شامل داده‌های غیرساخت یافته مثل متون آزاد و داده‌های نیمه ساخت یافته مثل اسناد HTML و یک نوع داده ساخت یافته تر مثل داده موجود در جداول یا پایگاه داده‌های تولید کننده صفحات HTML است. در

هر حال بیشتر محتوای وب داده غیرساخت یافته است. تحقیقات دربارهٔ بکار بردن تکنیکهای داده کاوی از متون غیر ساخت یافته با عنوان متن کاوی [3] نامیده می شود. بنابراین می توانیم متن کاوی را به عنوان نمونه ای از کاوش محتوای متن بدانیم. ما می توانیم تحقیقات انجام شده در زمینه کاوش محتوایی وب را از دو نقطه نظر متفاوت بررسی کنیم: دیدگاه های بازیابی اطلاعات و پایگاه های داده. هدف کاوش محتوایی وب از دیدگاه بازیابی اطلاعات اساساً کمک به بهبود جستجوی اطلاعات یا فیلتر کردن اطلاعات برای کاربران و معمولاً بر اساس پروفایل های استنباط شده یا تقاضا شده کاربران است. در حالیکه هدف کاوش محتوایی وب از نقطه نظر پایگاه داده اساساً تلاش برای مدل کردن داده روی وب و استفاده از آن به نحوی است که بتوان به جای جستجوهای ساده که تنها براساس کلمات کلیدی هستند پرس و جوهای پیشرفته تری داشت [4].

کاوش ساختاری وب

کاوش ساختاری وب به دنبال کشف مدلی است که در زیرساختار پیوندهای وب وجود دارد. کاوش ساختاری وب می تواند برای کشف اعتبار صفحات وب استفاده شود. تعیین اعتبار صفحات وب در پیوندهایی که به آن صفحه وجود دارد، پنهان است. وب تنها شامل یک سری صفحات جدا از هم نیست بلکه مهمترین ویژگی آن پیوندهایی است که صفحات را به هم متصل می کنند. این پیوندها دارای مقدار زیادی اطلاعات نهفته هستند که در ایجاد این تصور که اعتبارسنجی صفحات به صورت خودکار انجام می شود، کمک زیادی می کنند. در واقع وقتی نویسنده یک صفحه وب پیوندی به صفحه دیگری در صفحه اش می گذارد این پیوند به منزله مهر تأییدی است که این نویسنده روی آن صفحه زده است. هر چه یک صفحه وب دارای پیوندهای بیشتری از جاهای مختلف باشد، این نشان دهنده اهمیت بیشتر آن صفحه است و طبعاً اعتبار بیشتر آن صفحه را نشان می دهد. بنابراین پیوندهای وب اطلاعات غنی و کاملی در رابطه با مرتبط بودن، کیفیت و ساختار محتوایی وب ایجاد می کند و یک منبع غنی برای وب کاوی محسوب می شود. اما بر خلاف اظهار نظرهایی که ارائه می شود، پیوندهای وب ویژگی های خاصی دارند که استفاده از آنها را برای تعیین اعتبار صفحات با مشکل مواجه می کند.

اول اینکه هر پیوند نشان‌دهنده تأییدی که ما به دنبال آن هستیم نیست. در واقع بعضی از این پیوندها به هدف دیگری در نظر گرفته شده‌اند. دوم اینکه ممکن است به دلایل تجاری یا رقابتی یک صفحه دارای پیوندی از طرف نویسندگان رقیب نباشد که این مسئله سبب پایین آمدن اعتبار آن صفحه می‌شود [4].

کاوش کاربردی وب

در کاوش کاربردی وب رکوردهای گزارشهای وب برای کشف الگوهای دسترسی کاربر به صفحات وب بررسی می‌شوند. در حالیکه کاوش محتوایی وب و کاوش ساختاری وب از داده اولیه و واقعی وب استفاده می‌کنند، کاوش کاربردی وب داده ثانویه مشتق شده از فعل و انفعالات کاربران با وب را می‌کاود. تحلیل و پیدا کردن نظم موجود در رکوردهای گزارشهای وب می‌تواند مشتریهای احتمالی موجود در تجارت الکترونیکی را مشخص کند، کیفیت سرویس‌های اطلاعاتی اینترنت را برای کاربران افزایش دهد و کارایی سیستم سرویس دهنده وب را بهتر کند [4].

فعالیت های موتورهای جستجوگر	کاوش محتوایی وب	کاوش ساختاری وب	کاوش کاربردی وب
دسترسی به اطلاعات مورد نظر کاربر از طریق موتورهای جستجو با وجود بزرگی و پویایی وب	روشهای متن کاوی ، بخش بندی اسناد وب، روشهای بازیابی اطلاعات		
افزایش کیفیت سرویسهای اطلاعاتی اینترنت		کشف اعتبار صفحات از طریق ساختار پیوندهای هر صفحه	کاوش گزارشهای دسترسی به سرویس دهنده های وب و متا اطلاعات
کاربر فقط بخشی از وب را که به مربوط اوست ببیند.	بخش بندی اسناد وب		کاوش گزارشهای دسترسی به سرویس دهنده های وب و متا اطلاعات
مدل کردن کاربر بر اساس علایق و اهداف او			کاوش گزارشهای دسترسی به سرویس دهنده های وب و متا اطلاعات
اعتبار دهی صفحات وب		کشف اعتبار صفحات از طریق ساختار پیوندهای هر صفحه	
سازماندهی مناسب وب	طبقه بندی صفحات وب	بخش بندی وب از طریق بررسی ساختار پیوندها بین صفحات مختلف	کاوش گزارشهای دسترسی به سرویس دهنده های وب و متا داده

جدول (1-2): نحوه تشخیص عملیات مناسب برای حل مشکلات موتورهای جستجو [4]

ابزارها و روش‌های مناسب برای انجام هر یک از عملیات وب کاوی نیز در جدول (2-2) نشان داده شده است.

روشها و ابزار وب کاوی	تکنیکها	کاربرد
روشهای وب کاوی	روشهای متن کاوی	کاوش محتوای متون غیر ساختیافته وب برای بهبود کارایی موتورهای جستجو
	روشهای بازیابی اطلاعات	کاوش محتوای اسناد وب برای بهبود کارایی موتورهای جستجو، پیدا کردن الگوهایی در متون اسناد وب
	روشهای بخش بندی متن	کاوش محتوای اسناد وب برای بهبود کارایی موتورهای جستجو
	کاوش قوانین انجمنی	مدل کردن کاربر، مشخص کردن شمای وب سایت
	روشهای طبقه بندی	طبقه بندی اسناد وب
	روشهای خوشه بندی	خوشه بندی اسناد وب و تعیین بخشهای اولیه به منظور بخش بندی
	روشهای یادگیری آماری	بخش بندی اسناد وب، مدل کردن کاربر
تجارت	شبکه های عصبی	بخش بندی اسناد وب

جدول (2-2): ابزارها و روش‌های وب کاوی [4]

2-2-2 روش های دسترسی به اطلاعات وب

پس از تعریف محیط وب و انواع کاوش‌های ممکن در آن، نوبت به معرفی ابزارها و راه‌های دسترسی به اطلاعات در این بستر است. بدون شک آسان‌ترین روش در این زمینه مراجعه به صفحه حاوی اطلاعات مورد نظر با علم بر آدرس آن است. اما به دلیل حجم بالا اطلاعات و تغییرات متناوب صفحات، امکان اطلاع از تمام آنها وجود ندارد، در نتیجه باید به خدمات جستجو در وب مراجعه نمود. خدمات جستجویی که در وب ارائه می‌شود را می‌توان در سه گروه اصلی فهرست (دایرکتوری)ها و موتورهای جستجوگر و دروازه‌های وب دسته‌بندی نمود.

تفاوت اصلی این گروه‌ها در نحوه تشکیل پایگاه داده و جمع‌آوری اطلاعات آنهاست. در فهرست‌ها، اینکار به عهده انسان است اما در موتورهای جستجوگر جمع‌آوری اطلاعات پایگاه داده را نرم افزارها انجام می‌دهند.

دایرکتوری^۲

ساختار دایرکتوری‌ها در واقع نوعی رده‌بندی عمومی را در اختیار می‌گذارند که از مفاهیم کلی و انتزاعی در سطوح بالای رده‌بندی شروع می‌شود و به تدریج با مفاهیم خاص گسترش پیدا می‌کند. برای هر یک از گره‌های این ساختار، تعدادی صفحه‌ی وب مرتبط با آن مفهوم، به آن گره نسبت داده شده است. همچنین به ازای هر گره، یک یا چند جمله وجود دارد که آن را توصیف می‌کند. از جمله معروف ترین این ساختارها می‌توان دایرکتوری‌های وب Dmoz و Yahoo را نام برد. این ساختارها توسط انسان ایجاد شده‌اند و بطور مکرر مثلاً هر هفته یکبار به روز می‌شوند. فهرست هرگز از وجود سایت شما اطلاع نمی‌یابد مگر زمانی که شخصی آن را معرفی نماید. بعد از معرفی، ویراستار فهرست به سایت شما مراجعه نموده، در صورت رعایت قوانین فهرست و انتخاب گروه مناسب، سایت شما را به پایگاه داده فهرست اضافه می‌نماید. در واقع شما باید سایت خود را با عنوان و توضیحی مناسب به فهرست‌ها معرفی نمایید و بهترین گروه ممکن را برای سایت خود در نظر بگیرید. این کار بسیار مهم است زیرا عموماً فهرست‌ها همین عنوان و توضیح را به همراه آدرس صفحه اول سایت‌تان در پایگاه داده خود قرار می‌دهند.

یکی از مشکلات این ساختارها این است که معمولاً بصورت متوازن نیستند. به این معنی که گره‌های خاصی بسیار گسترده شده‌اند و گره‌های مربوط به برخی مفاهیم کمتر گسترش یافته‌اند. در نتیجه در بسیاری از موارد این ساختارها نمی‌توانند کمک چندانی داشته باشند.

مشکل عمده‌ی این ساختارها جهت استفاده از آن‌ها به عنوان یک منبع معنایی لغوی آن است که هر گره اطلاعات بسیار کمی در مورد مفهوم نسبت داده شده به آن در اختیار می‌گذارد و به منظور

^۲ Directory

دستیابی به اطلاعات بیشتر باید عمل خزش وب را روی ساختار انجام داد که زمانبر است و به پردازش صفحات وب نیاز دارد. حتی در صورت انجام این کار، به دلیل نویز فراوان موجود در صفحات وب، دریافت اطلاعات در مورد آن مفهوم از وبسایت‌های نسبت داده شده به آن مشکل است. به علاوه مشکل پوشش نیز در این ساختار وجود دارد، به این معنی که لیست سایت های هر گره ممکن است تمامی جنبه های آن مفهوم را پوشش ندهند [5].

دروازه وب³

ترجمه کلمه پرتال به فارسی در فرهنگ‌های لغت «دریچه»، «درگاه» و «مدخل» ذکر شده اما در تکنولوژی اطلاعات، پورتال صفحه وب واسطی است که امکان دسترسی آسان را به هر چیزی که کاربر، برای انجام وظیفه یا خواسته‌اش نیاز دارد - بدون توجه به اینکه محل فیزیکی آن کجاست-، فراهم می‌کند. دروازه‌ای که به تبعیت از ذات اصلی خود یعنی همان دسترسی آسان، دسترسی ما را به وب و در نهایت اینترنت راحت‌تر از پیش می‌کند و محلی است برای به اشتراک گذاری خدمات و محتوا توسط چندین وب سایت قدرتمند. در واقع دروازه وب به پایگاه وبی گفته می‌شود که اطلاعات را با توجه به موضوع آنها دسته‌بندی می‌کند تا به این صورت به یافتن آنچه به دنبال آن هستیم، کمک کند [6].

وب پورتال یک محصول خاص نیست و همانند وب سایت یک ماهیت استاندارد و تشکیل شده از دو زیرساخت خدمات و محتویات است که این ماهیت، خروجی چندین استاندارد، دیتا و موتور خاص، زیر نظر یک یا چند مجموعه نظارتی مشخص و همگام است.

ویژگی‌های اصلی یک پورتال عبارت است از تجمع اطلاعات، هدف دار بودن اطلاعات، در دسترس بودن اطلاعات و ایجاد دریچه ورود منحصر به فرد. هدف وب پورتال‌ها طبقه‌بندی کردن اطلاعات و نیز تعریف دسترسی آسان به آنهاست و برای جلوگیری از پراکندگی در چگونگی یافتن، دسترسی و

^۳ Portal

نگهداری اطلاعات توسط کاربر و به صورت خلاصه جلوگیری از سردرگمی آنها، به صورت همزمان

سرویس‌های اصلی و جانبی را در اختیار کاربران قرار می‌دهند

پورتال همواره ما را به سایت‌ها یا پورتال‌های دیگر راهنمایی می‌کند و به خودی خود تنها یک راهنما

است. برای همین است که در بعضی از موارد به پورتال‌ها، صفحات زرد اینترنتی⁴ نیز می‌گویند. آنها

عموماً "حاوی مطالبی هستند که جنبه اطلاعات عمومی دارد. داده‌هایی که از منابع مختلف بر روی

یک پورتال جمع‌آوری می‌شوند، معمولاً دارای پراکندگی فراوانی هستند. به همین علت، در بسیاری از

پورتال‌ها، ابزارهایی مانند دایرکتوری قرار داده می‌شود تا این اطلاعات را طبقه‌بندی نماید. از سوی

دیگر داده‌های قرار داده شده بر روی یک وب سایت، اولاً "از منابع محدودتری تامین می‌شوند و ثانیاً"

دارای پراکندگی زیادی نبوده و حول یک محور و موضوع مشخص می‌باشد. پورتال یک سیستم کاربر

محور می‌باشد. تامین اطلاعات پورتال را می‌توان یک مرکز ارائه خدمات و اطلاعات اینترنتی دانست

که بر چهار پایه اصلی انطباق‌پذیری⁵، اختصاصی کردن⁶، توصیه و معرفی کردن محتوا⁷ و خلاصه

کردن مطالب⁸ استوار است.

موتورهای جستجوگر⁹

موتور جستجو ابزاری است که امکان جستجوی اطلاعات در وب را فراهم می‌کند [7]. موتور جستجوگر

می‌تواند از وجود سایت شما اطلاع یابد، اگر راه ورود آن فراهم شده باشد. در واقع نرم افزار موتور

جستجوگر هر لحظه در حال وبگردی و به روز رسانی اطلاع قدیمی و همینطور افزودن اطلاعات جدید

به پایگاه داده موتور جستجوگر است.

⁴ Internet yellow pages

⁵ Customization

⁶ Personalization

⁷ Recommendation

⁸ Excerpt content

⁹ Search engines

اولین نسل موتور جستجوگر در سال 1994 با WebCrawler and Lycos کار خود را آغاز کرده و تنها به نمایش یک صفحه اطلاعات مرتبط از لحاظ محتوی و فرمت می‌پرداختند و بیشتر درخواست‌های اطلاعاتی را پوشش می‌دادند.

نسل دوم موتور جستجوگر در سال 1998 با شرکت گوگل شروع به کار کرده و عمل اولویت‌بندی نتایج را به سیستم اضافه نموده است. در این سیستم‌ها از اطلاعات خاص خارج از محتویات صفحه مانند آنالیز لینک‌ها، جریان داده مبتنی بر کلیک و ... نیز استفاده شده است. تمامی درخواست‌های اطلاعاتی و دنباله روی¹⁰ در این نسل پوشش داده شده است.

نسل سوم در اوایل سال 2000 کار خود را با تلاش برای ادغام منابع چندگانه از شواهد و دربرگیری تمام درخواست‌های ممکن آغاز نموده است. انتظار می‌رود که در نسل آینده شاهد موتورهای جستجوگر دسته‌بند و معناگرا بوده و بتوانیم با استفاده عمومی از آنها، نحوه و قدرت دسترسی کاربران به اطلاعات مورد نیازشان را افزایش دهیم.

بازیابی مدرن اطلاعات وب بیشترین تلاش خود را معطوف به استفاده از نتایج کلاسیک موجود نموده و بیشتر دست به توسعه ابتکاری و ادغام روش‌های سنتی نموده است. همه موتورهای جستجو در زمان پاسخ‌گویی به جستجوهای کاربران، تنها در پایگاه داده‌ای که در اختیار دارند به جستجو می‌پردازند و نه در وب! موتور جستجوگر به کمک بخش‌های متفاوت خود، اطلاعات مورد نیاز را قبلاً جمع‌آوری، تجزیه و تحلیل می‌کند، آنرا در پایگاه داده‌اش ذخیره می‌نماید و به هنگام جستجوی کاربر تنها در همین پایگاه داده می‌گردد.

اجزاء موتورهای جستجوگر

بخش‌های مجزای یک موتور جستجوگر شامل عنکبوت¹¹، خزنده¹²، بایگانی‌کننده¹³، پایگاه داده¹⁴ و سیستم رتبه‌بندی¹⁵ می‌گردد، که در ادامه به شرح مختصر آنها می‌پردازیم.

¹⁰ Navigation

¹¹ Spider

¹² Crawler

¹³ Indexer

عنکبوت

اسپایدر یا روبات، به صفحات مختلف سر می‌زند، محتوای آنها را می‌خواند، لینکها را دنبال می‌کند، اطلاعات مورد نیاز را جمع‌آوری می‌کند و آنها را در اختیار سایر بخش‌های موتور جستجوگر قرار می‌دهد. اسپایدر کدهای HTML صفحات را می‌بیند.

اسپایدر، به هنگام مشاهده صفحات، بر روی سرورها رد پا برجای می‌گذارد. شما اگر اجازه دسترسی به آمار دید و بازدیدهای صورت گرفته از یک سایت و اتفاقات انجام شده در آن را داشته باشید، می‌توانید مشخص کنید که اسپایدر کدام یک از موتورهای جستجوگر صفحات سایت را مورد بازدید قرار داده است.

کار یک اسپایدر، بسیار شبیه کار کاربران وب است. همانطور که کاربران، صفحات مختلف را بازدید می‌کنند، اسپایدر هم درست این کار را انجام می‌دهد با این تفاوت که اسپایدر کدهای HTML صفحات را می‌بیند اما کاربران نتیجه حاصل از کنار هم قرار گرفتن این کدها را، اسپایدرها کاربردهای دیگری نیز دارند، به عنوان مثال عده‌ای از آنها به سایت‌های مختلف مراجعه می‌کنند و فقط به بررسی فعال بودن لینک‌های آنها می‌پردازند و یا به دنبال آدرس ایمیل می‌گردند.

خزنده

کراولر، نرم افزاری است که به عنوان یک فرمانده برای اسپایدر عمل می‌کند. آن مشخص می‌کند که اسپایدر کدام صفحات را مورد بازدید قرار دهد. در واقع کراولر تصمیم می‌گیرد که کدام یک از لینک‌های صفحه‌ای که اسپایدر در حال حاضر در آن قرار دارد، دنبال شود. ممکن است همه آنها را دنبال کند، بعضی‌ها را دنبال کند و یا هیچ کدام را دنبال نکند. کراولر، ممکن است قبلاً "برنامه‌ریزی شده باشد که آدرس‌های خاصی را طبق برنامه، در اختیار اسپایدر قرار دهد تا از آنها دیدن کند. دنبال کردن لینک‌های یک صفحه به این بستگی دارد که موتور جستجوگر چه حجمی از اطلاعات یک

^{۱۴} Database

^{۱۵} Ranker

سایت را می‌تواند (می‌خواهد) در پایگاه داده‌اش ذخیره کند. همچنین ممکن است اجازه دسترسی به بعضی از صفحات به موتورهای جستجوگر داده نشده باشد.

شما به عنوان دارنده سایت، همان طور که دوست دارید موتورهای جستجوگر اطلاعات سایت شما را با خود ببرند، می‌توانید آنها را از بعضی صفحات سایت‌تان دور کنید و اجازه دسترسی به محتوای آن صفحات را به آنها ندهید. موتور جستجو، به صورت معمول قبل از ورود به هر سایتی ابتدا قوانین دسترسی به محتوای سایت را (در صورت وجود) در فایلی خاص بررسی می‌کند و از حقوق دسترسی خود اطلاع می‌یابد. به عمل کراولر، خزش نیز می‌گویند.

بایگانی کننده

تمام اطلاعات جمع‌آوری شده توسط اسپایدر در اختیار ایندکسر قرار می‌گیرد. در این بخش اطلاعات ارسالی مورد تجزیه و تحلیل قرار می‌گیرند و به بخش‌های متفاوتی تقسیم می‌شوند. تجزیه و تحلیل بدین معنی است که مشخص می‌شود اطلاعات از کدام صفحه ارسال شده است، چه حجمی دارد، کلمات موجود در آن کدامند، و در حقیقت ایندکسر، صفحه را به پارامترهای آن خرد می‌کند و تمام این پارامترها را به یک مقیاس عددی تبدیل می‌کند تا سیستم رتبه‌بندی بتواند پارامترهای صفحات مختلف را با هم مقایسه کند. در زمان تجزیه و تحلیل اطلاعات، ایندکسر برای کاهش حجم داده‌ها از بعضی کلمات که بسیار رایج هستند صرف‌نظر می‌کند. کلماتی نظیر a، an، the، is، WWW و ... از این گونه کلمات هستند.

پایگاه داده

تمام داده‌های تجزیه و تحلیل شده در ایندکسر، به پایگاه داده ارسال می‌گردد. در این بخش داده‌ها گروه‌بندی، کدگذاری و ذخیره می‌شود. همچنین داده‌ها قبل از آنکه ذخیره شوند، طبق تکنیک‌های خاصی فشرده می‌شوند تا حجم کمی از پایگاه داده را اشغال کنند. یک موتور جستجوگر باید پایگاه داده عظیمی داشته باشد و به طور مداوم حجم محتوای آنرا گسترش دهد و البته اطلاعات قدیمی را هم به روز رسانی نماید. بزرگی و به روز بودن پایگاه داده یک موتور جستجوگر برای آن امتیاز

محسوب می‌گردد. یکی از تفاوت‌های اصلی موتورهای جستجوگر در حجم پایگاه داده آنها و همچنین روش ذخیره‌سازی داده‌ها در پایگاه داده است.

سیستم رتبه بندی

بعد از آنکه تمام مراحل قبل انجام شد، موتور جستجوگر آماده پاسخ‌گویی به سوالات کاربران است. کاربران چند کلمه را در جعبه جستجوی آن وارد می‌کنند و سپس با فشردن اینتر منتظر پاسخ می‌مانند. برای پاسخ‌گویی به درخواست کاربر، ابتدا تمام صفحات موجود در پایگاه داده که به موضوع جستجو شده، مرتبط هستند، مشخص می‌شوند. پس از آن سیستم رتبه‌بندی وارد عمل شده، آنها را از بیشترین ارتباط تا کمترین ارتباط مرتب می‌کند و به عنوان نتایج جستجو به کاربر نمایش می‌دهد. حتی اگر موتور جستجوگر بهترین و کامل‌ترین پایگاه داده را داشته باشد اما نتواند پاسخ‌های مرتبطی را ارایه کند، یک موتور جستجوگر ضعیف خواهد بود. در حقیقت سیستم رتبه بندی قلب تپنده یک موتور جستجوگر است و تفاوت اصلی موتورهای جستجوگر در این بخش قرار دارد. سیستم رتبه‌بندی برای پاسخ‌گویی به سوالات کاربران، پارامترهای بسیاری را در نظر می‌گیرد تا بتواند بهترین پاسخ‌ها را در اختیار آنها قرار دهد. الگوریتم، مجموعه‌ای از دستورالعمل‌ها است که موتور جستجوگر با اعمال آنها بر پارامترهای صفحات موجود در پایگاه داده اش، تصمیم می‌گیرد که صفحات مرتبط را چگونه در نتایج جستجو مرتب کند. در حال حاضر قدرتمندترین سیستم رتبه‌بندی در اختیار گوگل می‌باشد.

نحوه عملکرد موتور جستجوگر

طرز کار موتور جستجوگر به طور خلاصه بدین گونه است که ابتدا یک آدرس وب را می‌یابد، آن را دنبال می‌کند و به صفحه‌ای جدید می‌رسد. محتوای آن صفحه را می‌خواند و پارامترهای آن را مشخص می‌کند. به عنوان مثال تعداد کلمات متن آن صفحه، حجم و تاریخ به روزرسانی آن، بعضی از پارامترهای آن صفحه است. سپس پارامترهای تعیین شده را به همراه آدرس آن صفحه به بایگانی موتور جستجوگر ارسال می‌کند و این اطلاعات در آنجا پس از فشرده‌سازی، ذخیره می‌گردد.

حال اگر کاربر کلماتی را جستجو کند، موتور جستجوگر در پایگاه داده‌ای که تشکیل داده است ابتدا تمام صفحات مرتبط با موضوع جستجو شده را یافته و سپس مرتبط‌ترین اسناد را به عنوان اولین نتیجه جستجو و بقیه صفحات را بر اساس میزان ارتباط بعد از آن در اختیار کاربر قرار می‌دهد. به عبارت دیگر اگر تعداد نتایج جستجو 1000 مورد باشد، سایت رده اول مرتبط‌ترین و سایت رده 1000 کم ارتباط‌ترین سایت به موضوع جستجو شده می‌باشد.

اگر عبارت یکسانی در تمام موتورهای جستجوگر، جستجو شود هیچ کدام از آنها نتایج یکسانی را ارائه نمی‌دهند و با نتایج کاملاً متفاوتی روبرو می‌شویم. تفاوت در ارائه نتایج جستجو در موتورهای جستجوگر از تفاوت آنها در الگوریتم (سیستم رتبه‌بندی) و بایگانی داده‌هایشان ناشی می‌شود. حتی اگر همه آنها از بایگانی داده یکسانی نیز استفاده کنند، بازهم نتایج جستجویشان متفاوت خواهد بود. هر موتور جستجوگری برای رده‌بندی صفحات وب، از الگوریتم خاصی استفاده می‌کند که منحصر به خودش بوده و فوق‌العاده محرمانه می‌باشد. الگوریتم نیز مجموعه‌ای از دستورالعمل‌ها است که موتور جستجوگر به کمک آن تصمیم می‌گیرد که سایت‌ها را چگونه در خروجی‌اش مرتب کند.

ضریب نفوذ مناسب به معنای حضور در موتورهای جستجوگر مهم و عمده، بایگانی شدن هر چه بیشتر صفحات سایت در پایگاه داده آنها و قرار گرفتن در صفحه‌های اول تا پنجم نتایج جستجوی آنهاست.

انواع موتورهای جستجو

انواع مختلفی از موتورهای جستجوگر وجود دارند که در ادامه به معرفی و شرح مختصری از آنها خواهیم پرداخت.

موتورهای جستجو عمومی

همانطور که پیش از این اشاره شد، حجم اطلاعات و داده‌های اینترنتی آنقدر زیاد و رو به افزایش است که کاربران برای یافتن اطلاعات دلخواه خود، نیازمند ابزارهایی (موتورهای جستجوگر) هستند که به آنها در پیدا کردن اطلاعات کمک کنند؛ موتورهای جستجو با ارائه خدمات جستجو علاوه بر

تسهیل کار کاربران، رقابتی میان وب سایتها و وبلاگها ایجاد کرده‌اند. تا جایی که مدیران سایتها، وب سایتها و وبلاگهای خود را بر اساس عملکرد موتورهای جستجوگر بهینه‌سازی می‌کنند. از معروفترین نمونه‌های خارجی این نوع موتور جستجو می‌توان به گوگل و یاهو اشاره کرد و از نمونه‌های داخلی، که برای فارسی زبانان برنامه‌ریزی شده است، می‌توان جس جو را نام برد.

موتورهای جستجوگر چندگانه¹⁶

با وجود محبوبیت بسیار بالا موتورهای جستجوگر نظیر گوگل، یاهو و ... به دلیل تفاوت الگوریتم‌های مورد استفاده، محدوده پوششی این صفحات همپوشانی نداشته و هر کدام تنها قادر به پوشش بخشی از فضای وب هستند، لذا برای ایجاد دیدی وسیع و همه جانبه نیاز به ابزاری داریم تا در حد امکان به تمام زوایای یک موضوع بپردازد. این جستجوگرها به هنگام دریافت درخواست کاربر، نتایج موتورهای جستجوی مختلف را برای آن موضوع بررسی می‌کنند و بهترین نتایج را در اختیار کاربر قرار می‌دهند. در این نوع عمل پیمایش و ایندکس‌گذاری صفحات صورت نگرفته و بیشتر به تجمیع نتایج جستجوگرهای دیگر پرداخته می‌شود. در نهایت با حذف اطلاعات کاملاً یکسان (محدوده پوششی مشترک)، به رتبه‌بندی مجدد اطلاعات پرداخته و آنها را به کاربر نمایش می‌دهند.

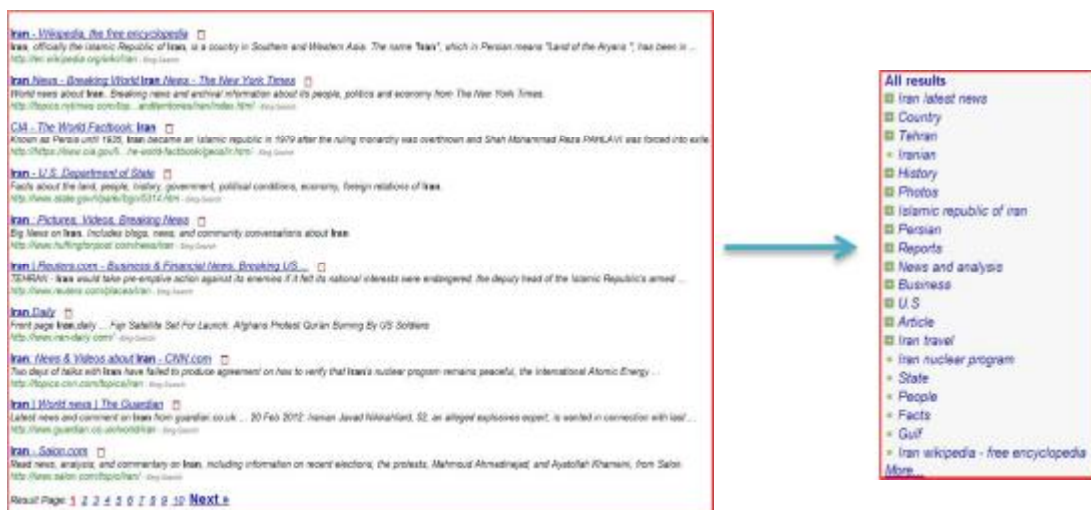
موتورهای جستجوگر خاص

این نوع موتورهای جستجو بر موضوعی خاص تمرکز دارند و تنها به جستجو به وب سایتها و وبلاگ‌های مرتبط با آن موضوع می‌پردازند. یکی از نمونه‌های این نوع موتور جستجو Chemical Search می‌باشد. در این موارد حجم کلی فضای مورد جستجو را با آموزش دادن به جستجوگر کاهش داده و تنها مطالب مرتبط را به نمایش می‌گذاریم. از آنجا که در هر حوزه تعداد وب سایت‌های معتبر محدود هستند، در ابتدا به صورت پیش فرض از آنها استفاده می‌کنیم.

¹⁶ Meta search engine

موتورهای دسته‌بندی نتایج وب

با وجود اینکه موتورهای جستجوگر چندگانه تا حد زیادی مشکل کاربران در دسترسی به منابع اطلاعاتی را کاهش داده‌اند، اما همچنان وجود تعداد نتیجه بازگردانده شده بالا موجب سردرگمی و ابهام در بدست‌آوردن اطلاعات است. از طرف دیگر با ادغام چندین موتور حجم بالایی از اطلاعات تکراری و مشابه در لیست نهایی مشاهده می‌گردد. عمل دسته‌بندی در واقع ساختار نمایش و ارائه اطلاعات را تغییر داده و هیچ‌گونه دخالتی در محتوای اسناد ندارد. تفاوت آرایه نتایج به صورت لیست ترتیبی و گروه‌های سازماندهی شده در شکل (1-2) مشاهده می‌گردد. به عنوان نمونه ای از این دسته می‌توان به موتور iboogie اشاره نمود.



شکل (1-2): تغییر ساختار ارائه اسناد وب از حالت لیست رتبه‌بندی شده به دسته‌های مجزا

{http://www.iboogie.com/searchtree.asp?name_query=iran&name_tab=0&name_do_search=1}

هدف از طراحی موتورهای دسته‌بند نتایج وب ارائه بهتر و منسجم‌تر داده‌ها به گونه‌ای است که سرعت کاربر در یافتن اطلاعات مورد نیاز و مورد علاقه‌اش را افزایش دهد. این امر علاوه بر صرفه‌جویی در هزینه و زمان، میزان رضایت کاربر را نیز بالاتر می‌برد. در مورد این موتورها به صورت مفصل در بخش سوم توضیح داده شده است.

