



وزارت علوم، تحقیقات، و فناوری
پژوهشگاه علوم و فناوری اطلاعات ایران
(ایرانداک)

طرح پژوهشی

**بررسی موتور جستجوی گنج به منظور بهبود بازیابی
اطلاعات در زبان فارسی**

مجری

مرضیه زرین بال ماسوله

شهریور ۱۳۹۸

بنام خداوند جان و

حقوق: پژوهشگاه علوم و فناوری اطلاعات ایران

طرح پژوهشی « بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی » پیرو قرارداد شماره ۳۳/۶۷۵۰ مورخ ۱۳۹۷/۰۶/۱۲ میان پژوهشگاه علوم و فناوری اطلاعات ایران (کارفرما) و خانم مرضیه زرین بال ماسوله (مجری) اجرا شده است.

این گزارش و تمامی حقوق مادی آن بر اساس «قانون حمایت حقوق مؤلفان و مصنفان و هنرمندان، مصوب سال ۱۳۸۴ و اصلاحیه‌های بعدی آن و همچنین آیین‌نامه‌های اجرایی این قانون متعلق به پژوهشگاه علوم و فناوری اطلاعات ایران است و هر گونه استفاده از تمامی یا پاره‌ای از آن، شامل: نقل قول، تکثیر، انتشار، کاربرد نتایج، تکمیل و مانند آنها به صورت چاپی، الکترونیکی یا وسایل دیگر فقط با اجازه کتبی پژوهشگاه امکانپذیر است. نقل قول در حد هزار واژه در انتشارات علمی مانند کتاب و مدرک با درج اطلاعات کامل کتابشناختی، نیازی به مجوز پژوهشگاه ندارد.

صحت مندرجات گزارش برعهده مجری طرح پژوهشی است.

بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی

مدیر طرح پژوهشی: مرضیه زرین بال ماسوله، پژوهشگاه علوم و فناوری اطلاعات ایران

نشانی: تهران، خیابان انقلاب، چهارراه فلسطین، شماره ۱۰۹۰

صندوق پستی: ۱۳۱۵۷۷۳۳۱۴، تلفن: ۰۲۱۶۶۴۹۴۹۸۰، دورنگار: ۰۲۱۶۶۴۹۴۹۸۰

وب گاه: www.irandoc.ac.ir

رایانامه: zarinbal@irandoc.ac.ir



برنام‌نما

آغاز نیم‌قرن دوم خدمات ارزشمند ایرانداک به علم، فناوری، و نوآوری کرامی باد

۱۳۹۸ - ۱۳۴۷

گواهی

انجام طرح پژوهشی:

◇ عنوان: بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی

◇ مجری: دکتر مرضیه زرین‌بال ماسوله

◇ همکار(ان): -

◇ شماره قرارداد: ۳۳/۶۷۵۰

◇ تاریخ قرارداد: ۱۳۹۷/۶/۱۲

◇ تاریخ شروع: ۱۳۹۷/۶/۱۲

◇ تاریخ پایان: ۱۳۹۸/۶/۶

◇ ناظر: دکتر آزاده محبی

در پژوهشگاه علوم و فناوری اطلاعات ایران گواهی می‌شود. گزارش پایانی این طرح، بر پایه داوری در نشست ۳۴۴ (تاریخ ۱۳۹۸/۷/۲) شورای پژوهشی پژوهشگاه به تصویب رسیده است. لازم می‌دانم مراتب قدردانی خود را به مجری و ناظر این طرح تقدیم دارم.

با آرزوی پیروزی
پرویز شهریاری
معاون پژوهش و آموزش

چکیده

بازیابی اطلاعات به معنای یافتن محتوایی با طبیعت غیر ساخت یافته (معمولاً متن) از مجموعه‌هایی به منظور برآورده ساختن نیاز(های) اطلاعاتی کاربران است. با بهره‌گیری از روش‌های بازیابی اطلاعات موجود در ادبیات پلتفرم‌های متن‌باز متعددی طراحی شده‌اند. کتابخانه آپاچی لوسن و موتور جستجوی آپاچی سولر از این نوع پلتفرم‌ها هستند. سامانه گنج نیز به منظور جستجو در اسناد و مدارک ذخیره شده در پایگاه داده خود، از این کتابخانه و موتور جستجو به عنوان هسته اصلی فرآیند بازیابی اطلاعات استفاده می‌کند. علی‌رغم کاربردهای گسترده، سولر دارای کاستی‌هایی بوده، تاکنون نیز هیچ‌گونه پژوهشی در زمینه شناسایی و واکاوی این کتابخانه و موتور جستجو در پژوهشگاه انجام نشده و تنها عملکرد رابط کاربری آن و مشکلات موجود بررسی شده است. لذا علی‌رغم ضرورت بررسی کارایی رابط کاربری گنج، تازمانی که هسته اصلی آن (خصوصاً با تمرکز بر زبان فارسی) بررسی و اصلاح نشود نمی‌توان انتظار داشت عملکرد سامانه مطلوب باشد. لذا ضروری است تا با بررسی محدودیت‌ها و مشکلات موجود در روش‌های مورد استفاده در موتور جستجوی آپاچی سولر (از منظر بازیابی اطلاعات) دیدگاهی کلان نسبت به آن در پژوهشگاه ایجاد شده تا از این طریق بتوان پیشنهادهایی به منظور ارتقاء عملکرد هسته اصلی فرآیند بازیابی اطلاعات سامانه گنج ارائه داد. براین اساس، طرح پژوهشی با عنوان «**بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی**» پیشنهاد شد. در این طرح دانش فنی درمورد این موتور جستجو و روش‌های مورد استفاده آن کسب شده و راهکارهایی به منظور ارتقاء نتایج بازیابی در زبان فارسی و در سامانه گنج ارائه شد. راهکارهای پیشنهاد شده در قالب دو دسته راهکارهای نیازمند پژوهش و راهکارهای فنی و عملیاتی دسته‌بندی شدند.

کلید واژه‌ها: بازیابی اطلاعات؛ آپاچی لوسن؛ آپاچی سولر؛ سامانه گنج

فهرست مطالب

عنوان	شماره صفحه
چکیده.....	ب
فهرست مطالب.....	ب
فهرست جدول ها.....	ج
فهرست شکل ها.....	خ
فصل اول: بیان مسئله.....	۲
۱-۱ بیان مسئله.....	۲
۱-۲ هدف ها.....	۴
۱-۳ پرسش ها.....	۴
۱-۴ روش پژوهش.....	۵
فصل دوم: کتابخانه لوسن و موتور جستجوی سولر.....	۷
۲-۱ مقدمه.....	۷
۲-۲ بازیابی اطلاعات.....	۸
۲-۳ کتابخانه لوسن.....	۱۰
۲-۴ موتور جستجوی سولر.....	۱۷
فصل سوم: مباحث پایه‌ای در موتور جستجوی سولر.....	۲۱
۳-۱ مقدمه.....	۲۱
۳-۲ مفاهیم کلیدی در سولر.....	۲۳
۳-۲-۱ تنظیمات در سولر.....	۲۳
۳-۲-۲ سند.....	۲۹
۳-۲-۳ تطابق فازی.....	۳۰
۳-۲-۴ ارتباط.....	۳۲

۳۵.....	۳-۲-۵- معیارهای سنجش بازیابی اطلاعات
۳۶.....	۳-۳- زیرسیستم مدیریت اسناد
۳۷.....	۳-۳-۱- تنظیمات ورود و بروزرسانی اسناد
۳۹.....	۳-۴- زیرسیستم تحلیل متن
۴۲.....	۳-۴-۱- ریشه‌یابی
۴۶.....	۳-۴-۲- زبان‌های پیش‌فرض در سولر
۴۸.....	۳-۵- زیرسیستم پردازش پرس‌وجو
۵۱.....	۳-۵-۱- مولفه Query
۵۸.....	۳-۵-۲- مولفه Facet
۵۸.....	۳-۵-۳- مولفه More Like This
۵۹.....	۳-۵-۴- مولفه Highlight
۶۲.....	۳-۵-۵- مولفه Last-Component
۶۵.....	۳-۵-۶- مولفه Response Writer
۶۶.....	۳-۶- جمع‌بندی
۷۰.....	فصل چهارم: مباحث پیشرفته در موتور جستجوی سولر
۷۰.....	۴-۱- مقدمه
۷۰.....	۴-۲- استفاده از داده‌های منابع بیرونی
۷۲.....	۴-۳- عملیات جستجوی پیچیده
۷۳.....	۴-۴- ارتقاء ارتباط سند با عبارت پرس‌وجو
۷۶.....	۴-۵- شخصی‌سازی جستجو
۸۱.....	۴-۶- جمع‌بندی
۸۳.....	فصل پنجم: موتور جستجوی گنج
۸۳.....	۵-۱- مقدمه
۸۳.....	۵-۲- سند Solrconfig سامانه گنج
۸۳.....	۵-۳- سند Schema سامانه گنج
۸۴.....	۵-۴- جمع‌بندی
۸۶.....	فصل ششم: راهکارها و روش‌های بهبود
۸۶.....	۶-۱- مقدمه
۸۶.....	۶-۲- راهکارهای بهبود بازیابی اطلاعات در سامانه گنج

۸۹.....	۱-۲-۶- راهکارهای نیازمند پژوهش
۹۱.....	۲-۲-۶- راهکارهای فنی و عملیاتی
۹۴.....	جمع‌بندی و نتیجه‌گیری
۹۹.....	ضمائم
۹۹.....	۱-۷- توابع
۹۹.....	۱-۱-۷- توابع ریاضی
۱۰۰.....	۲-۱-۷- توابع فاصله
۱۰۱.....	۳-۱-۷- توابع منطق بولی
۱۰۱.....	۲-۷- عوامل محاسبه ارتباط
۱۰۲.....	۳-۷- کلاس‌های محاسبه ارتباط
۱۰۳.....	۴-۷- خصوصیات Update Handler
۱۰۳.....	۵-۷- پارامترهای مولفه Highlight
۱۰۴.....	۶-۷- پارامترهای مولفه Spellcheck
۱۰۵.....	منابع
109.....	ABSATRACT

فهرست جدول‌ها

جدول ۱-۲- فرمت‌های ورودی قابلیت پردازش در لوسن (McCandless, Hatcher, and Gospodnetic 2010).....	۱۱
جدول ۱-۳- جستجوی Wildcard: عبارت پرس‌وجوی کاربر و عبارتی که سولر جستجو می‌کند.....	۳۰
جدول ۲-۳- جستجوی نزدیکی: عبارت پرس‌وجوی کاربر و توضیحات آن	۳۱
جدول ۳-۳- جستجوی بازه‌ای: عبارت پرس‌وجوی کاربر و عبارتی که سولر برای جستجو استفاده می‌کند.....	۳۲
جدول ۴-۳- در نظر گیری وزن برای فیلدهای مختلف	۳۴
جدول ۵-۳- متن اولیه مثال شکل ۳-۷، تغییرات اعمال شده و ابزارهای سولر برای تغییرات	۴۰
جدول ۶-۳- عملکرد ریشه‌یاب‌های KStemFilter،PorterStemFilter و EnglishMinimalStemFilter	۴۳
جدول ۷-۳- پارامتر نوع نتایج خروجی	۶۵
جدول ۱-۴- توابع محاسبه میزان ارتباط سند و عبارت پرس‌وجو	۷۲
جدول ۱-۷- توابع ریاضی در سولر	۹۹
جدول ۲-۷- توابع فاصله در سولر	۱۰۰
جدول ۳-۷- توابع منطق بولی در سولر	۱۰۱
جدول ۴-۷- عوامل محاسبه ارتباط در سولر	۱۰۱
جدول ۵-۷- کلاسهای محاسبه ارتباط در سولر	۱۰۲
جدول ۶-۷- خصوصیات Update Handler	۱۰۳

جدول ۷-۷- برخی از پارامترهای مولفه highlight ۱۰۳

جدول ۷-۸- برخی از پارامترهای مولفه Spellcheck ۱۰۴

فهرست شکل ها

- شکل ۱-۲- نحوه ارتباط بین آپاچی لوسن و آپاچی سولر (Grainger and Potter 2014) ۱۰
- شکل ۲-۲- کلاس های مورد استفاده در لوسن برای نمایه سازی متن (McCandless, Hatcher, and Gospodnetic 2010) ۱۱
- شکل ۳-۲- منوی مرور کلی در Luke (McCandless, Hatcher, and Gospodnetic 2010) ۱۳
- شکل ۴-۲- منوی اسناد در Luke (McCandless, Hatcher, and Gospodnetic 2010) ۱۴
- شکل ۵-۲- منوی جستجو در Luke (McCandless, Hatcher, and Gospodnetic 2010) ۱۵
- شکل ۶-۲- توضیحات مربوط به نحوه محاسبه وزن در Luke (McCandless, Hatcher, and Luke 2010) ۱۵
- شکل ۷-۲- منوی موارد افزوده در Luke (McCandless, Hatcher, and Gospodnetic 2010) ۱۶
- شکل ۱-۳- اجزای سولر ۲۲
- شکل ۲-۳- نمونه ای از سند Solrconfig ۲۴
- شکل ۳-۳- نتایج جستجوی عبارت «text analysis» برای سطح ریزدانگی هر کتاب و هر فصل هر کتاب ۲۷
- شکل ۴-۳- نمودار کلاس انواع فیلدها در Schema ۲۹
- شکل ۵-۳- شباهت کسینوسی بین بردارهای واژگان ۳۳

شکل ۳-۶- زیرسیستم مدیریت اسناد: فرآیند نمایه‌سازی و ذخیره‌سازی اسناد.....	۳۷
شکل ۳-۷- مثال تحلیل متن	۴۰
شکل ۳-۸- مثالی از StandardTokenizer	۴۱
شکل ۳-۹- اعمال دو فیلتر StopFilterFactory و LowerCaseFilterFactory در مثال شکل ۳-۸	۴۱
شکل ۳-۱۰- نمونه‌هایی از نتایج متفاوت بکارگیری روش‌های ریشه‌یابی و بُن‌واژه‌سازی	۴۵
شکل ۳-۱۱- درنظرگیری فیلدهای جداگانه در تحلیل‌های چند زبانی	۴۶
شکل ۳-۱۲- دسته‌بندی اسناد براساس زبان در تحلیل‌های چند زبانی	۴۷
شکل ۳-۱۳- تشخیص خودکار زبان سند در تحلیل‌های چند زبانی	۴۷
شکل ۳-۱۴- نمای کلی مدیریت درخواست جستجو	۴۸
شکل ۳-۱۵- کلاس‌های مدیریت درخواست در سولر	۴۹
شکل ۳-۱۶- نمای درختی کلاس‌های SearchHandler	۴۹
شکل ۳-۱۷- فرآیند کلی SearchHandler	۵۰
شکل ۳-۱۸- نمونه‌ای از کاربرد مولفه‌های بخش ۳ شکل ۳-۱۷	۵۱
شکل ۳-۱۹- گام‌های پردازش q و fq	۵۳
شکل ۳-۲۰- نمایش نمادین عبارت پرس و جو در eDisMax در حالت تفاوت وزن فیلدها	۵۵
شکل ۳-۲۱- نمونه‌ای از نتایج جستجو پس از فعال کردن مولفه hithighlight	۵۹
شکل ۳-۲۲- فرآیند نمایش snippet و پررنگ کردن بخشهایی از سند	۶۰
شکل ۳-۲۳- فرآیند پردازش اسناد و پرس و جو در سولر	۶۷
شکل ۳-۲۴- فرآیند پردازش اسناد و پرس و جو در سولر در قالب مثال	۶۸
شکل ۴-۱- سیستم پیشنهاد دهنده آمازون برای کالاهای مشابه	۸۰
شکل ۴-۲- سیستم پیشنهاد دهنده آمازون برای کالاهای مکمل	۸۰
شکل ۴-۳- جدول ارتباط بین کاربران و اسناد	۸۱

فصل اول

بیان مسئله

۱-۱ بیان مسئله

بازیابی اطلاعات شامل سازماندهی، ذخیره سازی، بازنمایی (نمایش) و دسترسی به اقسام اطلاعاتی است. به عبارت دیگر به فناوری و دانش پیچیده جستجو و استخراج اطلاعات، داده‌ها و فراداده‌ها در انواع گوناگون منابع اطلاعاتی مانند بانک اسناد، مجموعه‌های تصاویر، وب، و...، و در نهایت نمایش آنها برای کاربر آغاز کننده عملیات جستجو، بازیابی اطلاعات گفته می‌شود. سیستم‌های بازیابی اطلاعات جهت خود کارسازی فرآیند بازیابی بکار می‌روند. مراکز اطلاعاتی بسیاری همچون دانشگاه‌ها و کتابخانه‌های عمومی از سیستم‌های بازیابی اطلاعات بمنظور فراهم آوردن امکان دستیابی به کتب، روزنامه‌ها، و اسناد مختلف استفاده می‌کنند. سامانه گنج نیز به منظور جستجو در اسناد^۱ ذخیره شده در پایگاه داده خود، از کتابخانه آپاچی لوسن^۲ و موتور جستجوی آپاچی سولر^۳ استفاده می‌کند. این کتابخانه، کتابخانه‌ای متن‌باز^۴ برای راه‌اندازی موتورهای جستجو برپایه متن بوده و با زبان جاوا نوشته شده است. سولر نیز موتور جستجوی این کتابخانه بوده و مانند لوسن با زبان جاوا پیاده‌سازی شده است. به بیانی دیگر هسته اصلی سامانه گنج سولر و لوسن است

^۱ Document

^۲ Apache Lucene

^۳ Apache Solr

^۴ Open source

که هر دو متن باز بوده و با زبان جاوا پیاده‌سازی شده‌اند. برای تعامل با این هسته نیز رابط‌های کاربری متعددی توسعه داده شده‌اند، رابط کاربر سامانه گنج^۱ پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک) نمونه‌ای از این موارد است.