3-2-2 کدهای وضعیت¹⁷ HTTP و خطاهای سرور

هنگامی که در بستر پروتکل HTTP در وب در حال گشت و گذار و مرور صفحات مختلف هستیم، با وارد نمودن هر آدرسی که از طریق مرورگر خود از سرور سایت‌های مختلف درخواست می‌کنیم، کدهایی در پس زمینه بین واسط کاربری ما (مرورگر) و سرور رد و بدل می‌شوند که به آنها در اصطلاح کدهای وضعیت¹⁸ HTTP می‌گویند، این کدها توسط کنسرسیوم‌ها و بنیادهای جهانی استاندارد سازی وب از جمله¹⁹ IETF و²⁰ W3C تعریف و سازمان‌دهی شده‌اند و امروزه تقریباً سرور یا مرورگری وجود ندارد که از اصول آنها پیروی نکند. آشنایی با این کدها و مدیریت درست و به‌جا آنها باعث جلوگیری از ایجاد خطا در روند اجرا برنامه خواهد شد. این کدها به پنج دسته کلی تقسیم می‌شوند که در ادامه به آنها اشاره نموده و در فصل چهارم به تفصیل پیرامون کدهای چالش‌زا بحث خواهیم نمود.

کدهای سری 100، مربوط به اطلاعات

اولین سری از کدهای HTTP، با عدد 100 شروع می‌شود که در مورد نقل و انتقال بسته‌های اطلاعات مثل ارسال و دریافت فایل، کاربرد دارند و حالت موقت پاسخ سرور را نشان می‌دهند، به فرض وقتی از متد POST در فرم‌های وب استفاده می‌کنیم، دریافت کد 100 به معنی این است که سرور درخواست ما را پذیرفته و فرایند پردازش اطلاعات ادامه دارد، البته بدون ارسال کد 100 نیز این فرایند ادامه می‌یابد لذا ارسال آن از طرف سرور ضروری نیست و حتی در مرورگرهایی که از نسخه HTTP/1.0 استفاده می‌کنند، این کد قابل فهم و پردازش نیست.

^{۱۷} Hypertext Transfer Protocol

^{۱۸} HTTP response status codes

^{۱۹} Internet Engineering Task Force

^{۲۰} World Wide Web Consortium

کدهای سری 200، درخواست موفق

کدهای سری 200، به این معنی است که درخواست ارسالی واسط کاربری (که می‌تواند مرورگر یا ابزار دیگری باشد)، با موفقیت دریافت، موافقت، پردازش و پاسخ داده شده است، کدهای سری 200 معمولا به معنی بی نقص بودن درخواست و عملکرد صحیح سرور است.

کدهای سری 300، انتقال

کدهای سری 300 مربوط به مواردی هستند که پاسخ به درخواست واسط کاربری از سرور، باید با انجام اعمال دیگری (در سمت کاربر) کامل شود، این عملیات معمولا توسط واسط کاربری (مثلا مرورگر) و بدون دخالت کاربر (به صورت خودکار) انجام می‌شود، به فرض عمل انتقال خودکار از یک آدرس به آدرس دیگر، با ارسال کدهای سری 300 انجام می‌شود، نکته مهم در اینجا این مسئله است که انتقال‌ها نباید در یک درخواست، بیش از 5 بار تکرار شوند، در غیر اینصورت در اکثر مرورگرها، فرض بر حلقه (Loop) بی‌انتهای شده و ارتباط قطع خواهد شد.

کدهای سری 400، خطای سمت کاربر

کدهای سری 400 مربوط به رویداد خطایی از جانب کاربر (سمت کاربر) در ارائه درخواست به سرور است، در پاسخ، سرور معمولا و به طور پیش فرض، به همراه کدهای HTTP عباراتی در توضیح خطای رخ داده ارسال می‌کند و دائمی یا موقتی بودن مشکل به وجود آمده را نیز تعیین خواهد کرد.

کدهای سری 500، خطای سمت سرور

کدهای سری 500 به معنی نقص داخلی سرور است، با این حال سرور در مجموع سالم بوده و احتمالا به طور موقت در حال انجام به روزرسانی یا تغییراتی است و در ساعات آینده مشکل رفع خواهد شد.

2-3 دسته‌بندی اسناد

فرآیند کشف زیر مجموعه‌هایی از اسناد ورودی در راستایی که اسناد داخل یک گروه بیشترین میزان شباهت را با یکدیگر داشته و در عین حال دارای کمترین تشابه با دیگر گروه‌ها باشند را دسته‌بندی می‌نامیم. میزان این تشابه بر مبنای معیارهای اندازه‌گیری عددی سنجیده می‌شود. در مورد اسناد متنی، تشابه باید میان عناصر تشکیل‌دهنده یا به عبارتی محتوای هر سند استوار باشد که عملاً شامل کلمات و عبارات تشکیل‌دهنده آنها است. متن‌کاوی، به فرآیند کشف الگو در مجموعه متون غیر-ساخت‌یافته بر مبنای تکنولوژی‌های پردازش زبان طبیعی²¹، بازیابی اطلاعات²²، استخراج اطلاعات²³ و داده‌کاوی²⁴ گفته می‌شود [8]، که در آن دو متن با تعداد بالای کلمات یکسان متشابه نامیده می‌شوند. با توجه به نحوه اختصاص اسناد به دسته‌ها و امکان ایجاد همپوشانی میان آنها، روش‌های دسته‌بندی را به دو گروه سخت²⁵ با قابلیت تعلق تنها به یک دسته و نرم²⁶ به منظور اختصاص به چندین گروه دسته‌بندی می‌کنیم. هر چند اکثر الگوریتم‌ها در دسته‌بندی اسناد بر مبنای روش اول عمل می‌کنند. از نقطه نظر نظارتی فرآیند گروه‌بندی را به دو مدل دسته‌بندی و خوشه‌بندی طبقه‌بندی می‌کنیم.

2-3-1 خوشه‌بندی²⁷

خوشه‌بندی به عمل یادگیری بدون ناظر گفته می‌شود. در خوشه‌بندی هیچ نظارت یا اطلاعاتی به جز داده‌های آموزشی در اختیار یادگیرنده قرار ندارد و این یادگیرنده است که باید در داده‌ها به دنبال ساختار یا الگویی خاص بگردد. نتیجه خوشه‌بندی به نحوه نمایش آبجکت‌ها، معیار تشابه و الگوریتم آن بستگی دارد. در این روش به صورت خودکار ویژگی‌های متمایز کننده زیرگروه‌ها تعریف می‌گردد. انتظار می‌رود در نهایت به دسته‌هایی با عناصر درونی مشابه و متمایز با عناصر دیگر گروه‌ها دست یابیم [9,10,11].

²¹ Natural Language Processing

²² Information Retrieval

²³ Information Extraction

²⁴ Data Mining

²⁵ Hard/ crisp

²⁶ Soft/fuzzy

²⁷ Clustering

در خوشه بندی نه تعداد گروه‌ها و نه خصوصیت یا ویژگی مورد نظر برای تمایز گروه‌ها از قبل مشخص نمی‌باشند. در نتیجه پس از انجام خوشه‌بندی لازم است یک فرد خبره، خوشه‌های ایجاد شده را تفسیر کند و در بعضی مواقع نیاز است که پس از بررسی خوشه‌ها، بعضی از پارامترهایی که در خوشه‌بندی در نظر گرفته شده‌اند ولی نامرتبط بوده و یا اهمیت چندانی ندارند، حذف گردیده و جریان خوشه‌بندی مجدداً از اول صورت گیرد. پس از این که داده‌ها به چند گروه منطقی و توجیه‌پذیر تقسیم شدند، از این تقسیم‌بندی می‌توان برای کسب اطلاعات در مورد داده‌ها یا تقسیم داده‌های جدید استفاده نمود. انواع الگوریتم‌های خوشه‌بندی را می‌توان به دو روش چند بخشی²⁸ و سلسله مراتبی²⁹ تقسیم نمود. با توجه به این که روش‌های خوشه‌بندی سلسله مراتبی، اطلاعات بیشتر و دقیق‌تری تولید می‌کنند، برای تحلیل داده‌های با جرئیات بیشتر، پیشنهاد می‌شوند، ولی از طرفی چون پیچیدگی محاسباتی بالایی دارند، برای مجموعه داده‌های بزرگ، روش‌های خوشه‌بندی مسطح پیشنهاد می‌شود. تمامی الگوریتم‌های خوشه‌بندی به صورت خلاصه در جدول (2-3) آورده شده‌اند. در ادامه به شرح هر کدام در حد نیاز پرداخته شده است.

²⁸ Partitional
²⁹ Hierarchical

Algorithm	Time complexity	Overlap	Clustering type	Description
Single linkage	$O(n^2)$	Hard	Hierarchical	n = number of documents
Group average	$O(n^2)$	Hard	Hierarchical	n = number of documents
Complete linkage	$O(n^2)$	Hard	Hierarchical	n = number of documents
Centroid/ Median HAC	$O(n^2)$	Hard	Hierarchical	n = number of documents
k-Means	$O(nkt)$	Hard	Partitional	n = number of documents k = number of initial clusters t = number of iterations
Single pass	$O(n \log n)$	Hard	Partitional	n = number of documents
Scatter/ gather	$O(kn)/O(nm)$	Hard	Partitional/ Hierarchical	n = number of documents k = number of initial clusters t = number of iterations
STC	$O(n)$	Soft	Hierarchical	n = length of all documents
SHOC[12]	$O(n)$	Soft	Hierarchical	n = length of all documents
Lingo[13,14]	$O(m^2n + n^3)$	Soft	Hierarchical	n = number of documents m = number of features

جدول (3-2): خلاصه ای از روش‌های مورد استفاده در خوشه‌بندی اسناد

روش چند بخشی

روش متداول در مدل مسطح kmeans [15] است که مرکز دسته‌ها را به عنوان centroid و نقطه مرکزی را به عنوان نماینده مناسبی از گروه در نظر گرفته و با انتخاب k مرکز دسته و مقایسه فاصله cosine یا هر معیار تشابه دیگری (در ادامه معیارهای تشابه توضیح داده شده است) میان اسناد و این مراکز، تعلق آنها را مشخص می‌کند. به روز رسانی مراکز پس از انتساب‌های جدید صورت گرفته و تکرار روند، تا زمانی که دیگر هیچ تغییری در گروه‌ها رخ ندهد ادامه می‌یابد. از معایب این روش می‌توان به وابستگی بالا آن به انتخاب خوشه‌های اولیه و تعداد آنها و عدم وجود روال مشخص برای محاسبه اولیه مراکز خوشه‌ها نام برد. هرچند در ادامه انواع مختلفی از روش‌ها با تغییرات جزئی در

الگوریتم اصلی توسط محققان پیشنهاد شد که در هر یک سعی در کم کردن وابستگی به مقداردهی اولیه است [16].

روش دیگر استفاده از سیستم پراکنده سازی- جمع آوری³⁰ است که از دو الگوریتم چند بخشی با پیچیدگی زمانی خطی استفاده می‌کند. شکستن هر دسته و جمع‌آوری مجدد زیر دسته‌ها عمل پالایش را در هر دسته انجام داده و از روش‌های بالابه پایین، یا پایین به بالا بسته به نوع داده به منظور انتخاب مراکز اولیه خوشه‌ها استفاده می‌گردد. هدف اصلی این روش انتخاب بهینه مراکز و سپس انتساب اسناد به آنها است.

آخرین الگوریتم مورد استفاده، روش مسیر ساده³¹ نام دارد که در آن هر سند به مرکزی اختصاص می‌یابد که میزان تشابه میان آنها بیشتر از یک حد آستانه باشد. دسته‌ها با یک بار مرور اسناد ایجاد شده و هیچ تکراری صورت نمی‌گیرد.

مزیت اصلی این روش‌ها سهولت و پیچیدگی زمانی پایین آنها است. در حالی که عیب اصلی دسته-بندی نیاز به تعریف پارامترهای بسیاری مانند تعداد نهایی دسته‌ها، انتخاب اولیه مراکز دسته‌ها و ... است.

روش سلسله مراتبی

خروجی متدهای سلسله مراتبی، ساختاری درخت مانند دارد. این روش می‌تواند به دو صورت متراکم- ساز (هر سند یک دسته) و تفرقه‌ساز (تمام اسناد یک دسته منفرد) عمل کند. در روش تفرقه‌ساز، ابتدا تمام داده‌ها به عنوان یک خوشه در نظر گرفته می‌شوند و پس از طی یک فرایند تکراری در هر مرحله، داده‌هایی که شباهت کمتری به هم دارند، به خوشه‌های مجزایی شکسته می‌شوند و این روال تا رسیدن به خوشه‌هایی که دارای یک عضو هستند، ادامه پیدا می‌کند. اما در روش متراکم ساز ابتدا هر داده به عنوان خوشه‌ای مجزا در نظر گرفته می‌شود و در طی فرآیند تکراری در هر مرحله،

³⁰ Scatter/Gather

³¹ Single Pass

خوشه هایی که شباهت بیشتری با یکدیگر دارند، ترکیب می‌شوند تا در نهایت یک خوشه و یا تعداد معینی خوشه حاصل شود. برخی از این الگوریتم‌ها عبارتند از [17]:

اتصال منفرد³²: در این متد تشابه میان دو گروه برابر است با ماکزیمم تشابه میان تمام زوج اسناد، به گونه ای که هر عضو از دوتایی متعلق به یک گروه مجزا باشد. هر عضو از مجموعه باید حداقل به یکی از اعضا دسته شباهتی بیشتر نسبت به اسناد هر گروه دیگری داشته باشد. این روش دارای کمترین پیچیدگی زمانی است اما در ارائه نتایج مفید ضعیف عمل می‌کند.

اتصال کامل³³: در این روش تشابه میان دو گروه برابر است با مینیمم تشابه میان تمام زوج اسناد، به گونه ای که هر عضو از دوتایی متعلق به یک گروه مجزا باشد. هر عضو از مجموعه حداقل به یکی از غیر مشابه ترین اعضا دسته شباهتی بیشتر نسبت به غیر مشابه ترین سند هر گروه دیگری دارد. در این روش خوشه‌هایی منسجم تر از روش فوق ایجاد می‌گردد.

میانگین اتصال گروهی³⁴: در این متد تشابه میان دو گروه برابر است با میانگین تشابه میان تمام زوج اسناد، به گونه ای که هر عضو از دوتایی متعلق به یک گروه مجزا باشد و دارای بهترین نتیجه نسبت به دو مورد قبل می‌باشد.

روش‌های میانگین/ میانه³⁵: هر گروه ایجاد شده توسط میانگین/ میانه اعضا نمایش داده می‌شود، در هر مرحله دسته‌هایی که تشابه بیشتری در میانگین/ میانه دارند، با هم ادغام می‌شوند. تفاوت دو معیار بررسی شده در آن است که مقدار میانه وابستگی به سائز دسته نداشته و به صورت مستقل از تعداد اعضا عمل می‌نماید.

روش درخت پسوندی³⁶

این الگوریتم [18] با هدف دسته‌بندی اسناد متنی بازگردانده شده از موتورهای جستجوی وب بر اساس عبارات مشترک میان آنها طراحی شده است. ایده اصلی استفاده از عبارات به جای کلمات

³² Single link

³³ Complete link

³⁴ Group average link

³⁵ Centroid/median methods

³⁶ Suffix Tree Clustering

نشأت گرفته از میزان قدرت توصیفی بالا در آنها است. در این روش تمام اسناد وب به عنوان رشته متنی در نظر گرفته می‌شوند. دو مرحله کلی در این الگوریتم عبارتند از فاز تشخیص دسته پایه فاز ادغام دسته‌های پایه. در مرحله اول از یک درخت پیشوندی عمومی با استفاده از کلمات به عنوان المان‌های جملات استفاده شده است. گره‌های درخت نشان‌دهنده اسناد با عبارات یکسان بوده و مواردی که از حد آستانه مشخص بیشتر باشند، با هم ادغام می‌گردند. مهم‌ترین مزیت این روش استفاده از عبارات به جای کلمات برای ساخت برچسب‌های مفهومی‌تر است. نیاز به تعریف حد آستانه مانند تمام موارد دیگر مشکل عمده این روش می‌باشد، از طرف دیگر اگر سندی شامل عین عبارات تعیین شده نباشد، حتی در صورت تشابه با گروه هم امکان ادغام با آنها را ندارد.

روش آنلاین سلسله مراتبی مفهومی³⁷

این الگوریتم که به روش SHOC شناخته شده است نیز برای دسته بندی نتایج وب و بر اساس روش آنالیز معنایی پنهان³⁸ برای زبان چینی طراحی شده است. تحلیل معنایی پنهان اولین بار در سال 1988 [19] معرفی گردید. این روش سعی در کاهش بعد و بهبود معنای مدل برداری دارد و در آن، بردارهای عبارات به مفاهیم سطح بالاتر و در فضای با بعد کمتر نگاشت می‌شوند.

در تحلیل معنایی پنهان، همه‌ی مستندات و کلمات کلیدی، یک ماتریس دو بعدی مستند-کلمه را تشکیل می‌دهند. سپس روش تجزیه مقدار منفرد³⁹ (SVD) برای به دست آوردن ویژگی‌های معنایی، ماتریس مستند-کلمه را تجزیه می‌کند. روش SVD به منظور حذف اطلاعات نامربوط ماتریس مستند-کلمه با ابعاد زیاد و تبدیل آن به یک ماتریس مستند-کلمه با ابعاد کم استفاده می‌گردد. این روش بر خلاف STC از ساختار سلسله مراتبی و آرایه ای به جای درخت استفاده می‌کند. دو مفهوم کلی در این روش معرفی شده است که عبارتند از عبارت کامل و تعریف خوشه مستمر. در مورد اول به صورت کلی به دنبال ماکزیمم نمودن طول دنباله کلمات متوالی کشف شده در متن و اجتناب از

³⁷ Semantic hierarchical online clustering

³⁸ Latent Semantic Indexing

³⁹ Singular Value Decomposition

انتخاب عبارات کوتاه که به سختی قابل درک هستند، می‌باشیم. برای این کار از ساختار آرایه ای با پیچیدگی زمانی $O(n)$ که در آن n طول کل اسناد قابل پردازش است استفاده می‌کنیم. مورد دوم به اسناد اجازه می‌دهد تا با درجه عضویت متفاوت به چندین گروه تعلق داشته باشند. این مورد مانند روش فازی عمل نموده و مشکل هم‌پوشانی داده‌ها در گروه‌ها را برطرف می‌سازد. الگوریتم در سه فاز تشخیص عبارات کامل، کشف خوشه‌های پایه و همچنین ادغام خوشه‌های مشابه به کمک ساختار سلسله مراتبی عمل می‌کند. در فاز دوم از روش تجزیه مقدار منفرد به کمک برداهای متعامد و تعیین مقادیر ویژه استفاده می‌شود. وجود ابهام در تعیین مقادیر آستانه و مشاهده چندین نتیجه که تا حدی به صورت رندوم به تولید عبارات پرداخته بودند، از جمله معایب این روش محسوب می‌گردد.

روش لینگو⁴⁰

معنای لغوی Lingo عبارت است از طیف وسیعی از کلمات یا سبکی از زبان که در موقعیت خاص یا توسط گروهی مشخص از مردم مورد استفاده قرار می‌گیرد. هر دفعه که کاربر پرس‌وجویی را در سایت مطرح می‌کند، یک زبان جدید با ویژگی‌های خاص خود مانند لغات، اصطلاحات و عبارات مختص آن ایجاد می‌شود. کاربران به احتمال زیاد نسبت به توصیفات بیش از حد طولانی و مبهم بی‌اعتنا خواهند بود. هرچند اطلاع داشته باشند که محتوای آنها ممکن است ارزشمند باشد. در نتیجه باید سعی در ایجاد برچسب‌هایی معنادار، مجزا و دقیق برای دسته‌ها نماییم. تعیین برچسب در روشهای متداول و مرسوم که بعد از تعیین گروه‌ها و بر اساس محتوای اسناد صورت می‌گیرد نسبتاً دشوار است. روش‌های عددی تا حد بالایی "می‌دانند" که اسناد مشابه یکدیگر هستند، اما قادر به شرح این تشابه نیستند. در روش لینگو ابتدا توصیفات مد نظر قرار گرفته و انتخاب دقیق برچسب‌های نامزد اهمیت بسیاری دارد. از طرف دیگر باید این اطمینان حاصل گردد که برچسب‌ها به اندازه کافی نسبت به یکدیگر متفاوت بوده و تمام فضای نتایج جستجو را پوشش داده‌اند. برای یافتن این کاندیدها از مدل

⁴⁰ Lingo

فضای بردار⁴¹ (VSM) و تجزیه مقدار منفرد استفاده شده است. یعنی در مرحله اول به ایجاد لیبل پرداخته و در ادامه تنها قطعات به دست آمده را به هر برچسب نسبت می‌دهیم. این الگوریتم در دو فاز کلی استقرای برچسب خوشه به کمک تکنیک کاهش ابعاد و فاز تخصیص محتوای خوشه توسط مدل فضای بردار عمل دسته‌بندی را انجام می‌دهد. متد VSM یکی از روشهای موجود در بازیابی اطلاعات است که از عملیات جبری خطی برای مقایسه داده‌های متنی استفاده می‌کند. این روش به هر یک از اسناد موجود در مجموعه، برداری چند بعدی اختصاص می‌دهد و هر قسمت از بردار به یکی از لغات موجود در سند اشاره دارد. در نتیجه هر سند به کمک ماتریس (لغت - سند) منحصر به فرد نمایش داده می‌شود. مقدار هر یک از اجزا در ماتریس بستگی به نوع و میزان وابستگی لغت و سند متناظر دارد.

برخلاف VSM، LSI مجموعه ورودی را به کمک مفاهیم موجود در اسناد به جای معنی واقعی کلمات موجود نمایش می‌دهد. برای این کار با استفاده از تعداد معدودی از عوامل متعامد ماتریس (لغت - سند) تخمین زده می‌شود. این فاکتورها به نمایش مجموعه‌ای از مفاهیم انتزاعی می‌پردازند، که هر کدام بخشی از ایده مشترک میان زیر مجموعه‌ای از اسناد کلی را بیان می‌دارد. از دیدگاه lingo این مفاهیم، بهترین کاندیدهای برچسب دسته هستند، هر چند فهم نمایش ماتریسی آنها برای انسان سخت است. برای بدست آوردن برچسب‌های خلاصه و حاوی اطلاعات مفید، باید بار معنایی و مفهومی لغات مورد پردازش در LSI را تا حد امکان افزایش دهیم. در بین تمام رخدادهای لغوی در سند، آنهایی که به تعداد دفعات بیشتر از حد آستانه در سند تکرار شده‌اند، ارزش معنایی بیشتری برای کاربر داشته و در عین حال پیدا کردن آنها آسان است. این مجموعه اغلب شامل موجودیت‌های اسمی است. الگوریتم lingo از این عبارات تکرار شونده به عنوان برچسب دسته‌های خود استفاده می‌کند. سپس توسط روش SVM اسناد را به دسته‌های مربوطه الحاق می‌کنیم. پیچیدگی زمانی این الگوریتم برابر با $O(m^2n + n^3)$ است [20].

⁴¹ Vector space model

2-3-2 دسته‌بندی

دسته‌بندی یا خوشه‌بندی بانظارت [9]، یک تکنیک داده‌کاوی است که از طریق ارزیابی ویژگی‌های مجموعه‌ای از داده‌ها، آنها را به تعدادی دسته‌های از پیش تعیین شده، اختصاص می‌دهد. این عمل گامی بلند در بازیابی، مدیریت و کاوش داده‌های موجود در سیستم‌های اطلاعاتی برداشته است. خوشه‌بندی اسناد متنی به خوشه‌های مختلف، یک گام مهم در ایندکس کردن، بازیابی، مدیریت و کاوش داده‌های متنی فراوان موجود روی وب یا سیستم اطلاعات یک شرکت است. در این روش، عامل یادگیر با استفاده از یک سری از داده‌های آموزشی یاد می‌گیرد که در هر موقعیت چه کاری انجام دهد.

مسائل دسته‌بندی به شناسایی خصوصیات منجر می‌گردد که همزمان به درک مفاهیم موجود و پیش‌بینی نمونه‌های جدید کمک می‌کند. این روش اسناد را بر طبق یک ساختار شناخته شده مرتب می‌سازد. از مهم‌ترین روش‌های دسته‌بندی می‌توان به بیزین ساده، k -نزدیک‌ترین همسایه و ماشین بردار پشتیبان اشاره نمود.

روش بیزین

این روش متداول‌ترین روش در دسته‌بندی متون است که در آن متن به صورت مجموعه‌ای از کلمات مستقل از یکدیگر و مستقل از محل قرار گرفتن در متن دیده می‌شود. تعریف تابع احتمال هر متن از حاصلضرب احتمال کلمات آن و احتمال رخداد متنی با آن طول بدست می‌آید (2-1). احتمال هر دسته نیز از تعداد متن‌های متعلق به آن دسته نسبت به کل متن‌ها بدست می‌آید (2-2).

$$p(d_i) = p(|d_i|) * \prod_{k=1}^{|d_i|} p(w_k^i) \quad (1-2)$$

$$p(c_j) = \frac{\#(d_i \in c_j)}{\#(d_i)} \quad (2-2)$$

که در اینجا d_i متن i ام مجموعه آموزشی و w_k^i کلمه k ام در متن i ام بوده و تابع $\#(.)$ تعداد آیتم‌ها است. حال با توجه به تئوری بیز، احتمال تعلق هر متن به هر دسته طبق (2-3) محاسبه می‌گردد.

$$p(c_j|d_i) = \frac{p(c_j) * p(d_i|c_i)}{\sum_{c_j \in C} p(c_j) * p(d_i|c_i)} \quad (3-2)$$

از مهم‌ترین ویژگی‌های این روش سادگی آن و مقاوم‌بودن در برابر خصیصه‌های نامرتب است. این روش مبتنی بر قانون احتمالات بوده و نسبت به شبکه عصبی و درخت تصمیم‌گیری کارایی بالایی دارد، دقت و سرعت بالا در مجموعه داده‌های بزرگ، مستقل بودن از خصیصه‌ها، عدم تحمل و انطباق-پذیری با داده حقیقی و بلادرنگ نیز از دیگر خصیصه‌های آن محسوب می‌شود.

در بسیاری از موارد از مدل توسعه یافته این روش موسوم به π -گرام‌ها استفاده می‌گردد. این روش برخلاف بیز محل قرارگیری کلمات نیز در احتمال رخداد متن اهمیت دارد [21,22,23].

روش k-نزدیک‌ترین همسایه

در این روش به جای استفاده از تمام اسناد، از k سند در نزدیک‌ترین همسایگی مورد نظر استفاده می‌کند. در این متد از اندازه کسینوس میان دو بردار سند برای تعیین میزان شباهت استفاده می‌شود. در این روش، همسایه‌های نزدیک انتخاب شده و به وسیله امتیازدهی دسته همسایه برنده به سند جدید اختصاص می‌یابد. در صورتی که سند به بیش از یک همسایه شباهت داشته باشد، از مجموع امتیازها به عنوان وزن نهایی استفاده می‌کنیم. مهم‌ترین مشکل در این روش تعیین تعداد همسایه مجاور برای پردازش و حد آستانه‌ای برای انتساب سند به گروه‌ها است [23].

روش ماشین بردار پشتیبان⁴²

این الگوریتم که به روش SVM معروف است، مبتنی بر روش‌های آماری بوده و اطلاعات را از فضای برداری حاضر به فضای برداری دیگری (عموما در ابعاد بالاتر) می‌برد. این عمل توسط توابع کرنل صورت گرفته و پس از نگاشت، حداکثر حاشیه ممکن میان نمونه‌های مثبت و منفی از طریق بهینه-سازی مرز داده‌ها انجام می‌گیرد. از مزایای این روش می‌توان به عدم وابستگی آن به تعداد نمونه‌های آموزشی و کارایی بالا در تعداد مشخصه زیاد و نمونه کم نام برد. بسته به فضای کار، داده‌ها توسط

⁴² Support Vector Machine

خطوط و یا صفحات از هم جدا می‌شوند. این الگوریتم در حال حاضر کاربرد بسیار بالایی در اکثر موارد و خصوصا در متن‌کاوی پیدا کرده است [24,25].

3-3-2 دسته‌بندی اسناد متنی

دسته‌بندی متن به معنای انتساب اسناد متنی بر اساس محتوای آنها به یک یا چند طبقه از قبل تعیین شده است که یکی از مهمترین مسائل حوزه متن‌کاوی می‌باشد. متن‌کاوی، فرایند کشف نیمه اتوماتیک الگوها و مسیرها در مجموعه عظیمی از متن‌های غیرساخت‌یافته است. متن‌کاوی به عنوان روشی در استخراج دانش از متون، یکی از موضوعات مهم در گستره ای از اعمال مدیریت اطلاعات است.

یکی از روش‌های مورد استفاده در متن‌کاوی، استفاده از تکنیک‌های یادگیری ماشین است. مسأله قابل تأمل این است که این تکنیک‌ها ابتدا در مورد داده‌های ساخت‌یافته به کار گرفته شدند و علمی به نام داده‌کاوی را به وجود آوردند. داده‌های ساخت‌یافته به داده‌هایی اطلاق می‌گردد که به طور کاملاً مستقل از یکدیگر ولی یکسان از لحاظ ساختاری در یک محل گردآوری شده‌اند [26].

بنابراین تفاوت بارز داده‌کاوی و متن‌کاوی در این است که تمرکز اصلی داده‌کاوی بر داده‌های ساخت‌یافته و تمرکز متن‌کاوی بر داده‌های نیمه ساخت‌یافته و غیرساخت‌یافته می‌باشد. داده‌کاوی تلاش می‌کند تا از میان مجموعه داده‌های فراوان، الگوها و روابطی را در اختیار تحلیل‌گرها و تصمیم‌گیرنده‌ها قرار دهد [27]. به هر حال استفاده از داده‌کاوی در مورد متن، بنیان‌گذار شاخه‌ای جدید در علوم هوش مصنوعی گردید که امروزه متن‌کاوی نامیده می‌شود. در کاربرد خوشه‌بندی متن، هر سند به عنوان یک داده آموزشی در نظر گرفته می‌شود. هر خوشه در این کاربرد، حاوی اسناد مربوط به یک موضوع خاص است. تعداد خوشه‌های آموزشی بر اساس تنوع موضوعی و توسط فردی که داده‌های آموزشی را فراهم می‌کند، مشخص می‌شود. برچسب داده‌های آموزشی، شماره و یا نام خوشه و یا خوشه‌هایی است که هر داده آموزشی به آن خوشه و یا خوشه‌ها تعلق دارد. عامل یادگیر در فاز آموزشی با استفاده از داده‌های آموزشی (داده‌های برچسب دار)، آموزش داده شده و یک تعمیم از هر

خوشه، عامل را قادر به شناسایی اسناد مربوط به هر خوشه می‌کند. در فاز آزمایش، عامل یادگیر، اسناد متنی را که در فاز آموزشی ندیده است و متعلق به یکی از خوشه های آموزشی می‌باشد، دریافت کرده و احتمال تعلق این سند را به هر کدام از خوشه های آموزشی به دست می‌آورد. در این-جا خوشه بندی برای هر سند، به معنی انتخاب خوشه ای با بالاترین احتمال تعلق سند به آن خوشه می‌باشد.

2-3-3-1 آماده سازی و پیش پردازش داده ها

داده های خام جمع آوری شده معمولا دارای حجم زیاد و بسیار ناهمگن می‌باشند. این داده ها باید به داده های سازگار و یکپارچه تبدیل شوند تا بتوانند برای مرحله ی کشف الگو مفید واقع باشند. به فرایند پردازش داده های خام به منظور تبدیل به فرمت مناسب برای آموزش، پیش پردازش گفته می-شود. همانند بیشتر کاربردهای داده کاوی، پیش پردازش و آماده سازی داده شامل پاکسازی و حذف نویز، تبدیل فرمت داده، یکپارچه سازی و همچنین برطرف کردن ناسازگاری ها می‌باشد [28,29,30]. کارهای عمده ای که اغلب در آماده سازی یا پیش پردازش داده ها صورت می‌گیرد، از این قرارند:

ریشه یابی

عمل ریشه یابی [31,32,33,34] با هدف زدودن الحاقات و یافتن جوهره اصلی واژه به کمک جداول راهنما و حذف پیشوند و پسوندهای مرسوم در هر زبان انجام می‌گیرد. ریشه یابی لغات به معنای زبان شناسی آن نبوده بلکه صرفا به معنای دسته بندی کلمات در گروه های معنایی یکسان است. بدین صورت شیوه نوشتار تمام کلمات هم ریشه یکسان خواهد گشت. برای این کار نیاز به شناخت کافی از دستور زبان داشته و باید ماشین حالت مناسب و متناسب با ساختار زبان تهیه نماییم. الگوریتم های موجود در این زمینه مبتنی بر دیکشنری و قانون بوده که مورد اول به دلیل عدم قابلیت گسترش پذیری، کارایی بالایی ندارد. معروف ترین ریشه یاب در زبان انگلیسی Porter [35] بوده و تاکنون چندین نمونه فارسی بر اساس آن نوشته شده است. در زیر به بیان کلی تفاوت های میان این دو دسته می‌پردازیم.

هر دو ریشه یاب به دنبال پسوندهای خاصی جستجو می‌کنند و مراحل مختلفی را بر طبق لیست قوانین پسوندی پشته‌گذاری شده طی می‌کنند. با این حال تفاوت‌های مهمی بین این دو الگوریتم وجود دارد. برای مثال الگوریتم ریشه‌یاب پورتر به منظور تخمین محتوای اطلاعات، الگوی حروف صدادار و بی‌صدا را تشخیص می‌دهد؛ اما در فارسی بسیاری از حروف صدادار نوشته نمی‌شوند، بنابراین ریشه‌یاب فارسی از طول رشته برای تعریف کران پائین محتوای ریشه استفاده می‌کند. (در حال حاضر مینیمم طول ریشه 3 است). این محدودیت در بعضی موارد باعث خطا می‌گردد بخصوص زمانی که یک زیررشته که قسمتی از یک کلمه کوتاه است، به اشتباه به عنوان یک پسوند در نظر گرفته شود. تفاوت دیگر این دو الگوریتم آن است که این الگوریتم بر خلاف الگوریتم پورتر، پیشنوندها را هم شناسایی می‌کند.

حذف کلمات عمومی

در این مرحله به تشخیص، علامت گذاری و حذف کلمات عمومی [36,37,38] می‌پردازیم. به صورت کلی می‌توان کلمات موجود در یک متن را به سه دسته عمومی، خاص و کلیدی تقسیم نمود. کلمه عمومی، کلمه‌ای است که کلیه اسناد تمام کلاس‌ها بارها و بارها ظاهر شده و اهمیت و مفهوم چندانی ندارند و معمولاً در فاز پیش پردازش، این کلمات حذف می‌شوند. به عنوان نمونه می‌توان از ضمائر، قیود، حروف ربط و اضافه، افعال و ... نام برد. این کلمات در محتوای سند هیچ‌گونه ارزشی ندارند. کلمات خاص، کلماتی با فراوانی زیاد فقط در کلیه اسناد یک کلاس و فراوانی کم در سایر کلاس‌ها بوده و هم‌پوشانی حداقلی آنها، قابلیت مجزا سازی بهینه کلاس‌ها را به میزان زیادی افزایش می‌دهد. در نهایت کلمات کلیدی به عبارات یا کلمات مهم برای تشریح محتویات داخل سند گفته شده و پس از وزن‌دهی به کلمات کاندید اولیه انتخاب می‌گردند. این گروه دارای فراوانی زیاد در سند خاص بوده و در ایجاد خلاصه متن اهمیت بسیار بالایی دارند. در اکثر زبان‌ها لیستی از کلمات عمومی پرتکرار وجود دارد، که در مرحله پیش پردازش آنها را حذف می‌نماییم. عملیات حذف در دو مرحله قبل و بعد از عمل ریشه‌یابی صورت گرفته و نقش بسزایی در کاهش تعداد کلمات قابل پردازش دارد.

تقسیم‌بندی

هدف از این کار تشخیص محدوده جملات، عبارات و یا لغات قابل پردازش در یک متن است. به عنوان مثال در ساده‌ترین روش می‌توان از کاراکترهای خالی و علائم و نشانه‌های نقطه‌گذاری در متن استفاده نمود. در زبان فارسی با مدل‌های مختلف فضای خالی کامل و نصفه مواجه هستیم که به ترتیب برای جداسازی کلمات تک واژه‌ای و کلمات چند واژه‌ای به کار برده می‌شوند. هرچند در اکثر موارد هنوز استفاده از نیم فاصله مرسوم نبوده و برای بالا بردن دقت باید جدولی شامل کلمات چند بخشی ایجاد گردد.

لازم به ذکر است که تمامی موارد مورد بحث در قسمت پیش‌پردازش در صورتی کارایی مورد نظر را خواهند داشت که مطابق با زبان سیستم طراحی شده باشند، چرا که ویژگی‌ها و قوانین زبان‌شناسی، دستوری و نگارشی برای هر زبان متفاوت بوده و از استاندارد مشخصی پیروی نمی‌کنند.

2-3-3-2 وزن دهی به کلمات و ایجاد ماتریس

وزن‌دهی ویژگی به عنوان معیاری عددی برای محاسبه میزان تاثیر ویژگی خاص در مجموعه می‌باشد. به عنوان مثال اگر مجموعه کلی ما دارای i سند مجزا و k ویژگی متمایز کننده باشد، در نتیجه برای هر سند، بردار وزن $(1, k)$ مانند رابطه زیر ایجاد خواهد گشت.

$$\text{اگر } d_i \in D \rightarrow d_i = (w_{1i}, \dots, w_{ki})$$

هدف از وزن‌دهی به کلمات تعیین ارزش و اهمیت آنها در یک سند است. در بسیاری از موارد کلمات کلیدی نهایی بر اساس وزن به دست آمده از میان کلمات کاندید انتخاب می‌شوند. در این زمینه تا-کنون روش‌های بسیاری به کار برده شده است که اکثراً نتیجه ترکیب یک یا چند مورد از موارد زیر است که بسته به کاربرد و نوع داده ورودی از آن استفاده نموده‌اند. در ساده‌ترین شکل وزن‌دهی به صورت منطق $(1/0)$ عمل کرده و وجود کلمه در سند را با میزان عددی "یک" و عدم وجود آن را با مقدار "صفر" نمایش می‌دهیم [39]. در نهایت پس از وزن‌دهی به عبارات و لغات و تعیین ارزش و نقش آنها در اسناد، داده‌ها را به فضای محاسباتی برده و برای این منظور ماتریسی شامل تمام اسناد و

لغات تعیین شده رسم می‌کنیم. به گونه‌ای که هر سطر نشان‌دهنده یک متن، هر ستون بیان‌کننده یک لغت و سلول متناظر با هر سطر و ستون نشان‌دهنده وزن لغت در سطر باشد. بدیهی است عدم وجود عبارت در یک سند معادل با مقدار صفر خواهد بود. نمایش ماتریس به صورت شکل (2-2) است.

$$D = \{d_1, d_2, \dots, d_n\}$$

$$C = \{c_1, c_2, \dots, c_m\}$$

$$a_{ij} = \begin{cases} w_{ij} & d_j \in c_i \\ 0 & d_j \notin c_i \end{cases}$$

	d_1	...	...	d_j	...	...	d_n
c_1	a_{11}	...	...	a_{1j}	...	...	a_{1n}
...	...	...	...	...	...	...	...
c_i	a_{i1}	...	...	a_{ij}	...	...	a_{in}
...	...	...	...	...	...	...	...
c_m	a_{m1}	...	...	a_{mj}	...	...	a_{mn}

شکل (2-2): ماتریس سند-لغت، نمایانگر وزن هر لغت در سند متناظر

به صورت کلی روش‌های متداول وزن‌دهی ویژگی را بر حسب توابع توزیع آنها می‌توان به سه گروه مختلف دسته‌بندی نمود. در ادامه لیستی از مهم‌ترین این روش‌ها آورده شده است.

فرکانس تکرار

اساس این روش توجه به کلمات یا عبارات کاندید به عنوان موجودیتی مستقل بدون توجه به بار معنایی و نحوه قرارگیری در متن است. بدین منظور که صرف وجود یا روئیت کلمه در متن، به آن اعتبار خواهد داد.

⁴³ فرکانس تکرار لغت در سند

این متد به صورت محلی به اهمیت و میزان تکرار لغت یا عبارت در یک سند مجزا می‌پردازد. انواع مختلف وزن‌دهی بر اساس فرکانس لغت در سند در جدول (2-4) آورده شده است. در تمام موارد t_k بیانگر k امین واژه یا عبارت و d_i نشان‌دهنده i امین سند قابل پردازش است [40,41].

⁴³ Term frequency

نام روش	رابطه	توضیحات
TF	$w_{ki} = tf(t_k, d_i) = \begin{cases} \#(t_k, d_i) & t_k \in \text{vector of } d_i \\ 0 & t_k \notin \text{vector of } d_i \end{cases}$	$\#(t_k, d_i)$ تعداد تکرار ویژگی t_k در مستند d_i
normTF	$w_{ki} = \frac{tf(t_k, d_i)}{\sqrt{\sum_k (tf(t_k, d_i))^2}}$	فرمول نرمال شده
logTF	$w_{ki} = \log(tf(t_k, d_i))$	لگاریتم فرمول
ITF	$w_{ki} = ITF(t_k, d_i) = 1 - \frac{1}{1 + tf(t_k, d_i)}$	$r=1$
Sparck	$w_{ki} = Sparck(t_k, d_i) = tf(t_k, d_i) * (k - \log(p_k))$ $p_k = \sum_D tf(t_k, d_i)$	k تعداد ویژگی های منحصر به فرد در مجموعه D

جدول (2-4): روش های وزندهی بر اساس فرکانس تکرار لغت در سند

معکوس فرکانس تکرار لغت در اسناد⁴⁴

این شیوه به صورت عمومی عمل کرده و هدف اصلی آن تلاش برای کاهش وزن ویژگی با افزایش فرکانس ویژگی در مجموعه مستندات است. چرا که هر چه کلمه ای بیشتر در کل متون تکرار شده باشد، جنبه عمومی تری گرفته و برای جداسازی حداکثری دسته ها از هم قابلیت کمتری خواهد داشت. در اکثر موارد از ترکیب این روش و روش فوق استفاده می گردد. جدول (2-5) فرمول اصلی و موارد ترکیبی را نشان می دهد. در تمام موارد t_k بیانگر k امین واژه یا عبارت و d_i نشان دهنده i امین سند قابل پردازش است [41,42].

⁴⁴ Inverse document frequency

نام روش	رابطه	توضیحات
IDF	$w_{ki} = idf(t_k, d_i) = \log \frac{ D }{ D(t_k) }$	$ D $: تعداد کل اسناد مجموعه D $ D(t_k) $: تعداد مستندات شامل ویژگی t_k
TFIDF	$w_{ki} = tfidf(t_k, d_i) = tf(t_k, d_i) * idf(t_k, d_i)$	ترکیبی از فرکانس تکرار محلی و عمومی
normTFIDF	$w_{ki} = \frac{tfidf(t_k, d_i)}{\sqrt{\sum_k (tfidf(t_k, d_i))^2}}$	فرم نرمال شده

جدول (5-2): روش‌های وزندهی بر اساس معکوس فرکانس تکرار لغت در اسناد و ترکیب آن با روش محلی

فرکانس تکرار در دسته⁴⁵

همان‌گونه که پیشتر بیان شد هدف اصلی در فرآیند دسته‌بندی، ایجاد گروه‌ها و دسته‌های متمایز از یکدیگر است. هر چه تفاوت میان دسته‌ها بیشتر باشد، انتساب اسناد به گروه‌های مربوطه به صورت دقیق‌تری صورت می‌گیرد. در نتیجه باید تا حد امکان ارزش کلمات پرکاربرد عمومی کاهش یابد تا کلمات کلیدی مهم‌تر و خاص‌تر استخراج گردند. در تمام موارد t_k بیانگر k امین واژه یا عبارت، d_i نشان دهنده i امین سند و C_j بیانگر j امین دسته قابل پردازش است [41,43]. دو فرمول مهم در این زمینه در جدول (6-2) آورده شده است.

⁴⁵ Cluster frequency

نام روش	رابطه	توضیحات
TFRF	$w_{ki} = TFRF(t_k, d_i) = \frac{tf(t_k, d_i) * rf(t_k, c_{d_i})}{\sqrt{\sum_k (tf(t_k, d_i))^2 * (rf(t_k, c_{d_i}))^2}}$ $rf(t_k, c_j) = \log \left(2 + \frac{ D(t_k, c_j) }{\sum_{m=1, m \neq j}^{ C } D(t_k, c_m) } \right)$	$ D(t_k, c_j) $: تعداد مستندات مجموعه D و طبقه c_j شامل ویژگی t_k مجموع تعداد: $\sum_{m=1, m \neq j}^{ C } D(t_k, c_m) $ مستندات D و طبقه غیر از c_j شامل ویژگی t_k
TFCRF	$w_{ki} = TFCRF(t_k, d_i) = \frac{\log(tf(t_k, d_i) * crfValue(t_k, c_{d_i}))}{\sqrt{\sum_k (\log(tf(t_k, d_i) * crfValue(t_k, c_{d_i})))^2}}$ $negativeRF(t_k, c_i) = \frac{\sum_{m=1, m \neq j}^{ C } D(t_k, c_m) }{\sum_{m=1, m \neq j}^{ C } D(c_m) }$ $positiveRF(t_k, c_i) = \frac{ D(t_k, c_i) }{ D(c_i) }$ $crfValue(t_k, c_i) = \frac{positiveRF(t_k, c_i)}{negativeRF(t_k, c_i)}$	$ D(c_j) $: تعداد مستندات طبقه c_j

جدول (2-6): روش‌های وزن‌دهی بر اساس فرکانس تکرار لغت در دسته

46 مجاورت معنایی

در این روش به کمک مفاهیم آنتولوژی و زبان‌شناسی کلمات مترادف با کلمه کلیدی تطابقت داده شده و بدین وسیله امتیاز کلمه در سند افزایش خواهد یافت. بدین منظور باید از دیکشنری‌های لغت برای یافتن کلمات هم‌معنی و افزودن آنها به متن استفاده نماییم. از مهم‌ترین دیکشنری‌ها در این زمینه می‌توان به wordnet و farsnet اشاره نمود [44,45].

47 موقعیت کلمات

یکی دیگر از عوامل موثر در وزن‌دهی به کلمات، موقعیت مکانی آنها خواه به صورت مستقل و خواه بر حسب دیگر لغات کلیدی است. روش‌های بسیار زیادی در این قسمت نیز وجود دارد که بسته به کاربرد با افزایش ضریب خاص، ارزش کلمه را زیاد می‌کنند. علاوه بر موقعیت، نوع نگارشی (حروف بزرگ/کوچک)، نوع فونت (افزایش)، نوع فرمت (ضخیم/کج) و همچنین طول عبارت (شامل چندین

⁴⁶ Semantic proximity

⁴⁷ Term position

کلمه/طول یک کلمه) نیز در وزن‌دهی اهمیت بسیار بالایی دارد. به عنوان مثال قرار گرفتن کلمه در عنوان مطلب یا فرمت نوشتاری خاص وزن بیشتری به کلمات اختصاص خواهد داد. همچنین اسنادی که در آنها کلمات کلیدی به یکدیگر نزدیک‌تر هستند، به احتمال زیاد اسناد مشابه‌تری برای بازیابی می‌باشند.

2-3-3 کاهش داده و کاهش بعد

کار با اسناد متنی به دلیل تنوع و حجم بالای اطلاعات اغلب زمانبر است. هدف از این مرحله، دست یافتن به حجم کوچک‌تری از داده‌ها جهت بالا بردن زمان پردازش و کاهش حافظه مصرفی است. نکته مهم در این مرحله از آماده‌سازی داده‌ها آن است که دستیابی به نتایج تحلیلی مشابه با اصل و تمام داده‌ها تضمین گردد، چرا که در غیر این صورت این کاهش، اثر مثبتی برای ما در پی نخواهد داشت. هر چند در بسیاری از موارد از این تکنیک‌ها صرف‌نظر شده و تنها سعی در یافتن و بهینه‌سازی کلمات کلیدی حداقلی است.

2-3-3-4 معیارهای شباهت

معیارهای سنجش شباهت، ابزاری ضروری در بسیاری از کاربردها مانند مسائل تصمیم‌گیری، کلاس‌بندی الگوها و داده‌کاوی هستند. بیشترین کاربرد در این زمینه مباحث مربوط به کلاس‌بندی داده‌ها برای یافتن کلاس مربوط به یک داده جدید و وارد کردن آن به کلاس در مسائل طبقه‌بندی است. از این‌رو می‌توان بیان داشت که معیار شباهت نقش مهمی را در پیدا کردن ارتباطات میان داده‌ها و در نتیجه میزان عملکرد سیستم دارد. انتخاب یک معیار شباهت درست و کارا بستگی مستقیم به نوع داده ورودی داشته و باید در هر مسئله از روش مناسب استفاده نمود. وابسته به کاربرد مورد نظر ممکن است یک معیار شباهت k بعدی برای سنجش شباهت مورد نیاز باشد. در حال حاضر معیارهای شباهت بسیاری برای اندازه‌گیری میزان شباهت وجود دارند. این معیارها در سه گروه مختلف طبقه

بندی شده‌اند: میزان فاصله⁴⁸، میزان نزدیکی⁴⁹ و میزان باینری⁵⁰. مهم ترین معیارها به همراه فرمول-ها و پارامترهای مربوط به دو گروه اول به ترتیب در جداول (7-2) و (8-2) آورده شده است.

در معیارهای مبتنی بر فاصله، معیارها و مشخصه‌های اختصاصی یک داده متقابلاً با مشخصه‌های یک داده دیگر مورد مقایسه قرار می‌گیرد. به عبارت دیگر در این گروه فقط داده‌های اولیه و ناخالص دو مجموعه به منظور تشخیص میزان شباهت مورد بررسی قرار می‌گیرند. در معیارهای مبتنی بر نزدیکی اطلاعات اضافه دیگری در مورد داده‌ها مانند مقدار میانگین، انحراف استاندارد نیز برای سنجش میزان شباهت استفاده می‌شوند. معیارهای باینری نیز برای اندازه‌گیری میزان شباهت در داده‌های باینری مورد استفاده قرار می‌گیرند. در تمام موارد X ، Y به دو سند متفاوت از مجموعه اسناد که قصد بررسی میزان تشابه میان آنها را داریم، اشاره دارند [46].

Row num.	Distance Measure Methods	Calculation Formula
1	Euclidian	$Euclid(x, y) = \sqrt{\sum (x_i - y_i)^2}$
2	Chebychev	$Chebychev(x, y) = \max_i x_i - y_i $
3	Block(Hamming)	$Block(x, y) = \sum x_i - y_i $

جدول (7-2): روش‌های اندازه‌گیری معیار شباهت بر مبنای فاصله

⁴⁸ Distance Measures

⁴⁹ proximity measures

⁵⁰ binary measures

Row num.	Proximity Measure Methods	Calculation Formula
1	Correlation	$Correlation(x, y) = \frac{\sum(Z_{xi}Z_{yi})}{N}$
2	Cosine	$Cosine(x, y) = \frac{\sum(x_i y_i)}{\sqrt{((\sum x_i^2)(\sum y_i^2))}}$
3	Chi-Square	$Chisq(x, y) = \sqrt{\sum \frac{(x_i - E(x_i))^2}{E(x_i)} + \sum \frac{(y_i - E(y_i))^2}{E(y_i)}}$
4	Jensen	$Jensen(x, y) = \frac{1}{2} \sum_{i=1}^k \{x_i \log_2 x_i + y_i \log_2 y_i - (x_i + y_i) \log_2 (\frac{x_i + y_i}{2})\}$ $where \ x_i = \frac{x_i}{\sum_{i=1}^k x_i} \ , \ y_i = \frac{y_i}{\sum_{i=1}^k y_i}$

جدول (8-2): روش‌های اندازه‌گیری معیار شباهت بر مبنای نزدیکی

در انتها عمل دسته‌بندی بر اساس یکی از روش‌های مطرح شده در بخش مربوطه و به کمک یکی از معیارهای اندازه‌گیری بسته به نوع و ذات داده‌های ورودی انجام می‌پذیرد.

4-2 دسته‌بندی اسناد وب

با افزایش رو به رشد اطلاعات در بستر وب، یافتن سریع اطلاعات مورد نظر درخواستی از موتورهای جستجو دشوارتر گشته است. یکی از راهکارهای حل این معضل استفاده از تکنیک‌های دسته‌بندی اسناد و گروه‌بندی آنها در رده‌های مشابه به منزله تسهیل در نمایش و ارائه نتایج به صورت فشرده است. دسته‌بندی به معنای روند شکل‌دهی گروه‌هایی از اشیاء مشابه از یک مجموعه داده ورودی است. به عنوان یکی از مشخصه‌های دسته‌بندی خوب باید تمام عناصر درون گروهی با یکدیگر بیشترین میزان شباهت را داشته و با عناصر دیگر گروه‌ها دارای کمترین وجه تشابه باشند. هر چند ایده اصلی دسته‌بندی از محاسبات آماری روی داده‌های عددی سرچشمه گرفته است، اما نیازهای موجود در علوم کامپیوتر و به طور خاص داده‌کاوی، این مفهوم را در دیگر حوزه‌ها مانند داده‌های متنی و چند

رسانه‌ای نیز گسترش داده است. دسته‌بندی اسناد وب در ابتدا در سایت‌هایی مانند dmoz.org و یاهو تا سال 1998 به صورت دستی به کمک افراد خبره و کارشناسان صورت می‌گرفت. هرچند این نوع دسته‌بندی دارای درجه بالا دقت بود، اما ورود تمایلات و تصمیمات شخصی و عدم تشابه آن میان افراد مختلف و همچنین حجم بالا اطلاعات از کارایی سیستم‌ها کاسته و متعاقباً نیاز به تعداد نیروی انسانی زیاد و صرف زمان و هزینه بالا وجود داشت. لذا این وظیفه بر عهده ماشین‌های اتوماتیک گذاشته شد.

در مبحث دسته‌بندی نتایج جستجو وب سعی در بسط این ایده در اسناد بازگردانده شده از موتورهای جستجو (برحسب درخواست یا پیش فرض) داریم، به گونه‌ای که اسناد دسته‌بندی شده به کاربر در نحوه و زمان دسترسی به اطلاعات مورد نیاز و علاقه خود کمک رسانده و عمل دسترسی به اطلاعات را تسهیل بخشد. در این سیستم هر کدام از قطعات⁵¹ خبری را به یک گروه موضوعی معنادار نسبت می‌دهیم. این اقدام به معنای پس پردازش نتایج نهایی به دست آمده از موتورهای جستجوگر به منزله نمایش بهینه و سهولت دسترسی به اطلاعات است. چرا که بدون وجود موتورهای جستجوگر، اینترنت به حجم انبوهی از اطلاعات آشفته و بدون سازمان تبدیل می‌شود که عملاً استفاده از آن برای هیچ کس مقدور نخواهد بود. بر خلاف باور عمومی مردم موتورهای جستجوگر به هیچ عنوان جواب سوالات را ندارند، بلکه صرفاً راهی برای دسترسی سریع به جواب و اطلاعاتی که مردم دیگر بر روی وب قرار داده‌اند را پیش‌روی کاربر قرار می‌دهند و اغلب به جای پاسخ دادن، تنها صفحات متناظر با درخواست کاربر را باز می‌گردانند و رضایت کاربر را در یافتن اطلاعات مورد نظر بدست می‌آورند. مشکل اصلی زمانی رخ می‌دهد که موضوع درخواست عبارتی مبهم، دارای گستردگی مفهومی زیاد و یا عبارتی ضعیف از لحاظ معنایی باشد. در این هنگام با طیف وسیعی از موضوعات نامربوط، گیج کننده و بی-ارزش برای کاربر مواجه خواهیم شد. هدف موتور دسته‌بند مطالب، ارائه بهتر و فشرده‌تر موارد مشابه با هم به کاربر است.

⁵¹ Snippet

چهار معیار اساسی در ایجاد دسته‌ها عبارتند از: ایجاد عناوین مختصر، دقیق و مناسب، مشخص و مورد علاقه کاربر انسانی است، به طوری که مشخص نشود این عبارت توسط یک ماشین ایجاد شده است. یکی از مشخصه‌های موتورهای دسته‌بند اسناد، عدم حفظ عنوان و یا خود داده‌ها است، بلکه در اکثر موارد از نتایج بدست آمده از موتورهای جستجوگر دیگر استفاده می‌کنند. دقت پایین موتورهای لیست‌های طولانی نمایش نتایج کار یافتن اطلاعات مورد نظر را برای کاربران بسیار دشوار و زمانبر کرده است. در مواقعی حتی درخواست مشابه در چند موتور نتایج متفاوتی ارائه می‌دهد که بیشتر آنها نسبتی با مورد درخواستی ندارند. در نتیجه استفاده از موتور دسته‌بند قطعات وب، برای ایجاد مکملی در نمایش لیست نتایج مفید به نظر می‌رسد. این موتورها به نام موتور دسته‌بند نتایج وب نیز معروف هستند. در این روش قطعات خبری بدست آمده از موتورهای جستجوگر در گروه‌های موضوعی معنادار دسته‌بندی می‌شوند. این مقوله در هر دو زمینه تجاری و علمی تاکنون پیشرفت‌های بسیاری داشته است.

هرچند دسته‌بندی نتایج وب وابستگی بسیار شدیدی به دسته‌بندی اسناد دارد، اما چالش‌های بسیاری در این زمینه دیده می‌شود که اثر بخشی و بهره‌وری الگوریتم‌های پایه دسته‌بندی اسناد را تحت تاثیر قرار می‌دهد. اولین تفاوت تاکید بر روی کیفیت برچسب خوشه‌ها است تا جایی که مقالات بسیاری توجه خود را بیشتر بر روی این موضوع معطوف کرده اند.

دسته‌بندی متون عمدتاً به عنوان روشی برای افزایش اثربخشی مرتب‌سازی اسناد به کار می‌رود، چرا که به صورت منطقی و واضح، اسناد مشابه به شدت مربوط به درخواست‌های مشابه می‌باشند. این دسته‌بندی پایه و مبنای بسیاری از اقدامات از جمله دسته‌بندی اسناد وب و طراحی موتورهای دسته‌بند است. الگوریتم‌های دسته‌بندی نتایج وب بنا بر نظر [18] باید دارای ویژگی‌های زیر باشند.

1: ارتباط داشتن⁵². روش باید دسته‌هایی ایجاد نماید که قادر به گروه‌بندی اسناد مربوطه با درخواست

کاربر به صورت مجزا از اسناد غیر مربوطه ایجاد نماید.

⁵² Relevance

2: قابلیت مرور خلاصه⁵³: کاربر باید در یک نگاه کوتاه قادر به تشخیص مفید بودن یا نبودن دسته برای کارایی موردنظر باشد. هدف تنها جایگزینی دسته‌ها با نتایج مرتب شده نیست، در نتیجه الگوریتم باید قادر به ارائه توضیح مختصر و مفید از خوشه باشد.

3: هم پوشانی⁵⁴: از آنجا که هر سند دارای چندین موضوع است، نباید هیچ محدودیتی در اختصاص اسناد به چندین گروه وجود داشته باشد.

4: میزان تحمل پذیری در استفاده از قطعات⁵⁵: الگوریتم باید قادر به ایجاد دسته‌هایی با کیفیت بالا باشد، حتی اگر برای این کار از قطعات خبری برگرفته از موتورهای جستجوگر استفاده شود. چرا که در اکثر موارد کاربر تحمل انتظار برای بارگزاری کامل متن را ندارد.

5: سرعت⁵⁶: کاربران بسیاری وجود دارند که بنا به حساسیت بر روی جستجو، تمایل بررسی بیش از 100 قطعه اولیه نتایج را دارند، یکی از دلایل دسته بندی، اجازه دادن به کاربر برای مطالعه موضوعات مشخص و در نتیجه صرفه جویی در زمان وی است، در نتیجه الگوریتم دسته‌یابی باید قادر به گروه-بندی هزاران قطعه در کمترین زمان ممکن باشد. برای کاربران حساس، حتی ثانیه‌ها نیز ارزشمند هستند.

6: قابلیت افزایشی⁵⁷: به منزله صرفه جویی در زمان روش مورد نظر باید هر قطعه را به محض ورود به وب مورد پردازش قرار دهد.

2-4-1 مراحل خوشه بندی نتایج وب

مراحل خوشه بندی نتایج وب در شکل (2-3) آورده شده و در ادامه توضیح هر کدام با جزئیات بیشتر بیان گشته است. در این قسمت از دوباره‌گویی مطالب مشترک با تکنیک‌های پردازش متن چشم-پوشی کرده و تنها در مواردی که تفاوت خاصی وجود دارد، به آنها اشاره نموده‌ایم.

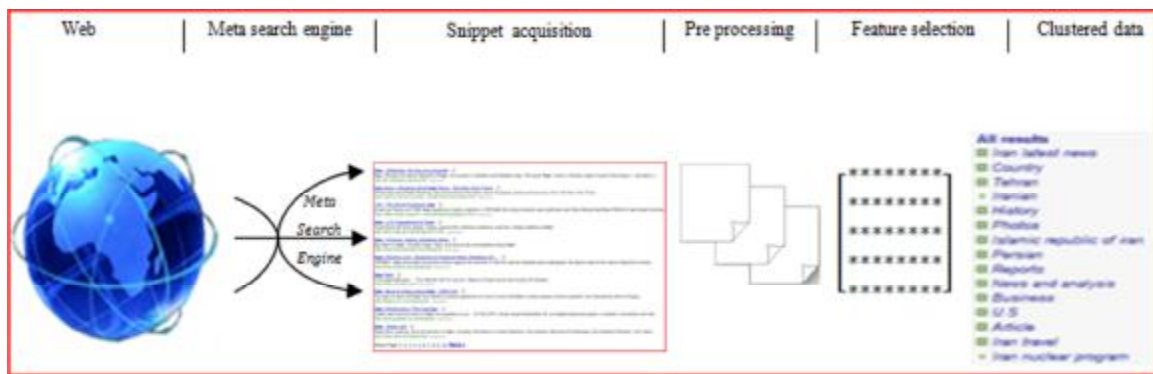
⁵³ Browsable summaries

⁵⁴ Overlap

⁵⁵ Snippet- Tolerance

⁵⁶ Speed

⁵⁷ Incrementality



شکل (2-3): مراحل خوشه‌بندی نتایج وب

وب: در حال حاضر وب به عنوان بزرگترین منبع اطلاعات درهم و غیر ساخت یافته شناخته شده است. برای بدست آوردن اطلاعات، تعدادی عامل توسط موتورهای جستجوگر در صفحات پخش شده و اطلاعات مشابه با درخواست مطرح شده و یا اطلاعات خام مورد نیاز را بر می گردانند.

موتورهای جستجوگر چندگانه: به دلیل عدم پوشانندگی کامل وب توسط یک موتور جستجوگر منفرد، برای دستیابی کامل و همه جانبه به اطاعات پیرامون یک موضوع خاص نیاز به ادغام چندین موتور وجود دارد. در همین راستا از موتورهای جستجوگر چندگانه استفاده می نماییم که به جای پوشش و ایندکس گذاری صفحات وب و جمع آوری و ذخیره نتایج در یک پایگاه، از نتایج بدست آمده از دیگر موتورها استفاده می نمایند. یکی از اقدامات مهم حذف منطقه همپوشانی موتورهای جستجو است بدین صورت که عناوین عینا مشابه که از سایت منبع و زمان یکسان برگردانده شده از میان لیست حذف شده و تنها قطعه مربوط به یکی از موتورها باقی می ماند.

تهیه قطعات: هدف این فاز جمع آوری قطعات متنی است که توسط موتورهای جستجو در جواب درخواست کاربر بازگردانده شده اند. تمام اطلاعات مورد نیاز در این قسمت از صفحات وب استخراج شده و تگ های اضافه و HTML آن جدا می شوند. در این قسمت تنها عنوان، محتوای قطعه، منبع اولیه و تاریخ ثبت نگهداری می شود. در این قسمت صفحه ای مانند صفحه نتیجه جستجو گوگل با لیستی ترتیبی از اطلاعات مشابه خواهیم داشت. این اطلاعات می توانند بر حسب درخواست کاربر و یا

بر حسب پیش‌فرض سیستم بازگردانده شوند که در مورد اول باید به پیش پردازش درخواست کاربر و اعمال قوانین جبری منطقی احتمالی تعبیه شده در صفحه رابط کاربر پرداخته شود.

پیش‌پردازش: عمل پیش‌پردازش به منزله یکدست‌سازی اطلاعات، حذف نویز و آماده‌سازی نهایی داده‌ها برای پردازش انجام می‌گیرد. در این قسمت عملیات فیلتر سازی، حذف کلمات عمومی، ریشه‌یابی و تقسیم نمودن سند به بخش‌های قابل پردازش انجام می‌پذیرد.

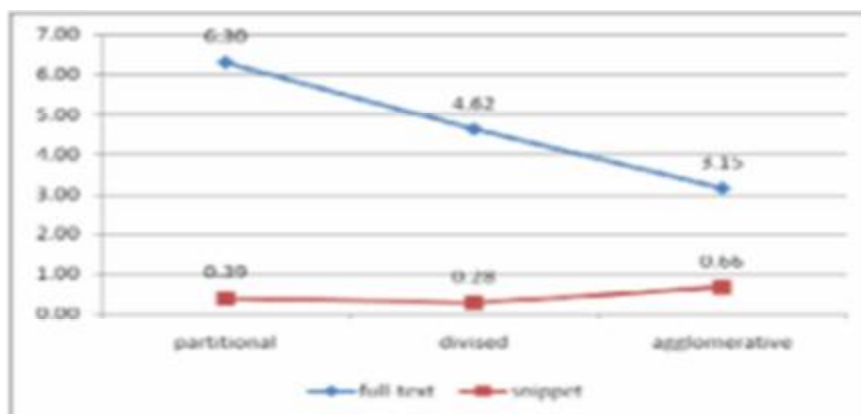
فیلتر سازی: در این مرحله تلفیق اطلاعات مشابه موتورهای جستجوگر و حذف عناوین عینا مشابه و حذف کاراکترهای غیر مجاز، اعمال عملیات منطقی و جبری تحت پوشش (و، یا) صورت می‌گیرد. از طرف دیگر در صورتی که کاربر فیلترینگ بازه زمانی، سایت‌ها یا گروه‌های پایه را انتخاب کرده باشد، این موارد نیز اعمال می‌گردند.