لذا عملکرد سامانه گنج را می‌توان از دو منظر بررسی کرد؛ (۱) عملکرد رابط کاربری طراحی شده و تعامل آن با کاربران و (۲) عملکرد هسته اصلی سامانه. در مورد عملکرد رابط کاربری مطالعات بسیاری در پژوهشگاه صورت پذیرفته که از این میان می‌توان به (ارشادی و همکاران، ۱۳۹۴)، (فتاحی، ۱۳۹۶) و (زرین‌بال، ۱۳۹۷) اشاره نمود. اما در مورد عملکرد هسته سامانه تاکنون نیز هیچ‌گونه پژوهشی در زمینه شناسایی و واکاوی این کتابخانه و موتور جستجو در پژوهشگاه انجام نشده و تنها به استفاده از تنظیمات پیش فرض اکتفا شده است. این هسته علی‌رغم کاربردهای وسیع و گسترده، کاستی‌هایی نیز دارد؛ به عنوان نمونه، در سولر از رویکردهای بهینه‌سازی موتور جستجو استفاده نشده است و ارتباط بین اسناد در بازیابی‌ها در نظر گرفته نمی‌شود؛ در صورتی که داده‌ای بین چند مدرک و یا فیلد اطلاعاتی مشترک باشد، این داده باید برای هر مدرک تکرار شود؛ در سولر نمی‌توان به سادگی اطلاعات مربوط به یک مدرک را برورسانی نمود و باید از ابتدا فرآیند نمایه‌سازی را انجام داد و سولر برای جستجوهای بهینه است که عبارات پرس‌وجو از تعداد کمی از کلمات تشکیل شده باشد. علاوه بر این، روش‌های مورد استفاده روش‌های پایه‌ای بازیابی اطلاعات هستند درحالی‌که در طی سال‌های گذشته توسعه‌های بسیاری برای آنها ارائه شده است. با این وجود، متن باز بودن موتور جستجوی سولر امکان تغییر ماژول‌ها و روش‌های بکاررفته در آن را برای طراحان و پیاده‌سازان فراهم آورده است.

بدیهی است که اکتفا به استفاده از تنظیمات پیش فرض سولر در سامانه گنج و عدم توسعه آن، با توجه به نیازهای کاربران سامانه گنج و خصوصاً زبان فارسی، سبب خواهد شد تا مشکلاتی در عملکرد رابط کاربری این سامانه ایجاد شده و در تعامل آن با کاربران مشکلاتی ایجاد شود، به برخی از این مشکلات در (ارشادی و همکاران، ۱۳۹۴) اشاره شده است. همچنین به منظور بهبود و ارتقاء عملکرد این سامانه ضروری است تا هسته اصلی بازیابی اطلاعات آن بهبود یابد. بنابراین بایستی کتابخانه آپاچی لوسن و موتور جستجوی آپاچی سولر و روش‌های مورد استفاده در برخی از مراحل و اجزای مختلف آن به صورت دقیق بررسی شده، مشکلات فعلی این روش‌ها در حوزه بازیابی اطلاعات و خصوصاً در زبان فارسی استخراج شده و راهکارهای بهبود پیشنهاد گردد.

^۱ <https://ganj-beta.irandoc.ac.ir/>

براین اساس، انجام طرح پژوهشی حاضر با هدف بررسی هسته سامانه گنج و با عنوان «بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی» پیشنهاد می‌شود. انتقال دانش فنی این کتابخانه و موتور جستجو و معرفی روش‌های مورد استفاده در هریک از ماژول‌های آن به پژوهشگاه و ارائه پیشنهادهایی به منظور بکارگیری روش‌های جایگزین معرفی شده در ادبیات در برخی از ماژول‌ها با تمرکز بر زبان فارسی از خروجی‌های این طرح خواهد بود.

بنابراین، بررسی رابط کاربری گنج و عملکرد آن در حیطه این طرح نیست. همچنین سطح شناسایی این کتابخانه و موتور جستجو محدود به نمودار مفهومی^۱ سطح صفر و یک بوده و فعالیت‌هایی چون بررسی دقیق و تفسیر کدهای موجود، ارائه الگوریتم جدید، کدنویسی و پیاده‌سازی آن در حیطه این طرح قرار ندارند.

۱-۲ هدف‌ها

براساس موارد مطرح شده، اهداف اصلی طرح پژوهشی مذکور بشرح زیر است:

□ کسب و انتقال دانش فنی کتابخانه آپاچی لوسن و موتور جستجوی آپاچی سولر برای ارتقاء بازیابی اطلاعات در گنج

جهت دستیابی به اهداف فوق اهداف فرعی طرح پژوهشی به شرح زیر تعیین می‌گردند:

□ بررسی سطح صفر و یک ماژول‌ها و روش‌های بکاررفته در کتابخانه آپاچی لوسن و موتور جستجوی آپاچی سولر؛

□ شناسایی و بررسی مشکلات و محدودیت‌های روش‌های بکاررفته در کتابخانه لوسن و موتور جستجوی سولر (با تمرکز بر زبان فارسی) برای بازیابی اطلاعات؛

□ ارائه روش‌های پیشنهادی و یا جایگزین به منظور بهبود عملکرد بازیابی اطلاعات در زبان فارسی در موتور جستجوی سولر.

۱-۳ پرسش‌ها

این طرح به دنبال یافتن پاسخ به پرسش‌های زیر است:

^۱ Context diagram

بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی □ ۵

- ساختار کتابخانه آپاچی لوسن و موتور جستجوی آپاچی سولر چگونه است و از چه ماژول‌هایی تشکیل شده است؟
- چه روش‌ها و یا الگوریتم‌هایی در این ماژول‌ها بکار رفته‌اند؟
- این روش‌ها از چه مشکلات و یا محدودیت‌هایی رنج می‌برند؟
- به منظور بهبود نتایج بازیابی شده در زبان فارسی از چه روش‌ها و یا الگوریتم‌های دیگری می‌توان استفاده کرد؟

۴-۱ روش پژوهش

پژوهش حاضر در قالب دو مرحله اصلی زیر صورت پذیرفته است:

مرحله اول: شناسایی و بررسی کتابخانه آپاچی لوسن و موتور جستجوی آپاچی سولر

- شناسایی ماژول‌های آپاچی لوسن و آپاچی سولر
- بررسی روش‌های بکاررفته در هر یک از ماژول‌ها

مرحله دوم: ارائه راهکارها و روش‌های بهبود

- شناسایی مشکلات فعلی سامانه گنج از منظر عملکرد هسته آن (ماژول‌های آپاچی لوسن و آپاچی سولر) خصوصاً در زبان فارسی
- ارائه پیشنهادهایی به منظور بهبود عملکرد بازیابی اطلاعات در زبان فارسی (شامل پیشنهادهایی برای جایگزینی روش‌های موجود با روش‌های بهبود یافته معرفی شده در ادبیات با تمرکز بر زبان فارسی. این پیشنهادها شامل تفسیر کدهای موجود، ارائه الگوریتم جدید، کدنویسی و پیاده‌سازی نخواهند بود.)

فصل دوم

کتابخانه لوسن

و موتور جستجوی سولر

۲-۱ مقدمه

اطلاع‌یابی یکی از بنیادی‌ترین نیازهای بشر برای توسعه و پیشرفت بوده و وجود منابع بزرگ اطلاعاتی و در راس آنها وب، اهمیت بازیابی اطلاعات را در سطح گسترده افزایش داده است. آنچه امروزه اهمیت بسیار دارد طراحی و بکارگیری روش‌هایی است که به کاربران کمک می‌کنند اطلاعات مورد نیاز خود را در انبوهی از اطلاعات غیرساخت‌یافته بیابند. مفهوم بازیابی اطلاعات^۱ در چنین بافتی تعریف می‌شود (Manning, Ragahvan, and Schutze 2009).

در سال‌های اخیر سامانه‌های بسیاری با هدف بازیابی اطلاعات طراحی شده‌اند و بسیاری از آنان از پلتفرم‌های متن‌باز برای این منظور بهره می‌برند. یکی از پرکاربردترین این پلتفرم‌ها، پلتفرم طراحی شده توسط آپاچی، کتابخانه لوسن و موتور جستجوی سولر، است. سامانه گنج پژوهشگاه ایرانداک نیز به منظور مدیریت و جستجو در اسناد ذخیره شده در پایگاه داده خود از این پلتفرم استفاده می‌کند.

این فصل با هدف بررسی بیشتر موتور جستجوی سولر نگاشته شده است هرچند نیم‌نگاهی نیز به کتابخانه لوسن دارد.

^۱ Information retrieval

۲-۲ بازیابی اطلاعات

جمع آوری مدارک و مستندات موجود در پایگاه داده و نمایه سازی^۱ آنها دو گامی است که عموماً پیش از شروع فرآیند بازیابی اطلاعات انجام می شود. پس از ورود نیاز اطلاعاتی کاربر در قالب عبارت پرس و جو^۲، عملیات متنی بر عبارت پرس و جوی اعمال شده و فرآیند بازیابی اسناد آغاز می شود. در انتها و پس از رتبه بندی نیز نتایج بازیابی شده به کاربر نمایش داده می شود. مقاله بوش با عنوان «As we may Think» (Bush 1945) را می توان نقطه آغازین فعالیت های بازیابی اطلاعات دانست و پس از آن روش ها و مدل های متعددی ارائه شده است که از این میان می توان به موارد زیر اشاره نمود:

- ۱۹۶۰- مارون و کوهن: بازیابی اطلاعات احتمالی (Maron and Kuhns 1960)
- ۱۹۷۱- روچیو: بازخورد در بازیابی اطلاعات (Rocchio 1971)
- ۱۹۷۱ و ۱۹۷۵- سالتون و همکاران: مدل برداری در بازیابی اطلاعات (Salton 1971) و (Salton, Wong, and Yang 1975)
- ۱۹۷۶- رابرتسون و اسپارک- جونز: وزن دهی به عبارت پرس و جو (Robertson and Sparck-Jones 1976)
- ۱۹۸۰- بوک اشتین: مدل بولین فازی در بازیابی اطلاعات (Bookstein 1980)
- ۱۹۸۶- ون ریجزبرگر: منطق در بازیابی اطلاعات (Van Rijsbergen 1986)
- ۱۹۸۸- دومیاس و همکاران: تحلیل معنایی پنهان در بازیابی اطلاعات (Dumais et al. 1988)
- ۱۹۸۹- کوک: شبکه های عصبی در بازیابی اطلاعات (Kwok 1989)
- ۱۹۹۵- رابرتسون و همکاران: BM25 (Robertson et al. 1995) و (Robertson, Walker, and Hancock-Beaulieu 1995)
- ۱۹۹۵- چن: یادگیری ماشین در بازیابی اطلاعات (Chen 1995)
- ۱۹۹۶- ریبیرو و مانتز: شبکه باور در بازیابی اطلاعات (Ribeiro and Muntz 1996)
- ۱۹۹۸- پونت و کرافت: مدل زبانی در بازیابی اطلاعات (Ponte and Croft 1998)
- ۱۹۹۸- اورد و دیکما: بازیابی اطلاعات بین زبانی (Oard and Diekema 1998)
- ۲۰۰۰- هیمسترا: ترکیب فراوانی کلمه- فراوانی معکوس سند^۳ و مدل زبانی (Hiemstra 2000)

^۱ Indexing

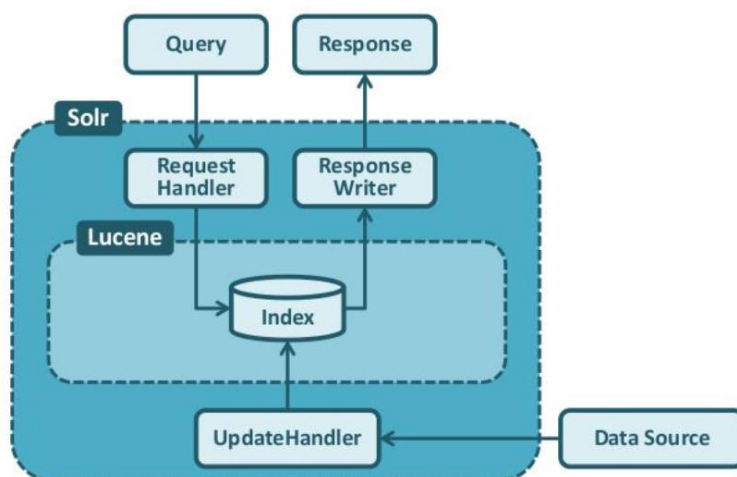
^۲ Query

^۳ Term frequency-inverse document frequency (TF-IDF)

- ۲۰۰۱- براون و جونز: بازیابی اطلاعات بافت آگاه (Brown and Jones 2001)
- ۲۰۰۲- آماتیو ون ریجزبرگر: بکارگیری معیار Divergence from Randomness (DFR) در بازیابی اطلاعات (Amati and Van Rijsbergen 2002)
- ۲۰۰۳- بروزا و سانگ و زاراگوزا و همکاران: بکارگیری احتمال در مدل‌های زبانی (Bruza and Song 2003) و (Zaragoza, Hiemstra, and Tipping 2003)
- ۲۰۰۴- متزلر و کرافت: شبکه‌های استنباطی در بازیابی اطلاعات (Metzler and Croft 2004)
- ۲۰۰۵- فنگ و ژای: رویکردهای معنایی در بازیابی اطلاعات (Fang and Zhai 2005)
- ۲۰۰۸- ژای: مدل‌های زبانی احتمالی در بازیابی اطلاعات (Zhai 2008)
- ۲۰۱۱- الای و دمیستر: مدل‌های احتمالی در بازیابی اطلاعات (Aly and Demeester 2011)
- ۲۰۱۲- گراوسمن و فریدر: روش‌های هیورستیک در بازیابی اطلاعات (Grossman and Frieder 2012)
- ۲۰۱۳- هافمن و همکاران: رتبه‌بندی نتایج بر اساس یادگیری برخط (Hofmann, Whiteson, and de Rijke 2013)
- ۲۰۱۴- شن و همکاران: تحلیل معنایی پنهان و یادگیری عمیق در بازیابی اطلاعات (Shen et al. 2014)
- ۲۰۱۵- لیو و همکاران: دسته‌بندی معنایی با کمک یادگیری عمیق در بازیابی اطلاعات (Liu et al. 2015)
- ۲۰۱۶- آی و همکاران: مدل برداری پاراگراف در بازیابی اطلاعات (Ai et al. 2016)
- ۲۰۱۷- کلوچ و همکاران: رتبه‌دهی بصری به نتایج بازیابی اطلاعات (Klouche et al. 2017)
- با بهره‌گیری از روش‌های ارائه شده، پلتفرم‌های بازیابی اطلاعات متعددی طراحی شده است و نکته اساسی در یک پلتفرم مناسب آن است که این پلتفرم برای جستجو باید بتواند نتایج مرتبط را در زمان کوتاهی بازیابی کند؛ اشتباهات تایپی کاربر را در عبارت پرس‌وجو اصلاح نماید؛ عبارت پرس‌وجوی مشابه را به صورت خودکار پیشنهاد دهد؛ کلمات جستجوی هم‌معنی و مشابه را شناسایی کرده و در جستجو از آن استفاده کند؛ قادر باشد عبارت‌های پرس‌وجوی حاوی کلمات پرکاربرد را مدیریت کند؛ و کاربر بتواند با کمک ابزارهای طراحی شده در رابط کاربر نتایج مورد نیاز خود را در سریع‌ترین زمان مشاهده کند (Grainger and Potter 2014).