حذف کلمات عمومی: این قسمت به تشخیص، علامت‌گذاری و حذف کلمات عمومی (بسته به کاربرد) می‌پردازد. در اکثر زبان‌ها لیستی از کلمات عمومی پر تکرار وجود دارد، که در مرحله پیش پردازش آنها را حذف می‌نماییم. عملیات حذف در دو مرحله قبل و بعد از عمل ریشه‌یابی صورت گرفته و نقش بسزایی در کاهش تعداد کلمات قابل پردازش دارد. در مواردی که قصد ایجاد لیبل یا برچسب برای گروه‌ها داریم، جهت بالابردن خوانایی و وضوح برچسب در ابتدا کلمات عمومی را علامت‌گذاری نموده و در مرحله آخر عباراتی که تنها شامل کلمات عمومی هستند یا ابتدا و انتهای آنها با این کلمات است را حذف می‌کنیم. به عنوان مثال کلمه "برای" جزء کلمات عمومی محسوب می‌گردد اما بدون شک خوانایی عبارت "تلاش برای آزادی" از عبارت فاقد کلمه عمومی آن یعنی "تلاش آزادی" بسیار بیشتر است.

تقسیم بندی: هدف از این کار تشخیص محدوده جملات، عبارات و یا لغات قابل پردازش در یک متن است. استفاده از علائم نشانه‌گذاری در زبان HTML کاربرد فراوانی دارد. به عنوان مثال مطالب درون تگ title به راحتی قابل تشخیص و جداسازی هستند.

انتخاب کلمات کلیدی: پس از انجام پیش‌پردازش، فضای محاسباتی و لغات قابل پردازش به میزان چشمگیری کاهش پیدا کرده و تنها تعدادی کلمه کاندید باقی می‌مانند. حال باید به کمک یکی از روش‌های وزندهی به لغات و یا ترکیبی از آنها، ارزش هر کلمه را محاسبه نموده و کلمات نهایی را تشخیص دهیم. مهم‌ترین روش‌های وزندهی در بخش قبلی آورده شده است، حال با توجه به نوع داده‌های موجود باید به انتخاب یکی از متدها پرداخته و با قرار دادن حد آستانه مناسب با داده‌ها، تعدادی از آنها را به عنوان لیبل‌های اصلی انتخاب نماییم.

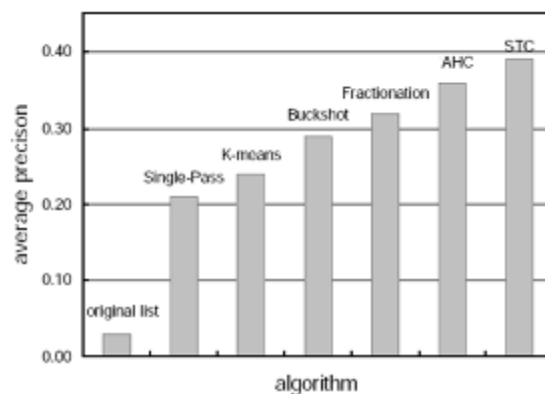
هر چند تمام روشها از مقبولیت و کارایی بالایی در پروژه‌های دسته‌بندی اسناد برخوردارند، اما همچنان ذات داده‌های مورد استفاده اهمیت بالایی دارند. به عنوان مثال در مورد استفاده از قطعات خبری به دلیل کوتاه بودن طول عبارات قابل پردازش فرکانس تکرار پایین است، و یا در مورد اسناد وب، بدون شک اهمیت واژه‌های آورده شده در عنوان و یا با فونت و حالت نوشتاری متفاوت اهمیت بیشتری دارد. عدم توجه به این امر باعث شده است که میزان دقت و کارایی در نمونه‌هایی که کل متن را پردازش می‌کنند نسبت به مواردی که بر روی قطعات کار می‌کنند، در زمانی قابل اغماض بالاتر باشد. این مطلب در شکل (2-4) نشان داده شده است.



شکل (2-4): مقایسه میانگین زمان اجرا و دقت برای تعداد دسته و درخواست مشابه در کل متن و قطعات وب [47]

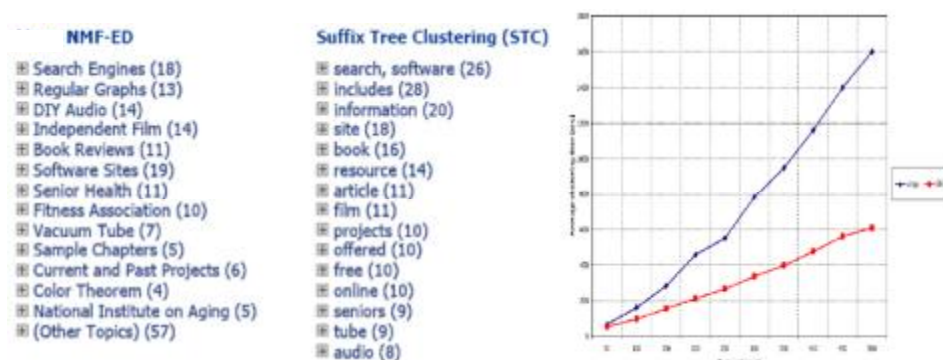
در خوانایی لیبل‌ها، اکثر موارد عبارات با طول بیشتر (بیش از یک کلمه) برای کاربر جذابیت و مفهوم واضح‌تری دارند. در نتیجه اولویت انتخاب با عبارات بوده و باید ضریبی برای افزایش وزن آن‌ها در نظر گرفته شود.

دسته‌بندی داده‌ها: پس از انتخاب کلمات کلیدی نوبت به دسته‌بندی اسناد در گروه‌های مجزا می‌رسد. برای این کار از دو دیدگاه کلی روش‌های مبتنی بر یادگیری ماشین و دسته‌بند می‌توان استفاده کرد. هر چند هر دو مورد به صورت مفصل در بخش قبلی بیان شده‌اند، اما در این بخش مقایسه اجمالی و کلی از موارد دسته‌بندها داشته و نتایج نهایی را نمایش خواهیم داد. لازم به ذکر است که در اکثر موارد از ترکیب روش‌های دسته‌بندی و یادگیری ماشین استفاده شده و این دو به صورت مکملی برای هم عمل نموده‌اند. نمودار مقایسه‌ای میانگین دقت الگوریتم دسته‌بندی برای 10 مجموعه سند در شکل (5-2) آورده شده است.



شکل (5-2): میانگین دقت الگوریتم‌های دسته‌بندی برای 10 مجموعه سند [18]

از میان تمام روش‌هایی که در بخش قبل توضیح داده شد، دقت الگوریتم STC در مورد دسته‌بندی داده‌های وب از تمام موارد بالاتر است. استفاده از ساختار درختی، توجه به عبارات و پیچیدگی زمانی کم از مهم‌ترین مزایای این روش به حساب می‌آید.



شکل (6-2): مقایسه روش‌های Lingo , STC راست: پیچیدگی زمانی، چپ: کیفیت و خوانایی لیبل‌ها [12]

شکل (6-2) به صورت خاص دو روش STC و Lingo را در دو معیار پیچیدگی زمانی و کیفیت برچسب‌های تولیدی مقایسه نموده است. همانطور که ملاحظه می‌گردد STC در زمان کوتاه‌تری قادر به دسته‌بندی تعداد داده مشابه است، در حالی که کیفیت، خوانایی و وضوح مطلب در Lingo بسیار بالاتر بوده و از بیان مباحث کلی در آن پرهیز شده است.

در قسمت دسته‌بندی قطعات با دو رویکرد متفاوت مواجه هستیم. در اولین مورد هر سند یا دسته اسناد جدید به تنهایی مورد پردازش قرار گرفته و با انتخاب کلمات کلیدی به دسته‌های جدیدی افزای می‌گردند، در ادامه دسته‌های جدید با موارد قبلی مقایسه شده و موارد مشابه با هم ادغام می‌گردند. روش دیگر انتساب اسناد جدید به گروه‌های قبلی است. که هر کدام از این موارد در جایگاه خود کاربرد دارند.

در انتها نیز عناوین به صورت دسته‌بندی شده به کاربر نمایش داده می‌شوند. برای مرتب‌سازی نمایش می‌توان از الگوریتم‌های دسته‌بندی و یا اولویت‌دهی بر حسب آخرین زمان به روز رسانی یا ایجاد دسته و یا بر حسب تعداد قطعات مرتبط با هر دسته یا برچسب استفاده نمود.

2-4-2 معرفی موتورهای دسته‌بند

در ادامه جدول (2-9)، شامل موتورهای دسته‌بند معروف به همراه الگوریتم‌های به کار رفته در آنها از مرجع [20] آورده شده است (به غیر از موارد مشخص شده). در این بخش تنها به ذکر موتورهای با بیشینه علمی پژوهشی اکتفا کرده و از موارد تجاری صرف نظر کرده‌ایم، چرا که در این موارد هیچ

گونه معیاری برای مقایسه و ارزیابی نهایی وجود ندارد. هرچند به دلیل اهمیت و اعتبار بالایی که موتور **vivisimo** دارد، در انتها به آن نیز اشاره نموده‌ایم. در ادامه شرح مختصری از نحوه عملکرد چند موتور به عنوان نمونه آورده شده و تصویری از خروجی **carrot2** نیز در شکل (2-7) نمایش داده شده است. باید به این موضوع توجه داشت که تمام موتورهای زیر بر اساس درخواست کاربر و ورود کلمه پرس‌وجو عمل می‌کنند.

System name (Alg. alias)	Year	Text features	Cluster labels	Clustering method	On- line demo	Cluster structure
Grouper (STC)	1998	Single words, phrases	phrases	STC	Yes (dead)	Flat
Lassi	2000	Lexical affinities	Lexical affinities	AHC	No	Hierarchy
CIIRarchies	2001	Single words	Word sets	Graph analysis	Yes (dead)	Hierarchy
WICE(SHOC)	2002	Single words, phrases	Phrases	SHOC	Yes (dead)	Hierarchy
Carrot ² (Lingo)	2003	Frequent phrases	Phrases	Lingo	Yes	Flat
WebCat	2003	Single words	Words	k-Means	Yes (dead)	Flat
AOsearch	2004	Single words	Word sets	AHC	Yes (dead)	Hierarchy
SnakeT [48]	2004	Approximate sentences	phrases	Approx. sent. coverage	Yes (dead)	Hierarchy
EigenCluster	2005	Single words	3 salient terms	Divide merge	No	Flat
WhatsOnWeb	2006	Single words	phrases	Edge connectivity	No	Graph
NeSReC[49]	2006	Single words, phrases	phrases	STC	No	Hierarchy

جدول (2-9): خلاصه‌ای از سیستم‌های پژوهشی دسته‌بندی نتایج جستجو [20]

Grouper

این موتور جستجوگر رابط دسته‌بند اسناد در سرویس جستجوگر چندگانه HuskySearch می‌باشد که در سال 1995 پیاده‌سازی شده است. HuskySearch، نتایج مورد نظر را از چندین موتور جستجوگر معروف دریافت کرده و Grouper، نتایج را به محض دریافت توسط روش دسته‌بند درخت پسوندی (STC) دسته‌بندی می‌نماید. در ادامه رابط کاربری این سیستم به صورت خلاصه شرح داده شده است. کار با ورود درخواست توسط کاربر آغاز می‌گردد. کاربر می‌تواند در این مرحله نحوه برخورد با درخواست خود را مشخص نموده (به عنوان عبارت، کلمات مجزا و ..) و همچنین تعداد و میزان شرکت پذیری هر کدام از موتورهای جستجوگر پایه و تعداد نهایی نتایج جستجو (10-200) را تعیین نماید. به صورت پیش فرض سیستم از 10 موتور استفاده کرده و متعاقباً بین 70 الی 1000 سند را پس از حذف اسناد مشابه دریافت می‌کند. پس از آن نتایج اصلی در قالب تعداد دسته‌های یافته شده و اسناد مرتبط با آنها در صفحه نمایش داده می‌شود. دسته‌ها در یک جدول بزرگ نشان داده شده و هر دسته یک ردیف مجزا را به همراه خلاصه‌ای سند به خود اختصاص می‌دهد. دسته‌ها بر اساس میزان ارتباط تخمین زده شده با درخواست مرتب می‌شوند. خلاصه هر دسته شامل سائز آن (تعداد اسناد) بوده و تلاش می‌کند برای انتقال بهینه محتوای دسته عبارات مشترک میان اسناد و عناوین نمونه را به نمایش گذارد. اعدادی که در پرانتز مقابل هر عبارت آورده شده است نمایانگر درصد اسنادی از دسته است که شامل آن عبارت خاص می‌باشند [50].

WICE (SHOC)

موتور خوشه‌بند اطلاعات وب⁵⁸ از الگوریتمی به نام دسته‌بند آنلاین سلسله مراتبی مفهومی⁵⁹ در سیستم خود استفاده نموده است که به صورت محلی داده‌ها را اداره کرده و با موفقیت با سیستم‌های آلفابت بزرگ کنار می‌آید. برای این کار از ساختار آرایه پس‌وندی به جای درخت پسوندی برای استخراج تعداد تکرار عبارات استفاده شده است [51].

⁵⁸ Web Information Clustering Engine

⁵⁹ Semantic Hierarchical Online Clustering

WebCAT

مبنای اصلی این روش که در سال 2003 معرفی شده است، استفاده از نوعی الگوریتم $kmeans$ برای دسته‌بندی داده‌های طبقه‌بندی شده می‌باشد. الگوریتم در اصل برای کار با پایگاه‌های داده توسعه داده شده است و در آن از پردازش تعاملی و تعاریف دقیق از تشابه/عدم تشابه میان اسناد به کمک کلمات طبقه‌بندی شده استفاده شده است. پیچیدگی زمانی این الگوریتم نسبت به تعداد اسناد، خطی بوده و دارای تعداد تکرار ثابت است [52].

SnakeT

این نام متعلق به سیستم دسته‌بند و الگوریتم مورد استفاده در موتور می‌باشد. ویژگی جالب آن ایجاد سلسله مراتبی از پوشه‌هایی است که احتمالاً با یکدیگر همپوشانی دارند. این سیستم ویژگی جدیدی به نام جملات تقریبی شامل عبارات غیر پیوسته (جملات تشکیل شده از عباراتی هستند که در میان آنها فاصله وجود دارد) را معرفی نموده است. پیچیدگی زمانی این الگوریتم برابر با $O(n \log 2n + m \log 2mp)$ است که در آن n , m , p به ترتیب از چپ به راست نمایانگر تعداد اسناد، تعداد ویژگی‌ها و تعداد لیبل‌ها است [53].

Carrot2

این سیستم در سال 2003 معرفی شده و ترکیبی از چندین الگوریتم دسته‌بندی نتایج جستجو از جمله STC، Lingo، ، دسته‌بندی مبتنی بر هوش ازدحامی (بر مبنای نظریه کلونی مورچگان) و تکنیک‌های ساده سلسله مراتبی می‌باشد. اساس کار Lingo استفاده از تکنیک SVD به عنوان مکانیزم اولیه‌ای برای ایجاد خوشه‌ها است. این موتور اغلب تمایل به ایجاد تعداد پوشه بیش از تعداد کل قطعات دارد که نتیجه آن تاثیر منفی بر روی قابلیت و کارایی سیستم است. از طرف دیگر این سیستم قادر به ادغام لیبل‌های مشابه مانند "دانش و کشف دانش" نبوده و هر کدام را به عنوان یک دسته منفرد در نظر می‌گیرد و این مسئله به صورت سلسله مراتبی اطلاعات اضافی هنگام مرور نتایج به وجود می‌آورد [20].



شکل (7-2): تصویری از خروجی موتور دسته بند carrot2 در مورد درخواست Iran
 { <http://search.carrot2.org/stable/search/> }

همان گونه که در شکل (7-2) مشاهده می گردد، نمایش خروجی به سه شیوه متفاوت پوشه ای، دایره-ای و ساختار درختی فوم آورده شده است. این کار علاوه بر دید زیبایی شناسی، دیدی عینی از نتایج به کاربر خواهد داد. به عنوان مثال در مورد آخر هر چه قسمت رنگی مختص به یک دسته بزرگتر باشد، بیشک اهمیت آن نیز بیشتر بوده و وسعت داده بالاتری را پوشش داده است. در نتیجه کاربر به آسانی می تواند عناوین مهم و مورد علاقه خود را شناسایی نموده و با مراجعه مستقیم به آنها در زمان خود صرفه جویی نماید.

Vivisimo

آخرین موتور مورد بحث، یک مورد تجاری است که هر چند در جدول مقایسه ای آورده نشده اما از کارایی بالایی برخوردار می باشد. به احتمال زیاد اولین موتور خوشه بند تجاری به نام Northern Light متعلق به اواخر دهه 1990 می باشد. این موتور شامل تعدادی دسته از پیش تعیین شده بوده و نتایج به دست آمده به هر کدام از این دسته ها انتساب داده می شود. موفقیت بزرگ دیگر در این زمینه موتور vivisimo است که در آن دسته ها و لیبل آنها به صورت اتوماتیک استخراج و تولید می شوند. این موتور توسط محققان دپارتمان کامپیوتر دانشگاه کارنگی ملون زیر نظر و تحت حمایت های مالی

بنیاد ملی علوم در سال 2000 تاسیس شده است. این موتور در سال‌های 2001 الی 2003 جایزه بهترین موتور جستجوگر چندگانه را توسط Search Engine Watch به خود اختصاص داده بود. اساس کار بر مبنای نظریه قدیمی هوش مصنوعی مبنی بر اینکه "یک سیستم دسته‌بند خوب، سیستمی است که دارای توضیحات خوب و خوانا باشد" پایه‌گذاری شده است. نرم افزارهای دسته‌بندی اسناد و موتور جستجوگر چندگانه به صورت اتوماتیک نتایج را به محض دریافت (on-the-fly) درون دسته‌های سلسله مراتبی قرار می‌دهند. سرعت این موتور با استفاده از مزایای استانداردهای XML و XSL از ساختاری مدرن بهره می‌گیرد. تا جایی که vivisimo قابلیت رسیدگی به اطلاعات ترابایتی با حداقل هزینه زیر ساختی نسبت به موارد مشابه در بازار تجاری را دارد. از دیگر سیستم‌های تجاری می‌توان به WebClust, KartOO و Mooter اشاره نمود [20].

2-4-3 سیستم دسته‌بندی اسناد در مقابل موتورهای دسته‌بند وب

الگوریتم‌های دسته‌بندی برای استفاده در سیستم‌های مدیریت پایگاه داده و سازماندهی صفحات وب در موتورهای جستجو توسعه داده شده‌اند. با توجه به اینکه اسناد مرتبط با یکدیگر، با درخواست‌های مشابه تطابق دارند، عمل دسته‌بندی اسناد در وب به منزله بهبود اثربخشی و کارایی در رتبه‌بندی اسناد صورت گرفته است. بر این اساس اسناد به گروه‌های مجزا دسته‌بندی شده و گروه‌های با بیشترین میزان تشابه به عنوان خروجی در رتبه‌های بالاتر به صورت سلسله مراتبی به کاربر ارائه داده می‌شود. موتورهای دسته‌بند اسناد به کاربر این اجازه را می‌دهند که تنها دسته‌های مورد نظرش را انتخاب نموده و به مطالعه آن بپردازد. ورودی این سیستم عبارت است از تعدادی نتیجه برگشته از موتورهای جستجو شامل URL، عنوان، تاریخ، سایت منبع و قطعه متنی (متنی کوتاه که به صورت خلاصه محتوای سند را نشان می‌دهد) و خروجی آن نمایش تعدادی دسته لیبل‌گذاری شده در ساختاری سلسله‌مراتبی، درخت‌مانند یا پارتیشن‌های مسطح است. در جدول (2-10) عمده تفاوت‌های میان این سیستم‌ها آورده شده است.

Clustering type	Cluster label	Cluster computation	Input data	Cluster intersection	GUI
Search results clustering	Pre Defined	Online	Snippets	Overlapping	Yes
Document clustering	Centroid	Offline	Documents	Disjoint	No

جدول (2-10): دسته‌بندی نتایج جستجو در مقابل دسته‌بندی‌های عمومی و سنتی متون [54]

جدول فوق تفاوت‌های کلی میان دسته‌بندی‌های عمومی و سنتی متون و اسناد وب را نشان می‌دهد. منظور از دسته‌بند عمومی به عنوان مثال سیستم دسته‌بندی عناوین کتب درون یک کتابخانه الکترونیک می‌باشد. برچسب اصلی هر دسته که به گونه‌ای نماینده اعضا مجموعه محسوب می‌گردد در اسناد وب، نوعی برچسب مفهومی شامل یک عبارت، لغت و یا توصیفی به زبان طبیعی گفتاری است در حالی که در سیستم‌های دسته‌بندی متون اغلب از یک ویژگی آماری مانند میانگین نمایش عددی مفاهیم استفاده می‌گردد در نتیجه خوانایی و قابلیت درک برچسب در سیستم‌های اول اهمیت بالاتری دارند. به همین دلیل در صورت اضافه‌شدن یک سند جدید به اسناد در وب، برچسب همچنان بدون تغییر باقی می‌ماند اما در سیستم‌های متنی، میانگین داده‌های دسته هرچند بسیار کم و نامحدود تغییر می‌یابد. در سیستم‌های عمومی داده‌ها از پیش در یک پایگاه ذخیره شده و برای پردازش آماده گشته‌اند در حالی که اسناد وب در هر لحظه بنا به درخواست کاربر و یا گذر زمان به صورت مداوم و پیوسته تغییر کرده و باید عملیات به صورت آنلاین بر روی آنها صورت گیرد. نوع داده‌های مورد پردازش نیز در وب، قطعات بدست آمده از موتورهای جستجوگر است که اغلب شامل عنوان، خلاصه‌ای کوتاه، تعداد تگ معنایی، آدرس و وب‌سایت مرجع است. در حالی که در دسته‌بندی متون از کل متن و محتوای موجود که اغلب طولانی است، استفاده می‌کنیم. تعداد دسته‌های نهایی در متون مشخص بوده و تنها باید انتساب اسناد به گروه‌ها صورت گیرد ولی در وب این تعداد نامشخص بوده و هر داده می‌تواند در زمان واحد به بیش از یک گروه با درصدهای متفاوت تعلق داشته باشد. در نهایت برای سیستم‌های وب نیاز به یک واسط کاربری برای دریافت اطلاعات و نمایش خروجی و به

صورت کلی تعامل با کاربر وجود دارد ولی در دسته‌بندی متون انتساب داده‌ای به منزله سهولت دسترسی خود سیستم بوده و نهایتاً نیاز به یک صفحه برای نمایش اطلاعات داریم.

2-4-4 ارزیابی عملکرد موتورهای دسته‌بندی

6 معیار کلی برای ارزیابی عملکرد سیستم‌های سیستم‌های بازیابی اطلاعات وجود دارد که عبارتند از: پوشانندگی درخواست⁶⁰، فراخوانی⁶¹، دقت⁶²، زمان پاسخ⁶³، تلاش کاربر⁶⁴ و فرمت خروجی⁶⁵. از آنجا که این معیارها در سال 1973 [55] معرفی گشته‌اند بسیاری از آنها ارزش خود را تا حدی کاهش داده‌اند، به عنوان مثال میزان تلاش کاربر برای دسترسی به نتایج و ارسال درخواست خود با توجه به طراحی رابط گرافیکی بسیار کاهش یافته است، اما همچنان برای سیستم‌های کنونی مورد استفاده قرار می‌گیرد. اما در عین حال توجه به جنبه‌های زیر نیز بسیار اهمیت دارد [54].

ترکیب شاخص‌های وب: زمانی که یک درخواست ایجاد می‌شود، شاخص نتایج به کمک ربات‌ها و عنکبوت‌ها ایجاد می‌گردد نه توسط خود صفحات. در نتیجه ترکیب شاخص‌های وب تاثیر مهمی بر عملکرد موتور جستجو دارد.

سه جزء کلی وجود دارد که در ساختار آرایش شاخص وب اهمیت دارد که عبارتند از پوشش وب، زمان به‌روزرسانی و بخش‌های شاخص‌گذاری (به عنوان مثال عناوین به انضمام چند خط اول، یا کل صفحه وب). از سوی دیگر محدوده پوشش بزرگتر، به روز رسانی مکرر و نمایه‌سازی کل متن لزوماً به معنی بهتر بودن موتور جستجوگر نیست.

قابلیت‌های جستجو: موتور جستجوگر باید امکانات اساسی که کاربر با آنها آشنایی دارد را داشته باشد که شامل عبارات منطقی بولین، جستجو عبارت، اختصارات و تسهیلات محدود ساز (محدودیت تعداد و حجم اطلاعات) باشد.

⁶⁰ Coverage

⁶¹ Recall

⁶² Precision

⁶³ Response time

⁶⁴ User effort

⁶⁵ Form of output

عملکرد میزان بازیابی: این عملکرد به صورت سنتی شامل سه پارامتر دقت، فراخوانی و زمان پاسخ می‌باشد.

مجموعه اسناد مرتبط با یک پرس وجو را با $\{Relevant\}$ نشان می‌دهیم، و مجموعه اسناد بازیابی شده را با $\{Retrieved\}$ نشان می‌دهیم. مجموعه اسنادی که هم مرتبط هستند و هم بازیابی شده-اند با $\{Relevant\} \cap \{Retrieved\}$ نشان داده می‌شوند.

دو مقیاس اصلی برای تشخیص کیفیت بازیابی متن وجود دارد:

- دقت: در صد اسناد بازیابی شده که در واقع به پرس‌وجو مرتبط هستند (برای مثال پاسخ‌های «درست»)

$$precision = \frac{|\{Relevant\} \cap \{Retrieved\}|}{|\{Retrieved\}|} \quad (4-2)$$

- فراخوانی: در صد اسنادی که به پرس‌وجو مرتبط هستند و در واقع بازیابی شده‌اند.

$$recall = \frac{|\{Relevant\} \cup \{Retrieved\}|}{|\{Retrieved\}|} \quad (5-2)$$

- معیار F : که ترکیبی از دو روش فوق بوده و به صورت زیر محاسبه می‌گردد.

$$F - measure = \frac{2 * Recall * Precision}{Recall + Precision} \quad (6 - 2)$$

فرمت و شمای خروجی: این جزء ارزیابی باید از دو دیدگاه مورد بررسی قرار گیرد. یکی از موارد تعداد گزینه‌های خروجی و دیگری توجه به محتوای حقیقی خروجی است.

تلاش کاربر: این گزینه بیشترین ارتباط را با نحوه آماده سازی مستندات و رابط کاربری تهیه شده دارد. از آنجا که تعداد بسیار زیادی موتور جستجوگر وجود دارد، باید به میزان سهولت کار برای کاربر به منزله جذاب نمودن هرچه بیشتر آن توجه نمود.

5-2 دسته‌بندی اسناد خبری در وب

با توجه به ماهیت پویای اطلاعات در صنعت خبرگزاری و به‌روزرسانی ثانیه‌ای اخبار در بستر وب در ساعات مختلف شبانه روز و همچنین عدم تطابق زمانی در مناطق مختلف جغرافیایی، نیاز به مراجعه مکرر به وبسایت‌ها برای مطلع‌شدن از جدیدترین اخبار وجود دارد. از سمت دیگر به دلیل عدم امکان جمع‌آوری اخبار در ایام تعطیلات و ساعات غیر اداری و وجود تعداد زیاد سازمان‌های خبری معتبر در جهان، طراحی سیستمی اتوماتیک برای این منظور لازم به نظر می‌رسد.

در برخی موارد شاهد انتشار اخبار ضد و نقیض بسیاری هستیم و متأسفانه به دلیل گستردگی شبکه نت و دسترسی تمام افراد به آن، امکان نظارت مراجع اجرایی بر روی تمام موارد وجود ندارد، لذا باید سیستمی جهت تطبیق و ارزش‌دهی به اخبار تعبیه گردد. از طرف دیگر در بسیاری موارد مشاهده می‌گردد که اخبار منتشر شده مکمل اخبار قدیمی بوده و باید به منظور شفاف‌سازی بیشتر خبر از ترکیب، ادغام و نهایتاً جمع‌بندی آنها استفاده نمود.

در نتیجه می‌توان این چنین بیان نمود که تضاد عمیقی میان حجم اخبار و رشد سریع منابع قابل جستجو و عدم توانایی در کنترل و زمانبندی ورود این اطلاعات به شبکه و توان، زمان و ساعات کاری نیروی انسانی وجود دارد و برای مقابله با آن نیاز به طراحی و پیاده‌سازی سیستمی برای دسته‌بندی اخبار انتشار یافته در وب وجود دارد. در این بخش به صورت خاص اسناد خبری را تحلیل و آنالیز کرده و قصد طراحی سیستمی به منزله دسته‌بندی قطعات خبری به دست آمده از موتورهای جستجو چندگانه داریم. بیشترین کاربرد این طرح، استفاده از آن در خبرگزاری‌ها و سازمان‌های دولتی به منزله ایجاد مزیت رقابتی در میان دیگر سازمان‌ها در زمان و اعتبار دسترسی و بازپخش اخبار و دسته‌بندی بدون ناظر آنها بسته به نمودار خبری مختص سازمان است، هر چند در ادامه می‌توان از این روش به منزله شخصی‌سازی اسناد خبری برای کاربران این صفحات نیز استفاده نمود.

2-5-1 تفاوت اسناد وب و اسناد خبری

دسته‌بندی اسناد و نتایج حاصل از جستجو در وب یکی از شاخه‌های دسته‌بندی اسناد متنی است با این تفاوت که در وب با حجم بالا و غیر قابل کنترل، نظارت و پیش‌بینی صفحات مواجهیم. علاوه بر اینکه احتمال تکراری بودن محتوای صفحات یا تکرار محتوای یک سند در وب وجود دارد، از سمت دیگر با تغییر مداوم مفهوم و متن درون صفحات وب مواجه هستیم. داده‌های موجود در صفحات وب اغلب به صورت بدون ساختار بوده و از ناهمگونی بالایی برخوردار بوده و در اکثر موارد شامل اطلاعات اضافی بسیاری مانند اتصالات درون و برون صفحه‌ای، ابر اطلاعات و ... می‌باشند.

بنا به دلایل فوق در سیستم‌های دسته‌بندی وب به سیستم‌هایی با سرعت پردازش بالا، قابلیت انعطاف‌پذیری در مقابل تغییرات مداوم محتوا و کاربرگرا نیازمندیم. تعداد دسته اولیه به ازاء هر خبرگزاری ثابت است، حال آنکه هیچ استاندارد کلی که تمام موسسات از آن پیروی کنند وجود ندارد. از طرف دیگر هر خبر می‌تواند در چندین حوزه تاثیرگذار باشد. بدلیل اهمیت اعتبار و صحت اخبار انتشاری، تنها باید به سایت‌های معتبر مراجعه نمود، هر چند تعیین و ردگیری سایت‌های پر بازدید و معتبر ثانویه و افزودن آنها به سیستم نیز امکان‌پذیر است.

هر عنوان خبری مربوط به یک شخصیت، مکان و زمان به خصوص است که این موضوع اهمیت تشخیص موجودیت‌های اسمی⁶⁶ را بیشتر می‌کند، از طرف دیگر هر حادثه بلافاصله پس از به وقوع پیوستن توسط اکثر خبرگزاری‌ها انتشار می‌یابد (حجم اخبار مشابه در بازه زمانی کوتاه) و در اکثر موارد از نوع نوشتار رسمی (زبان و قواعد) استفاده می‌شود.

⁶⁶ Name Entities

2-5-2 چالش‌های موجود در وب

برای انتشار اخبار صحیح باید اعتبار سنجی (به کمک سایت‌های اولیه معتبر) انجام گیرد. وجود اختلاف زمانی بدلیل موقعیت‌های جغرافیای متفاوت و عدم امکان پوشش 24 ساعته اخبار بدلیل ساعت کاری و نیرو و توان محدود انسانی، امکان دسترسی دستی به اخبار دقیق در کوتاه‌ترین زمان ممکن را دشوار می‌سازد که این امر می‌تواند در اخذ تصمیمات، مدیران را با چالش مواجه سازد. برخی از مهم‌ترین چالش‌های موجود در این زمینه در زیر لیست شده‌اند.

- تطبیق اخبار و اعتبار سنجی
- توان، زمان و ساعات کار محدود نیروی انسانی
- اطلاع دقیق و سریع از اخبار جهت اخذ تصمیم
- عدم وجود استاندارد خبری و وجود شاخه‌های غیر مشابه
- پراکندگی زمانی بدلیل تفاوت در مختصات جغرافیایی و میزان تاثیر پذیری

2-5-3 بررسی دیدگاه‌های متفاوت

ایجاد فضایی برای تشخیص گروه خبر

اولین دیدگاه در دسته‌بندی اخبار تعریف یک چارت سازمانی دقیق به همراه موجودیت‌های ارزشمند در هر کدام از دسته‌ها است. بدین صورت می‌توان با وارد نمودن متن و عنوان خبر میزان تعلق آن به گروه‌های مختلف استاندارد و از پیش تعیین شده را تشخیص داد. نتیجه این رویکرد در شکل (2-8) آورده شده است.

AUTOMATIC CLASSIFICATION OF NEWS ARTICLES IPTC

Title of the article (may be empty):

Content of the article:

The agreement was hailed as a step forward for Greece, but doubts immediately emerged as to whether it would do much more than deal with its most pressing debt problems.

Greece will need more help if it is to bring its debt down to the level envisaged in the bailout and will remain "accident prone" in coming years, according to a deeply pessimistic report by international experts obtained by Reuters.

After 15 hours of talks, ministers finalized measures to cut Greece's debt to 120.5 percent of gross domestic product by 2020, a fraction above the target, to secure its second rescue in less than ten years and mark a hard turnaround in March.

Classify only in:

(in specific)

Accept

Clear

Example: Spanish English Portuguese Catalan

شکل (2-8): ایجاد فضایی برای تشخیص گروه خبر توسط موتور دسته بند نتایج
چپ: عنوان یا متن خبری ورودی، راست: دسته های متعلق به خبر و درصد تعلق
{<http://showroom.daedalus.es/>}

Automatic classification:

Topic:

- 1- economy, businesses and finances - macroeconomics - international economic institutions (PTC 0400010) (100.0%)
- 2- economy, businesses and finances - macroeconomics - central banks (PTC 0400001) (94.0%)
- 3- labor - collective agreements - agreements/labor standards (PTC 0600003) (82.1%)
- 4- economy, businesses and finances (PTC 0400000) (75.8%)
- 5- economy, businesses and finances - financial and commercial services - banking services (PTC 0400002) (66.2%)
- 6- politics - elections (PTC 1100003) (55.1%)
- 7- economy, businesses and finances - macroeconomics - bonds (PTC 0400007) (54.3%)
- 8- catastrophes and accidents - rescue (PTC 0200003) (54.1%)
- 9- catastrophes and accidents (PTC 0200000) (49.8%)
- 10- economy, businesses and finances - macroeconomics - state debt (PTC 0400006) (46.4%)
- 11- economy, businesses and finances - information for companies - shopping (PTC 0401006) (44.7%)

دسته‌بندی موضوعات خبری

یکی دیگر از دیدگاه‌های مطرح دسته‌بندی موضوعات خبری به عناوین داغ و جدید روز، عناوین پر بازدید، عناوین اخیر و ... است که در این موارد بیشترین پردازش و آنالیز بر روی اطلاعات اضافی خبر (زمان انتشار، تعداد کلیک، آنالیز لینک های درونی و بیرونی و ...) بوده و کمتر به مفاهیم و محتویات خبر پرداخته می‌شود. در این زمینه سایت‌های زیادی مانند google news و yahoo news کار کرده‌اند. اما این شیوه نیز نمی‌تواند دیدی دسته‌بندی شده و آسان در اختیار کاربر قرار دهد. این نمونه در مرجع [54] نیز بیان شده و در شکل (2-9) نمایش داده شده است.



شکل (2-9): دسته بندی موضوعات خبری توسط موتور دسته بند نتایج
{<http://news.google.com/>}

ایجاد و ردیابی جریان خبری

در اغلب موارد مردم علاوه بر جستجو کلمه‌ای عبارات تمایل دارند از امور جاری سر در بیاورند. یکپارچه سازی اطلاعات جامع در مورد موضوعات خبری امری لازم است که آنرا "جریانات خبری" می‌نامیم و شامل پس‌زمینه، تاریخچه، پیشبرد جریانات، نقدها و نظرات مختلف پیرامون آن است. به طور سنتی این کار توسط ویرایشگران خبری در خبرگزاری‌ها و مراجعه به بایگانی و آرشیو خبری صورت می‌گیرد که بسیار زمانبر بوده و امکان بروز بودن و دسترسی آنی به آنها وجود ندارد. نمونه یک جریان خبری در شکل (2-10) آورده شده است.

این کار به ایجاد یک تاریخچه غیر مدون از اخبارها در بازه زمانی طولانی می‌پردازد. به عنوان مثال تمامی جزئیات و اخبار مربوط به مسابقات المپیک لندن از روز نخست تا امروز، این سیستم‌ها هنوز به صورت عمومی قابل استفاده نبوده و در حال حاضر به صورت پژوهش‌های علمی [57, 58] به آن پرداخته می‌شود.

Source Rank	Daily Rank	Weekly Rank	Monthly Rank
News with image: Jiao Liuyang, champion of women 200 meter butterfly stroke of the 6th City Games 13:49:00 Oct 26, 2007 from SINA Jiao Liuyang, an athlete from Harbin, left the competition terrain after the match on Oct 26. She won a golden medal with a 2'06"16 result in the final of women 200 meter butterfly stroke of the 6th City Games held in Wulumu. Photo by She dan, a reporter from Xinhua news agency. News with image: Xie Zhi in the final of men 200 meter butterfly stroke of the 6th City Games -SOHU- (3 related news) Result of the final of men 200 meter butterfly stroke of the 6th City Games -SOHU- (1 relate news) News with image: Lai Zhongxin in the final of men 200 meter butterfly stroke of the 6th City Games -SOHU- (1 related news) 4 related topics>>			
Amare Stoudemire is back and Grant Hill is leading, Suns defeat Nuggets: 116-113 13:06:00 Oct 26, 2007 from TOM News from TOM sports: Suns play at home with Nuggets, and win 116-113 on Oct 26, Beijing time. Amare Stoudemire is back from an operation and plays as main force. Suns use the complete main force, and Grant Hill substitutes Boris Diaw as the main force. Amare Stoudemire is back and Suns defeat Nuggets: 116-113 -ENORTH- (1 related news) Live records of the first period, Nuggets vs Suns: 32-22 -ENORTH- (1 relate news) "AI combination" unable to win for Nuggets, Amare Stoudemire helps Suns defeat Nuggets by 3 points -SPORTS.CN- (1 related news) 5 related topics>>			
News with image: [Preseason of NBA] Bucks defeated by Bulls, Dan Gadzuric is defending 12:33:00 Oct 26, 2007 from SOHU News from SOHU sports: The preseason of NBA is continuing at 8:30 on Oct 26, Beijing time. Milwaukee Bucks was defeated away by Chicago Bulls and have now lost two straight games. Jianlian Yi played for 21 minutes, made 2 of 7 shots and 3 of 4 penalty shots, won 7 marks and 3 rebounds, and made 2 turnovers and 3 personal fouls. Images: S1-97, Bucks defeated away by Bulls -SPORTS.PEOPLE.COM.CN- (4 relate news) [News with image] Jianlian Yi in the match between Bucks and Bulls -TOM- (4 related news) 6 related topics>>			

شکل (2-10) ایجاد جریان خبری توسط موتور دسته بند نتایج [58]

فصل سوم: مروری بر تحقیقات

3-1 مقدمه

در این فصل به معرفی مجموعه‌ای از پژوهش‌های انجام شده در حوزه‌های مورد بحث می‌پردازیم. بدین منظور، شماری از مطالعات صورت گرفته در زمینه پردازش و دسته‌بندی متون، موتورهای جستجوگر و دیگر فیلدهای مرتبط با پروژه را معرفی خواهیم نمود. لازم به ذکر است که در این بخش از بیان جزئیات و دوباره‌گویی عملکرد روش‌ها و الگوریتم‌هایی که در فصل قبل بیان گشته‌اند، چشم‌پوشی کرده و تنها به بیان کلیات خواهیم پرداخت.

نحوه استفاده از تکنیک خوشه بندی با نظارت برای دسته‌بندی نامه‌های دریافتی در سامانه 137 شهرداری تهران در [26] شرح داده شده است، که در آن به کمک الگوریتم بیزین عمل دسته‌بندی صورت گرفته است. در این پروژه بعد از پیش‌پردازش داده‌های خام به منظور تبدیل آنها به فرمت مناسب برای آموزش و وزندهی به عبارات، پیام‌های فارسی ارسالی از طرف شهروندان از طریق موبایل به گروه‌های از پیش تعیین شده نسبت داده شده‌اند.

در مقاله [30] از یک روش آماری ترکیبی، برای استخراج کلمات کلیدی اسناد فارسی استفاده شده است. در این روش پس از اعمال پیش‌پردازش‌های مورد نیاز ویژگی‌های آماری برای کلمات مختلف محاسبه شده و با استفاده از فازی‌سازی و اعمال قواعد فازی، کلمات کلیدی محتمل، انتخاب می‌شوند. گام بعدی محاسبه رخداد همزمان 4 پیشین و پسین کلمات کلیدی محتمل، با کلمات تکرار شونده، 5 در جملات سند است. در ادامه با اعمال حد آستانه مشخص وفقی روی رخداد همزمان کلمات، عبارات دو کلمه‌ای تشخیص داده می‌شوند.

در این پروژه بر روی سه گروه کلمه عمومی (کلماتی که در کلیه اسناد کلیه کلا سها پخش شده‌اند)، کلمه خاص (کلماتی با فراوانی زیاد فقط در کلیه اسناد یک کلاس و فراوانی کم در سایر کلا سها) و کلمات کلیدی (عبارت یا کلمه مهم و دارای فراوانی زیاد در سند خاص) مانور زیادی داده شده است. به همین منوال از روش‌های وزن‌دهی متنوعی نیز در سه گروه کلمه فوق استفاده شده است.

مقاله [44] با هدف یافتن و استخراج کلمات کلیدی مستندات فارسی، به منظور طبقه‌بندی کارآمد آنها در موتورهای جستجو به ترکیب چند تکنیک ریشه‌یابی و خلاصه سازی اسناد پرداخته است. پس از ورود هر متن عملیات حذف وندها و کلمات اضافه و افعال صورت گرفته و به هر جمله امتیاز مشخصی داده می‌شود. سپس به حذف کلمات هم‌معنا و تعمیم کلمات به کمک دیکشنری‌های جانبی پرداخته می‌شود. در تعریف فاکتور اهمیت برای کلمات با در نظر گرفتن محدوده‌ای مشخص میزان تکرار بین حد فرکانس پایین و بالا بر اساس تعداد کلمات موجود در متن را محاسبه نموده و منبع کلمات مهم را ایجاد می‌کنیم. در نهایت درجه اهمیت هر کلمه توسط حاصل ضرب فرکانس تکرار کلمه در سند در تعداد جملات شامل کلمه در سند محاسبه می‌شود. در مرحله بعد به منزله یافتن فاکتور اهمیت جمله، با تعیین حدی به عنوان حداکثر فاصله دو کلمه مهم، یک بخش در جمله را پیدا می‌کند که به وسیله کلمات مهمی که بیش از L کلمه غیرمهم بینشان قرار ندارد، جدا شده‌اند. سپس تعداد کلمات مهم موجود در این بخش را شمرده و جذر این تعداد را بر تعداد کل کلمات موجود در آن بخش تقسیم نموده و فاکتور اهمیت جمله را بدست می‌آورد. بعد از آن به حذف کلمات هم معنی با استفاده از لغت‌نامه و انتخاب کلمه جایگزین پرداخته می‌شود که نتیجه آن همگونی و ساده سازی متون می‌باشد. در این پروژه از تعداد 100 متن غیر استاندارد در 5 گروه و مقایسه آن با روش سنتی و دستی استفاده شده است.

مقاله [49] با گرفتن در خواست از کاربر، جمع‌آوری تکه‌های خبری توسط Altavista به مدیریت نتایج جستجو به منزله جلوگیری از مشاهده نتایج تکراری توسط کاربر پرداخته و در نهایت بر مبنای لیبل‌گذاری، عمل دسته بندی را انجام می‌دهد. در این روش پس از دریافت نتایج اولیه از موتور جستجوگر به استخراج عناوین کاندید و اولویت‌بندی عنوان‌ها پرداخته و در نهایت حاصل دسته‌بندی نمایش داده می‌شود. تمام کلمات غیر مهم بدون نیاز به ریشه‌یابی و همچنین عناوین با کمتر از طول 3 کلمه حذف گشته و با درجه‌بندی عناوین کاندید از ریزاسناد خبری، عنوان برتر انتخاب می‌گردد. در

نتیجه عمل دسته‌بندی به جای استفاده از تشابه میان متون، از تشابه لیبل‌های انتخابی استفاده می‌کند.

در مرجع [56] رویکردی جدید در یافتن و دسته‌بندی موارد مشابه به کمک روش SVM در اخبار فارسی مورد بررسی قرار گرفته است. از آنجا که مطالعه اخبار مشابه در وبسایت‌ها و اعتبار سنجی آنها برای کاربران بسیار زمانبر است، سیستم دسته‌بندی برای این منابع در نظر گرفته شده و از عنوان و توضیحات اخبار برای دسته‌بندی استفاده شده است. در نتیجه در پایگاه داده مربوطه آدرس، تاریخ و کلمات کلیدی ذخیره می‌شود. مرحله اول پروژه به استخراج ویژگی متن یک خبر و ذخیره‌سازی آن در پایگاه داده محلی پرداخته و مرحله دوم به استخراج اخبار مشابه بر اساس ویژگی‌های مشخص شده در دیگر صفحات می‌پردازد. این پروژه بر مبنای درخواست کاربر بوده و عمل پالایش و حذف کلمات عمومی در آن به دو صورت آفلاین هنگام پویش اخبار از صفحات وب و آنلاین به معنای حذف کلمات اضافی در هنگام دریافت درخواست کاربر صورت می‌گیرد. در مورد پیدا کردن قطعات مشابه اخبار نیز اولویت بندی قطعات اخبار به صورت تشابه کامل با تمام کلمات درخواست، تشابه با تعداد کلمات منهای یک و ... انجام می‌گیرد.

در اغلب موارد مردم علاوه بر جستجو کلمه‌ای، تمایل دارند از امور جاری با خبر شوند. یکپارچه سازی اطلاعات جامع در مورد موضوعات خبری امری است که آنرا "جریان‌ات خبری" می‌نامیم و شامل پس زمینه، تاریخچه، پیشبرد جریان‌ات، نقدها و نظرات مختلف پیرامون موضوع بوده و مقاله [58] به آن پرداخته است. به طور سنتی این کار توسط ویرایشگران خبری در خبرگزاری‌ها و مراجعه به بایگانی صورت می‌گیرد. در این مقاله الگوریتم سه مرحله‌ای تشخیص موضوع، ردیابی خبر و ادغام موضوعات با یکدیگر، برای ایجاد روند جریان‌ات اخبار در وب پیشنهاد شده است. این سیستم چندین جریان خبری شامل "خبرهای تازه"⁶⁷، "خبرهای داغ"⁶⁸ و "چه خبر"⁶⁹ را پوشش می‌دهد که پرسیدن آنها

⁶⁷ What's new

⁶⁸ What's hot

⁶⁹ What's going on

به صورت درخواست نتیجه‌ای در موتورهای جستجو نخواهد داشت. عناوین خبری اصولاً برای چند روز اهمیت داشته و سپس در مدتی کوتاه ارزش خود را از دست می‌دهند، در حالی که جریانات خبری در مدت زمان طولانی‌تری پایداری داشته و به رکوردهای قابل توجهی تبدیل می‌شوند.

مقاله [59] به بررسی روش‌های مختلف یادگیری ماشین برای دسته‌بندی اتوماتیک اسناد به زبان فارسی پرداخته و دو روش SVM و KNN را با یکدیگر مقایسه نموده است. در این مقاله نیز پس از انجام پیش پردازش و ایندکس گذاری داده‌ها به روش TFIDF و TFDRE، کلمات مهم انتخاب شده‌اند. از 800 سند موجود در همشهری متعلق به 8 دسته مجزا جهت آموزش و تست سیستم استفاده شده است. بهترین نتیجه در روش SVM با کرنل polynomial و در روش KNN مجاورت یک همسایه با معیار شباهت Cosine تشخیص داده شده است. این مقاله بیشتر بر مبنای مقایسه روش‌ها و توجه به ساختار کلمات در زبان فارسی عمل نموده است.

در مقاله [60] کار استخراج کلمات کلیدی بر روی جملات و عبارات غیر رسمی موجود در شبکه‌های اجتماعی به منزله مدیریت تبلیغات و علاقه‌مندی‌های کاربران است. این ریز اسناد اغلب شامل مجموعه کلمات بیانگر وضعیت، علایق و اخبار شخصی برای ارائه به دوستان است. به دلیل طول کوتاه اغلب TF برابر با 1 است، در نتیجه این تعداد تکرار کم، فرکانس کارایی مناسبی ندارد. در این روش تمرکز بر روی موقعیت نسبی کلمه در عبارت، طول عبارت، تعداد کل کلمات و ارزش کلمه، بر حسب نوع نوشتار کلمه با حروف بزرگ یا کوچک است. در نهایت کلمات با امتیاز بالاتر از حد آستانه مشخص، انتخاب می‌شوند. در این پروژه از 1830 عبارت موجود در شبکه اجتماعی "فیس بوک" استفاده شده است.

در مرجع [61] یک دسته‌بند اسناد فارسی با استفاده از الگوریتم K-NN پیشنهاد شده و دلیل استفاده از آن نیز موفقیت زیاد در دسته‌بندی اسناد مشابه در زبان انگلیسی عنوان گشته است. در الگوریتم پیشنهادی، از 540 متن فارسی برگرفته از روزنامه‌های آنلاین ایران و جام‌جم، به عنوان داده آموزشی و از 60 متن مشابه به عنوان داده آزمایشی استفاده شده است. متن‌ها در 6 دسته از پیش

تعریف شده قرار گرفته و اندازه‌گیری مشابهت اسناد که به وسیله برداری از فرکانس خصیصه‌ها در هر سند نمایش داده می‌شوند، از طریق محاسبه شباهت کسینوسی صورت پذیرفته است. برای ارزیابی کارایی، میزان دقت و یادآوری مورد بررسی قرار گرفته و نتایج به دست آمده حاکی از آن است که K-NN قادر است با تعداد کم داده آموزشی به خوبی کار نماید.

یک روش سریع برای خوشه بندی نتایج جستجوی وب نیز در مقاله [62] معرفی شده که در آن از خوشه‌بندی به عنوان روشی هوشمند برای بهبود نمایش نتایج جستجو برگشتی از موتورهای جستجوی وب استفاده شده است. در ابتدا به کمک روش برداری، ماتریس عبارت-متن ایجاد شده و در نهایت با تکیه بر مفهوم دریافتی از موضوع خوشه بندی صورت می‌گیرد. این پروژه نیز مبتنی بر درخواست کاربر بوده و در آن اطلاعات مورد نیاز از موتورهای جستجوگر بدست می‌آیند که هر نتیجه جستجو شامل عنوان، توضیحات، خلاصه و آدرس با ضرایب اهمیت متفاوت است. وزن‌دهی نهایی لغات نیز بر اساس محل قرار گرفتن لغت و تعداد دفعات تکرار آنها صورت می‌گیرد. در این پروژه به کلمه ای لغت کاندید گفته می‌شود که در بخش عنوان یا توضیحات بوده و در بیشتر از دو بخش نتایج دیده شود. در ادامه مباحث بدست آمده از جستجو خوشه بندی شده به ترتیب تعداد سند موجود در هر کدام که نشانگر اهمیت دسته است، نمایش داده می‌شوند. عدم وابستگی به مقادیر آستانه جز مزایای این روش محسوب می‌گردد.

[63] به دسته‌بندی نتایج جستجو بر مبنای ویژگی‌های مدارک و امکان‌سنجی استفاده از الگوریتم-های خوشه‌بندی مختلف در سطح وب پرداخته است. از آنجا که بازآرایی و سازماندهی مدارک همواره بر اساس ویژگی‌های هر مدرک صورت می‌پذیرد، بازسازی نتایج جستجو نیز منوط به ویژگی‌های مذکور است. ویژگی‌های مدارک در این پروژه به دو گروه ردیف اول و ردیف دوم دسته‌بندی شده است. در این مقاله، با اشاره به انواع ویژگی‌های مطرح برای مدارک و اشاره به این نکته که خوشه‌بندی یکی از روشهای رایج در دسته‌بندی نتایج جستجو است، تلاش شده تا فصل مشترک میان این ویژگی-ها و نیز روشهای مختلف خوشه‌بندی مشخص گردد. این در حالی است در سطح وب، امکان محدود

ساختن نتایج جستجو به تاریخ، عنوان، دامنه، قالب مدرک، در نظر گرفتن پیوندهای فرامتن در نتایج جستجو و نیز حتی جستجو در پی ناشر اثر، همگی تاکنون در ابزارهای بازایی اطلاعات مورد توجه واقع شده است و نه در الگوریتم‌های دسته‌بندی. در نهایت، با توجه به اینکه خوشه‌بندی نتایج جستجو از فناوریهای نسبتاً رایج در سطوح وب بوده و طراحان موتورهای جستجو برای بهبود نتایج از این ویژگیها در بازآرایی نتایج خود استفاده نمی‌کنند، در پروژه مقایسه‌ای جهت نشان دادن این مهم صورت گرفته است.

مرجع [63]، روشی ارائه نموده که با توجه به خصوصیات متن قرآن کریم قادر است ابتدا آیات را درون موضوعاتی از پیش تعریف شده دسته‌بندی کند، سپس آیات همه سوره‌ها را با توجه به وزن کلمات آن برای هر موضوع، درون دسته‌ها طبقه‌بندی نماید. الگوریتم مورد استفاده از خانواده الگوریتم‌های بیزین بوده و مبتنی بر TFIDF می‌باشد. در این روش برای بالا بردن دقت از درخت ارتباطات و جدول کلمات هم معنی استفاده شده است. تعداد بالا داده‌های آزمایشی، یکی از نقاط قوت موجود در این پروژه بوده است.

در مقاله [64] سیستم جدیدی برای دسته‌بندی خودکار متون فارسی ارائه شده است. این سیستم شامل دو مرحله اصلی است: مرحله پردازش و مرحله دسته‌بندی. در مرحله اول با پردازش داده‌های آموزشی بهترین ویژگی‌های نماینده هر کلاس استخراج شده و برای آموزش دسته‌بندی کننده مبتنی بر ماشین بردار پشتیبان استفاده می‌شوند و سپس در فاز دسته‌بندی، ماشین بردار پشتیبان قادر خواهد بود داده‌های تست را به یکی از کلاس‌های آموزش داده شده نسبت دهد. در روش ارائه شده در این مقاله برای افزایش دقت دسته‌بندی کننده از دانش معنایی موجود در گنجوازه بهره گرفته شده است. به این ترتیب که ویژگی‌های مربوط به هر کلاس می‌تواند با استفاده از گنجوازه گسترش یابد.

در مقاله [65] الگوریتمی جهت دسته‌بندی نتایج جستجو مبتنی بر درخواست کاربر به منظور افزایش کارایی موتورهای جستجوگر و کاهش تلاش کاربر در یافتن نتیجه مورد نظر در لیست نتایج ارائه شده

است. ابتدا صفحه مرورگر بارگزاری شده و پیشنهادهایی برای کاربر ارائه می‌شود، برخی پارامترها مانند تعداد موتور جستجوگر برای جستجو، اولویت زبانی، تعداد دسته ها و ... به او پیشنهاد می‌گردد. پس از بازبازی درخواست کاربر که توسط واسط کاربری مطرح شده است، متن درخواست به یکی از موتورهای جستجو ارجاع داده شده و پس از دریافت نتایج مورد نظر، اسناد به دست آمده از وب پردازش شده و نتایج دسته‌بندی شده نمایش داده می‌شوند. از ترکیب وزن‌دهی TFIDF و الگوریتم K-means برای امر دسته‌بندی استفاده شده است.

مقاله [66] رویکرد دسته‌بندی معنایی مبتنی بر سیستم‌های چند عامله در بازبازی اطلاعات در وب را برای بهبود روند جستجو پیش‌رو گرفته است. ایده اولیه این روش از آنجا نشأت گرفته است که بیشتر موتورهای جستجو سنتی تلاشی برای درک مفهوم درخواست ها نکرده و تنها به یافتن عینی (شمول/عدم شمول) کلمات و عبارات می‌پردازند. بدین منظور از پایگاه هستی‌شناسی برای افزودن معنا و محاسبات زبان‌شناسی به عبارات درخواستی کاربر استفاده شده است. در این برنامه از یک سیستم شامل چندین عامل اطلاعاتی هوشمند متعامل جهت حل مسایلی که موجودیت‌های منفرد و روشهای قدیمی از پس آنها بر نمی‌آیند بهره گرفته شده است. عامل اطلاعاتی، عاملی است برای بهبود سرویس‌های جستجو و امکان دسترسی به یک یا چندین منبع اطلاعاتی که قادر به تطبیق و دستکاری اطلاعات منابع برای پاسخگویی به درخواست کاربر می‌باشد. هر کدام از متغیرهای درخواست به یک عامل داده شده و هر کدام 100 نتیجه اول را از سایت‌های پایه موجود بازبازی می‌کنند. این نتایج بر اساس معیارهای مشخصی دوباره اولویت‌دهی می‌شوند. سپس مناسب‌ترین کلمات برای شرح بهتر مفهوم انتخاب می‌شوند. در ادامه ابعاد ماتریس حاصل به کمک روش SVD کاهش یافته و لیبل خوشه‌ها شناسایی می‌شود. سپس هر سند به لیبل مربوطه نسبت داده شده و خروجی به صورت گروه‌های مرتب‌سازی و اولویت‌بندی شده نمایش داده می‌شود. معیار ارزیابی به کار گرفته شده در این سیستم، میزان رضایت کاربر در معنادار بودن لیبل و تناسب لیبل با محتویات اسناد است.

برچسب‌زنی موضوعی که در مقاله [68] به آن پرداخته شده است، محتمل‌ترین موضوعی که محتوای متن بدان اشاره دارد را مشخص می‌کند. برای رسیدن به این هدف، از روش فاصله‌یابی در فضای بردار بسامدی بهره گرفته شده است. در این روش فاصله بین بردارهای بسامد وزن‌دار کلمات کلیدی محاسبه می‌شود. کلمات کلیدی شامل کلمات با بیشترین حاصلضرب بسامد پیکره‌ای و فاکتور وابستگی به طبقه هستند که مجموعه این کلمات از پیکره متنی زبان فارسی استخراج شده‌اند. در مرحله هرس، علاوه بر کلمات دستوری مانند حروف اضافه و ربط، کلمات کلیدی زاید نیز به صورت دستی بر اساس دیدگاه‌های زبان‌شناختی و معنی‌شناختی حذف می‌شوند. در گام بعد، برای هریک از طبقات مورد نظر یک بردار بسامد وزن‌دار کلمات کلیدی تشکیل می‌گردد. از بین این بردارها، برداری که کمترین فاصله را با بردار بسامد وزن‌دار متن آزمون دارد، طبقه مربوط به متن آزمون را مشخص می‌نماید. پایگاه داده مورد استفاده شامل 8486 کلمه در 10 دسته بوده و پایگاه کلمات زائد بیش از 60 هزار کلمه را در خود جای داده است.

از [69] به عنوان یکی دیگر از مقالاتی که به دسته‌بندی نتایج جستجو مبتنی بر درخواست کاربر به منزله سهولت کاربران در مشاهده سریع و بهینه نتایج موتورهای جستجو پرداخته‌اند، می‌توان نام برد. در این روش پس از اعلان درخواست از سمت کاربر و دریافت لیست ترتیبی اسناد از موتورهای جستجو، متد پیشنهادی به استخراج و اولویت‌دهی عبارات ممتاز به عنوان اسامی کاندید دسته‌ها بر مبنای مدل آموزشی رگرسیون خطی می‌پردازد. سپس هریک از اسناد به عبارت نزدیکتر اختصاص یافته و دسته‌های نهایی با ادغام کاندیدها ایجاد می‌گردد. الگوریتم پیشنهادی ابتدا کل عنوان و قطعه خبری بدست آمده از موتورهای جستجوگر پایه را تجزیه کرده، تمام عبارات مهم را در یک π -گرام ایجاد کرده و برای هر کدام چندین پارامتر مشخص را محاسبه می‌نماید. سپس به کمک یک مدل رگرسیون بر اساس داده‌های آموزشی قبلی و اعمال ترکیب این خصوصیات، امتیاز برجسته‌سازی عبارت را محاسبه می‌کند و تعدادی از عبارات با بالاترین امتیاز به عنوان نام و شرح دسته انتخاب می‌شوند. روش وزن‌دهی به عنوان و قطعه متفاوت است چرا که به احتمال زیاد عناوین شامل مباحث

مهم‌تری هستند، طول یا تعداد لغات تشکیل عبارت، فرکانس تکرار هر کلمه، آنتروپی دسته و استقلال عبارت جهت وزندهی و معیار تشابه cosine برای دسته‌بندی اسناد استفاده شده است.

در این بخش با تعدادی از پژوهش‌های کلی انجام شده در زمینه دسته‌بندی اسناد متنی و اسناد وب آشنا شدیم. همچنین نگاهی کلی به مقالات مرتبط با موتورهای دسته‌بند اسناد وب داشتیم. هرچند این نمونه‌ها تمام موارد پژوهشی را در بر نمی‌گیرند، اما نشان‌دهنده خط مشی و رویکردی کلی از پروژه پیش‌رو است.

فصل چهارم: روش تحقیق

1-4 مقدمه

امروزه گسترش چشمگیر تعداد وبسایت‌های خبری و حجم بالای اطلاعات موجود در آنها و تمایل برای کسب دقیق و به موقع اطلاعات پیرامون حوادث جاری از اهمیت بسیار بالایی برای مردم برخوردار است.

سرعت بالای گسترش اخبار در اینترنت، امکان پیگیری مداوم اخبار به صورت دستی را غیر ممکن نموده است. این حجم بالا نه تنها در بدست آوردن اطلاعات مورد نیاز به کاربران کمک نمی‌کند، بلکه باعث سردرگمی و ابهام بیشتر آن‌ها نیز می‌گردد، تا آنجا که نیاز به سیستمی برای مدیریت و در نهایت متمایز ساختن اسناد خبری از یکدیگر و دسته‌بندی آن‌ها در گروه‌های متشابه در این بستر به چشم می‌خورد. اجزای یک سیستم یادگیری نظارتی شامل عامل یادگیر، مجموعه داده‌های آموزشی و داده‌های آزمایشی (داده‌های برچسب‌دار و غیر برچسب‌دار) و تعداد مشخصی دسته می‌باشد. در کاربرد خوشه‌بندی متن، هر سند به عنوان یک داده آموزشی در نظر گرفته می‌شود. هر خوشه در این کاربرد، حاوی اسناد مربوط به یک موضوع خاص است. تعداد خوشه‌های آموزشی بر اساس تنوع موضوعی تعیین می‌گردد. این تعداد توسط کاربری که داده‌های آموزشی را فراهم می‌کند، مشخص می‌شود. برچسب داده‌های آموزشی، شماره و یا نام خوشه و یا خوشه‌هایی است که هر داده آموزشی به آن خوشه و یا خوشه‌ها تعلق دارد. عامل یادگیر در فاز آموزشی با استفاده از داده‌های آموزشی، آموزش داده شده و یک تعمیم از هر خوشه یاد می‌گیرد که عامل را قادر به شناسایی اسناد مربوط به هر خوشه می‌کند. در فاز آزمایش، عامل یادگیر، اسناد متنی را که در فاز آموزشی ندیده است و متعلق به یکی از خوشه‌های آموزشی می‌باشد، دریافت کرده و احتمال تعلق این سند را به هر کدام از خوشه‌های آموزشی به دست می‌آورد. در این‌جا خوشه‌بندی برای هر سند، به معنی انتخاب خوشه‌ای با بالاترین احتمال تعلق سند به آن خوشه می‌باشد.

در این بخش راهکار پیشنهادی برای دسته‌بندی و نمایش بهینه اخبار جهت سهولت زمانی و دسترسی کاربران بیان شده است. برای رفع این معضل، از تلفیق روش‌های وب‌کاوی و دسته‌بندی داده‌های

متنی برای مدیریت بهینه اسناد خبری بهره گرفته شده است. قسمتی از متن اسناد خبری در ابتدا به کمک یک خدمتگذار وب از سایت‌های خبری مورد تایید و مشخص برداشته شده و پس از آموزش سیستم، پیش‌پردازش‌های مورد نیاز بر روی اسناد قابل تست صورت گرفته، هر سند خبری به دسته‌های مربوطه و انتخابی کاربر ارجاع داده شده و نتیجه به صورت لیستی از اخبار دسته‌بندی شده نمایش داده می‌شود. در ادامه به معرفی و تشریح اجزای تشکیل دهنده معماری پیشنهادی می‌پردازیم.

در فصل‌های قبل با مراحل روند دسته‌بندی اسناد در وب و مقدمات فنی آن آشنا شدیم. در این فصل به صورت متمرکز کار بر روی اسناد خبری حقیقی را در پیش گرفته و از بیان مجدد کلیات جلوگیری می‌کنیم.

2-4- اخبار موجود در وب

به صورت کلی وظیفه انتشار اخبار معتبر و رسمی کشور در اینترنت به عهده سه دسته سازمان می‌باشد. دسته اول مربوط به نهادهای خبرگزاری کشور می‌باشد. تهیه و انتشار خبرهای رسمی کشور تنها در خبرگزاری و توسط خبرنگاران مربوط به هر واحد صورت می‌پذیرد. از جمله خبرگزاری‌های معتبر می‌توان به خبرگزاری رسمی جمهوری اسلامی ایران (ایرنا)، خبرگزاری دانشجویان ایران (ایسنا)، خبرگزاری مهر، ایلنا، ایتنا و ... اشاره نمود. نحوه دسته‌بندی اسناد خبری در این سازمان‌ها بدین گونه است که پس از ترسیم و تعریف نمودار خبری در هر سازمان (بسته به سیاست و خط مشی سازمان)، گروه‌های نهایی انتخاب شده و هر واحد، تعدادی خبرنگار مشخص و انحصاری را استخدام می‌کند. هر خبرنگار با توجه به گرایش کاری خود اخبار مربوطه را ردیابی نموده و سند خبری را تهیه می‌کند. این کار برای هر دو گروه اخبار داخلی (مراجعه مستقیم و تهیه گزارش) و اخبار خارجی (ترجمه اخبار معتبر دیگر کشورها یا تهیه گزارش حضوری) صورت گرفته و پس از تأیید مسئول واحد، کارشناس خبری آن را جهت انتشار در سایت قرار می‌دهد. گزارشات مربوطه به سایر استان‌ها نیز پس از تأیید فرمانداری و کارشناس خبری در گروه استانی انتشار می‌یابد. در حقیقت دسته‌بندی اسناد در

دنیای حقیقی در اولین فاز ممکن (تهیه خبر) صورت گرفته و هیچ گونه عمل جداسازی و تفکیک خبری و یا پردازش و گروه‌بندی اخبار صورت نمی‌پذیرد. قرار دادن اخبار بر روی سایت، در محدوده زمانی یا تعدادی مشخصی صورت نگرفته و هر کدام از خبرگزاری‌ها بسته به تعداد و مهارت پرسنل و بازه زمانی به صورت 24 ساعته به نشر اخبار جدید می‌پردازند. میانگین نرخ به روز رسانی اخبار سه خبرگزاری در جدول (4-1) آورده شده است.