برخی از پلتفرم‌های طراحی شده متن‌باز هستند مانند آپاچی سولر^۱، Lemur^۲، MeTA^۳ و Terrier^۴ و برخی نیز به صورت اختصاصی توسعه داده شده‌اند مانند Google، Bing و Amazon.

از آنجا که سامانه گنج از پلتفرم آپاچی، لوسن و سولر، به منظور مدیریت و جستجو در اسناد ذخیره شده در پایگاه داده خود استفاده می‌کند، تمرکز این طرح بر این پلتفرم خواهد بود. کتابخانه آپاچی لوسن و موتور جستجوی آپاچی سولر با زبان جاوا نوشته شده‌اند اما در عین حال قابلیت پیاده‌سازی با زبان‌هایی چون Python، Ruby، C++، C#، Perl، Delphi و PHP را دارند. نحوه ارتباط کتابخانه لوسن و موتور جستجوی سولر در شکل ۱-۲ نمایش داده شده است. لوسن وظیفه نمایه‌سازی داده‌های ورودی را برعهده داشته و وظایف سولر نیز تحلیل عبارت پرس‌وجو، اجرای جستجو و نمایش اطلاعات بازیابی شده به کاربر است (Grainger and Potter 2014).



شکل ۱-۲- نحوه ارتباط بین آپاچی لوسن و آپاچی سولر (Grainger and Potter 2014)

۲-۳ کتابخانه لوسن

لوسن کتابخانه بازیابی اطلاعات مقیاس‌پذیر^۵ با کارایی بالا بوده و امکان اضافه کردن جستجو را برای طراحان سامانه‌ها فراهم می‌کند. به بیانی دیگر، لوسن هسته پردازشی برای نمایه‌سازی داده‌های

^۱ <http://lucene.apache.org/solr/>

^۲ <http://www.lemurproject.org/>

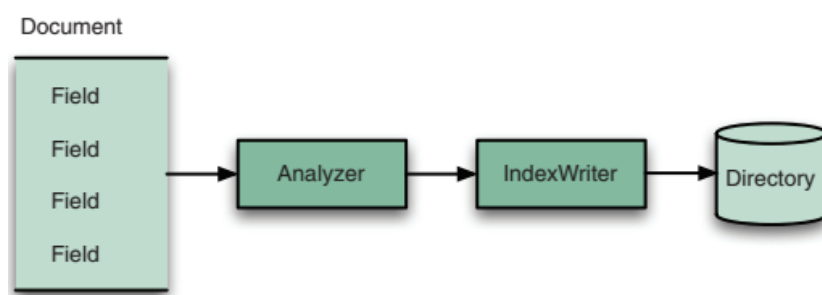
^۳ <https://meta-toolkit.org/>

^۴ <http://terrier.org/>

^۵ Scalable

ورودی در قالب متن و تا سطح ریزدانگی کلمه بوده و برای بکارگیری به عنوان موتور جستجو نیازمند ابزارهای دیگری است (McCandless, Hatcher, and Gospodnetic 2010).

کلاس‌های مورد استفاده در لوسن برای نمایه‌سازی متن به صورت شماتیک در شکل ۲-۲ نشان داده شده‌اند. در ابتدا هر سند (متشکل از یک یا چند فیلد) توسط کلاس Analyzer تحلیل شده و به واژه^۱ شکسته می‌شود و سپس نمایه‌های^۲ استخراج شده توسط IndexWriter در مکانی (Directory) ذخیره می‌شوند (McCandless, Hatcher, and Gospodnetic 2010).



شکل ۲-۲- کلاس‌های مورد استفاده در لوسن برای نمایه‌سازی متن (McCandless, Hatcher, and Gospodnetic 2010)

لوسن قابلیت پردازش ورودی‌ها با فرمت‌های متعددی را دارد. این موارد به همراه کتابخانه مورد استفاده در جدول ۱-۲ گزارش شده‌اند (McCandless, Hatcher, and Gospodnetic 2010).

جدول ۱-۲- فرمت‌های ورودی قابلیت پردازش در لوسن (McCandless, Hatcher, and Gospodnetic 2010)

کتابخانه مورد استفاده	فرمت ورودی
Apache POI	Microsoft's OLE2 Compound Document Format (Excel, Word, PowerPoint, Visio, Outlook)
PDFBox	Adobe Portable Document Format (PDF)
Java Swing API (RTFEditorKit)	Rich Text Format (RTF)
ICU4J library	Plain text character set detection
CyberNeko library	HTML
Java's javax.xml classes	XML
Java's built-in zip classes, Apache Commons Compress	ZIP Archives
Apache Ant, Apache Commons Compress	TAR Archives
Apache Commons Compress	BZIP2 compression
	AR Archives
	CPIO Archives

^۱ Token

^۲ Index

Java's built-in support (GZIPInputStream) , Apache Commons Compress ava's <code>javax.imageio</code> classes ASM library (JCR-1522) Java's built-in zip classes and ASM library, Apache Commons Compress Implemented directly Parses XML directly Parses metadata from FLV files directly Java's built-in support (<code>javax.audio.midi.*</code>)	GZIP compression Image formats (metadata only) Java class files Java JAR files MP3 audio (ID3v1 tags) OpenDocument Adobe Flash MIDI files (embedded text, eg song lyrics) WAVE Audio (sampling metadata)
---	---

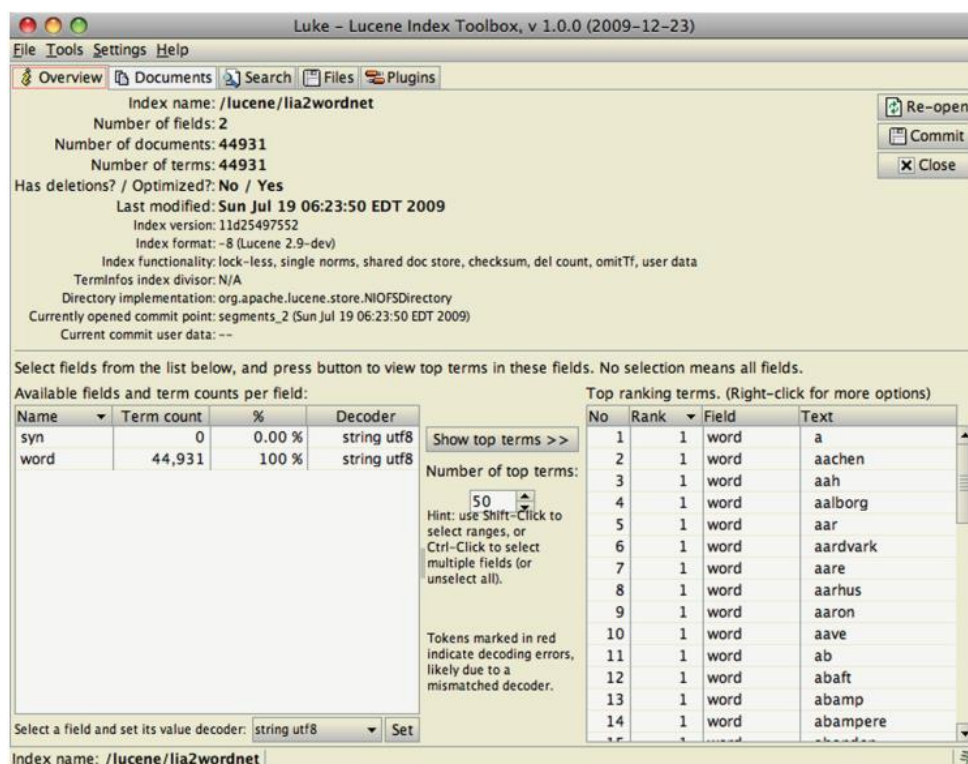
پس از ورود داده به لوسن متن باید تحلیل گردد که این امر توسط کلاس Analyzer انجام می‌شود. به جهت اهمیت، به این کلاس به صورت مفصل در بخش **Error! Reference source not found.** زیرسیستم مدیریت اسناد، پرداخته می‌شود.

یکی از ابزارهای مهم لوسن برای مدیریت نمایه‌ها، مرور اسناد و واژگان بر اساس نمایه‌ها، گزارش آمارهای مربوط به نمایه‌ها و ساختاردهی مجدد به اسناد، ابزار ^۱ Luke است که برای اجرای آن بهره‌مندی از نسخه ۱.۵ به بعد (JRE) Java Runtime Environment ضروری است (McCandless, Hatcher, and Gospodnetic 2010).

این رابط از پنج منو تشکیل شده است: نمای کلی (Overview)، اسناد (Documents)، جستجو (Search)، فایل‌ها (Files) و موارد افزوده (Plugins). در منوی نمای کلی، کاربر یا مدیرسیستم قادر است تصویر کلان نمایه‌های لوسن را مشاهده کند. این نما شامل تعداد فیلدها، تعداد اسناد و تعداد واژگان است (شکل ۲-۳). همچنین واژگان پرکاربرد در یک یا چند فیلد انتخابی در پنل Top Ranking Terms نشان داده شده است و با انتخاب هر واژه، اسناد حاوی آن نمایش داده می‌شود (McCandless, Hatcher, and Gospodnetic 2010).

^۱ <http://code.google.com/p/luke>

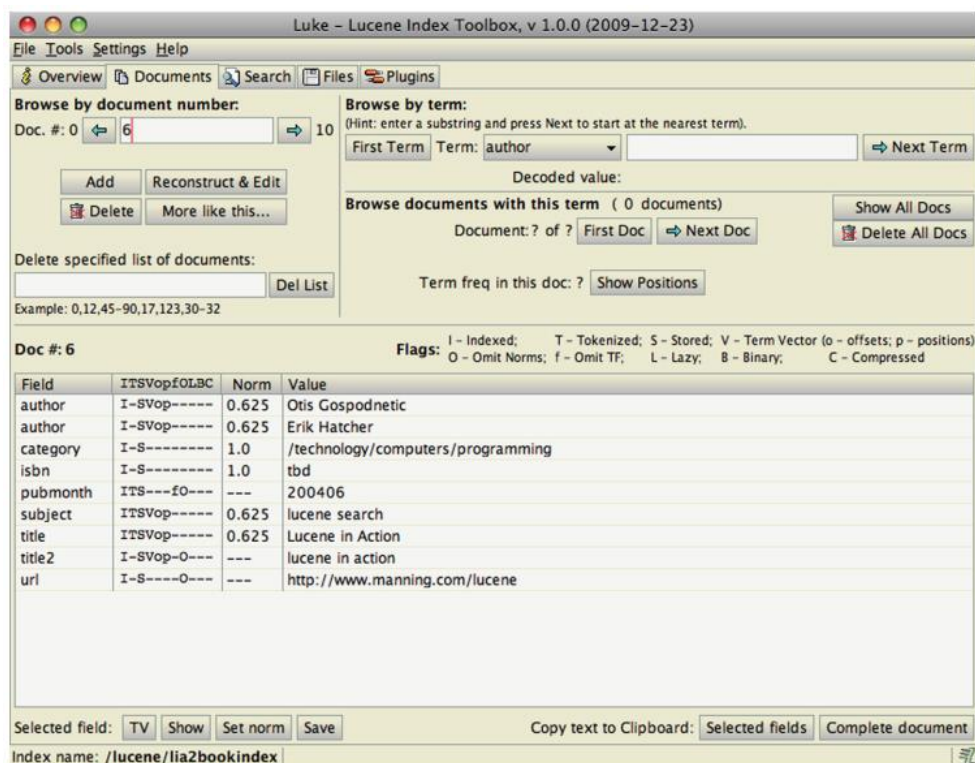
بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی □۱۳



شکل ۲-۳- منوی مرور کلی در Luke (McCandless, Hatcher, and Gospodnetic 2010)

در منوی اسناد کاربر قادر است اسناد را با کمک دو روش مشاهده و مرور نماید: (۱) با انتخاب شماره سند (Browse by document number)، که در این حالت فیلدهای مربوط به سند به همراه مشخصات و اطلاعات هر فیلد را در بخش پایین صفحه نمایش داده می شود (شکل ۲-۴) و یا (۲) با انتخاب واژه (Browse by term)، که در این حالت یا با انتخاب پرکاربردترین واژه در فیلد منتخب (در شکل ۲-۴ Frist Term در فیلد author فعال است) و یا با تایپ واژه مورد نظر در محل مربوطه می توان اسناد حاوی چنین واژگانی را مشاهده نمود (McCandless, Hatcher, and Gospodnetic 2010).