دسته دوم که به عنوان پورتال‌های خبری کشور شناخته شده‌اند، مربوط به سازمان‌ها، نهادها و ارگان‌های دولتی و خصوصی بوده و تنها اخبار درون سازمانی را برای اطلاع پرسنل و یا بخش محدودی از افراد جامعه درون وب سایت خود قرار می‌دهند. از جمله پورتال‌های سازمانی می‌توان به سایت مربوط به دانشگاه‌ها و وزارتخانه‌ها اشاره کرد. عمده اخبار انتشاری بر روی آنها شامل عزل و نسب پرسنل، بخش‌نامه‌ها و جلسات برگزار شده و تصمیمات اتخاذ شده است. بدون شک اخبار بین سازمانی و یا خبرهایی که تاثیر مستقیم بر روند امور جامعه دارند، توسط روابط عمومی سازمان و یا با تشکیل کنفرانس خبری و دعوت خبرنگاران و در برخی موارد بنا بر پیگیری خبرنگاران به اطلاع عموم مردم رسانده می‌شود.

آخرین دسته مربوط به سایت‌های خبری - تحلیلی می‌گردد که رویکردی غیر دولتی داشته و هر کدام به جریان، گروه یا فردی خاص منتسب بوده و یا مستقل عمل می‌کنند. در این دسته بیشتر با تحلیل، مقایسه و بازپخش اخبار با تاکید بر گردش و نشر آزاد افکار و اندیشه‌های مردم و تلاش در نشر دیدگاه‌ها، نقدها و نظرهای آنان در چارچوب حفظ هویت دینی و ملی و حدود قانون و باورهای مردم مواجه هستیم. عملاً کار تهیه خبر بر عهده این سازمان نبوده و هیچ دخل و تصرفی در اصل اخبار ندارند و در اکثر موارد شاهد تقسیم‌بندی‌های شخصی‌سازی شده بر اساس نظر مدیر سایت هستیم و هر کدام بر روی موضوعات خاصی تمرکز دارند. یکی از اهداف این سایت‌ها که تعداد آنها نیز بسیار زیاد است جمع‌آوری مجموعه اخبار و سهولت دسترسی کاربران به مجموعه اخبار گزینشی است. از این دسته می‌توان به سایت‌های خبری تحلیلی کلمه، فردا، بامداد نیوز و ... اشاره کرد.

تمرکز و توجه پروژه کنونی بر روی اخبار انتشار یافته در خبرگزاری‌های رسمی بوده و در این میان سه خبرگزاری ایرنا، ایسنا و مهر انتخاب شده‌اند. در عین حال سیستم نهایی متعلق به هیچ کدام از گروه‌بندی‌های فوق نبوده و تنها به جمع‌آوری و گروه‌بندی اخبار دور از دسترس و نظارت کاربر انسانی می‌پردازد. به همین دلیل از عنوان موتور دسته بند اسناد خبری در بستر وب برای آن استفاده نموده-ایم.

16-15	15-14	14-13	13-12	12-11	11-10	10-9	9-8	8-7	7-6	خبرگزاری
13	17	15	14	19	26	16	23	31	13	ایرنا
30	31	27	23	25	31	26	34	23	17	ایسنا
12	5	6	4	8	17	9	11	18	1	مهر
6-1	1-24	24-23	23-22	22-21	21-20	20-19	19-18	18-17	17-16	خبرگزاری
8	6	6	15	13	18	30	11	10	10	ایرنا
3	7	0	1	8	8	17	12	26	26	ایسنا
2	1	1	0	3	4	8	2	6	12	مهر

جدول (1-4): نرخ به روز رسانی اخبار در سایت‌های پایه

3-4 پوشش صفحات و جستجو لینک‌های خبری

بدون شک وجود یک پایگاه با داده‌هایی متناسب، برای هر پروژه‌ای الزامی است. از آنجا که تمرکز اصلی کار بر روی نمونه‌های به‌روز موجود در سایت‌های خبری می‌باشد، لذا ایجاد پایگاهی از داده‌های مشابه نمونه‌های حقیقی برای آموزش سیستم یکی از نخستین اقدامات است، در نتیجه باید به جستجوی اسناد خبری در سایت‌های متعلق به خبرگزاری‌های معتبر پرداخته و سعی در جمع‌آوری داده‌های اولیه نماییم. به دلیل عدم پوشانندگی تمام اخبار توسط یک سایت خبر رسانی منفرد، برای دستیابی کامل و همه جانبه به حوادث نیاز به ادغام اخبار موجود در چندین سایت خبری وجود دارد. پیشتر بیان شد که به دلیل اهمیت صحت اسناد، نمی‌توان به هر صفحه‌ای در دنیای وب مراجعه نمود، در نتیجه برای این کار از میان خبرگزاری‌های رسمی و مورد تایید در کشور، چند خبرگزاری را انتخاب نموده و عملیات آموزش و تست را بر روی آنها انجام دادیم. خبرگزاری‌های مورد بحث عبارتند از: خبرگزاری دانشجویان ایران (ایسنا)، خبرگزاری جمهوری اسلامی ایران (ایرنا) و همچنین خبرگزاری

مهر، بدیهی است سیستم با توجه به نیاز و درخواست کاربر قادر به گسترش و انطباق‌پذیری می‌باشد. بعد از ورود به پورتال هر یک از خبرگزاری‌ها، به استخراج بخشی از اسناد خبری می‌پردازیم.

لازم به ذکر است که بنا بر روند کار، موتورهای جستجوگر می‌توانند میتنی بر درخواست کاربر و یا بر اساس پیش‌فرض‌های تعریف شده برای سیستم عمل پویش صفحات را انجام دهند. در همین راستا اولین اقدام ما پویش صفحات خبری خاص در بازه‌های زمانی تعریف شده به منزله دسترسی به لینک خبرها است، برای این کار از یک خدمت‌گزار وب استفاده کرده‌ایم. بدین صورت که پس از انتخاب و تایید وبسایت‌های پایه، به آدرس صفحات مربوطه رفته و کد HTML ایجاد کننده صفحه را خوانده و در آرایه‌ای ذخیره می‌کنیم. پس از انجام این کار به دنبال تگ مربوط به لینک‌های درونی صفحه رفته و آدرس آنها را به عنوان مشخصه کلیدی و انحصاری هر صفحه استخراج می‌کنیم. در ادامه با ارسال یک خدمت‌گزار وب دیگر اطلاعات صفحات لینک را خوانده و بخش مورد نظر از هر صفحه را در قالب یک فایل به فرمت xml ذخیره می‌کنیم.

از جمله مشکلات پیش‌رو برای افزودن خبرگزاری‌های جدید، عدم وجود استاندارد در فرمت و محتوای صفحات می‌باشد. از آنجا که هیچ گونه ساختار مشخص، قالب و همچنین زبان کدینگ واحدی برای سایت‌ها در نظر گرفته نشده و هر خبرگزاری صرفاً بر مبنای تمایل و علم متخصصان درون یا برون سازمانی به صورت مجزا، اقدام به طراحی و ارائه سایت نموده، نتیجتاً هنگام فراخوانی و ایجاد پایگاه نمونه‌های اولیه با تنوع بالایی از فرمت چه در سطح آرایش اخبار و چه در سطح دسته‌بندی‌های نهایی مواجه می‌شویم. نمودار خبری سه خبرگزاری مورد بحث -بدون توجه به زیر گروه‌های متعدد تعریف شده برای هر کدام در جدول (2-4) آورده شده است.

خبرگزاری	نمودار خبری
ایرنا	سیاسی، اقتصادی، اجتماعی، فرهنگی، علمی، ورزشی، خارجی، استانها، پژوهش‌های خبری
ایسنا	اندیشه امام و رهبری، علمی و فناوری، اجتماعی، دانشگاه و حوزه، اقتصادی، سیاسی، بین‌الملل، فرهنگی و هنری، معارف و حقوق، فرهنگ حماسه، ورزشی، عکس، استان‌ها
مهر	فرهنگ و هنر، فرهنگ و ادب، دین و اندیشه، حوزه و دانشگاه، فناوری‌های نوین، اجتماعی، اقتصادی، سیاسی، بین‌الملل، ورزشی، انرژی هسته‌ای، عکس، استان‌ها، دفاع مقدس

جدول (2-4): نمودار خبری مربوط به سه سازمان خبرگزاری پایه

1-3-4 مدیریت خطاهای احتمالی به کمک کدهای وضعیت HTTP

پویایی، تغییر و حذف اطلاعات در بستر نت به دلیل دسترسی همگانی به آن امری بدیهی بوده و روزانه تعداد لینک‌های تهی یا خراب بسیاری ایجاد می‌گردد. در صورت عدم وجود نظارت و مدیریت لینک‌ها احتمال ورود آدرس اشتباه به لیست لینک‌های قابل پردازش وجود دارد، که در نهایت به دلیل عدم توانایی دسترسی عامل خزنده، برنامه با خطا مواجه شده و عملکرد آن متوقف می‌گردد. از جمله خطاهایی که در مرحله تست به دفعات مکرر با آن مواجه گشتیم می‌توان به نمونه‌های زیر اشاره کرد. در صورت عدم وقوع این خطاها، کد وضعیت 200 (سالم) را به خود اختصاص داده و عملیات ارجاع به آن صورت می‌گیرد.

کد 400: درخواست بد⁷⁰

کد 400 به دلیل درک نشدن شیوه نگارش⁷¹ درخواست واسط کاربری از سرور رخ می‌دهد، در این حالت مفهوم تقاضای کاربر برای سرور روشن نیست و درخواست قابل پردازش نمی‌باشد، این خطا ممکن است به دلایل دیگر، از جمله نقص در انتقال داده‌ها (به فرض به دلیل قطع یا افت سرعت ارتباط) نیز رخ دهد.

⁷⁰ Bad Request

⁷¹ Syntax

کد 403: دسترسی غیر مجاز⁷²

کد 403 مربوط به مواقعی است که کاربر درخواست منبعی را از سرور دارد که دسترسی به آن برای همه کاربران محدود شده است، در اینجا حتی با ورود نام کاربری و کلمه عبور نیز امکان دسترسی مقدور نخواهد بود، معمولاً مدیران سایت‌ها، دسترسی مستقیم به فولدرها و نمایش فایل‌ها به صورت لیست را غیر فعال می‌کنند، در نتیجه وقتی آدرس یک فولدر را از آن سرور درخواست می‌کنیم، با خطای 403 مواجه خواهیم شد.

کد 404: منبع درخواستی پیدا نشد⁷³

کد 404 در مواقعی رخ می‌دهد که واسط کاربری تقاضای منبعی (به طور مثال یک فایل یا صفحه) را از سرور دارد که در حال حاضر موجود نبوده یا حذف شده است (و یا ممکن است نام آن تغییر کرده باشد)، البته احتمال دارد در آینده مجدداً آن منبع ایجاد شده و در دسترس قرار گیرد.

کد 410: محذوف⁷⁴

کد 410 به معنی حذف همیشگی منبع درخواستی از سرور است، بر خلاف خطای 404، کد 410 به واسط کاربری یا موتورهای جستجو می‌گوید که نباید مجدداً آن منبع را درخواست کنند، چرا که برای همیشه حذف شده است، البته در عمل موارد استفاده از این کد خیلی محدود است و تنظیم خطای 404 بهتر و اصولی‌تر است.

کد 500: خطای داخلی سرور⁷⁵

کد 500 به معنی وقوع یک خطای داخلی در سرور است و معمولاً به دلیل نقص تنظیمات یا انجام به روزرسانی نرم افزاری یا سخت افزاری رخ می‌دهد، تنظیم این کد در مواقعی که می‌خواهیم در سایت، تغییراتی اعمال کنیم که باعث از دسترس خارج شدن آن می‌شود، می‌تواند مفید باشد.

⁷² Forbidden

⁷³ Not Found

⁷⁴ Gone

⁷⁵ Internal Server Error

2-3-4 انتخاب قطعه خبری

همانگونه که بیان شد، داده‌های مورد پردازش شامل اسناد خبری انتشار یافته در اینترنت می‌باشد. یکی از مهم‌ترین تفاوت‌های موجود میان اسناد متنی اداری و اسناد متنی در وب، اندازه اسناد است، بدین‌منظور که در سیستم‌های مبتنی بر وب به دلیل اهمیت حجم و زمان پردازش داده تمرکز اصلی بر روی ریز اسناد بوده و پردازش مستقیم بر روی کل متن صورت نمی‌گیرد. داده‌های نهایی ما ریز اسنادی شامل: آدرس دسترسی مستقیم به خبر، شرح متن و زمان استخراج خبر می‌باشد، که در ادامه با خصیصه‌های هر کدام آشنا می‌شویم.

آدرس دسترسی مستقیم: همانطور که قبلاً بیان شد، در هر صفحه تعداد زیادی لینک مشاهده می‌گردد، اما تنها لینک‌های درون صفحه‌ای، نشان‌دهنده ارجاع به اخبار انتشاری داشته و دیگر لینک‌ها که آنها را برون صفحه‌ای می‌نامیم، کار برقراری ارتباط میان خبرگزاری با دیگر سازمان‌ها و نهادها را بر عهده دارند. آدرس دسترسی مستقیم، به مجموع آدرس وب سایت پایه (خبرگزاری) و آدرس لینک درون صفحه‌ای که خزنده از آن استفاده نموده است گفته می‌شود. از آنجا که هر لینک تنها به یک سند خبری اشاره دارد، لذا از آن به عنوان کلمه مشخصه هر سند استفاده نموده و برای جلوگیری از ورود اسناد تکراری به پایگاه داده، پس از مشاهده متن هر کدام از اسناد، لینک دسترسی مستقیم آنرا ذخیره می‌نماییم. در انتها نیز از همین آدرس جهت رساندن کاربر به خبر مورد نظر استفاده خواهیم نمود.

شرح متن: هدف اصلی این بخش استخراج و گزینش بخشی از متن خبر برای انجام پردازش می‌باشد. جهت استخراج شرح ریزاسناد روش‌های بسیاری وجود دارد. در بسیاری از پروژه‌ها مشاهده می‌گردد که داده آموزشی و گاهی داده تست توسط کاربر و به کمک پیش‌فرض‌هایی جهت سهولت بخشیدن به کار پردازش آماده می‌گردد. هر چند این عمل میزان دقت و نتیجه نهایی را تا حدی افزایش خواهد داد، اما بدون شک برای کار با داده‌ها و فضای حقیقی مشکلات بسیاری به وجود خواهد آورد. در صورتی که بخواهیم کار استخراج را به صورت دستی انجام دهیم، بدون شک مجموع متون موجود در

عنوان، توضیحات و چند جمله اول که طبیعتاً چکیده کل مطلب را در بردارد، بهترین گزینه جهت فهم کلی مطلب بوده و به اندازه کافی قادر به تفکیک‌سازی عناصر مهم متون از یکدیگر می‌باشند. اما از آنجا که نمونه‌های نهایی به صورت خودکار از صفحات استخراج می‌گردد، جهت بالا بردن تطابق داده‌های تست و آموزش، در فاز آموزش نیز داده‌ها به صورت اتوماتیک توسط عامل خزنده تهیه شده‌اند. یکی از روش‌های اولیه در این زمینه استفاده از نشانه‌های موجود در کد HTML اسناد می‌باشد. بدین صورت که پس از مراجعه به لینک‌های درون صفحه‌ای، متون موجود در نشانه‌های "عنوان"⁷⁶، "توضیحات"⁷⁷ و "کلمات کلیدی"⁷⁸ را که نویسنده و یا اپراتور هنگام درج مطلب تکمیل نموده را به عنوان ریزسند انتخاب نماییم. با نگاهی به صفحه منبع سایت‌های موجود و نشانه‌گذاری‌های متفاوت (به دلیل عدم وجود استاندارد در فرمت داده‌ها) این روش کارآمد نبوده و در بسیاری موارد این نشانه‌ها خالی می‌باشند. جدول (3-4) نشان‌دهنده تفاوت‌های نوشتاری در صفحه خبرگزاری‌های مورد بحث است. روش مورد استفاده در این پروژه، استخراج سه پاراگراف ابتدایی موجود در هر متن به دلیل پوشانندگی نسبی و قابل قبول محتوای خبری در تمام صفحات است.

زمان: آخرین پارامتر مهم در جمع‌آوری، انتشار و باز پخش اسناد خبری، زمان است. در سیستم طراحی شده نتایج مربوط به دسته‌بندی در آخرین زمان بازبینی اسناد نمایش داده می‌شود. از آنجا که در مورد پارامتر زمان نیز هیچ گونه فرمت مشخصی در سایت‌ها وجود ندارد، لذا زمان خبر را برابر با زمان جمع‌آوری داده از صفحات قرار داده‌ایم. زمان بازبینی صفحات توسط کاربر و یا با زمان پیش فرض 30 دقیقه تنظیم شده است.

⁷⁶ Title

⁷⁷ Description

⁷⁸ Keywords

خبرگزاری	ایسنا	ایرنا	مهر	نهایی
<URL	inner link	inner link	inner link	Base website+ inner link
<title	*	* + <h1	وب سایت	-
<p1	Title	<h2	P3	*
<p2	P1	P3	P4	*
<p3	P2	P4	P5	*
<time, date	<div class="newsPubDate">	تاریخ انتشار	-	System time
Keywords	-	*	-	-
Description	* 1 sen.	*	-	-
Start point	<p1	<h1	<Span	<p

جدول (3-4): چکیده ساختار صفحه منبع کد سایت‌های پایه و سیستم نهایی

4-4 فاز آموزش

روش دسته‌بندی اسناد در سیستم طراحی شده مبتنی بر نظارت بوده و بر آموزش تکیه دارد. به این معنی که در این روش به یک مجموعه اولیه از اسناد که قبلاً تحت دسته‌هایی مشخص دسته‌بندی شده‌اند نیاز میباشد. در انتها نیز ارزیابی به کمک مجموعه آزمایشی صورت می‌پذیرد.

4-4-1 جمع‌آوری اسناد آموزشی

هدف این فاز جمع‌آوری قطعات متنی است که توسط موتورهای جستجو در جواب نیاز برنامه بازگردانده شده‌اند. تمام اطلاعات مورد نیاز از صفحات وب استخراج شده و تگ‌های اضافه و HTML آن جدا می‌شوند. در نتیجه تنها عنوان، محتوای قطعه، منبع اولیه و تاریخ ثبت نگهداری می‌شود. در این قسمت صفحه‌ای مانند صفحه نتیجه جستجو گوگل با لیستی ترتیبی از اطلاعات مشابه خواهیم داشت. این اطلاعات می‌توانند بر حسب درخواست کاربر و یا بر حسب پیش فرض سیستم بازگردانده شوند که در مورد اول باید به پیش‌پردازش درخواست کاربر و اعمال قوانین جبری منطقی احتمالی تعبیه شده در صفحه رابط کاربر پرداخته شود. اما از آنجا که سیستم مورد بحث ما مبتنی بر پیش-

فرض سیستم بوده و درخواست کاربر نقشی در نوع یا تعیین کلید واژه مشخص اسناد استخراجی ندارد، در نتیجه کار را با دنبال گیری تمام لینک های خبری موجود در خبرگزاری ها ادامه می دهیم. برای انجام نمونه گیری و آماده سازی داده ها از آرشیو خبری روزانه (بازه زمانی، تعداد اخبار) در دسته های مختلف استفاده شده است. بدین صورت که هر کدام از اسناد خبری به صورت مجزا در یک فایل ذخیره شده و با علم بر طبقه مربوط به آن در پایگاه داده جمع آوری شده است.

به صورت عمده می توان چنین بیان داشت که استفاده از تعدادی سند رسمی برای آموزش این چنین سیستم هایی کافی به نظر می رسد، اما از آنجایی که اسناد مورد بررسی در زمینه اطلاعات و سازمان های خبری از هیچ گونه نظم و روند مشخصی در زمان و نحوه اعلام بر خوردار نیستند، در نتیجه می توانیم از دو متد متفاوت استفاده نماییم. در مورد اول می توانیم با استفاده از آرشیو خبری مربوط به بازه زمانی مشخص پروسه آموزش سیستم را انجام دهیم. اما در برخی موارد احتمال آن می رود که در طی بازه مورد بحث تعداد بسیار اندکی سند در رابطه با شاخه ای خاص وجود داشته باشد، که این امر نه تنها به روند آموزشی سیستم ضربه وارد می کند، بلکه میزان اعتماد و ضریب کارکرد سیستم را نیز پایین خواهد آورد. در روش دوم آموزش را با حداقل تعداد مشخص خبر از هر زیرشاخه بدون در نظر گرفتن محدودیت زمانی انجام می دهیم. به کمک تلفیق دو روش فوق می توان به جامعیت مورد نظر دست یافت. در این پروژه نیز از تعداد حداقل 30 سند در هر شاخه استفاده شده و اخبار در چند بازه زمانی جمع آوری شده اند. تعداد 808 سند نهایی به کار برده شده به تفکیک به ازای هر گروه در جدول (4-4) آورده شده است.

گروه	علمی	استان ها	بین الملل	اجتماعی	اقتصادی	دانشگاهی	مذهبی	فرهنگی	ورزشی	سیاسی
تعداد	78	95	78	87	82	49	57	111	84	87

جدول (4-4): تعداد اسناد آموزشی به تفکیک گروه ها

2-4-4 تطبیق و تعیین گروه‌ها

مشکل عدم وجود استاندارد در سایت‌های خبری در این مرحله نیز برای ما مشکل آفرین است؛ چراکه هیچ الگویی برای دسته‌های خبری وجود نداشته و هر خبرگزاری بسته به نظر شخصی مسئولان و سیاست‌های خاص سازمانی گزینش دسته‌های نهایی را انجام داده است. در این برنامه نیز مجموعه‌ای از برآیند دسته‌های خبری انتخاب شده است و در قسمت آموزش اخبار مربوط به دسته‌های متفرقه، به صورت گزینشی به دسته‌های نهایی نسبت داده شده‌اند. این مشکل در سطح بین‌المللی نیز در صنعت خبری مشاهده می‌گردد. جدول (4-5) نشان‌دهنده نمودار دسته‌های خبری مربوط به خبرگزاری‌های مورد بحث و دسته‌های نهایی انتخابی می‌باشد. پس از تکمیل فرآیند نمونه‌گیری که در بخش قبل آورده شد، اخبار به گروه‌های نهایی انتساب داده شده و جهت آموزش به سیستم، لیبل گذاری می‌شوند.

خبرگزاری	نهایی	ایسنا	ایرنا	مهر
1	علمی، فناوری	علمی و فناوری	علمی	فناوری های نوین
2	استان ها	استان ها	استان ها	استان ها
3	بین الملل	بین المللی	خارجی	بین الملل
4	اجتماعی	اجتماعی	اجتماعی	اجتماعی
5	اقتصادی	اقتصادی	اقتصادی	اقتصادی
6	حوزه، دانشگاه	دانشگاه و حوزه		حوزه و دانشگاه
7	دینی، مذهبی	معارف و حقوق		دین و اندیشه
8	فرهنگی	فرهنگی و هنری، فرهنگ حماسه، اندیشه امام و رهبری	فرهنگی	فرهنگ و هنر، فرهنگ و ادب، دفاع مقدس
9	ورزشی	ورزشی	ورزشی	ورزشی
10	سیاسی	سیاسی، بیداری اسلامی	سیاسی	سیاسی، انرژی هسته ای
			پژوهش های خبری	

جدول (4-5): نمودار گروه بندی سازمان‌های خبرگزاری و نمودار خبری نهایی

3-4-4 آماده‌سازی اسناد آموزشی

تا این قسمت عملیات استخراج داده‌های مورد نیاز از وبسایت‌های پایه تعریف شده در برنامه صورت گرفته است. در ادامه به کمک برخی روش‌های مربوط به متن کاوی، اسناد را جهت پردازش نهایی آماده ساخته و عملیات پیش‌پردازش را بر روی آنها اعمال می‌کنیم. داده‌های آموزشی، مجموعه اطلاعاتی هستند که دانشی را به صورت پیش‌فرض به عامل یادگیر منتقل می‌کنند. در واقع تفاوت عمده دو روش خوشه‌بندی با یادگیری با نظارت یا بدون نظارت از تفاوت در نوع آموزش آنها نشأت می‌گیرد. هدف نهایی این مرحله یکسان‌سازی اسناد و حذف داده‌های ناهمگون جهت سهولت در امر پردازش است.

نرمال سازی

در دوران گذشته پیش از اختراع دستگاه‌های تحریر، هر انسان بسته به ذوق و سلیقه خود سعی در نوشتار مرتب و منظم متون با دست‌خط خویش می‌نمود. این تفاوت در شیوه‌های ظاهری نوشتاری، میزان خوانایی و درک متون توسط دیگر انسان‌ها را تا حد زیادی کاهش می‌داد. وجود تکنولوژی‌های امروزه، مشکل ناخوانایی متون را کاهش داده و نوشتار را مستقل از نویسنده نموده است. چراکه همگی کاربران با استفاده از استانداردها و کدهای تعریف شده متون را تایپ نموده و خروجی یکسانی ایجاد می‌نمایند. مشکل جدید در این زمینه وجود تعداد و تنوع بالای فونت‌ها است که به آنها قلم نگارش کامپیوتری نیز گفته می‌شود. علاوه بر شکل ظاهری متفاوت، وجود فونت‌های عربی فارسی با کدهای اختصاصی برای هر کدام و عدم پوشش برخی کاراکترها در زبان عربی، تفاوت‌های پنهان بسیاری را در مقایسه کلمات با ظاهر مشابه ایجاد نموده است. استفاده از شیوه‌های نگارشی متنوع و عدم امکان استانداردسازی نوشتاری را نیز باید به مشکلات بالا افزود. به عنوان مثال استفاده از انواع شکل همزه در میان و پایان کلمات را می‌توان نام برد.

تشخیص این تفاوت‌ها به دلیل شباهت ظاهری، برای کاربران انسانی دشوار بوده و از طرف دیگر کامپیوتر به دلیل استفاده کدهای مختلف آنها را جدا از هم و به صورت کلمات مجزا در نظر می‌گیرد.

به همین دلیل در ابتدای کار قبل از انجام هرگونه پردازشی، کد فونت‌ها را برای حروف مشابه یکسان نموده و همزه انتهایی کلمات را نیز حذف نموده‌ایم. این مرحله باید برای تمام داده‌های آموزشی، تست، کلمات اضافی، فهرست‌های مکمل و به صورت کلی برای هر متن ورودی انجام پذیرد. با این کار نه تنها افزونگی داده از بین می‌رود، بلکه قادر به ایجاد پایگاه لغات غیر تکراری نیز خواهیم شد. در جدول (4-6) کد اختصاصی به حروف مشابه و نمونه‌هایی از هر مورد آورده شده است. از آنجا که اسناد قابل پردازش، متون رسمی خبرگزاری‌ها می‌باشند، در نتیجه احتمال وجود غلط‌های املایی در نظر گرفته نشده است.

حرف	کدهای متفاوت حرف	مثال	کد نهایی
و	1572 و 1608	مودب... مؤدب	1608
ه	1577 و 1607 و 1726 و 1729	تذکره... تذکره... تذکره	1607
ی	1574 و 1609 و 1610 و 1740	مسایل... مسائل و کتبی... کتبی	1610
ک	1603 و 1705	سرکش دار عربی	1603
ء	1569	امضا... امضاء	32

جدول (4-6): یکسان‌سازی حروف مشابه

در این مرحله همچنین کاراکترهای نگارشی به کاربرده شده در متن و فضاهای خالی میان کلمات حذف شده و با تعریف حد آستانه طول کلمه، کلمات خالص متن بدست می‌آید. در بسیاری از موارد به دلیل عدم رعایت نشانه گذاری یا فاصله، دو کلمه به هم متصل شده و عملاً به عنوان داده خراب وارد متن می‌شود. از جهت دیگر طول بیشتر وندهای اضافه شده به کلمات (فعل، اسم) کمتر از حد مشخصی است. هرچند این وندها در قسمت ریشه‌یابی پس از تطبیق با جدول وندها حذف می‌گردند، اما لازمه این کار مقایسه تک تک کلمات با جدول بوده و بدون شک دارای بار محاسباتی و زمانی برای برنامه خواهد بود. لذا در این مرحله از پیش‌پردازش، تمامی لغاتی که طول کمتر از سه کاراکتر و بیش از هشت کاراکتر دارند را از مجموعه اصلی حذف می‌کنیم. با این عمل طول سند قابل پردازش به حد چشمگیری کاهش یافته و از ورود لغات با ارزش معنایی پایین جلوگیری به عمل خواهد آمد. در مورد

با انجام دادن این کار و حذف وندها از کلمات، کار حذف کلمات عمومی و افعال نیز به سادگی بیشتری صورت خواهد پذیرفت. باید توجه داشت که در موارد بسیاری با حذف وندها، به ساختار اصلی کلمه آسیب وارد شده و کلمه شکل ناقصی به خود می‌گیرد. اما از آنجایی که اولاً تعداد این کلمات خیلی زیاد نیست و ثانیاً عمل ریشه‌یابی بر روی تمام داده‌های ورودی به شیوه یکسان انجام می‌پذیرد، در نتیجه تأثیر منفی آن قابل اغماض و چشم‌پوشی است.

حذف کلمات عمومی و افعال

انواع کلمه در زبان فارسی عبارتند از: فعل، اسم، صفت، ضمیر، قید، حرف و شبه جمله، که در این میان قیود، ضمائر و حروف تنها برای شکل‌دهی و بالا بردن فهم جملات کاربرد داشته و به صورت مستقل حامل پیام خاصی نیستند. لذا تعداد تکرار این کلمات در اسناد بسیار بالا بوده اما نقش خاصی در تمایز و طبقه‌بندی اسناد ندارند. در بخش قبل با انجام عمل ریشه‌یابی تمامی کلمات را به شکل واحد بن آنها درآورده و در این قسمت به تشخیص و حذف کلمات عمومی (مستقل از کاربرد) می‌پردازیم. در اکثر زبان‌ها لیستی از کلمات عمومی پر تکرار وجود دارد، که در مرحله پیش پردازش آنها را حذف می‌نماییم. عملیات حذف در دو مرحله قبل و بعد از عمل ریشه‌یابی صورت گرفته و نقش بسزایی در کاهش تعداد کلمات قابل پردازش دارد. در این قسمت تمام ضمائر، انواع حروف ربط و اضافه، بن افعال و برخی کلمات پر تکرار عمومی حذف می‌گردند.

تقسیم بندی

هدف از این کار تشخیص محدوده لغات قابل پردازش در یک متن است. استفاده از علائم نشانه‌گذاری در این قسمت کاربرد فراوانی دارد. اما از آنجا که در قسمت قبل تمام کاراکترهای نشانه‌گذار و جداساز کلمات تعیین و حذف شده‌اند، در نتیجه از کد مربوط به کاراکتر "فضای خالی" جهت تعیین محدوده کلمات استفاده نموده‌ایم. در صورتی که متن اصلی استخراجی به صورت خام وارد برنامه میشد، می‌توانستیم برای جداسازی ریز سند نیز از کدهای موجود در `html` استفاده نماییم. اما در این پروژه جداسازی متن ریز سند در اولین مرحله صورت گرفته است.

آماده‌سازی اسناد آموزشی، علاوه بر یکسان‌سازی نوشتاری لغات، باعث کاهش 32.6 درصدی، معادل 33132 کلمه پرارزش قابل استفاده در بخش بعد شده است. این کاهش تعداد لغات، تسریع زمان و قدرت پردازش را به همراه داشته و بالطبع کارایی سیستم را بهبود داده است.

4-4-4 ایجاد پایگاه لغت

پس از تکمیل فرآیند پیش پردازش، متن اسناد به لغات کلیدی تشکیل دهنده آنها تقسیم‌بندی شده و جهت آموزش سیستم تمام لغات وارد یک پایگاه متمرکز می‌شوند. بر خلاف بسیاری از پروژه‌های داده‌کاوی امکان قرار دادن حد آستانه برای لغات با مرتبه تکرار بسیار بالا و یا بسیار پایین وجود ندارد. مهم‌ترین دلیل این امر پویایی اسناد خبری است. به عنوان مثال امر انتخابات ریاست جمهوری هر چهار سال یکبار در ایران رخ می‌دهد. در صورتی که لغات مربوط به اولین سند منتشره در مورد انتخابات، به دلیل تعداد پایین دفعات تکرار در محدوده لغات قابل پذیرش قرار نگیرند، بدون شک هیچ‌گاه این خبر وارد پایگاه داده نخواهد شد، لذا در این سیستم هر کلمه به صورت منفرد ارزش داشته و از ورود هیچ یک از کلمات به پایگاه داده کلی نمی‌توان صرف‌نظر کرد.

4-4-4-1 ادغام پایگاه‌های لغت کمکی

در قسمت قبل بیان شد که هدف کلی از ساخت دیکشنری، ایجاد یک پایگاه داده منظم و دسته‌بندی‌شده از اسناد آموزشی جهت تطبیق، مقایسه و تعیین گروه نهایی اسناد تست می‌باشد. این امر به صورت دریافت اسناد مجزا از گروه‌بندی‌های مختلف در سایت‌های پایه بوده و پس از پاکسازی، مجموعه لغات هر گروه در آرایه‌ای متناظر گردآوری شده‌اند. به کمک این روش به جای مقایسه هر سند ورودی با کل اسناد، هر سند را با مجموعه لغات نهایی تطابق می‌دهیم. از آنجایی که تعداد اسناد آموزشی محدود می‌باشد (از نظر تعداد و زمان)، در نتیجه امکان پوشش کلی دایره لغات مربوط به هر دسته با تکیه بر مجموعه آموزشی امکان‌پذیر نمی‌باشد.

راهکاری که برای رفع این مشکل در نظر گرفته شده است، استفاده از فهرست لغات مکمل می‌باشد. در این فهرست سعی بر آن شده است که تا حد توانایی لغات کلیدی متمایز کننده مربوط به هر دسته

شناسایی شده و به صورت یک پایگاه کمکی در اختیار سیستم قرار داده شود. به عنوان مثال در دسته استان‌ها، نام تمامی شهرهای ایران، در بخش بین‌الملل، نام کشورهای مختلف، در قسمت حوزه و دانشگاه، لیستی از نام تمامی دانشگاه‌ها و رشته‌های تحصیلی و در بخش هنری سبک‌های مختلف موسیقی، شعر و نام نویسندگان و شعرا و آلات موسیقی آورده شده است. در دیگر بخش‌ها نیز تلاش شده است تا تخصصی‌ترین کلمات موجود اضافه گردند. بدون شک این فهرست تمامی لغات را در بر نداشته و هیچ گونه ادعایی بر این موضوع نخواهیم داشت. یکی از مهم‌ترین این دلایل پویایی زبان و حجم بالای لغات جدید و اصطلاحات و ترجمه‌های گوناگون از کلمات دیگر زبان‌ها می‌باشد.

در تهیه فهرست تکمیلی باید توجه داشت که لغات تا بالاترین حد ممکن ویژگی تمایزی گروه‌ها را پوشش دهند. به عنوان مثال در نام‌گذاری دانشگاه‌ها یکی از مرسوم‌ترین روش‌ها اختصاص نام شهر یا استان در ادامه نوع دانشگاه می‌باشد. حال آنکه از این لغات در فهرست استان‌ها به عنوان نمونه‌ای انحصاری نام برده شده است. در نتیجه در معرفی دانشگاه‌ها اسامی خاص دیگر گروه‌ها (مانند نام شهر، استان و کلماتی چون امام، شهید و ...) حذف شده است. برای ایجاد این فهرست‌ها سعی در استفاده از آخرین نسخه‌های ویرایش شده در دانشنامه آزاد ویکی پدیا نموده‌ایم.

2-4-4-4 تعیین و حذف کلمات پرکاربرد

هر چند در ابتدا تعدادی از لغات به عنوان کلمات عمومی و پرکاربرد از مجموعه اصلی هر سند حذف شدند، اما به عنوان آخرین اقدام در فاز آموزش سیستم، کار حذف کلمات پرکاربرد مربوط به حوزه خبرگزاری را انجام می‌دهیم. بدین منظور با تحلیل لیست لغات پایگاه‌های گروهی، کلماتی را که در تمام گروه‌ها به دفعات مکرر استفاده شده‌اند را، فارغ از توجه به معنی و ارزش کلمه، از پایگاه کلمات حذف نموده و به لیست کلمات عمومی، جهت حذف شدن آنها در فاز تست، اضافه می‌کنیم. از جمله می‌توان به کلماتی همچون خبر، خبرگزاری، بلامانع، ذکر و نام خبرگزاری‌ها اشاره کرد. هرچند این لغات بار معنایی منحصر به فرد دارند، اما تکرار بیش از اندازه آنها، اهمیت متمایز کنندگی‌شان را در

این پروژه از بین برده است. مجموعه این لغات به انتهای لیست کلمات عمومی اضافه شده و در فاز تست از این لیست نهایی به جای لیست اولیه استفاده شده است.

3-4-4-4 تعلق کلمات به گروه‌ها

در این بخش علاوه بر ایجاد پایگاه داده اصلی که پیشتر به آن پرداختیم، به تعداد گروه‌های نهایی نیز پایگاه منفرد ایجاد می‌گردد. اما تفاوت عمده پایگاه لغت گروهی با پایگاه اصلی، کلمات تشکیل دهنده آنها است. همانطور که بیان شد، در پایگاه اصلی هر کلمه به محض مشاهده برای اولین بار به لیست اضافه شده است، این عمل تنها وجود یا عدم وجود کلمه را در بردارد، اما در مورد گروه‌ها، تمام لغات موجود در پایگاه اصلی به انضمام تعداد تکرار هر کدام در دسته مربوطه آورده شده است. بدون شک، عدم وجود یک کلمه در گروه، مقدار عددی صفر را در آرایه متناظر در بر خواهد داشت.

در نتیجه این قسمت، ماتریسی با ابعاد [تعداد دسته، تعداد کل لغات غیر تکراری] ایجاد می‌گردد. در ادامه از این ماتریس جهت میزان تعلق لغات برای گروه‌بندی استفاده خواهیم نمود. در اکثر پروژه‌های دیگر از ماتریسی مشابه برای اسناد استفاده می‌شود. به منزله یافتن میزان تعلق کلمه به گروه، از روش احتمال پیشین⁷⁹ فرمول دسته‌بندی بیزین در که فصل دوم توضیح داده شد، با اندکی تغییرات جهت بهینه‌سازی جواب استفاده نموده‌ایم.

برای این منظور با توجه به گسسته بودن مقادیر برای محاسبه $P(X_k|C_i)$ از فرمول (1-4) استفاده می‌کنیم که در آن احتمال تعلق لغت k ام (X_k) به دسته i ام (C_i) بدست می‌آید. تعداد کل دسته‌ها در پروژه ده دسته تعریف شده است. در نتیجه این عمل پارامتری به نام (NGF) به ازاء هر لغت در هر گروه بدست می‌آید.

$$P(X_k|C_i) = \frac{C_i \text{ فرکانس تکرار هر کلمه برای موضوع}}{\text{فرکانس تکرار هر کلمه برای کل موضوعات}} \quad (1-4)$$

⁷⁹ Prior Probability

5-4 فاز تست

عملیات دسته‌بندی اسناد خبری به صورت دستی، بدین شکل انجام می‌گیرد که هر خبرنگار، بسته به نوع گرایش خبری که دارد، هر یک از گزارشات خبری مربوط به حوزه خود را تهیه و تدوین نموده و برای انتشار در اختیار مسئول واحد مربوطه قرار می‌دهد. از آنجا که این افراد قبلاً توسط متخصصان آموزش داده شده‌اند، اشراف کلی بر روی موضوعات داشته و بالطبع هیچ گونه مشکلی برای دسته‌بندی نهایی وجود نخواهد داشت.

اما در ادامه باید هر کدام از اسناد خبری آزمایشی در حوزه‌های تعریف شده پوشش داده شوند. در نتیجه پس از اتمام فاز آموزش، هر سند به یک یا بیش از یک گروه خبری نسبت داده می‌شود.

5-4-1 تهیه قطعات

در این قسمت به تهیه اسناد خبری برای تست سیستم خواهیم پرداخت. یکی از مسایل مهم در روند اجرای برنامه، تلاش برای جلوگیری از ورود و ذخیره اطلاعات مشابه و تکراری است، چرا که سیستم بر اساس آموزه‌های اسناد تکراری را به درستی دسته‌بندی نموده و این کار میزان اعتمادپذیری و کارایی آن را در مورد داده‌های حقیقی کاهش خواهد داد. به همین منظور آدرس تمام صفحات پیمایش شده، ذخیره شده و در صورت عدم تشابه، صفحه مربوط به لینک خبری جدید به عنوان نمونه تست پذیرفته می‌شود. داده‌هایی که به منظور آزمایش سیستم پیشنهادی مورد استفاده قرار گرفته‌اند، برگرفته از مجموعه اخبار انتشار یافته در سایت‌های پایه با در نظر گرفتن حداقل 10 سند خبری در هر گروه، در چندین بازه زمانی متفاوت است که توسط سیستم پویسگر از سایت‌های پایه بدست آمده، تعداد اسناد آموزشی نهایی مورد استفاده در هر گروه جهت ارزیابی سیستم در جدول (7-4) آورده شده است.

گروه	علمی	استانی	خارجی	اجتماعی	اقتصادی	دانشگاهی	مذهبی	فرهنگی	ورزشی	سیاسی	رسانه‌های دیگر	پژوهش خبری
تست	22	30	57	27	27	18	14	49	51	56	20	1

جدول (7-4): تعداد اسناد آزمایشی به تفکیک گروه‌ها

4-5-2 آماده‌سازی اسناد آزمایشی

در این مرحله تمام داده‌های تست جهت یکسان شدن با اسناد آموزشی و سهولت در امر پردازش، مانند قسمت قبلی آماده شده و عملیات نرمال سازی، ریشه یابی، تعیین و حذف کلمات و افعال عمومی و تقسیم بندی متن به کلمات قابل پردازش با تمام قوانین و شرایط ذکر شده پیشین بر روی آنها اعمال می‌گردد. در نهایت به دلیل آماده‌سازی مشابه گروه آموزشی و آزمایشی، به داده‌هایی با فرمت و ظاهر مشابه دست خواهیم یافت. از آنجا که در نهایت متن اصلی خبر به انضمام دسته نهایی برای پذیرش و اعتبارسنجی به کاربر نمایش داده خواهد شد، لذا باید متن اولیه بدون اعمال دیگر پیش‌پردازش‌ها ذخیره گردد.

4-5-3 وزندهی به کلمات

پس از اتمام مرحله پیش پردازش و آماده سازی داده‌ها، به تعداد حداقل کلمات با مفهوم درون یک متن دست یافتیم که می‌تواند بیانگر خلاصه متن بوده و به ما در درک کلی مفهوم متن یاری رساند. عمل وزندهی به کلمات به منزله یافتن میزان اهمیت و نقش کلمه در عبارات و اسناد صورت می‌پذیرد. بدون شک کلمه با تعداد تکرار زیاد نسبت به طول متن، نشانگر تعلق متن به دسته خاص خواهد بود. این تعلق با مقایسه وزن کلمات سند با فرکانس تکرار کلمات مشابه، در پایگاه داده متعلق به گروه‌های نهایی که در قسمت آموزش ایجاد شده بودند، مشخص می‌گردد. برای این کار ابتدا لیستی شامل تمام لغات تشکیل دهنده متن و فرکانس تکرار آنها در سند اولیه بدست می‌آید. برای نرمال نمودن مقادیر بدست آمده، بعد از شمارش تعداد تکرار به ازای هر کلمه، حاصل را بر تعداد کل کلمات تشکیل دهنده متن تقسیم می‌نماییم. با انجام این عمل پارامتر (TF) برای هر لغت موجود در متن تست بدست خواهد آمد. در ادامه با مقایسه هر کلمه با کلمات موجود در پایگاه لغات اصلی، احتمال حضور کلمات در هر یک از گروه‌ها را استخراج نموده و به ماتریس قبلی اضافه می‌کنیم. در نتیجه این عمل ماتریسی در ابعاد [تعداد کلمات کاندید سند تست $\times$ 12] خواهیم داشت که ستون اول و دوم به ترتیب مربوط به لغات کاندید متن تست و تعداد تکرار نرمال شده آنها و 10 ستون دیگر

متعلق به احتمال حضور کلمه در هر گروه است. سپس به محاسبه میزان تعلق سند به گروه‌ها به کمک فرمول (2-4) می‌پردازیم:

$$Similarity(D, G) = \sum_{w=1}^{\# \text{ of } D \text{ words}} \sum_{g=1}^{10} \frac{TF_w \times TF_w NGF_g}{\sum_{k=1}^{10} TF_w NGF_k - TF_w NGF_g} \quad (2-4)$$

در فرمول فوق w معادل تعداد کلمات غیر تکراری سند آزمایشی پس از اعمال پیش‌پردازش‌ها و g و k بیانگر تعداد دسته‌ها است. در مخرج سعی بر آن شده است که تا حد امکان، میزان تمایز نتایج نهایی با ارجاع چندگانه به پارامترهای اصلی، افزایش یابد. TF_w فرکانس تکرار کلمه در متن و $TF_w NGF_{g/k}$ فرکانس تکرار در گروه‌ها را نشان می‌دهد. در نهایت وزن هر سند (D) در هر گروه (G) بدست خواهد آمد. نتایج در نهایت به صورت درصد تشابه نمایش داده می‌شوند.

4-5-4- انتساب اسناد و تعیین مقدار آستانه

بعد از یافتن میزان تعلق سند به هر دسته، بدون شک گروه با بالاترین مقدار تشابه، به عنوان گروه برنده معرفی شده و لیبل مربوط به سند را به خود اختصاص خواهد داد. در موارد بسیاری شاهد آن هستیم که یک سند خبری به بیش از یک گروه خبری اختصاص دارد، دلیل این امر همپوشانی موضوعات خبری و عدم تفکیک جنبه‌های مختلف رخدادها و روزمره زندگی است. به عنوان مثال حادثه مربوط به آتش‌سوزی در یکی از مدارس شهرستان‌های کشور، علاوه بر اینکه یک خبر مربوط به استان‌ها می‌باشد، اما در مقیاس کشوری به مسایل اجتماعی، سیاسی نیز بستگی دارد. همچنین نقد و همدردی ادبی مربوط به این حادثه در دسته فرهنگی نیز جای خواهد داشت. به همین دلیل نمی‌توان در تمام موارد تنها گروه با بالاترین میزان تشابه را به عنوان گروه برنده معرفی نمود، بلکه با در نظر گرفتن مقدار آستانه، گروه‌های با تشابه بیش از مقدار مشخص را به عنوان گروه‌های برنده با درجه تشابه متفاوت نشانه‌گذاری خواهیم کرد. روش مورد استفاده برای تعیین گروه برنده، استفاده از الگوریتم k -نزدیک‌ترین همسایه است که در این مورد نزدیک بودن به معنای مشابه بودن متن ریز سند و ریز سند ایجاد شده به کمک کلمات تشکیل دهنده متن تست در هر دسته است. از آنجا که انتخاب عدد k وابسته به متن و ارزش کلمات آن در گروه‌های مجزا است، لذا انتخاب یک عدد ثابت

روش بهینه‌ای محسوب نمی‌شود. به همین دلیل به جای انتخاب عدد معین، تعداد گروه‌ها تا زمان رسیدن به حداقل میزان تشابه را با یک آستانه‌گذاری مشخص نموده‌ایم. تعیین مقدار آستانه در اغلب موارد به صورت تجربی و پس از تحلیل نتایج به ازای هر مقدار اولیه انجام می‌گردد. در این پروژه نیز از ترکیب روش تجربی تحلیل نتایج فنی و درک انسانی استفاده شده است. بدون شک قرار دادن یک عدد ثابت به عنوان حد آستانه نتیجه درستی در پی نخواهد داشت، چرا که تمام اسناد از توزیع یکسانی برای تعلق به دسته‌ها استفاده نمی‌کنند. تفاوت در تعداد و محتوای لغات مورد استفاده در متون نیز به این دلیل دامن می‌زند. فرمول (3-4) به محاسبه حد آستانه در محدوده صفر (عدم تطابق کامل) و بیش از 100 درصد (دسته‌بندی نرم)، به ازای هر سند به صورت اختصاصی می‌زند.

$$Threshold = \frac{\max(Similarity) + \min(Similarity) + 4 \times \frac{\max(Similarity) + \min(Similarity)}{2}}{4} + \min(Similarity) \quad (3-4)$$

اساس و پایه این روش بر مبنای برآورد سه نقطه‌ای⁸⁰ تخمین زمان و هزینه در مدیریت پروژه‌ها می‌باشد که در این شرایط شما باید یک برآورد در شرایط خوش‌بینانه، یک برآورد در شرایط بدبینانه و یک برآورد در شرایط عادی داشته باشید. در این حالت با یک فرمولی که شامل وزن‌دهی به این شرایط می‌باشد، می‌توان زمان و هزینه فعالیت‌های پروژه را تخمین زد. این فرمول معمولاً به این صورت است که برآورد فعالیت برابر است با مجموع برآوردهای خوشبینانه و بدبینانه به علاوه چهار برابر برآورد در شرایط نرمال تقسیم بر مجموع ضرایب به کار رفته در صورت. در این پروژه به ترتیب حداکثر و حداقل درصد تعلق سند به گروه را معادل با شرایط خوش‌بینانه و بدبینانه در نظر گرفته و میانگین زمانی را شرایط عادی می‌نامیم. بدین صورت فرمول به شکل بالا درآمده است. در نهایت با جمع نمودن مجدد حداقل مقدار، نتیجه بدست آمده را با میزان درک کاربر انسانی تطابق داده‌ایم. در برخی موارد مشاهده شده است که به دلیل نزدیکی میان درصد تشابه گروه‌ها هیچ دسته‌ای مقدار بیش از حد آستانه را ندارد که در این مواقع از بالاترین نسبت با چشم‌پوشی از مقدار آستانه استفاده

⁸⁰ Three point Estimates

می‌کنیم. یکی از موارد قابل بحث در این بخش تعیین گروه‌بندی پایه و مرجع به منزله مقایسه نتایج خروجی سیستم با آن است. هرچند روند آموزش سیستم بر مبنای گروه‌بندی اخبار در وبسایت‌های پایه بوده است، اما به دلیل وجود تعداد تناقض زیاد بین دسته‌های معرفی‌شده و دسته‌های حقیقی و وجود ریزمتن‌های خبری تست متعلق به دسته‌های تعریف نشده اولیه در سیستم (در قسمت نتیجه-گیری به تفصیل بیان شده است)، گروه‌بندی پایه دیگری اضافه نموده‌ایم. در این گروه‌بندی جدید داده‌های تست توسط یک کاربر انسانی نشانه‌گذاری شده و عمل مقایسه و ارزیابی سیستم بر مبنای آن صورت گرفته است. جدول (4-8) میزان تشابه/عدم تشابه گروه‌های خبرگزاری، انتسابی (سیستم) و انتخابی (کاربر) را در مورد داده‌های تست نشان می‌دهد.

میزان تشابه		خبرگزاری - انتسابی		انتسابی - انتخابی		خبرگزاری - انتخابی		هر سه گروه	
عدم تشابه	0	21%	77	10%	36	16%	56	2%	8
متشابه	حداقل	79%	44	90%	61	84%	46	74%	48
	حداکثر	230	230	254	249	211	211	211	211

جدول (4-8) تشابه/عدم تشابه گروه‌های خبرگزاری، انتسابی و انتخابی

در مورد گروه انتخابی از جانب کاربر انسانی، قطعات خبری بدست آمده به ده گروه 37 تایی به انضمام نام گروه‌های نهایی در اختیار 10 داوطلب قرار داده شده است. به منزله انتخاب سند از هر کدام از افراد خواسته شد که ریز متن را مطالعه نموده و بنا به سلیقه و درک شخصی سند را با حداقل یک و حداکثر تعداد مورد تائیدی خود نشانه‌گذاری نمایند. در هر دسته خبر سعی شده است تا تنوع گروه خبری لحاظ گردد. افراد تکمیل کننده فرم‌ها در بازه سنی 24 الی 60 سال با حداکثر تمایز در دیدگاه، علایق و سن، به منزله پوشش حداکثری دیدگاه‌های متفاوت در آن انتخاب شده‌اند.

در جدول (4-9) نتایج بدست آمده از سیستم برای چندین نمونه آزمایشی به همراه گروه‌بندی اصلی سند خبری و گروه انتسابی از جانب کاربر انسانی آورده شده است.

شماره خبر	متن ریز خبر دریافتی	عنوان گروه اصلی	عنوان گروه انتسابی سیستم	عنوان انتخابی کاربر
1	به گزارش خبرنگار اقتصادی ایرنا ، اردشیر سعد محمدی روز دوشنبه در نخستین نشست مطبوعاتی با اصحاب رسانه در مسند شرکت یادشده اظهار داشت تامین نیاز فولاد کشور یک ضرورت است و دستیابی به تولید میلیون تن فولاد در دستور کار قرار دارد وی با بیان اینکه تولید فولاد کشور در سال گذشته به میلیون تن رسید ، گفت ظرفیت ما در این بخش اکنون میلیون تن است و انتظار می رود تا پایان سال جاری بیش از میلیون تن تولید داشته باشیم و رسیدن به رقم میلیون تن تولید نیز پیش بینی می شود سعد محمدی خاطرنشان ساخت آمار عملکرد تولید هشت ماهه فولاد حاکی از رشد درصد نسبت به مدت مشابه سال گذشته است و میلیون و یکصد هزار تن تولید داشتیم وی، طرحهای فولادی را قائلان ، چهارمحال وبختیاری ، بافت ، سبزوار ، شادگان ، میانه و نی ریز ذکر کرد و افزود ...	اقتصادی	اقتصادی	اقتصادی
2	رییس دستگاه دیپلماسی بامداد دوشنبه در گفت وگویی اختصاصی با خبرنگار ایرنا در باره اظهارات اخیر وزیر امور خارجه آمریکا مبنی بر تمایل واشنگتن برای مذاکره مستقیم با تهران افزود مذاکره موضوعی در موارد خاص تاکنون با آمریکا شده که از آن جمله مذاکره در مورد افغانستان و عراق است وی افزود همچنین مذاکره در باره پرونده هسته ای نیز در قالب مذاکرات ایران با گروه + که آمریکا نیز در آن حضور دارد، انجام می شود، اما اگر در اذهان عمومی مذاکره فراگیر سیاسی مطرح شده باشد، این مساله در حدود اختیارات رهبر معظم انقلاب است و معظم له تشخیص می دهند که مثلا چنین اقدامی شود یا نشود وزیر امور خارجه یادآور شد مذاکره موضوعی تاکنون در چند مرحله شده است ...	سیاسی	سیاسی	سیاسی خارجی
3	به نوشته این منابع خبری ، کشتی های حامل سامانه های موشکی پاتریوت آلمان ، امروز بندر ' تراومونده ' در استان ' لوبک ' آلمان را به سوی بندر ' اسکندرون ' ترکیه ترک کردند براساس این گزارش ، دو سامانه موشکی پاتریوت آلمان روز ژانویه دوم بهمن برای استقرار در ترکیه وارد بندر اسکندرون خواهند شد به گزارش منابع خبری ترکیه ، یک هیات نظامی متشکل از تن از نظامیان هلند و نفر از نظامیان آلمان نیز جهت فراهم کردن زمینه برای استقرار سیستم های موشکی ...	خارجی	سیاسی خارجی	سیاسی خارجی
3	به گزارش ایرنا، 'نسبیل دوفلو' روز دوشنبه در باره نامه درخواست کمک خود از کلیسا، گفت از اسقف اعظمی پاریس خواسته ام 'ساختمان های تقریبا' خالی در اختیار خود را برای اسکان بی خانمانها، افراد و خانواده های فاقد مسکن مناسب در دسترس این وزارتخانه قرار دهد وی تاکید کرد که تصرف و تملک ساختمان های خالی عمومی و در اختیار شخصیت های معنوی، ازهم اکنون تا پایان سالجاری میلادی محقق خواهد شد دوفلو با بیان اینکه امیدواراست دراین زمینه نیازی به اعمال قانون و قدرت از سوی دولت نباشد، افزود این قابل درک نخواهد بود اگر کلیسا نخواهد، برای محقق شدن اهداف و برنامه دولت در زمینه اسکان بی سرپناهان، همراهی و همیاری نکند دوفلو طرحی را به شورای وزیران فرانسه ارائه کرده است که براساس ...	خارجی	سیاسی	خارجی
4	به گزارش ایرنا از سازمان لیگ حرفه ای فوتبال، تیم فوتبال دسته یکی برق شیراز نیز از حضور در جام حذفی آزاد سازی خرمشهر انصراف داد بر اساس این گزارش، پیش از این نیز تیم آلومینیوم هرمزگان، شهرداری تبریز، سنگ آهن بافق و نفت مسجد سلیمان نیز از حضور در جام حذفی انصراف داده بودند بدین ترتیب مسابقات ...	ورزشی	ورزشی	ورزشی
5	به گزارش خبرگزاری مهر، دانشمندان تا کنون به طور کامل نمی دانستند چگونه حیوانات رنگها را می بینند چرا که رنگدانه های بصری در چشم های آنها شامل کروموفورهای کاملا شبیه به هم (بخش جذب نور مولکولها) هستند اما با این حال می توانند طول موج های مختلف نور را جذب کنند.همه حیوانات دارای کروموفور شبکیه ای (الدئید ویتامین آ) هستند و بسته به پروتئین های گیرنده نور (اوپسین ها) مرتبط با این کروموفورها، همان مولکول می تواند طیفی از رنگها را از آبی و حتی فرابنفش تا قرمز جذب کند.اینکه چگونه یک تک مولکول می تواند این کار را انجام دهند تا کنون یک راز بود ...	علمی	علمی	علمی
6	برنامه 24 ساعته برای درمان سرماخوردگی،تصور کنید که یک روز از خواب بلند می شوید و احساس می کنید سردرد دارید، بدنتان درد می کند و بینی تان گرفته است ابتدا به سرماخوردگی اولین حدسی است که می زنید با این حال کارهای زیادی است که باید انجام دهید در این شرایط چه تصمیمی می گیرید آیا با همان کسالت و درد لباستان را می پوشید و به محل کار می روید یا تصمیم می گیرید یک ...	رسانه دیگر	اجتماعی	اجتماعی

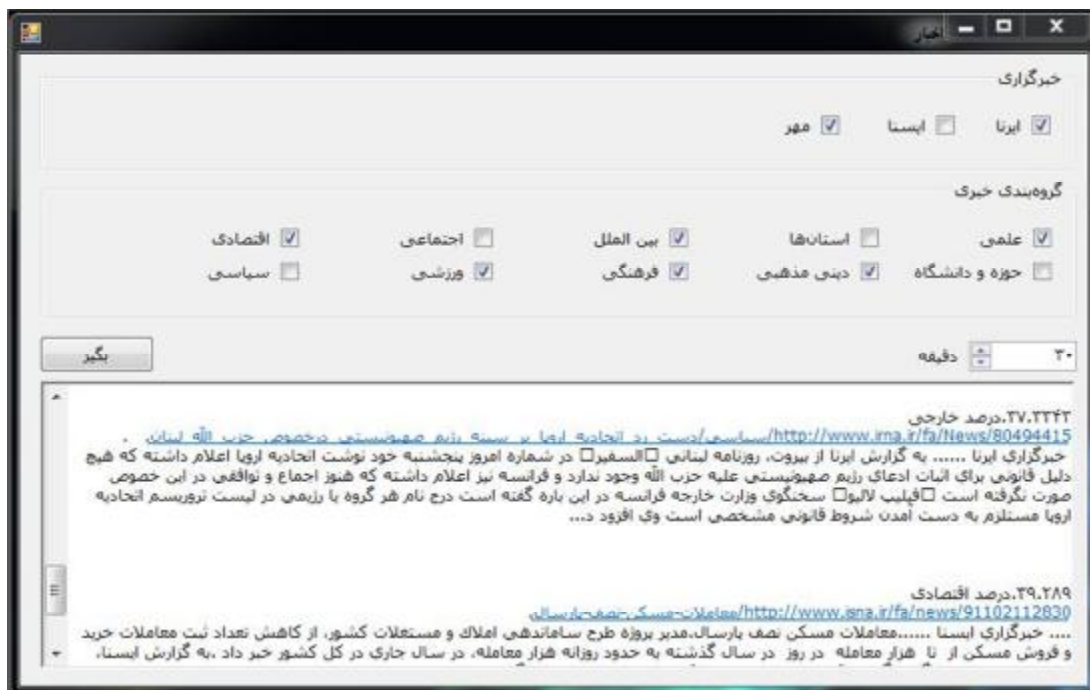
جدول (4-9): نمونه ریز متون تست و دسته های نسبت داده شده به آنها

4-5-5 بازسازی مجدد پایگاه لغات

پس از انتساب سند به گروه یا گروه‌های مربوطه، باید عملیات بازسازی مجدد پایگاه لغات و وزندهی کلمات در گروه‌ها صورت پذیرد. این روش به آموزش و یادگیری سیستم در حین عملکرد کمک بسیاری خواهد نمود. در موارد زیادی مشاهده می‌گردد که به عنوان مثال با یک رویداد سیاسی یا اقتصادی جدید مواجه می‌شویم که بالطبع اسامی خاص و تازه‌ای نیز در خود خواهند داشت. با این دامنه لغات گروه با کلمات خاص تعریف شده آموزش می‌بیند و در آینده با ضریب تشابه بالاتری اسناد مشابه را دسته‌بندی خواهد کرد. در این قسمت پس از مقایسه کلمه با پایگاه لغات اصلی در صورت غیر تکراری بودن کلمه اضافه شده و یک واحد به گروه/ گروه‌های مربوطه اضافه می‌شود و نتایج کار مستقیماً برای سند تست بعدی قابل استفاده است. یکی دیگر از دلایل عدم تمایل به ورود سند تکراری در مرحله تست نیز همین قسمت است.

4-6 نمای کلی برنامه

آخرین مرحله برنامه، به نمایش اسناد خبری در قالب دسته‌های نهایی و میزان تعلق خبر به دسته‌ها دارد. هر کاربر در لحظه آغاز قادر به گزینش یک یا چندین شاخه مورد نظر می‌باشد. این عمل به منزله شخصی سازی و تطابق برنامه با نیازها و علاقه مندی‌های کاربر صورت گرفته است. در انتهای کار، متن خبری کامل (متن قبل از ریشه‌یابی و حذف کلمات تکراری و عمومی) به همراه لیبل متعلق به گروه‌های نشانه‌گذاری شده و درصد تعلق به آنها با فرمت مورد نظر در یک فایل متنی ذخیره می‌گردد. استفاده از متن خام و اولیه در خوانایی و تطابق نتیجه با نظر کاربر، افزایش جذابیت و وضوح مفهوم تاثیر بسیاری دارد. در ابتدا قسمت جمع‌آوری و پویش اسناد خبری به کمک محیط برنامه نویسی سی‌شارپ انجام می‌گیرد، در ادامه پردازش متن در محیط مطلب صورت گرفته و در نهایت مجدداً به کمک برنامه سی‌شارپ نتایج در رابط کاربری طراحی شده نمایش داده می‌شود. تصویر (1-4) نشانگر نمای کلی برنامه و خروجی نهایی است.



شکل (4-1): نمای کلی برنامه

7-4 ارزیابی

تا این قسمت به معرفی داده‌های مورد استفاده و روند آموزش، تست و عملکرد سیستم پرداختیم. حال به بررسی و تحلیل نتایج بدست آمده و ارزیابی سیستم خواهیم پرداخت. اولین و متداول‌ترین روش برای تعیین میزان کارایی الگوریتم، معیارهای دقت (P)، فراخوانی (R)، دقیق بودن (A) و معیار شباهت F است. هرچند این معیارها در فصل قبل توضیح داده شده‌اند، اما در این قسمت نیز به صورت خلاصه فرمول هرکدام آورده شده است. برای محاسبه این معیارها به چهار پارامتر کلی نیاز داریم که در جدول (4-10) آورده شده‌اند.

متن	نسبت داده شده به گروه G_g	نسبت داده نشده به گروه G_g
متعلق به گروه G_g	True Positive (TP)	False Negative (FN)
متعلق به غیر گروه G_g	False Positive (FP)	True Negative (TN)

جدول (4-10): پارامترهای مربوط به معیارهای ارزیابی

معیار دقیق بودن، معمول‌ترین معیار مورد استفاده برای تعیین میزان دقت در الگوریتم‌های یادگیری

ماشین بوده و به کمک رابطه (4-4) محاسبه می‌گردد:

$$A = \frac{TP + TN}{TP + TN + FP + FN} \quad (4 - 4)$$

معیار دقت، میزان دقیق بودن یک جستجو را به کمک رابطه (5-4) بدست می‌آورد:

$$P = \frac{TP}{TP + FP} \quad (5 - 4)$$

و در نهایت معیار فراخوانی، میزان کامل بودن یک جستجو را از طریق فرمول (6-4) اندازه‌گیری می‌کند.

$$R = \frac{TP}{TP + FN} \quad (6 - 4)$$

در روش آخر نیز با ترکیب دو مقدار دقت و فراخوانی، معیار تشابه F (7-4) را بدست می‌آوریم:

$$F - \text{measure} = \frac{2 * R * P}{R + P} \quad (7 - 4)$$

نتایج حاصل از ارزیابی سیستم نسبت به گروه اصلی به ازای 372 سند آزمایشی از سه وبسایت پایه نسبت به گروه‌بندی انتسابی و انتخابی در جدول (11-4) آورده شده است. با توجه به اطلاعات فوق مقادیر زیر به دست آمده است.