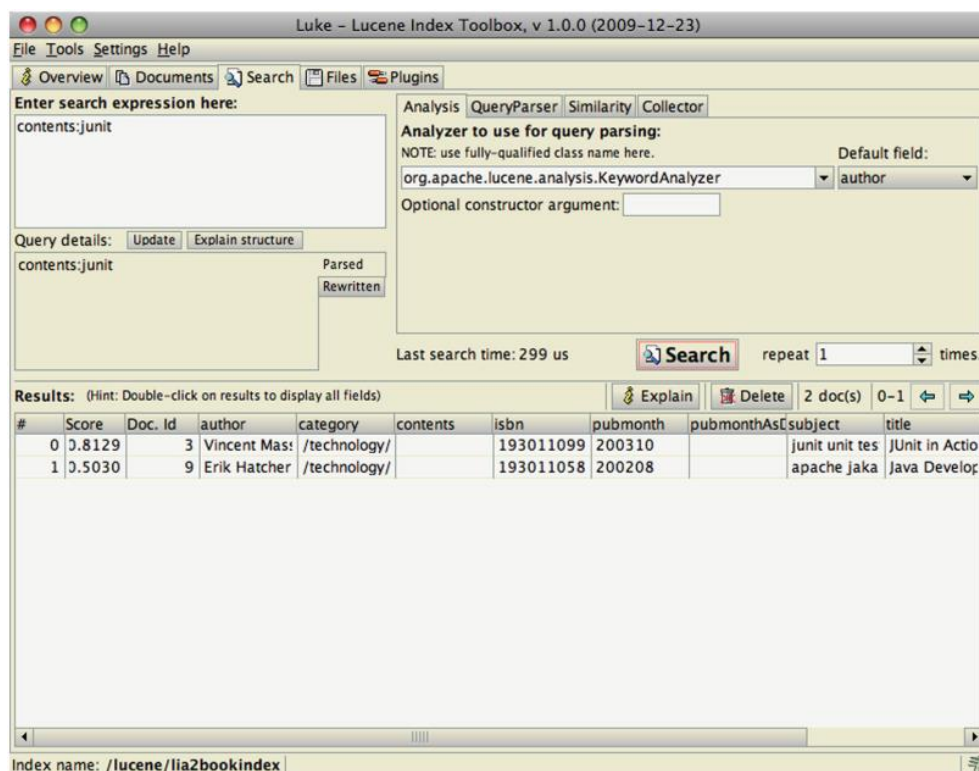
بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی □ ۱۴



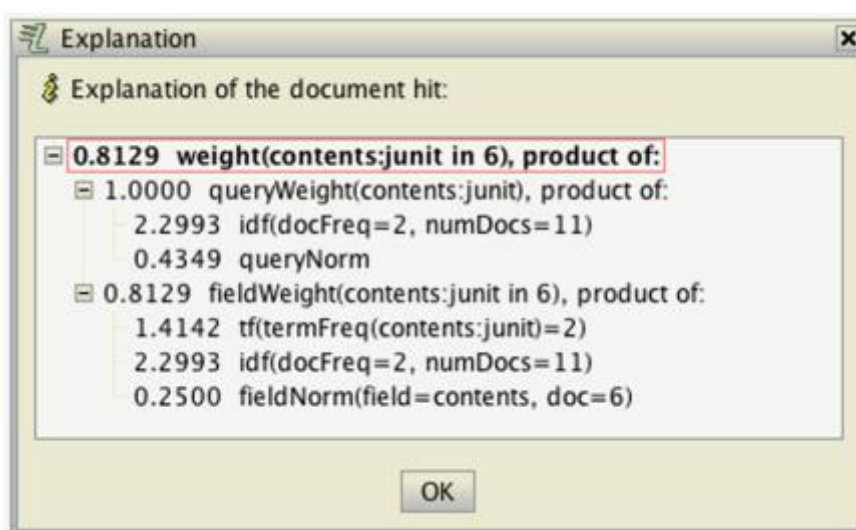
شکل ۲-۴- منوی اسناد در Luke (McCandless, Hatcher, and Gospodnetic 2010)

در منوی جستجو، کاربر می‌توان عبارت پرس‌وجوی خود را وارد کرده و از منوهای Analysis و QueryParse روش تجزیه را انتخاب کرده و جستجو را انجام دهد. با انتخاب سند بازیابی شده کاربر مجدداً به منوی اسناد منتقل شده و اطلاعات جزئی‌تر سند به وی نمایش داده می‌شود (شکل ۲-۵). همچنین با انتخاب گزینه «Explaine» در بخش میانی این صفحه نحوه محاسبه وزن اسناد بازیابی شده برای کاربر نمایش داده می‌شود (شکل ۲-۶) (McCandless, Hatcher, and Gospodnetic 2010).

بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی □ ۱۵



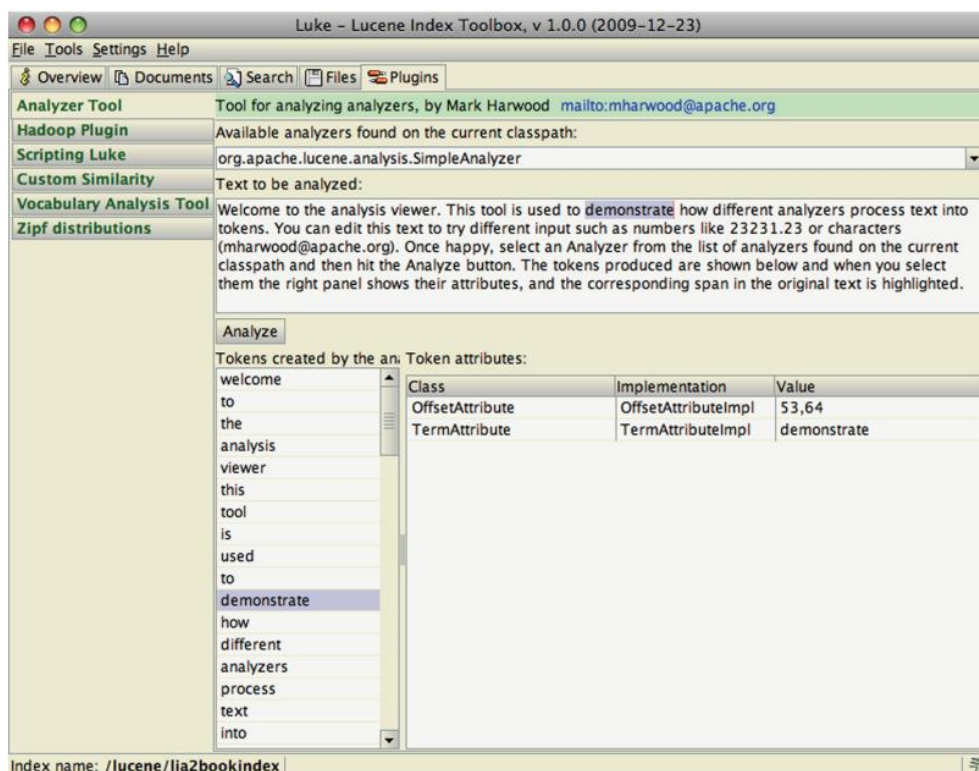
شکل ۲-۵- منوی جستجو در Luke (McCandless, Hatcher, and Gospodnetic 2010)



شکل ۲-۶- توضیحات مربوط به نحوه محاسبه وزن در Luke (McCandless, Hatcher, and Gospodnetic 2010)

در منوی موارد افزوده، توضیحات مربوط به ابزارهای توسعه داده شده برای luke قرار داده شده است و کاربر می‌تواند به صورت آزمایشی آنها را برای متن محدود بکار گیرد (.). همچنین می‌توان ابزارهای دیگری را نیز توسعه داد و به Luke افزود (برای مطالعه بیشتر در مورد این ابزار به (McCandless, Hatcher, and Gospodnetic 2010) مراجعه شود).

بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی □ ۱۶



شکل ۲-۷- منوی موارد افزوده در Luke (McCandless, Hatcher, and Gospodnetic 2010)

یکی دیگر از موارد مهم در تنظیمات لوسن، تعیین فضای مورد نیاز برای ذخیره‌سازی نمایه‌های ایجاد شده است. این فضا برای یک سند به عوامل مختلفی بستگی دارد؛ در بهترین حالات حدوداً یک سوم متن اصلی و در بدترین حالت (در نظرگیری بردار واژگان، اسناد حذف شده و مرورگرهای باز برای بررسی نمایه‌ها) حدوداً ده برابر متن اصلی است (به عنوان مثال برای نمایه‌سازی اسناد ویکی‌پدیا در نهایت ۱۴,۲ گیگابایت فضا نیاز است اما در طی پردازش فضای مورد استفاده لوسن به ۳۲,۴ گیگابایت هم می‌رسد). از معادله زیر برای تخمین فضای مورد نیاز استفاده می‌شود (McCandless, Hatcher, and Gospodnetic 2010):

$$\frac{1}{3} \text{ indexed} + \text{stored} + 2 \text{ term vectors} \quad (2-1)$$

به عنوان مثال اگر نیاز باشد علاوه بر نمایه و ذخیره شدن یک فیلد برای تمامی اسناد پایگاه داده، بردار واژگان آن نیز محاسبه شود، فضای مورد انتظار برابر خواهد بود با $\frac{10}{3}$ فضای مورد نیاز برای ذخیره‌سازی تمامی اسناد. این فضا را می‌توان با غیرفعال کردن برخی از ویژگی‌ها نظیر نگهداری (TF) فیلدهای غیرضروری، اطلاعات مکانی و offsetها در هنگام نمایه‌سازی بردارهای

واژگان و نمایه‌سازی و ذخیره‌سازی فیلدهای کمتر برای هر سند مدیریت نمود (McCandless, Hatcher, and Gospodnetic 2010).

۲-۴ موتور جستجوی سولر

موتور جستجوی سولر وظیفه جستجو براساس عبارت پرس و جوی کاربر و نمایش نتایج بازیابی شده را برعهده دارد. قابلیت‌های متعددی برای سولر معرفی شده است که می‌توان آنها را در قالب موارد زیر دسته‌بندی نمود (Grainger and Potter 2014):

- تجربه کاربر: برای فراهم آوردن تجربه مناسب سولر از سرویس^۱ REST برای استاندارد های XML، JSON و HTTP استفاده می‌کند تا کاربران قادر باشند از طریق زبان‌ها و وب‌سرویس‌های مختلف با آن ارتباط برقرار کنند. همچنین وجود قابلیت‌هایی چون مرتب‌سازی و صفحه‌بندی^۲ نتایج، جستجوی چندوجهی^۳، پیشنهاد خودکار^۴، املاي کلمات^۵، پررنگ کردن کلمات پرس و جو^۶ و جستجوی جغرافیایی از دیگر امکاناتی است که با پیاده‌سازی آنها سولر تلاش دارد تا تجربه مناسبی را برای کاربر فراهم آورد.
- مدلسازی داده: برای مدلسازی اسناد موجود در پایگاه داده امکاناتی در سولر طراحی شده است: تخصیص ویژگی‌هایی خاص به همه اسناد موجود در پایگاه داده، روش‌های پردازش نظیر عبارت‌های منطقی AND، OR و NOT، جستجوی Wildcard، جستجوی فازی و ...، خوشه‌بندی اسناد، پشتیبانی از فرمت‌های PDF و Word، انتقال داده‌ها از پایگاه‌های داده رابطه‌ای و پشتیبانی چندزبانه.

^۱ Representational State Transfer (REST)

^۲ Sorting and pagination

^۳ Facet Searching

^۴ Autosuggest

^۵ Spell checking

^۶ Hit highlighting

همچنین سولر به گونه‌ای طراحی شده است که قادر است حجم عظیمی از داده‌های متنی را با سرعت به صورت نزدیک به زمان واقعی^۱ پردازش کند. به بیانی دیگر، اسناد با فاصله زمانی اندکی از اضافه شدن به پایگاه داده قابلیت جستجو را دارند. با این حال (Grainger and Potter 2014):

- در سولر از رویکردهای بهینه‌سازی موتور جستجو استفاده نشده است؛
- به علت استفاده از طراحی^۲ NoSQL و ساختار غیر رابطه سولر، ارتباطی بین اسناد وجود ندارد و در صورتی که داده‌ای بین چند سند و یا چند فیلد اطلاعاتی مشترک باشد، این داده‌ها باید برای هر سند تکرار شوند؛
- سولر زمانی عملکرد خوبی دارد که داده‌ها غیر ساخت یافته باشند و در صورتی که داده‌ها ساخت یافته باشند بهتر است از پایگاه‌های داده رابطه‌ای استفاده شود؛
- در سولر می‌توان به راحتی سندی را اضافه و یا کم کرد اما نمی‌توان به سادگی اطلاعات مربوط به یک سند را بروز نمود و باید از ابتدا فرآیند نمایه‌سازی انجام شود؛
- سولر زمانی عملکرد بهینه دارد که عبارت پرس‌وجو از واژگان کمی تشکیل شده باشد.
- نسخه‌های فعلی سولر قابلیت مقیاس‌پذیری الاستیک را ندارد (در صورت افزایش بار محاسباتی سرورهای جدید به صورت کاملاً خود کار فعال نمی‌شوند). طراحی سولر Cloud با کمک آپاچی Zookeeper گامی در راستای بهبود این شرایط است.
- سولر به صورت پیش فرض از روش‌های پایه‌ای در بازیابی اطلاعات نظیر روش شاخص معکوس، بولین، مدل فضای برداری و TF-IDF استفاده می‌کند. در حالی که در طی سال‌های گذشته توسعه‌ها و اصلاحات بسیاری برای آنها ارائه شده است.

علی‌رغم موارد ذکر شده، متن باز بودن لوسن و سولر امکان تغییر مازول‌ها و روش‌های بکاررفته در آن را برای طراحان و پیاده‌سازان موتورهای جستجو فراهم آورده و این امر را می‌توان یکی از دلایل عمده شهرت و کاربردهای متعدد آنها دانست. با این وجود و تاکنون هیچ گونه پژوهشی در زمینه شناسایی و واکاوی این کتابخانه و موتور جستجو در پژوهشگاه انجام نشده است. به جهت تمرکز این طرح بر موتور جستجوی سولر، فصل سوم به آن تخصیص یافته است.

^۱ Near real-time (NRT)

^۲ Not Only SQL (NoSQL)

بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی □ ۱۹

فصل سوم

مباحث پایه‌ای در موتور جستجوی

سولر

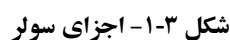
۳-۱ مقدمه

سولر موتور جستجوی متن‌باز و مقیاس‌پذیر بر مبنای تکنولوژی NoSQL^۱ است که برای جستجو در داده‌های متنی حجیم طراحی شده و نتایج بازیابی شده را براساس میزان ارتباط^۱ با عبارت پرس‌وجو نمایش می‌دهد. داده‌ها در این موتور جستجو ساختاری تخت^۲ داشته (داده‌ها به یکدیگر وابستگی نداشته و سلسله‌مراتبی نیستند) و نرخ بروزرسانی کمی دارند. همچنین از آنجا که سولر به صورت NoSQL طراحی شده است ضرورتی ندارد تمامی داده‌ها دارای ساختاری یکسان باشند. به بیانی دیگر، سولر برای مدیریت داده‌هایی بهینه شده است که مبتنی بر متن (Text-centric) بوده و ساختاری تخت (Document-Oriented) و منعطف (Flexible-schema) داشته و نرخ بروزرسانی در داده‌ها کمتر از نرخ خواندن آن‌ها (Read-dominant) باشد (Grainger and Potter ۲۰۱۴)^۳.

^۱ Relevancy

^۲ Flat

^۳ تمامی مطالب ارائه شده در این فصل از کتاب (Grainger and Potter ۲۰۱۴) استخراج شده‌اند که منبعی اصلی برای درک عملکرد سولر و ماژول‌های اصلی آن است (<http://lucene.apache.org/solr/resources.html>). لذا به جهت اختصار در ادامه بدان ارجاع داده نخواهد شد.



کاربران سولر باید قادر باشند از طریق نرم افزارها و زبان های مختلف با آن ارتباط برقرار کنند، بدین منظور وب سرویس REST برای استانداردهای XML، JSON و HTTP طراحی شده است. همچنین سولر می تواند هسته های پردازشی متعددی داشته باشد که امکان پردازش موازی داده ها و پشتیبانی از به اشتراک گذاری منابع^۲ را برای رایانش ابری فراهم می آورد.

' Java servlet

⁂ Multitenancy

زیر سیستم‌های (۱) مدیریت اسناد، (۲) تحلیل متن و (۳) پردازش پرس‌وجو سه زیر سیستم سولر بوده و هریک به صورت فرآیندهایی ماژولار طراحی شده‌اند. طراحی ماژولار زیرسیستم‌ها، آنها را قادر می‌سازد تا به صورت موازی فعالیت کرده و قابلیت توسعه و یا شخصی‌سازی را داشته باشند.

قابلیت مقیاس‌پذیری سولر توسط دو ویژگی تعداد پرس‌وجوهای پردازش شده در ثانیه و تعداد اسناد نمایه شده مشخص می‌شود. تعداد پرس‌وجوهای پردازش شده با قابلیت پردازشی سرورها مرتبط است و هرچه تعداد سرورها و قابلیت‌های پردازشی آنها بیشتر باشد، تعداد پرس‌وجوهای پردازش شده در ثانیه افزایش خواهد یافت. همچنین اگر حجم اسناد و یا تعداد آنها زیاد باشد، کارایی بازیابی اطلاعات کاهش خواهد یافت. بدین منظور ماتریس اصلی نمایه‌ها به بخش‌های کوچک‌تری بنام Shard تقسیم و شکسته شده و هر بخش به صورت موازی و جداگانه پردازش می‌شوند. در انتها نیز نتایج این بخش‌ها ترکیب شده و نتیجه نهایی برای کاربر نمایش داده می‌شود. پیش از پرداخت عمیق‌تر به این زیر سیستم‌ها ضروری است تا برخی از مفاهیم کلیدی بیان شده و درباره آنها توضیحاتی ارائه شود.

۳-۲ مفاهیم کلیدی در سولر

۳-۲-۱- تنظیمات در سولر

انجام تنظیمات در سولر در دو سند اصلی Solrconfig و Schema انجام می‌شود. تنظیمات اصلی سامانه‌ای که بر سولر ساخته شده است در Solrconfig انجام شده و Schema تعریف‌کننده ساختار نمایه‌سازی و نوع-فیلد^۱ و فیلدهای مورد استفاده است. به خصوصیات این دو سند در ادامه پرداخته می‌شود.

SOLRCONFIG

در زمان اجرا، سولر فایل اصلی تنظیمات را با نام SolrConfig در مسیر اصلی نصب خود فراخوانی می‌کند. چهار بخش اصلی این سند شامل تنظیمات مربوط به محل ذخیره‌سازی نمایه‌ها (Data directory location)، تنظیمات حافظه (Cache parameters)، مدیریت درخواست‌ها

^۱ Field type

(Request handlers) و مولفه‌های جستجو (Search components) است. نمونه‌ای از این سند در شکل ۲-۳ آورده شده است.



شکل ۲-۳- نمونه‌ای از سند Solrconfig

توضیحات مختصری از چهار بخش فوق به شرح زیر است (این بخش‌ها به صورت مفصل‌تر در زیرسیستم‌های مرتبط بررسی خواهند شد):

- محل ذخیره سازی نمایه‌ها: نمایه سازی یکی از بخش‌های اصلی سولر و لوسن است و با کمک این بخش محل ذخیره سازی نمایه‌های ساخته شده توسط «زیرسیستم تحلیل متن» و مسیر کتابخانه‌های مورد استفاده در «زیرسیستم مدیریت اسناد» و «زیرسیستم تحلیل متن» تعیین می‌شود.

□ مدیریت درخواست‌ها: در این بخش درخواست‌های دریافتی از کلاینت در طی فرآیند مدیریت و پردازش جستجو تحلیل شده و نتایج را برای نمایش در کلاینت تنظیم می‌شود. جزئیات مربوط به این بخش در «زیرسیستم پردازش پرس‌وجو» بیان خواهد شد.

□ تنظیمات حافظه cache: یکی از راهکارهای سولر برای بهبود کارایی پردازش پرس‌وجوها استفاده از Cache است. برای استفاده بهینه از این حافظه و مدیریت آن، توجه به موضوعات زیر ضروری است:

○ اندازه Cache: تعیین حداکثر تعداد موارد ذخیره شده در Cache سیستم موجب خواهد شد تا تعادلی بین حافظه اصلی و نهان ایجاد شده و سولر بتواند فعالیت اصلی خود را با کارایی بالاتری انجام دهد. زمانی که اندازه Cache به حد آستانه می‌رسد باید برخی از اسناد موجود پاک شود تا اسناد جدید بتوانند در آن ذخیره شوند. برای پاک کردن دو سیاست وجود دارد: سیاست Least Recently Used (LRU) (روش پیش‌فرض) که اسنادی قدیمی را پاک می‌کند (قدمت یک سند بر اساس آخرین زمانی که مورد استفاده قرار گرفته مشخص می‌شود)؛ و سیاست Least Frequently Used (LFU) که اسناد با کمترین استفاده (غیر جذاب‌ترین اسناد) را از Cache پاک می‌کند. از آنجا که این سیاست نیازمند نگهداری تعداد دفعات استفاده از اسناد است، حافظه و حجم اطلاعات بیشتری را به خود تخصیص می‌دهد.

○ سود حاصل از نگهداری و تعیین تعداد اسناد پاک شده^۱: محاسبه سود حاصل شده از نگهداری و تعیین تعداد اسناد پاک شده سبب خواهد شد تا اندازه بهینه‌ای برای Cache تعیین گردد.

○ ناکارایی اسناد در Cache: تعیین مصادیق ناکارایی اسناد موجود در Cache ارتباط مستقیمی با سیاست پاک کردن این حافظه و تعریف اسناد ناکارآمد در بافت سیستم دارد.

○ آماده‌سازی خودکار Cache جدید: همانطور که عنوان شد، سولر جستجوگری را به صورت فعال و جستجوگرهای دیگری را در حال آماده‌باش دارد که به محض بسته شدن جستجوگر فعلی، جستجوگر آماده شده بعدی بارگذاری

^۱ Hit ratio and eviction

خواهد شد. برخی از اسناد جستجوگر فعلی می توانند در Cache جستجوگر آماده شده بعدی ذخیره شوند. به این فرآیند آماده سازی خود کار Cache جدید گفته می شود. حداکثر تعداد اسناد جستجوگر فعلی برای ذخیره سازی و یا تعیین درصد اشغال Cache از پارامترهای مهم در مدیریت آن است.

□ مولفه های جستجو: پرس و جوها در سولر توسط مولفه ای به نام جستجوگر^۱ پردازش می شود. در هر لحظه از زمان تنها یک جستجوگر فعال وجود داشته و تمامی درخواستها توسط این جستجوگر پردازش می شوند. جستجوگر د ستر سی read-only به نمایه های لوسن داشته و اگر در زمان فعال بودن جستجوگر سند جدیدی اضافه شود، اطلاعات آن در جستجو دخالت داده نخواهد شد. برای این منظور جستجوگر باید بروزرسانی شود (جستجوگر فعلی بسته شده و جستجوگر جدیدی فعال شود). دلیل این امر آن است که اگر در لحظه اضافه شدن سند جدید، جستجوگر بروزرسانی شود، صفحه نتایج کاربر تغییر کرده و این اتفاق خوبی برای کاربر نیست (تجربه نامناسب کاربر). همچنین اگر به ازای اضافه شدن هر سند جدید، جستجوگر جدیدی از ابتدا بازخوانی و اجرا شود، بازیابی و نمایش نتایج کند خواهد شد (کاهش سرعت).

برای رفع این مشکل، سولر در پشت صحنه یک یا چند جستجوگر را در حال آماده شدن (warming) نگه داشته و پارامترهای جستجو در آن را براساس مقادیر پر کاربرد در سیستم بازیابی اطلاعات تنظیم می کند. جستجوگر فعلی، آنقدر فعال می ماند تا جستجوگر پشت صحنه کاملاً آماده شود و به محض آماده شدن، جستجوگر فعلی بسته شده و جستجوگر جدید جایگزین خواهد شد. تعداد این جستجوگرها توسط پارامتر maxWarmingSearchers با مقدار پیش فرض ۲ در Schema تنظیم می شود.

همچنین اگر درخواست جستجویی به سیستم داده شود و هیچ جستجوگر فعالی در سیستم نباشد شود با کمک دستور useColdSearcher می توان به سولر دستور داد که جستجو را آغاز کند چه جستجوگر به طور کامل آماده باشد چه نه (مقدار True) و یا آنکه صبر کند تا جستجوگر به طور کامل آماده شود (مقدار False).

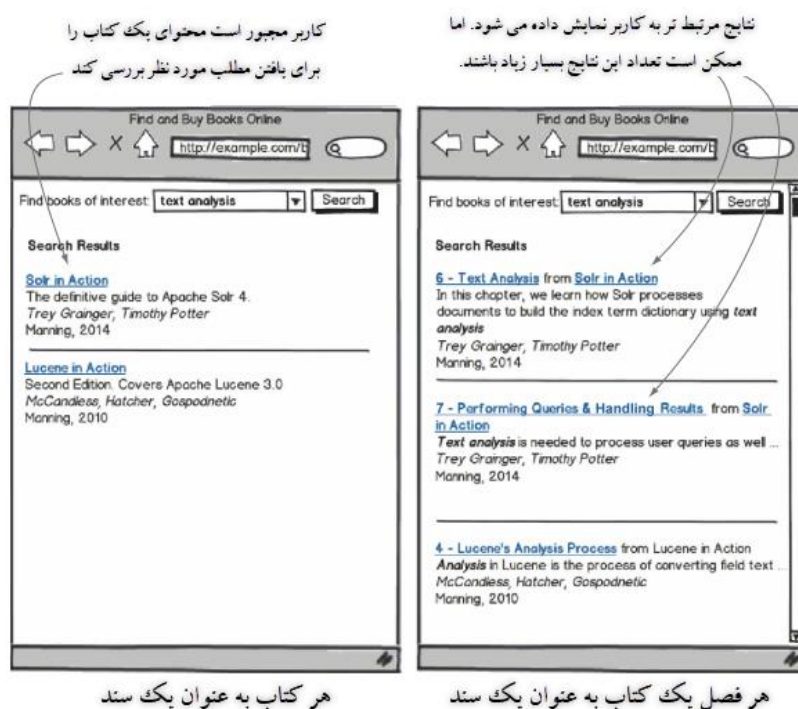
^۱ Searcher

SCHEMA

همانگونه که ذکر شد، ساختار نمایه‌سازی و نوع-فیلد و فیلدهای مورد استفاده در سندی به نام Schema انجام و ذخیره می‌شود. برای تنظیم Schema پاسخ به سوالات زیر می‌تواند کمک کننده باشد:

□ معنا و مفهوم سند در سیستم طراحی شده چیست؟

پاسخ به این سوال سطح ریزدانگی بررسی و پردازش اسناد را تعیین می‌کند. به عنوان مثال، سند می‌تواند به معنای یک کتاب کامل باشد و یا هر فصل از یک کتاب سندی جداگانه محسوب شود. در حالت نخست، کاربر باید خود به دنبال فصل/مطلب مورد نظر در کل یک کتاب (سند بازیابی شده) باشد اما در حالت دوم می‌توان فصل مرتبط با عبارت جستجوی کاربر را برای وی نمایش داد (که البته در مواردی سبب سردرگمی کاربر نیز خواهد شد چرا که ممکن است نتایج بسیار زیادی صرفاً از یک کتاب برای او نمایش داده شود). شکل ۳-۳ نمونه‌ای از جستجوی عبارت «text analysis» را برای سطح ریزدانگی کتاب و فصل‌های هر کتاب به صورت جداگانه نمایش می‌دهد.



شکل ۳-۳- نتایج جستجوی عبارت «text analysis» برای سطح ریزدانگی هر کتاب و هر فصل هر کتاب

□ مشخصه کلیدی^۱ هر سند چیست؟

به عنوان نمونه برای یک کتاب مشخصه کلیدی آن می تواند ISSN کتاب باشد. دلیل اهمیت تعیین این مشخصه آن است که اگر دو سند دارای مشخصه یکسان باشند، سولر آخرین سند اضافه شده را نگه داشته و بقیه اسناد را حذف خواهد کرد. همچنین اگر معماری سولر به گونه ای باشد که در آن از سرورهای متعدد برای پردازش استفاده شود، از مشخصه کلیدی هر سند برای تبادل اطلاعات و ادغام نتایج پردازش ها استفاده می شود.

□ فیلدهایی که نمایه می شوند کدامند؟

سولر جستجو را بر روی فیلدهای نمایه شده^۲ انجام می دهد. بنابراین تعیین دقیق این فیلدها بر اندازه شاخص معکوس، دقت، بازخوانی و سرعت جستجو اثر بسیار زیادی خواهد داشت.

□ چه فیلدهایی باید برای کاربر نمایش داده شوند؟

فیلدهایی که قابلیت نمایش برای کاربر را دارند در فیلدهای ذخیره شده^۳ تعیین می شوند. تعیین دقیق این فیلدها نیز تاثیر بسیاری بر سرعت و کفایت نمایش نتایج برای کاربر خواهد داشت.

یکی از اقدامات مهم دیگر در تنظیم Schema تعیین فیلدها ست. برخی از انواع این فیلدها در ادامه توضیح داده می شوند.

□ فیلدهای ضروری (Required fields): تعیین کننده ضروری بودن آن فیلد است و مقدار این پارامتر برای فیلد مشخصه کلیدی باید true باشد.

□ فیلدهای چند مقداره (Multivalued fields): یک فیلد می تواند شامل یک یا چند مقدار باشد و سولر هر یک از این مقادیر را به صورت جداگانه ذخیره می کند.

□ فیلدهای دینامیک (Dynamic fields): در مواردی که تعدادی از فیلدها مشخصات یکسانی دارند، می توان به جای تکرار هر مشخصه برای تک تک فیلدها، یک مشخصه کلان برای تمامی فیلدهای همسان تعریف کرد. به عبارت بهتر، اگر:

```
<dynamicFieldName="*_s" type="string" indexed="true" stored="true" />
```

^۱ Unique key

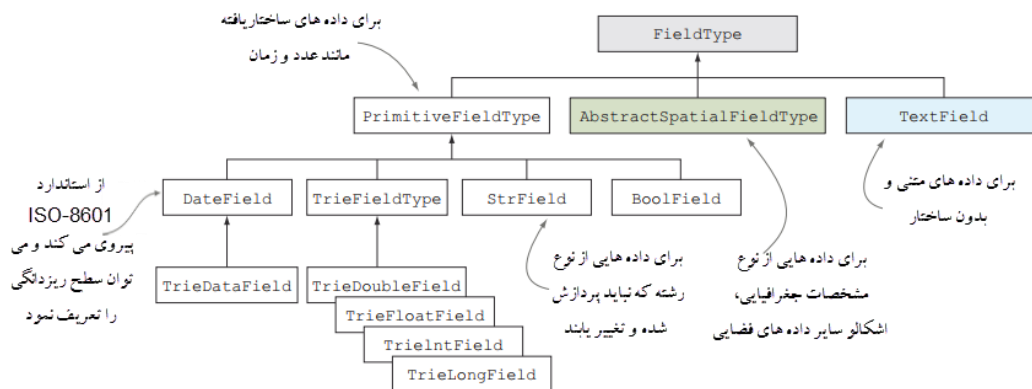
^۲ Indexed fields

^۳ Stored fields

بدین معناست که هر فیلدی که در انتهای نام آن "s_" وجود داشته باشد از نوع رشته بوده، نمایه شده و در صفحه نتایج برای کاربر نمایش داده می‌شود.