گروه انتسابی نسبت به گروه انتخابی				گروه انتسابی نسبت به گروه خبرگزاری				پارامتر گروه
دقیق بودن	دقت	فراخوانی	معیار F	دقیق بودن	دقت	فراخوانی	معیار F	
98.38	100	75	85.71	98.00	94.11	72.72	82.05	علمی، فناوری
96.77	71.42	55.55	62.5	92.30	61.53	26.66	37.21	استان ها
93.27	93.84	74.39	82.99	92.87	75.80	82.45	78.99	بین الملل
95.16	82.97	79.59	81.25	92.31	50	77.77	60.87	اجتماعی
97.84	88.88	82.75	85.71	95.73	73.07	70.37	71.69	اقتصادی
99.73	100	95.23	97.56	99.43	90	100	94.73	حوزه، دانشگاه
97.58	100	40	57.14	97.43	85.71	42.85	57.14	دینی، مذهبی
95.96	81.66	92.45	86.72	93.73	75.47	81.63	78.43	فرهنگی
99.46	100	96.42	98.18	98.57	92.59	98.03	95.23	ورزشی
87.90	64.86	92.30	76.19	82.33	47.11	87.5	61.25	سیاسی
96.20	88.36	78.37	81.39	94.27	74.54	74.00	71.76	میانگین

جدول (4-11): نتایج حاصل از ارزیابی سیستم نسبت به گروه انتخابی و خبرگزاری

همانگونه که پیشتر نیز بیان شد، معیار دقیق بودن یک جستجو بیان کننده مقدار داده متعلق به گروه غیر دسته مربوطه است، که الگوریتم از ورود آنها ممانعت به عمل آورده است. در نتیجه پایین بودن این عدد نشان می‌دهد که تعداد زیادی از داده‌های متعلق به دسته‌های دیگر در این دسته جای می‌گیرند. با مشاهده جدول فوق نفوذ همه جانبه اطلاعات در دسته سیاسی در هر دو گروه‌بندی مشاهده می‌گیرد، که این امر با توجه به ذات داده‌های مورد پردازش دور از ذهن نیست.

معیار فراخوانی نشان دهنده میزان اسنادی است که به صورت صحیح در دسته مربوطه نشانه‌گذاری شده و به دسته دیگری ارسال نشده‌اند. بالا بودن این معیار نشان‌دهنده میزان تفکیک‌پذیری دیگر دسته‌ها از دسته مربوطه است. این مسئله به وضوح در مقولات دانشگاهی و ورزشی نشان داده شده است.

با وجود تمام مشکلات میانگین میزان دقت دسته‌بندی بالای 90 درصد در هر دو دیدگاه، مهر تأییدی بر نحوه عملکرد صحیح الگوریتم انتخابی برای نوع داده‌های ورودی مورد بحث است.

یکی از دیگر مسائل مهم در ارزیابی سیستم، رعایت سه جزء کلی پوشش اخبار در وب، نحوه بروز-رسانی و شاخص گذاری اخبار به منزله نمایش آنها به کاربر است. سازمان های خبرگزاری پایه به نحوی انتخاب شده اند، که پوشش حداکثری اخبار مورد تأیید در کشور را دارا بوده و با توجه به اینکه به-صورت متوسط (با توجه به جدول 1-4) در هر ساعت 41 سند خبری است، در نظر گرفتن زمان معادل 60 دقیقه برای به استخراج داده های جدید منطقی به نظر می رسد. در هنگام نمایش نتیجه نیز، اسناد بنا بر سلیقه و انتخاب کاربر به ترتیب در گروه های مجزا ارائه می شوند.

فصل پنجم: نتیجه گیری و پیشنهادات

1-5 مقدمه

متن کاوی اسناد وب به عنوان یک رشته مشترک مابین داده کاوی و یادگیری ماشین قرار گرفته است، که سعی دارد با تکیه بر الگوریتم‌ها و ترکیب متدهای این دو فیلد علمی به استخراج، کشف و مدیریت این دسته از داده‌ها بپردازد. در این طرح پژوهشی با ادغام روش‌های موجود و گزینش روش بهینه بر مبنای داده‌های اولیه به تعریف سیستم دسته‌بندی اسناد خبری در بستر وب پرداخته شده است. بدین‌منظور پس از گزینش دو دسته نهایی، کار آموزش سیستم را با تعداد 808 سند آموزشی از وب-سایت مربوط به سه خبرگزاری رسمی کشور آغاز نموده و پس از پاک‌سازی و آماده‌سازی داده‌ها به انتخاب کلمات کلیدی آنها پرداخته و پایگاه آموزشی نهایی را ایجاد نموده‌ایم. در ادامه عمل تست سیستم را با 372 داده ادامه داده و در نهایت کارایی سیستم را بر مبنای عنوان بیان شده در خبرگزاری و گروه انتخابی کاربران انسانی محاسبه نموده‌ایم. لازم به ذکر است که تمامی داده‌های مورد استفاده به کمک یک خدمتگذار وب، از منبع صفحات خبری استخراج شده‌اند. در این سیستم بیشترین تاکید بر روی بالا بردن ارتباط میان کلمات کلیدی انتخابی و نفش آنها در دسته‌های مختلف است. در نهایت k نزدیک‌ترین همسایه با در نظر گرفتن مقدار آستانه حداقل تشابه به عنوان گروه‌های انتسابی سیستم معرفی شده و به انضمام درصد تعلق نمایش داده شده است.

پیاده‌سازی و ارزیابی سیستم حاکی از دقت بالای 90 درصد انتساب سند به گروه‌های خبری بر مبنای گروه پیشنهادی خبرگزاری پایه و گروه انتخابی کاربران است. در ادامه به بیان برخی موارد که تاثیر منفی در روند آموزش، تست و ارزیابی سیستم داشته‌اند پرداخته و در نهایت اقدامات مورد نظر برای ادامه کار بر روی سیستم پیشنهادی بیان شده است.

2-5 تحلیل نتایج

در این قسمت با مشاهده نتایج بدست آمده از سیستم و تحلیل گروه‌های انتسابی به مواردی که در کارایی سیستم اخلاص ایجاد نموده و همچنین موارد خاص استخراجی مانند لزوم ایجاد دسته‌های جدید و ادغام دسته‌های موجود اشاره نموده‌ایم. مباحث پایه تمامی این موارد در فصل قبل مطرح شده و از دوباره‌گویی و کپی مجدد جداول و نتایج در این بخش صرف‌نظر شده است.

عدم تطابق عناوین نهایی با نمودارهای خبری:

یکی از مشکلاتی که پیشتر نیز در مورد آن صحبت نمودیم، عدم تطابق و هماهنگی یک به یک میان عناوین نهایی گروه‌های انتخابی در برنامه و عناوین مربوط به دسته‌ها تعریف شده در نمودار خبری سازمان‌های خبرگزاری است. به عنوان مثال در خبرگزاری "ایرنا" اخبار مربوط به موضوع مذهبی به عنوان زیرشاخه‌ای در گروه فرهنگی جای داده شده است. تاثیر این امر، بالطبع منجر به ایجاد اشتباه در فرآیند آموزش گشته و در ادامه داده‌های تست با عنوان غیر مذهبی نشانه‌گذاری می‌شوند. در مورد اخبار دانشگاهی نیز در این سازمان هیچ‌گونه گروه یا زیرگروه مشخصی در نظر گرفته نشده است، حال آنکه صرف‌نظر از این دو دسته خبری در گروه‌بندی نهایی امکان‌پذیر نبوده و با جامعیت کار تضاد دارد. هرچند این عمل از دید مسئولین خبری بلامانع است اما در ارزیابی سیستم، به عنوان یک خطا محسوب می‌گردد. از دیگر نمونه‌ها در خبرگزاری "ایرنا" می‌توان به بخش "پژوهش‌های خبری" اشاره نمود که به دلیل عدم همخوانی با دیگر داده‌ها در بخش آموزش حذف گشته است، اما در داده‌های تست مشاهده شد. این مطلب در جدول (4-7) نیز قابل مشاهده است.

وجود دسته‌های از پیش تعریف نشده:

در قسمت تست با برخی داده‌ها با عنوان "رسانه‌های دیگر" مواجه گشتیم که در قسمت آموزش و همچنین در نمودار اولیه خبرگزاری "ایسنا" نامی از آن برده نشده بود. از آنجا که فرآیند آموزش بر اساس گروه‌های تعریف شده صورت گرفته است، لذا این اسناد خبری متعلق به هیچ گروهی نبوده و در نهایت امکان ارزیابی را با مشکل مواجه نموده است. در واقع در نهایت امر تعدادی داده نشانه

گذاری شده، بدون نشانه اصلی وجود دارد که تاثیر منفی بر نتایج بدست آمده وارد کرده است. جدول (4-7) حاکی از آن است که تعداد 20 سند خبری استخراجی تحت این عنوان نشانه گذاری شده اند، که بدلیل عدم وجود در دسته های نهایی در ارزیابی به عنوان داده های محسوب می گردند که به اشتباه نشانه گذاری شده اند. به همین دلیل این داده ها در ارزیابی سیستم در نظر گرفته نشده اند.

عدم دسته بندی صحیح اولیه:

در ابتدای فصل نحوه جمع آوری و نشانه گذاری سنتی اسناد در سازمان های خبرگزاری را بیان نمودیم. در هر سازمان، گروه بندی بسته به حوزه خبری و کاری هر خبرنگار معین گشته و به صورت معمول اخبار در بخش مربوطه انتشار می یابند. در موارد متعددی مشاهده شده است که خبر نه تنها به گروه تعلق ندارد، بلکه در زمینه انتشار خبر نیز کار به درستی صورت نگرفته است. به عنوان مثال تصویر (5-1) یک سند خبری است که برخلاف تهیه آن توسط خبرنگار علمی، به عنوان یک خبر اقتصادی منتشر شده است. از این زمینه ناهمگونی ها به شدت در صفحه اصلی سایت های پایه مشاهده می گردد که بی شک توجه بیشتر مسئولین امر را می طلبد.



شکل (5-1): نمونه ای از تناقض در گروه بندی پیشنهادی خبرگزاری ها

همراستای گروه‌ها:

با مطالعه و مشاهده نتایج بدست آمده، می‌توان چنین بیان داشت که مذهب در فرهنگ مردم و جامعه تا جایی نفوذ پیدا کرده است که امکان تفکیک این دو مورد از یکدیگر به سختی حاصل می‌شود. لذا شاید بتوان مذهب را به عنوان زیر شاخه گروه فرهنگی در نظر گرفت. از طرف دیگر عنوانی مانند "فرهنگ حماسه" هرچند در دید اول بیشتر به مقولات فرهنگی می‌پردازد، اما از آنجا که بیشترین عنصر در این زمینه فرهنگ جنگ و دفاع مقدس و پیامدهای آن می‌باشد، لذا اکثر اخبار در این زمینه همراستای اخبار و رویدادهای سیاسی کشور شکل گرفته است. سیستم نیز در نهایت تمام اخبار به دسته فرهنگ حماسه را به صورت ترکیبی فرهنگی و سیاسی نشانه‌گذاری کرده است.

در همین زمینه می‌توان به اخبار ارسالی از مراکز استان‌ها و یا شهرستان‌ها نیز اشاره کرد. به عنوان مثال سخنرانی یکی از سیاستمداران یا اختصاص بودجه و افتتاح یک طرح عمرانی در یک شهر، با وجود تعلق خبر به شهر مربوطه بدون شک در دید کشوری در دسته اخبار سیاسی اقتصادی نیز قرار می‌گیرند. در موارد بسیاری نیز مشاهده می‌گردد که به عنوان مثال ارزش فرهنگی ساخت یک فیلم یا مستند از محل فیلمبرداری بیشتر بوده و نام شهر موردنظر در مقایسه با فعالیت اصلی، قابل چشم‌پوشی است. این در حالی است که در وبسایت مربوط به خبرگزاری‌ها هر عنوان تنها به یک دسته تعلق داشته و این امر نیز در ارزیابی نهایی، سیستم را با مشکل مواجه نموده است.

از دیگر مواردی که پس از آزمایش تست مشاهده شد، می‌توان به گروه‌بندی جدید "سیاست خارجه" نام برد. در بسیاری موارد متن خبری، متن ترجمه و یا ارسالی از جانب خبرگزاران خارج از مرز بوده و در نهایت در دسته خارجی قرار داده شده‌اند. اما این اخبار در عین حال بر مسائل سیاسی متمرکز بوده و از آنجا که سیاست خارجه و علی‌الخصوص اخبار خاورمیانه تاثیر مستقیم بر سیاست داخلی ایران دارد، لذا قرار گرفتن آنها در بخش سیاسی نیز الزامی است. در همین راستا سیستم در حدود 29 مورد را با هر دو عنوان سیاسی خارجی نامگذاری کرده و مهمتر از آن اینکه از نقطه نظر کاربر انسانی نیز این نامگذاری دسته در بیش از 22 مورد مشاهده شد.

ردپای سیاست نیز (همانند روند معمول و متداول جامعه) در هر واقعه و رخدادی قابل مشاهده و پیگیری بوده و تمام جوانب کار کشور را تحت تاثیر قرار داده است، این امر در نحوه عملکرد سیستم نیز به وضوح مشاهده می‌گردد.

3-5 اقدامات آینده

بدون شک هدف نهایی انجام این پروژه، فراتر از معرفی یک سیستم دسته‌بندی اسناد بوده و سیستم طراحی شده، تنها به عنوان سیستم پایه برای اهدافی است که در آینده قصد دستیابی به آنها را داریم. آخرین بخش این گزارش به بیان اقدامات آینده در راستای تکمیل پروژه می‌پردازد.

گسترده‌گی داده‌های مورد پردازش:

در بخش قبل مشاهده شد که تنها از اطلاعات موجود در سه سازمان خبرگزاری در سیستم استفاده شد، هرچند این سه سازمان به دلیل اهمیت آنها در کشور، حجم بالایی از اخبار را پوشش می‌دهند، اما در ادامه قصد وارد نمودن دیگر سازمان‌های خبرگزاری به منزله دستیابی کامل به تمام موارد خبری جاری در کشور را داریم. این مسئله به منزله استفاده از وبسایت‌های تحلیلی خبری و یا پورتال‌های سازمانی و یا استفاده از موتورهای جستجوگر برای جمع‌آوری اخبار موجود در وبلاگ‌ها و سایر منابع خبری غیر رسمی نیست. چرا که با این کار نه تنها حجم اطلاعات غیر مفید بالا می‌رود، بلکه به دلیل عدم نظارت مراجع ذیربط بر کار این دسته صفحات امکان ورود اخبار غیر واقعی افزایش می‌یابد.

ادغام داده‌های خبری:

استفاده از چندین پایگاه خبری، علاوه بر پوشش حداکثری وقایع، مارا با حجم عظیمی از اسناد تکراری مواجه خواهد نمود. بیان این اطلاعات مشابه، نه تنها کاربر را بار دیگر دچار سردرگمی می‌کند، بلکه باعث به هدر رفتن زمان و فضای اختصاص یافته برای اخبار می‌گردد. به همین دلیل یکی از اقدامات مورد نظر تلاش به منزله ادغام اسناد خبری مشابه و ذکر یک خبر به انضمام لینک دسترسی به اخبار مشابه آن در دیگر وبسایت‌ها است. از آنجا که هر خبرگزاری بلافاصله پس از وقوع یک

حادثه سعی در آماده سازی و انتشار خبر دارد، در نتیجه در هر بار به روزرسانی سیستم (بازبینی صفحات) به احتمال زیاد با اخبار مشابه مواجه خواهیم شد. به همین دلیل ادغام این اخبار در درجه اول پیش از انتساب آنها به گروه‌ها صورت می‌گیرد. این امر در بالا بردن نقش کلمه در سند و همچنین کاهش تعداد کلمات قابل مقایسه با دسته‌های دیگر و نهایتاً کاهش زمان دسته‌بندی اسناد نقش بسزایی خواهد داشت.

ایجاد جریان‌ات خبری:

یکی از مواردی که این پروژه بر مبنای آن طراحی شده است، عدم استفاده از درخواست کاربر به منزله نمایش اسناد خبری است. یعنی برخلاف پروژه‌های بسیاری که در آن پس از اعلام درخواست کاربر با حجم محدود اطلاعات به منزله دسته‌بندی مواجه هستند، این سیستم مستقل از کاربر بوده و تلاش در پوشش تمام اخبار انتشار یافته را دارد. تنها نقش کاربر انتخاب دسته‌هایی است که تمایل دارد اخبار مربوط به آنها را مرور نماید. تا این قسمت از کار به بیان اخبار داغ در هر بازه بازبینی صفحات پرداخته‌ایم. خبر داغ به خبری گفته می‌شود که اخیراً به وقوع پیوسته و انتشار یافته شده است. گام بعدی برقراری ارتباط معنادار میان اخبار جدید و آرشیو خبری است. بدین صورت که اخبار مشابه با یکدیگر در هر گروه با یک عنوان جدید (به کمک الگوریتم‌هایی که در فصل دوم به آنها اشاره شد) نامگذاری شده و به کاربر نمایش داده شود. ایجاد این جریان‌ات خبری به مطالعه روند یک واقعه پرداخته و تاریخچه‌ای غیر مدون از واقعه را در اختیار کاربر قرار خواهد داد.

[1] Nasraoui O., Soliman M., Saka E., Badia A. and Germain R., (2008), "A Web Usage Mining Framework for Mining Evolving User Profiles in Dynamic Web Sites", IEEE TRANSACTIONS ON KNOWLEDGE AND DATA ENGINEERING, vol. 20, 1041-1047.

[2] Achananuparp P., Han H., Nasraoui O. and Johnson R., (2007), "Semantically Enhanced User Modeling", Proceedings of the 2007 ACM Symposium on Applied Computing, SAC '07. ACM, New York, NY, 1335-1339.

[3] Tan A.-H. (1999), "Text mining: The state of the art and the challenges", In Proc of the Pacific Asia Conf on Knowledge Discovery and Data Mining PAKDD'99 workshop on Knowledge Discovery from Advanced Databases, pages 65-70.

[4] صفی خانی ژ. و عبدالله زاده بارفروش ا.، (1384)، "طراحی SPAM: ابزاری برای کاوش در شبکه گسترده جهانی"، یازدهمین کنفرانس سالانه انجمن کامپیوتر ایران.

[5] قادریان م.، (1387)، پایان نامه ارشد: "بهبود مدل کاربر در وبسایت بصورت خودکار با استفاده از معناشناسی با مفاهیم خاص دامنه"، دانشکده کامپیوتر، دانشگاه صنعتی امیر کبیر، تهران.

[6] مطهری نژاد ح. ر.، (1382)، پایان نامه کارشناسی ارشد: "بازیابی کارا و مؤثر اطلاعات وب از طریق دستاوردهای یادگیری ماشین: طراحی و تکامل روشهای یادگیری تقویتی در کاوش متمرکز"، دانشکده کامپیوتر، دانشگاه صنعتی امیر کبیر، تهران.

[7] LookOFF E-book, (2002), Engine Basice,
<http://www.LookOff.com/tactics/engines.php3>.

[8] Uramoto, N., Matsuzawa, H., Nagano, T., Murakami, A., Takeuchi, H. and Takeda, K., (2004), "A text-mining system for knowledge discovery from biomedical documents" IBM Systems Journal., Vol.43: 516-533.

[9] Eick, C., F., Zeidat, N. and Zhao, Z., (2004), "Supervised Clustering- Algorithms and Benefits" 16th IEEE International Conference on Tools with Artificial Intelligence (ICTAI'04): 774-776.

- [10] Finley, T. and Joachims, T., (2005), "Supervised clustering with support vector machines" ICML '05 Proceedings of the 22nd international conference on Machine learning.
- [11] Berkhin, P., (2002), "Survey Of Clustering Data Mining Techniques" Springer. Vol. 12: 1-56.
- [12] Osinski S., (2004), "Dimensionality Reduction Techniques For Search Results Clustering", Msc.thesis,computer science department of sheffield university, UK.
- [13] Osinski S., Stefanowski J., Weiss D., (2004), "Lingo: Search Results Clustering Algorithm Based on Singular Value Decomposition", Springer, 359--368.
- [14] Osiński S. and Weiss D., (2005), "A Concept-Driven Algorithm for Clustering Search Results", [IEEE](#) Computer Society, Intelligent Systems.
- [15] Wu, X., Kumar, V., Quinlan, J., R., Ghosh, J., Yang, Q., Motoda, H., McLachlan, G., J., Ng, A., Liu, B., Yu, P., S., Zhou, Z., H., Steinbach, M., Hand, D., J. and Steinberg, D., (2007), "Top 10 algorithms in data mining", ICDM, Springer-Verlag London Limited. Vol. 14.
- [16] ANDERBERG, M. R. (1973), "Cluster Analysis for Applications", Academic Press, Inc., New York, NY.
- [17] JAIN A.K., MURTY M.N., FLYNN P.J., (1999), "Data Clustering: A Review", ACM Computing Surveys, Vol. 31, No. 3.
- [18] Zamir O. and Etzioni O., (1998), "Web Document Clustering: A Feasibility Demonstration", ACM New York, NY, USA.
- [19] Dumais S. T. et.al, (1988), "Using Latent Semantic Analysis to Improve Access to Textual Information", Proceedings of the Conference on Human Factors in Computing Systems.
- [20] Carpineto C., Osinski S., Romano G., Weiss D., (2009), "A Survey of Web Clustering Engines", ACM Computing Surveys (CSUR).
- [21] Pavlou, D., Balasubramanian, R., Dom, B., Kapur, S. and Parikh, J., (2004), "Document preprocessing for naive Bayes classification and clustering with

mixture of multinomials" Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining.

[22] Isa, D., Lee, L., H., Kallimani, V., P. and RajKumar, R., (2008), "Text Document Preprocessing with the Bayes Formula for Classification Using the Support Vector Machine", IEEE Trans, Knowledge and Data Engineering. Vol. 20: 1264-1272.

[23] guyon, I. and Elisseeff, A., (2003), "An Introduction to variable and feature selection", Machine learning research., Vol. 3: 1157-1182.

[24] Yuan, R., Li, Z., Guan, X. and Xu, L., (2008), "An SVM-based machine learning method for accurate internet traffic classification", Springer. Vol. 12: 149-156.

[25] Joachims, T., (1998), "Text Categorization with Support Vector Machines: Learning with Many Relevant Features", In European Conference on Machine Learning (ECML).

[26] حقیری م.، (1390) ، پایان نامه ارشد: "استفاده از تکنیک خوشه بندی با نظارت برای دسته بندی نامه های دریافتی در سامانه 137 شهرداری تهران"، دانشگاه شیراز.

[27] Qi, Y., Zhang. Y. and Song M., (2009), "Text Mining for Bioinformatics: State of the Art Review" iccsit, 2nd IEEE International Conference on Computer Science and Information Technology: 398- 401.

[28] احمدی م. ح. و منجمی ا. ح. و آیت س.، (1389) ، "دسته بندی متون فارسی با استفاده از قوانین انجمنی"، چهارمین کنفرانس داده کاوی ایران، تهران :دانشگاه صنعتی شریف.

[29] Han, J. and Kamber, M., (2006), "Data Mining: Concepts and Techniques" 2nd ed, SanFrancisco: Morgan Kaufmann.

[30] محمدی جنقرا م. و آنالوئی م.، (1386) ، "استخراج کلمات کلیدی اسناد فارسی"، سیزدهمین کنفرانس ملی انجمن کامپیوتر ایران.

[31] Taghva, K., Beckley, R., Sadeh, M., (2005), "A Stemming Algorithm for the Farsi Language", ITCC '05 Proceedings of the International Conference on Information Technology: Coding and Computing (ITCC'05) - Volume I - Volume 01 Pages 158-162

- [32] پورمعصوی آ. و کامیار ح. و کامیار م.، (1387)، "گزارش مطالعه و بررسی الگوریتم‌های موجود برای ریشه‌یابی کلمات"، دانشگاه فردوسی مشهد.
- [33] Rahimtoroghi E., Faili H., Shakery A., (2010), "A Structural Rule-based Stemmer for Persian", International Symposium on Telecommunications.
- [34] صفائی ع. و بدیع ک. و عباسیان ع.، (1387)، "یک ریشه‌یاب برای واژگان عاریتی زبان فارسی"، چهاردهمین کنفرانس ملی سالانه انجمن کامپیوتر ایران، دانشگاه صنعتی امیرکبیر.
- [35] Porter M.F., (1980), "An Algorithm for suffix stripping". Program, 14(3):130t137.
- [36] Sheykh Esmaili K., Rostami A., (2006), "List of Persian Stopwords", Technical Report No. 2006-03, Semantic Web Research Laboratory, Sharif University of Technology, Tehran, Iran, June.
- [37] Taghva K., Beckley R., Sadeh M., (2003), "A List of Farsi Stopwords", Technical Report, Information Science Research Institute, University of Nevada, Las Vegas.
- [38] Davarpanah M.R., Sanji M., Aramideh M., (2009) "Farsi lexical analysis and stop word list", Library Hi Tech, Vol. 27 Iss: 3, pp.435 – 449.
- [39] مقدم چرکری ن. و زمانیان م.، (1388)، "دسته‌بندی متون فارسی با استفاده از الگوریتم SVM و بررسی روشهای کاهش خصیصه"، سومین کنفرانس داده‌کاوی ایران، تهران: دانشگاه علم و صنعت ایران.
- [40] Yang, Y. and Pedersen, J. O., (1997), "A Comparative Study on Feature Selection in Text Categorization", Proc. ICML: 412-420.
- [41] ملکی م. و عبدالله زاده بارفروش ا.، (1385)، "TFCRF: روش جدید وزن دهی ویژگی مبتنی بر اطلاعات کلاس در حوزه طبقه‌بندی مستندات"، دوازدهمین کنفرانس بی‌نام‌لی انجمن کامپیوتر ایران.
- [42] Salton, G., Yang, C.S., (1973), "On the Specification of Term Values in Automatic Indexing", Journal of Documentation, Vol. 29, No. 4, pp. 351-357.
- [43] Lan, M., Sung, S.Y., Low, H.B., Tan, C.L., (2005), "A Comparative Study on Term Weighting Schemes for Text Categorization", IEEE International Conference on Neural Networks (IJCNN05), pp. 546-551.

[44] عربی نرئی س. و وحیدی اصل م. و مینایی بیدگلی ب.، (1386) ، "استخراج کلمات کلیدی جهت طبقه بندی متون فارسی"، دانشگاه صنعتی امیرکبیر.

[45] Shamsfard M., Hesabi A., Fadaie H., Famian A., Bagherbeigi S., Mansoory N., Fekri E., Monshizadeh M., Assi S.M., (2010), "**Semi-Automatic** Development of FarsNet; The Persian WordNet", Accepted at Global WordNet Conference (GWA2010), India.

[46] Hassanpour H., Darvishi A., (1387), "A New Similarity Measure For Signal Processing Applications", IEEE International Workshop on Signal Processing and Its Applications.

[47] Sánchez M. M., Gutiérrez E. B., Gallinas R. B., García A. M. F., Vidales M. A. S., (2008), "Clustering of Web Documents: Full-text or Snippets?", IADIS.

[48] Ferragina P., Gull A., (2004), "The Anatomy of a Hierarchical Clustering Engine for Web-page, News and Book Snippets", Data Mining, Fourth IEEE International Conference.

[49] Sayyadi H., Salehi S., AbolHassani H., (2006), "NeSReC: A News meta-Search Engines Result Clustering Tool", Advances and Innovations in Systems, Computing Sciences and Software Engineering.

[50] Zamir, O. AND Etzioni, O. (1999), "Grouper: A dynamic clustering interface to Web search results", Comput. Netw. 31, 11–16, 1361–1374.

[51] Zhang,D., Dong,Y., (2004), "Semantic, hierarchical, online clustering of Web search results". In Proceedings of 6th Asia-PacificWeb Conference (APWeb). Lecture Notes in Computer Science, vol. 3007. Springer, 69–78.

[52] Giannotti, F., Nanni, M., Pedreschi, D., Samaritani, F. (2003), "WebCat: Automatic categorization of Web search results". In Proceedings of the 11th Italian Symposium on Advanced Database Systems (SEBD), 507–518.

[53] Ferragina, P., Gulli, A. (2004), "The Anatomy of SnakeT: A Hierarchical Clustering Engine for Web- Page Snippets", In Proceedings of the 8th European Conference on Principles and Practice of Knowledge Discovery in Databases. Lecture Notes in Computer Science, vol. 3202. Springer, 506–508.

[54] Subhashini R., Jawahar Senthil Kumar V., (2010), "The Anatomy of Web Search Result Clustering and Search Engines" Indian Journal of Computer Science and Engineering Vol. 1 No. 4 392-401.

[55] Lancaster, F.W., and Fayen, E.G. (1973). "Information Retrieval On-Line" Los Angeles, CA: Melville Publishing Co. Chapter 6.

[56] Ezzati Jivan, N., Fazeli, M., Yousefi, K., (2010), "New Approach for Automated Categorizing and Finding Similarities in Online Persian News", 2nd international symposium on information management in a changing world.

[57] Lee S., Kim H., (2008), "News Keyword Extraction for Topic Tracking", Fourth International Conference on Networked Computing and Advanced Information Management, NCM.

[58] Wang C., Zhang M., Ma S., Ru L., (2008), "Automatic Online News Issue Construction in Web Environment", ACM, (IW3C2).

[59] Farhoodi M., Yari A., (2010), "Applying Machine Learning Algorithms for Automatic Persian Text Classification", [Advanced Information Management and Service \(IMS\)](#).

[60] Li Z., Zhou D., Juan Y., Han J., (2010), "Keyword Extraction for Social Snippets", WWW '10 Proceedings of the 19th international conference on World wide web, Raleigh, North Carolina, USA, 1143-1144.

[61] نعمتی ش. و بصیری م.، (1386)، "دسته‌بندی اسناد فارسی با استفاده از الگوریتم KNN"، همایش منطقهای برق و الکترونیک باشگاه پژوهشگران جوان واحد فسا.

[62] رحمتی م. و رضوی ر.، (1382)، "یک روش سریع برای خوشه بندی نتایج جستجوی وب"، اولین کنفرانس بین المللی فناوری اطلاعات و دانش، تهران، دانشگاه صنعتی امیرکبیر.

[63] ارسطوپور ش.، (1388)، "دسته بندی نتایج جستجو بر مبنای ویژگیهای مدارک و امکان سنجی استفاده از الگوریتمهای خوشه بندی مختلف در سطح وب"، فصلنامه کتابداری و اطلاع رسانی، شماره 46، صفحه 145 – 174.

[64] دقاقزاده ح. و عبداللهی ازگمی م. و دقاقزاده م.، (1388)، "آرائه روشی بر پایه کاوش معنایی به منظور بالا بردن دقت نتایج جستجو در نرمافزارهای قرآنی"، سومین کنفرانس داده‌کاوی ایران، تهران، دانشگاه علم و صنعت ایران.

[65] مقصودی ن. و همایون پور م. م.، (1388) ، "ارائه روشی جدید در طبقه‌بندی متون فارسی با استفاده از دانش معنایی"، پانزدهمین کنفرانس بین المللی سالانه انجمن کامپیوتر ایران.

[66] Lipai A., (2008), "Web Metasearch Result Clustering System", *Revista Informatica Economica* 4(48).

[67] Bassma S. Alsulami, Maysoon F. Abulkhair, Fathy A. Essa, (2012), "Semantic Clustering Approach Based Multi-agent System for Information Retrieval on Web", *IJCSNS International Journal of Computer Science and Network Security*, VOL.12 No.1.

[68] عبدی قویدل ه. و وزیرنژاد ب. و بحرانی م.، (1389) ، "برچسب‌زنی موضوعی متون فارسی"، دانشگاه صنعتی شریف، تهران، ایران.

[69] Zeng H., He Q., Chen Z., Ma W., Ma J., (2004), "Learning to Cluster Web Search Results", 27th annual international ACM SIGIR conference on Research and development in information retrieval, ACM New York, NY, USA.

Abstract

Nowadays the need and tendency of being announced about news, has an unavoidably and direct effect on making decision in personal, social and political fields. On the other hand, Uncontrolled growth of website numbers and their information has becomes one of the most important problems which cause many difficulties for users on the World Wide Web. By the way, wide speed of unsupervised news publishing and spreading on the web, and impressive increase of news websites in different fields, has reduced the ability for checking them. This huge amount of information has confused users instead of helping them to achieve their purpose. Because of it, design and implementation of an automatic news clustering system for managing the news on the web has explainable. Proposed method tries to resolve the problem by using data and web mining methods for clustering text news snippets.

At the first, a web client has defined for crawling the valid news agencies for acquiring news for learn and test levels. The clustered news in pre-defined groups has been showed as the result of the project. User can select both base news agency and beloved groups for final representation.

At last, too many researches have been done till now for clustering texts but all of these efforts use either term frequency or document frequency for weighing extracted keywords and then tried to make a comparison between different similarity measures, but recently a few numbers of researches have been focused on clustered features instead of documents and words. In this thesis we intend to use group's features by building a semi text with main text words in each group and calculating its similarity. The proposed architecture consists of a number of components namely news extraction, preprocessing, group database creation, keyword extraction, keyword

weighting, clustering, and database reconstruction. Our main contribution is that we introduce a mechanism to automatically catch and cluster news snippets on the web. At last, the groups which have passed the defined threshold are labeled as winner groups. Our evaluations show that the proposed method along with news snippet structures represents it's efficiently and effectiveness.

Keywords: *Clustering, News Documents, Data Mining, Web Mining*



Shahrood University of Technology

Faculty of Information Technology &
Computer Engineering

An Automatic Method for Clustering News on the Internet

Melika Yaghoubi

Supervisor: Dr. Hamid Hassanpour

Adviser: Dr. Morteza Zahedi

January ۲۰۱۳