□ فیلدهای کپی (Copy fields): در اغلب پایگاه‌های داده، اسناد در قالب فیلدهای مختلف ذخیره می‌شوند (فیلد نام نویسنده، فیلد عنوان، فیلد کلمات کلیدی و ...). در مواردی که کاربر بخواهد عبارت پرس و جوی خود را جستجو کند بدون آنکه فیلد جستجو را تعیین کرده باشد، سیستم بازیابی اطلاعات باید نتایج را مستقل از آنکه این عبارت در کجای سند قرار دارد نمایش دهد. اما اتفاقی که رخ می‌دهد آن است که موتور جستجو صرفاً در فیلدی که به عنوان فیلد پیش فرض برای جستجو طراحی شده است به دنبال نتایج مرتبط خواهد بود. برای رفع این مشکل دو راهکار وجود دارد: (۱) استفاده از فیلدهایی از نوع کپی که در واقع فیلدی سایه‌ای بوده و به صورت غیر رسمی تمامی فیلدهای یک سند در آن کپی شده‌اند و (۲) استفاده از تجزیه گرهای غیر از تجزیه گر پیش فرض لوسن (این موارد در بخش ۱-۵-۳- مولفه Query بررسی خواهند شد).

نمودار کلاس انواع فیلدها در Schema در شکل ۳-۴ نمایش داده شده است.



شکل ۳-۴- نمودار کلاس انواع فیلدها در Schema

۲-۲-۳- سند

سولر موتور جستجوی اسناد است و هر سند می‌تواند دارای یک یا چند فیلد اطلاعاتی از نوع رشته^۲، بولی^۳، زمان، مکان و غیره باشد. خصوصیات هریک از این فیلدها در سندی به نام Schema

^۱ Parser

^۲ String

^۳ Boolean

تعریف می‌شود و با کمک این سند سولر قادر است اسناد ورودی را مدیریت و پردازش نماید. ذکر مجدد این نکته نیز ضروری است که از آنجا که سولر به صورت NoSQL طراحی شده است ضرورتی ندارد همه فیلدها برای همه اسناد وجود داشته و یا دارای مقدار باشند.

همچنین در صورتی که سندی دارای فیلدی باشد که با نامی دیگر در Schema تعریف شده، سولر این قابلیت را دارد که به صورت خودکار نوع فیلد را حدس زده و آن را با ملاحظات به Schema اضافه کند.

۳-۲-۳ تطابق فازی

تطابق فازی^۱ به معنای توانایی یافتن عبارت‌ها و یا واژگان ناهمسان با عبارت پرس‌وجوی کاربر است. سولر این توانایی را با کمک چندین روش عملیاتی ساخته است:

□ جستجوی Wildcard: زمانی که کاربر املای درست کلمه پرس‌وجو را نمی‌داند و یا می‌خواهد تمام حالت‌های ممکن از کلمه پرس‌وجو را در پایگاه داده جستجو کند. مثال‌هایی از حالت‌های مختلف این جستجو در جدول ۳-۱ نمایش داده شده‌اند.

جدول ۳-۱- جستجوی Wildcard: عبارت پرس‌وجوی کاربر و عبارتی که سولر جستجو می‌کند

عبارتی که سولر جستجو می‌کند	عبارت پرس‌وجوی کاربر
office, officer, official, ...	offi*
offer, officer, officiador, ...	off*r
offer	off?r
“softwar* Eng?neering”	“softwar* Eng?neering”

وجود علامت * در کلمه پرس‌وجو به معنای یافتن و جستجوی همه کاراکترهای ممکن با معنا و وجود علامت ؟ به معنای یافتن و جستجوی تنها یک کاراکتر با معنا است (سولر برای شناسایی کلمات با معنا از نمایه ساخته شده برای پایگاه داده استفاده می‌کند). همچنین علامت “X” به معنای جستجوی دقیق عبارت X بوده و اعمال روش جستجوی Wildcard برای عبارت X ممکن نیست.

جستجوی اصلاح فاصله^۲: در جستجوی Wildcard حداکثر کاراکترهایی که باید یافته و جستجو شوند مشخص نمی‌شود. این امر می‌تواند باعث کاهش سرعت بازیابی اطلاعات

^۱ Fuzzy matching

^۲ Edit-distance search

شود. لذا در مواردی لازم است تعداد کاراکترها محدود شوند. در این حالت می توان از جستجوی اصلاح فاصله استفاده کرد. مقدار این فاصله در سولر بر اساس فاصله دامرا-لونا شتاین^۱ که تعمیم یافته فاصله لونا شتاین است (Damerau ۱۹۶۴) محاسبه می شود. البته توابع فاصله دیگری نیز تعریف شده اند که از جمله آنها می توان به موارد جدول ۷-۲ ضمايم اشاره کرد.

برای فراخوانی این جستجو از دستور $x \sim$ استفاده می شود: 1~ به معنای جستجو با یک فاصله اصلاح، 2~ جستجو با دو فاصله اصلاح (حالت پیش فرض) و $N \sim$ به معنای جستجو با N فاصله اصلاح است. به عبارت دیگر، اگر عبارت پرس و جوی کاربر به صورت $\sim administrator$ باشد، سولر از عبارت هایی مانند administrator، administrator، administrator نیز علاوه بر عبارت پرس و جوی کاربر استفاده می کند.

□ جستجوی نزدیکی^۲: حالت جستجوی اصلاح فاصله برای یافتن واژگان مشابه با واژه پرس و جوی کاربر بکار می رود، اما اگر کاربر بخواهد عبارات مشابه را بیابد (ترکیبی از واژگان)، جستجوی نزدیکی در سولر طراحی شده است. جدول ۳-۲ مثال هایی از این جستجو را نشان می دهد.

جدول ۳-۲- جستجوی نزدیکی: عبارت پرس و جوی کاربر و توضیحات آن

عبارت پرس و جوی کاربر	توضیح
"chief officer"~0	جستجوی دقیق (حداکثر فاصله دو واژه برابر ۰ است)
"chief officer"~1	chief و officer باید حداکثر فاصله ای برابر یک داشته باشند. مانند: "chief executive officer", "chief financial officer", ...
"chief officer"~2	chief و officer باید حداکثر فاصله ای برابر دو داشته باشند. مانند: "chief business development officer", "officer chief", ...
"chief officer"~N	chief و officer باید حداکثر فاصله ای برابر N داشته باشند.

□ جستجوی بازه ای^۳: جستجو در بازه های عددی و یا حرفی نیز در سولر میسر است. نمونه هایی از این جستجو در جدول ۳-۳ گزارش شده است.

^۱ Damerau Levenshtein

^۲ Proximity search

^۳ Range search

جدول ۳-۳- جستجوی بازه‌ای: عبارت پرس‌وجوی کاربر و عبارتی که سولر برای جستجو استفاده می‌کند

عبارتی که سولر جستجو می‌کند	عبارت پرس‌وجوی کاربر
Yearsold=18, 19, 20, 21	Yearsold: [18 To 21]
Yearsold=18, 19, 20	Yearsold: [18 To 21}
Yearsold= 19, 20	Yearsold: {18 To 21}
Title= boat, boil, book, boulder, ...	Title: [boat To boulder]
Price= 12.99, 13.00009, 14.01, ...	Price: [12.99 To 14.99]

۴-۲-۳- ارتباط

یافتن اسناد مرتبط با عبارت پرس‌وجوی کاربر یکی از اصلی‌ترین فعالیت‌ها در موتورهای جستجو است و تعیین میزان این ارتباط^۱ یکی دیگر از فعالیت‌ها مهم در این رابطه است. در سولر میزان ارتباط با کمک کلاس similarity که در Schema تعریف شده است، تعیین می‌شود و سولر به صورت پیش‌فرض از تابع DefaultSimilarity که متعلق به لوسن است استفاده می‌کند.

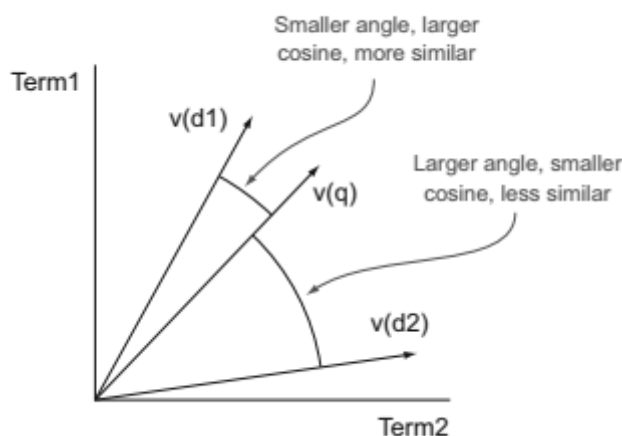
در این تابع، ابتدا با استفاده از منطق بولی اسناد مرتبط با عبارت پرس‌وجوی کاربر شناسایی شده و سپس از مدل فضای برداری^۲ و شباهت کسینوسی^۳ بین بردار سند و عبارت پرس‌وجو برای امتیازدهی^۴ (RS) و یافتن میزان ارتباط استفاده می‌شود. شکل ۳-۵- نمایی از شباهت کسینوسی را نمایش می‌دهد. در این شکل بردار عبارت پرس‌وجو $v(q)$ به بردار سند $v(d1)$ نزدیک‌تر از بردار سند $v(q2)$ است (این نزدیکی براساس زاویه بین دو بردار اندازه‌گیری می‌شود؛ هرچه زاویه کوچکتر باشد دو بردار به هم نزدیک‌تر بوده و یا به هم شباهت بیشتری دارند).

^۱ Relevancy

^۲ Vector space model (VSM)

^۳ cosine

^۴ Relevancy score



شکل ۳-۵- شباهت کسینوسی بین بردارهای واژگان

همچنین در این تابع برای محاسبه میزان ارتباط بین هر سند و عبارت پرس وجو از رابطه (3-1) استفاده می شود:

$$Score(q, d) = \sum_{t \text{ ind}} (tf(t \text{ ind}) \times idf(t)^2 \times t.getboost() \times norm(t, d)) \times coord(q, d) \times queryNorm(q) \quad (3-1)$$

در این رابطه t نشان دهنده واژه، q نشان دهنده عبارت پرس وجو و d نشان دهنده سند است. $tf(t \text{ ind})$ بیانگر مجذور میزان تکرار واژه t در سند d و $idf(t)$ معیاری برای تعیین تکرار واژه t در کل اسناد بوده و از طریق رابطه (3-2) محاسبه می شود.

$$idf(t) = 1 + \log\left(\frac{\text{numDocs}}{\text{docFreq} + 1}\right) \quad (3-2)$$

در رابطه (3-1)، $t.getboost()$ نشانگر وزنی^۱ است که برای فیلدهای مختلف (نظیر عنوان سند، چکیده، نام نویسنده و ...) در هنگام جستجو در نظر گرفته می شود. به تعبیری دیگر، اگر بدانیم فیلد خاصی در بافت جستجو اهمیت بیشتری نسبت به بقیه فیلدها دارند می توان در زمان پردازش و بازیابی اطلاعات وزن بیشتری را برای آن در نظر گرفت. به عنوان نمونه:

^۱ Boost

جدول ۳-۴- درنظرگیری وزن برای فیلدهای مختلف

اهمیت	عبارت پرس وجو	عبارتی که سولر جستجو می کند
عنوان مهم تر از توضیحات	Title: solr in action Description: solr in action	title:(solr in action)^2 & description:(solr in action)
جریمه توضیحات	Title: solr in action Description: solr in action	title:(solr in action) & description:(solr in action)^0.2
هر واژه دارای اهمیت خاص		title:(solr^2 in^0.1 action^1.5)^3

همچنین در رابطه (3-1)، $norm(t, d)$ وظیفه نرمال سازی فیلد^۱ را برعهده داشته و براساس اهمیت هر سند $d.getboost()$ ، اهمیت هر فیلد سند $f.getboost()$ و طول سند $lengthNorm()$ (رابطه (3-3)) محاسبه می شود. اگر فیلدهای یک سند دارای وزنهای مختلف باشند، $f.getboost()$ از حاصلضرب این وزن ها محاسبه می شود.

$$norm(t, d) = d.getboost() \times lengthNorm() \times f.getboost() \quad (3-3)$$

عامل $coord(q, d)$ در رابطه (3-1) نشاندهنده میزان ارتباط هر سند با عبارت پرس وجو است؛ هرچه سند دارای تعداد تکرارهای بیشتری از عبارت پرس وجو باشد، میزان ارتباط بیشتری با عبارت پرس وجو خواهد داشت. این عامل با کمک رابطه (3-4) محاسبه می شود.

$$coord(q, d) = \frac{\text{num Terms in Document from Query}}{\text{num Terms in Query}} \quad (3-4)$$

عامل $queryNorm(q)$ در در رابطه (3-1) نیز این قابلیت را به سولر می دهد تا امتیازهای عبارت های پرس وجو قابل مقایسه باشند. رابطه (3-5) برای محاسبه این عامل استفاده می شود.

$$queryNorm(q) = \frac{1}{\left(q.getboost()^2 \times \sum_{t \text{ in } q} (idf(t) \times t.getboost())^2 \right)^{1/2}} \quad (3-5)$$

از عامل های دیگری نیز می توان برای محاسبه ارتباط بین عبارت پرس وجوی کاربر و اسناد استفاده کرد. این عوامل در جدول ۷-۴ در بخش ضmann ارائه شده اند. همچنین کلاس های محاسبه ارتباط

^۱ Field normalization

دیگری نیز در سولر تعریف شده‌اند که می‌توان از آنها به جای DefaultSimilarity استفاده کرد. به این موارد در جدول ۷-۵ بخش ضمائم اشاره شده است.

۳-۲-۵- معیارهای سنجش بازیابی اطلاعات

معیارهای دقت^۱ و بازخوانی^۲ دو معیار مورد استفاده برای سنجش کارایی بازیابی اطلاعات در سولر هستند. دقت نتایج بازیابی شده توسط معیار دقت سنجیده می‌شود. این معیار بدنبال پاسخ دادن به پرسش «آیا اسناد بازیابی شده همان‌هایی هستند که کاربر به دنبال آن بوده؟» است. به بیانی دیگر، precision معیاری برای سنجش میزان ارتباط نتایج با عبارت پرس‌وجو است و توجهی به میزان دربرگیری کل اسناد مرتبط در پایگاه داده ندارد (not the comprehensiveness of the results). رابطه (3-6) نحوه محاسبه این معیار را نشان می‌دهد.

$$precision = \frac{\#Correct\ Matches}{\#Total\ Results\ Returned} = \frac{TP}{TP+FP} \quad (3-6)$$

معیار بازخوانی جامعیت نتایج بازیابی شده را سنجیده و بدنبال پاسخ دادن به پرسش «چه تعداد سند درست بازیابی شده؟» است. رابطه (3-7) نحوه محاسبه این معیار را نشان می‌دهد.

$$recall = \frac{\#Correct\ Matches}{\#Correct\ Matches + \#Missed\ Matches} = \frac{TP}{TP+FN} \quad (3-7)$$

هدف در سیستم‌های بازیابی اطلاعات بیشینه کردن دقت و بازخوانی است اما تمرکز اصلی در سولر بر بهبود بازخوانی است و سپس به دقت توجه می‌شود. به عبارت بهتر، ابتدا تمامی نتایج درست بازیابی شده (بازخوانی) و سپس نتایج به گونه‌ای مرتب می‌شوند که در صفحات نخست مرتبط‌ترین نتایج با عبارت پرس‌وجوی کاربر (دقت) نمایش داده شوند.

اما این دو معیار نمی‌توانند به صورت کامل کارایی سیستم‌های بازیابی اطلاعات را بسنجند؛ به عنوان نمونه کاربر ممکن است بخواهد نتایج نزدیک به محل سکونت خود را بیابد و یا پراکندگی نتایج بازیابی شده و نمایش تمامی حالت‌های ممکن در صفحه نخست مطلوب او باشد. به عبارت دیگر، به عنوان مثال اگر کاربر کلمه hamburger را جستجو کند آیا برای او مطلوب است صفحه نخست نمایش نتایج دربرگیرنده تمامی شعب Burger King باشد و یا اینکه تمامی رستوران‌های

^۱ Precision

^۲ Recall

دارای hamburger برای او در صفحه نخست نمایش داده شود. بنابراین محاسبه کارایی سیستم با کمک سنجش میزان ارتباط آن با کلیدواژگان^۱ کاربر تنها یکی از معیارها است و درک مفهوم ارتباط در بافت مورد بررسی و نظر کاربر نیز در این رابطه اهمیت بسیار دارد.

برخی مفاهیم کلیدی دیگر نیز در سولر وجود دارد که در ادامه صرفاً به آنها اشاره می‌شود:

□ استفاده از شاخص معکوس^۲ برای جستجو (در این ماتریس هر واژه به یک یا چند سند نگاشته شده و به صورت صعودی به ترتیب الفبایی مرتب شده است).

□ پشتیبانی از منطق بولی در پردازش عبارت‌های پرس‌وجو:

○ AND یا + در صورتی که کاربر بخواهد واژه‌ای حتماً در نتایج جستجو باشد؛

○ OR در صورتی که ضرورتی در وجود واژه در نتایج جستجو نباشد؛

○ NOT یا - در صورتی که کاربر بخواهد واژه‌ای حتماً در نتایج جستجو نباشد؛

○ "... در صورتی که کاربر به دنبال جستجوی عبارتی خاص باشد.

□ Deformalized بودن اسناد: از آنجا که سولر دارای طراحی NoSQL است، خصوصیات

اسناد به صورت یکسان و واحد تعریف نمی‌شود. این امر قابلیت مقیاس‌پذیر بودن سولر را به دنبال دارد.

□ جستجوی توزیع شده^۳: از آنجا که سولر دارای طراحی توزیع شده است، عبارت پرس‌وجو

به هسته‌های پردازشی مختلفی که هریک حاوی تعدادی از اسناد پایگاه داده هستند فرستاده شده، نتایج بازیابی و ادغام شده و براساس میزان ارتباط نهایی برای کاربر نمایش داده می‌شوند.

همانگونه که ذکر شد، سولر برای به انجام رساندن جستجو از سه زیرسیستم (۱) مدیریت اسناد، (۲) تحلیل متن و (۳) پردازش پرس‌وجو استفاده می‌کند. این زیرسیستم‌ها در ادامه بررسی خواهند شد.

۳-۳ زیرسیستم مدیریت اسناد

مدیریت اسناد در سولر را می‌توان به صورت فرآیند نمایه‌سازی و ذخیره‌سازی اسناد با گام‌های زیر تعریف کرد (شکل ۳-۶):

^۱ Keyword relevancy

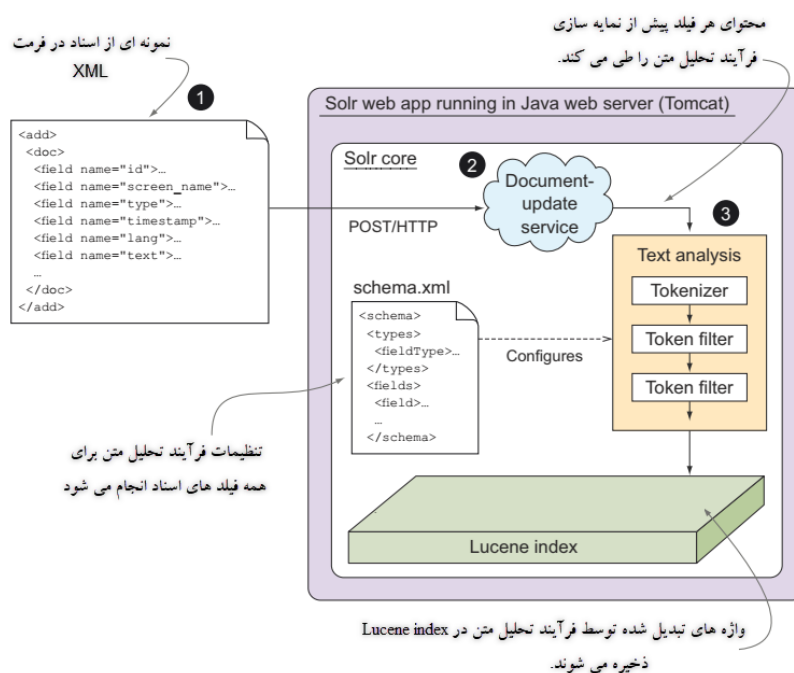
^۲ Inverted index

^۳ Distributed searching

گام ۱: تبدیل سند از فرمت اولیه به فرمت قابل پردازش در سولر مانند XML، CSV و یا JSON (۱)؛

گام ۲: اضافه کردن سند تبدیل شده به سولر از طریق واسطه‌هایی نظیر HTTP POST (۲)؛

گام ۳: اعمال تغییرات و تبدیل‌های مورد نیاز در سند در طی نمایه‌سازی (۳).



شکل ۳-۶- زیرسیستم مدیریت اسناد: فرآیند نمایه‌سازی و ذخیره‌سازی اسناد

برای اجرایی سازی گام‌های فوق نیاز است تا تنظیماتی در سولر انجام شود. به این تنظیمات در دو بخش تنظیمات Schema و تنظیمات ورود و بروزرسانی اسناد پرداخته می شود. فرآیند تحلیل متن نیز در بخش ۴-۳، تحلیل متن، بررسی می شود.

۳-۳-۱- تنظیمات ورود و بروزرسانی اسناد

همانگونه که در شکل ۳-۶ نشان داده شده است، در ابتدا باید فرمت اسناد ورودی به فرمت‌های قابل پردازش در سولر مانند XML، CSV و یا JSON تبدیل شده و سپس از طریق واسطه‌هایی به سولر اضافه شوند. علاوه بر افزودن، این واسطه‌ها وظیفه بروزرسانی اسناد موجود را نیز برعهده دارند. یکی از پرکاربردترین این واسطه‌ها، واسطه HTTP POST است که از طریق اپلیکیشن post.jar و کلاینت سولر وظیفه اضافه کردن و بروزرسانی اسناد را برعهده دارد. سه نمونه از سایر واسطه‌های پرکاربرد عبارتند از:

□ Data Import Handler (DHI): امکان اضافه کردن داده‌ها را از منابع بیرونی مانند وب‌سایت‌ها و یا پایگاه‌های داده رابطه‌ای پشتیبانی کننده از اتصال مبتنی بر جاوا^۱ (مانند MySQL، Postgres، Oracle) فراهم می‌آورد.

□ ExtractingRequestHandler (Solr Cell): امکان اضافه کردن داده‌های متنی با فرمت‌هایی چون PDF، Ms Office و OpenOffice را فراهم می‌آورد. در پشت صحنه Solr Cell از آپاچی Tika project برای قطعه‌بندی^۲ اسناد و استخراج متن و فراداده سند استفاده می‌کند.

□ Nutch: خزشگر وب مبتنی بر جاوا است و برای قابل جستجو کردن لینک‌های وب در حجم وسیع بکار می‌رود.

واسط‌های ذکر شده درخواست اضافه کردن و یا بروزرسانی را به سولر داده و سولر توسط Update Handler این وظیفه را مدیریت می‌کند. خصوصیات این handler در جدول ۶-۷ بخش ضمایم آمده است و به جهت اهمیت، خصوصیت سپردن^۳ اسناد برای نمایه‌سازی در ادامه بررسی بیشتری می‌شود.

به طور معمول، زمانی که یک سند به سولر اضافه می‌شود، تا وقتی نمایه نشده باشد در نتایج جستجو نشان داده نمی‌شود. در سولر 4.0 دو نوع سپردن وجود دارد: سپردن معمولی^۴ (سخت) و سپردن نرم^۵.

در سپردن سخت، سولر فرآیند نمایه‌سازی را به یکباره و برای تمامی اسناد جدید انجام داده و سپس اسناد نمایه شده را در محلی به صورت کامل در مکان ذخیره‌سازی طولانی مدت (durable storage) ذخیره می‌کند.

سپردن نرم فرآیندی متفاوت دارد؛ پشتیبانی سولر از جستجوی نزدیک به زمان واقعی ایجاب می‌کند تا به محض اضافه شدن یک سند، آن سند قابل جستجو باشد. به عبارت دیگر، اسناد باید در نزدیک به زمان واقعی نمایه شوند. لذا در فرآیند سپردن نرم از بخش‌های زمان‌بر نمایه‌سازی

^۱ Java database connectivity (JDBC)

^۲ Fragment

^۳ Commit

^۴ Normal (hard) commit

^۵ Soft commit

چشم‌پوشی شده و اسناد در durable storage ذخیره نمی‌شوند، بلکه در بازه‌های زمانی مشخص شده (مثلاً انتهای هر هفته) عملیات ذخیره‌سازی انجام می‌شود.

برای تعیین زمان آغاز فرآیند سپردن می‌توان از معیارهای زیر استفاده کرد: (۱) نمایه‌سازی هر سند در زمان اضافه شدن، (۲) نمایه‌سازی اسناد زمانی که تعداد اسناد نمایه نشده به حد مشخصی برسند و یا (۳) نمایه‌سازی اسناد در بازه‌های زمانی مشخص شده. همچنین میزان بازدید کاربران از سامانه (تعداد جستجوها) و یا میزان اضافه شدن سند جدید دو عامل کلیدی در تعیین نوع سپردن هستند.

پس از نمایه‌سازی، سولر هر دسته از اسناد را در یک بخش^۱ ذخیره می‌کند و این بخش‌ها ثابت می‌شوند، لذا، دسته جدید اسناد نمایه شده در بخش جدید دیگری ذخیره خواهند شد. در هنگام جستجو، هر دسته به صورت مجزا پردازش شده و نتیجه نهایی از ترکیب نتایج بخش‌ها بدست خواهد آمد. لذا تعداد دسته‌ها و حجم اسناد هر بخش عامل مهمی در تعیین سرعت بازیابی و نمایش اطلاعات است. همچنین اگر سندی حذف شود، صرفاً از بخش مربوطه حذف شده ولی همچنان در durable storage وجود دارد تا زمانی که بخش‌ها برای ذخیره‌سازی نهایی ادغام شده و durable storage بروزرسانی شود.

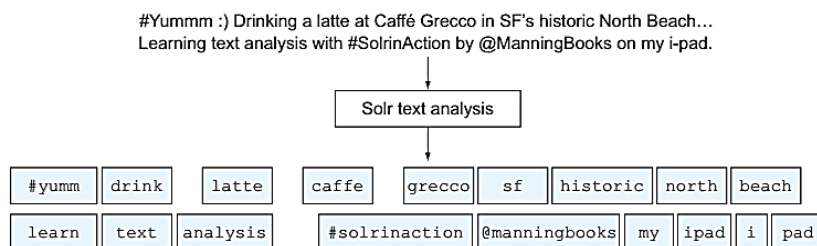
زیرسیستم تحلیل متن دومین زیرسیستم مورد بررسی در پلفرم سولر است. بخش بعد به بررسی این زیرسیستم اختصاص دارد.

۳-۴ زیرسیستم تحلیل متن

تحلیل اسناد موجود در پایگاه داده، اعمال تغییرات و تبدیل‌های مورد نیاز در سند در طی نمایه‌سازی و تحلیل عبارت‌های پرس‌وجو از جمله اهداف اجرای زیرسیستم تحلیل متن است. برخی از ابزارهای مورد استفاده سولر در این زیرسیستم در قالب یک مثال (شکل ۳-۷ و

جدول ۳-۵) نشان داده شده است.

^۱ Segment



شکل ۳-۷- مثال تحلیل متن

جدول ۳-۵- متن اولیه مثال شکل ۳-۷، تغییرات اعمال شده و ابزارهای سولر برای تغییرات

متن اولیه	تغییرات اعمال شده	ابزارهای تحلیل سولر
تمامی واژه‌ها	تبدیل به حروف کوچک	LowerCaseFilterFactory
a, at, in, with, by, on	حذف ایست واژه ^۱	StopFilterFactory
Drinking, learning	drink, learn (ریشه‌یابی) ^۲	KStemFilterFactory
SF's	sf, san francisco	WordDelimiterFilterFactory
Caffé	caffe	ASCIIFoldingFilterFactory
i-Pad	ipad, i pad	WordDelimiterFilterFactory
#Yummm	#yummm	PatternReplaceCharFilterFactory
#SolrInAction, @ManningBooks	#solrination, @manningbooks	WhitespaceTokenizer WordDelimiterFilterFactory

فرآیند تحلیل متن در سولر شامل دو مرحله است: (۱) واژه-واژه کردن^۳ (تجزیه) و (۲) فیلتر کردن واژه‌های تجزیه شده^۴:

□ مرحله ۱: واژه-واژه کردن

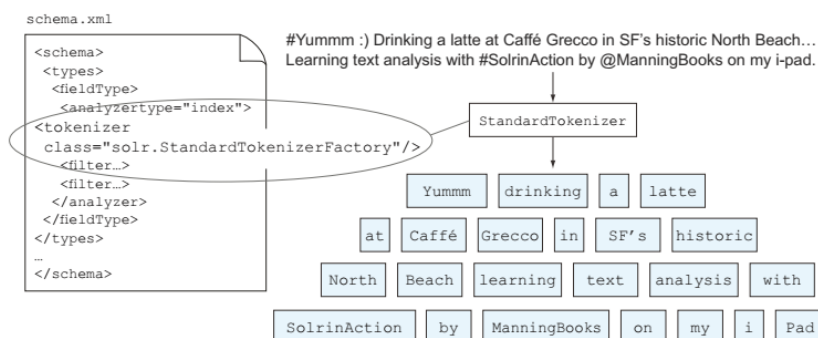
در این مرحله متن به مجموعه‌ای از واژه‌ها تجزیه می‌شود. StandardTokenizer یکی از معمولی‌ترین و پرکاربردترین ابزارها است که در طی آن واژه‌ها بر اساس فاصله و نشانه‌ها جدا شده و آدرس‌های اینترنتی، رایانامه‌ها و سرنام‌ها مدیریت می‌شوند. نمونه‌ای از این تغییرات به همراه نحوه فراخوانی آن در شکل ۳-۸ نمایش داده شده است.

^۱ Stop words

^۲ Stemming

^۳ Tokenization

^۴ Token filtering



شکل ۳-۸- مثالی از StandardTokenizer

□ مرحله ۲: فیلتر کردن واژه‌های تجزیه شده

پس از جداسازی واژه‌ها در متن گام بعدی فیلتر کردن واژه‌های تجزیه شده (حذف ایست‌واژه‌ها، تبدیل حروف بزرگ به کوچک و ...) است. برای این منظور فیلترهای مختلفی در سولر طراحی شده است، نظیر:

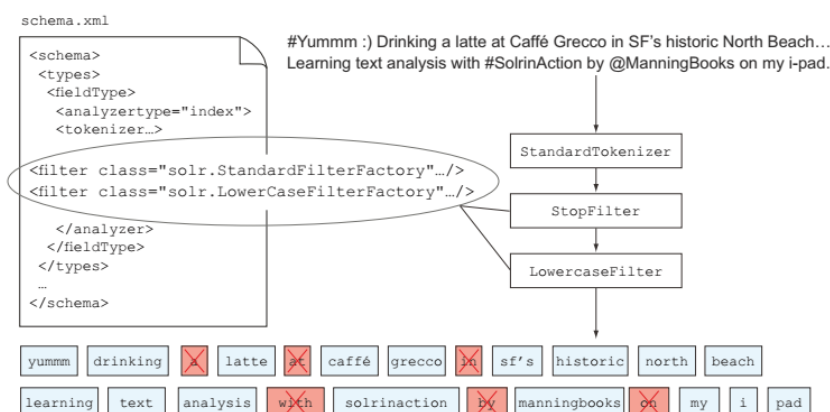
○ StopFilterFactory برای حذف ایست‌واژه‌ها بکار می‌رود. لیست این کلمات

باید در فایل txt جداگانه تعریف شده و در Schema فراخوانی شود.

○ LowerCaseFilterFactory برای تبدیل حروف بزرگ به کوچک بکار

می‌رود. نتیجه اعمال این دو فیلتر در مثال شکل ۳-۸ به همراه نحوه فراخوانی در

شکل ۳-۹ نمایش داده شده است.



شکل ۳-۹- اعمال دو فیلتر StopFilterFactory و LowerCaseFilterFactory در مثال شکل ۳-۸

○ WordDelimiterFilterFactory برای حذف نشانه‌ها (مانند ":", ";", ".") و مدیریت

کاراکترهای خاص (مانند "@", "#") بکار می‌رود (شکل ۳-۷).

- `ASCIIFoldingFilterFactory` برای تبدیل حروف به کدهای ASCII بکار می‌رود و بهتر است بعد از `LowerCaseFilterFactory` اعمال شود. از این فیلتر برای تبدیل حروف لاتین استفاده می‌شود.
- `KStemFilterFactory` برای ریشه‌یابی متون انگلیسی بکار می‌رود. سایر فیلترهای ریشه‌یابی در بخش بعد تشریح خواهند شد.
- `PatternReplaceCharFilterFactory` برای جایگزین کردن الگوهای نوشتاری بکار می‌رود (به عنوان مثال جایگزین کردن `#Yamm` با `#Yam` در شکل ۳-۷). بهتر است این فیلتر قبل از شروع فرآیند واژه-واژه کردن اعمال شود و درواقع نوعی پیش‌پردازش محسوب می‌شود.
- `HTMLStripCharFilterFactory` برای حذف کردن آدرس‌های اینترنتی از متن بکار می‌رود و بهتر است قبل از شروع فرآیند واژه-واژه کردن اعمال شود.
- `MappingCharFilterFactory` برای جایگزین کردن کاراکترهای متن با کاراکترهای از قبل تعیین شده بکار می‌رود.
- `SynonymFilterFactory` برای اضافه و یا جایگزین کردن کلمات مترادف بکار می‌رود و بهتر است در آخرین حلقه زنجیره پردازش متن و صرفاً در هنگام پردازش عبارت پرس‌وجو استفاده شود. فایل `txt` کلمات مترادف باید از قبل تعریف شده و در `Schema` فراخوانی شود. برای تعریف این فایل می‌توان از پایگاه داده `WordNet` نیز استفاده نمود.

یکی از اصلی‌ترین اقدامات در طی فرآیند تحلیل متن یافتن ریشه کلمات یا ریشه‌یابی است. به روش‌های مورد استفاده در سولر برای این منظور در ادامه پرداخته می‌شود.

۱-۴-۳ ریشه‌یابی

جزئی از واژه پس از حذف وندهای آن ریشه نامیده شده و ریشه‌یابی را می‌توان یکی از راهکارهای بهبود کارایی سیستم‌های بازیابی اطلاعات دانست. درواقع ریشه‌یابی باعث بهبود شاخص بازخوانی می‌شود و اما می‌تواند بر شاخص دقت اثر منفی داشته باشد.

فیلترهای ریشه‌یابی متعددی در سولر طراحی شده است که از آن میان می‌توان به ریشه‌یاب `Krovets` با دستور `KStemFilterFactory`، ریشه‌یاب `Porter` با دستور

PorterStemFilterFactory و ریشه‌یاب EnglishMinimalStemFilter اشاره کرد. مقایسه‌ای از عملکرد این ریشه‌یاب‌ها در جدول ۳-۶ آمده است.

جدول ۳-۶- عملکرد ریشه‌یاب‌های EnglishMinimalStemFilter و KStemFilter، PorterStemFilter

EnglishMinimalStemFilter	KStemFilter	PorterStemFilter	واژه اصلی
country	country	countri	country
country	country	countri	countries
dog	dogs	dog	dogs
run	runs	run	runs
running	running	run	running
association	association	associ	association
associate	associate	associ	associates
listing	list	list	listing
investing	invest	invest	investing
investment	investment	invest	investments
investor	invest	investor	investors
organization	organization	organ	organizations
organize	organize	organ	organize
organic	organic	organ	organic
generous	generous	gener	generous
general	general	gener	generals
generalization	generalization	gener	generalizations

ریشه‌یاب Porter نسخه دیگری نیز دارد که با دستور SnowballPorterFilter فراخوانی می‌شود. تفاوت این دو نسخه در سرعت و دقت آنهاست؛ نسخه اصلی سریع‌تر بوده و برای زبان انگلیسی پیشنهاد می‌شود اما نسخه دوم (Snowball) دقیق‌تر عمل می‌کند. اگر پارامتر زبان در نسخه دوم به صورت Lovins تنظیم شود، این دو نسخه مشابه هم عمل می‌کنند^۱ و اگر پارامتر زبان به صورت English تنظیم شود، نسخه دوم بر اساس الگوریتم بهبود یافته Porter عمل خواهد کرد^۲.

گاهی در ریشه‌یابی حالات خاصی رخ می‌دهد، به عنوان نمونه:

□ گاهی اسامی خاص و یا واژگان تخصصی وجود دارند که روش‌های ریشه‌یابی نباید تغییراتی در آنها اعمال کند. بدین منظور لیستی از اسامی محافظت شده ایجاد می‌شود. برای این منظور در زنجیره پردازش متن و قبل از فراخواندن ریشه‌یاب، دستور

^۱ <http://snowball.tartarus.org/algorithms/lovins/stemmer.html>

^۲ <http://snowball.tartarus.org/algorithms/english/stemmer.html>

KeywordMarkerFilterFactory اجرا شده و آدرس فایل حاوی اسامی محافظت

شده در قالب txt به سولر داده شود.

□ گاهی نیاز است ریشه‌یاب، برخی از ریشه‌های بدست آمده را به واژگانی خاص نگاشت

کند، این امر توسط دستور StemmerOverrideFilterFactory محقق می‌شود. این

دستور مانند دستور KeywordMarkerFilterFactory نیازمند لغت‌نامه‌ای در قالب

فایل txt است که نگاشت‌ها در آن تعریف شده باشد.

نمونه‌ای از نحوه بکارگیری این دو دستور در Schema به صورت زیر است:

```
<fieldType name="text_en_kstem_with_overrides"
class="solr.TextField"
positionIncrementGap="100">
<analyzer>
<tokenizer class="solr.StandardTokenizerFactory"/>
<filter class="solr.StopFilterFactory"
ignoreCase="true"
words="lang/stopwords_en.txt"/>
<filter class="solr.LowerCaseFilterFactory"/>
<filter class="solr.EnglishPossessiveFilterFactory"/>
<filter class="solr.KeywordMarkerFilterFactory"
protected="protwords.txt"/>
<filter class="solr.StemmerOverrideFilterFactory"
dictionary="stemdict_en.txt" />
<filter class="solr.KStemFilterFactory"/>
</analyzer>
</fieldType>
```

protwords.txt

listing

manning

stemdict_en.txt

dogs dog

runs run

investors investor

با استفاده از این دو دستور در زنجیره تحلیل متن می‌توان از بسیار از اشتباهات بوجود آمده در اثر

اجرای ریشه‌یابی جلوگیری نمود. با این حال مشکلاتی نیز وجود دارند و استفاده ناکارا از

روش‌های ریشه‌یابی می‌تواند تاثیرات نامطلوبی بر معیارهای دقت و بازخوانی داشته باشد. به عنوان

مثال:

□ اگر تمام حالت‌های یک واژه به یک ریشه ختم نشود (مانند نتایج روش Porter برای واژه‌های dry و dries که برای واژه اول ریشه dry و برای واژه دوم ریشه dri حاصل می‌شود) نتایج بازیابی کامل نبوده و مقدار بازخوانی کاهش می‌یابد.

مشکلی با نام overmatching، مانند نتایج روش Porter برای واژه‌های generals، generous و generalizations در

□ جدول ۳-۵، سبب می‌شود تا نتایج خوبی بازیابی نشده و دقت نتایج بازیابی شده کاهش یابد.

برای برطرف نمودن این کاستی‌ها دو راهکار کلی پیشنهاد شده است:

□ استفاده از ریشه‌یاب‌های کمتر تهاجمی مانند EnglishMinimalStemFilter بجای روش‌های تهاجمی‌تر مانند PorterStemFilter (روش‌های تهاجمی بازخوانی بالاتری داشته و روش‌های کمتر تهاجمی دقت بالاتری دارند).

□ استفاده از روش‌های بُن‌واژه‌سازی^۱. هدف در هر دو دسته روش‌های ریشه‌یابی و بُن‌واژه‌سازی دستیابی به ریشه واژه‌هاست با این تفاوت که روش‌های ریشه‌یابی، روش‌هایی مبتنی بر الگوریتم بوده و روش‌های بُن‌واژه‌سازی روش‌هایی مبتنی بر لغت‌نامه هستند (هرچند یافتن لغت‌نامه جامع دشوار و در برخی موارد غیرممکن است). شکل ۳-۱۰ نمونه‌هایی از نتایج متفاوت بکارگیری این دو دسته از روش‌ها را نمایش می‌دهد.

Original terms:

goose	geese	generations	generals	several	severity
-------	-------	-------------	----------	---------	----------

Stemmed terms:

goos	gees	gener	gener	sever	sever
------	------	-------	-------	-------	-------

Lemmatized terms:

goose	goose	generation	general	several	severe
-------	-------	------------	---------	---------	--------

شکل ۳-۱۰- نمونه‌هایی از نتایج متفاوت بکارگیری روش‌های ریشه‌یابی و بُن‌واژه‌سازی

تبادل بین دقت و بازخوانی می‌تواند به عنوان معیاری برای تصمیم‌گیری در مورد بکارگیری روش‌های ریشه‌یابی و یا بُن‌واژه‌سازی باشد. عدم استفاده از ریشه‌یابی سبب می‌شود تا دقت نتایج افزایش یابد و استفاده از آن باعث بهبود بازخوانی و کاهش دقت خواهد شد. اگر هر دو این

^۱ Lemmatization

معیارها در سیستم بازیابی مهم باشند بهتر است از روش‌های بُن واژه‌سازی و یا روش‌های پیشرفته‌تر نظیر پردازش زبان طبیعی استفاده کرد.

یکی دیگر از موارد موثر در انتخاب و در نتیجه نتایج روش‌های بازیابی اطلاعات زبان اسناد مورد بررسی است و توانایی تحلیل زبان‌های مختلفی در سولر طراحی شده است.

۲-۴-۳- زبان‌های پیش فرض در سولر

سولر به صورت پیش فرض توانایی تحلیل اسناد با زبان‌های مختلفی را دارد نظیر انگلیسی، عربی، لاتین، هندی، ژاپنی، چینی و دستورات زنجیره تحلیل متن به زبان فارسی در ادامه آمده است که این دستورات باید در Schema در بخش تحلیل متن و تحلیل عبارت پرس و جو فراخوانی شوند (برای توضیح بیشتر به بخش ۳-۵، سند Schema سامانه گنج مراجعه شود).

PersianCharFilterFactory

StandardTokenizerFactory

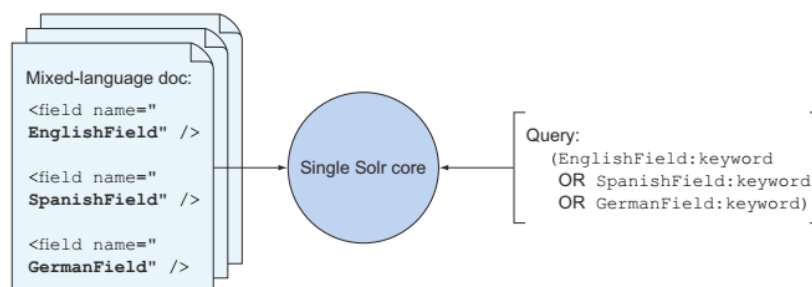
LowerCaseFilterFactory

ArabicNormalizationFilterFactory

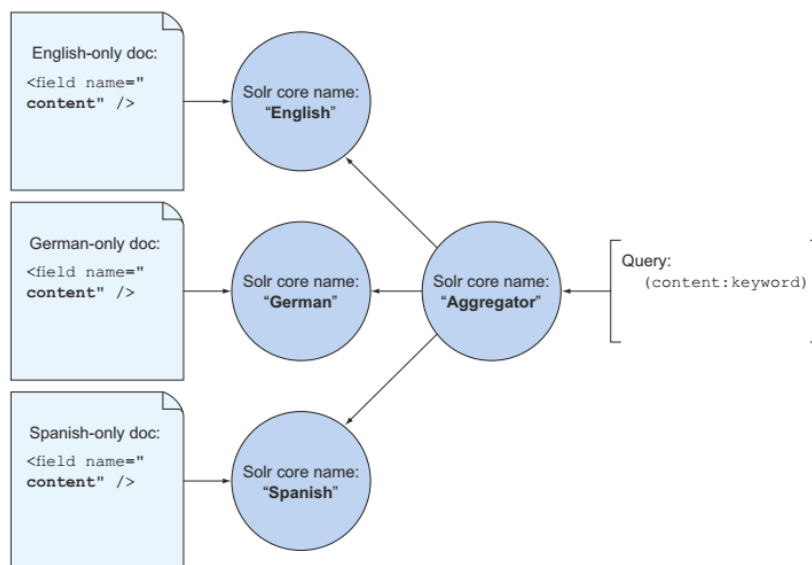
PersianNormalizationFilterFactory

StopFilterFactory [words="lang/stopwords_fa.txt"]

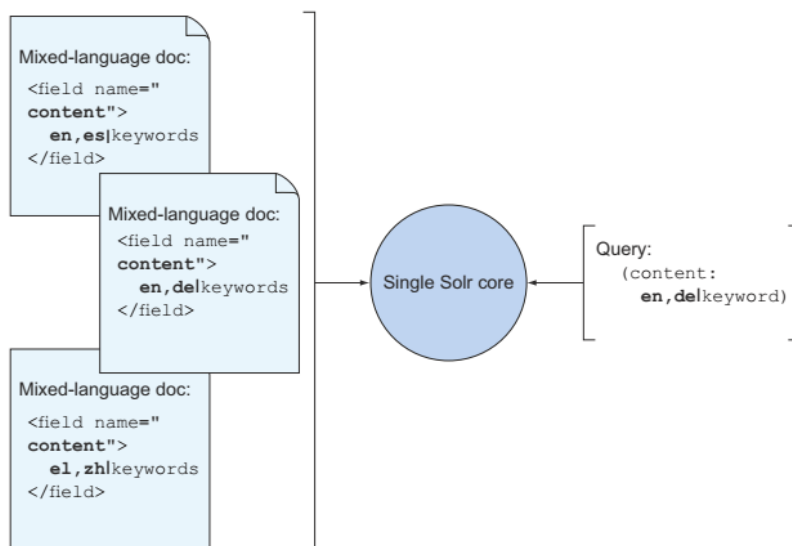
همچنین ممکن است کاربر عبارت پرس و جو خود را در زبان‌های مختلف بیان کند. برای پردازش چنین عبارت‌هایی راهکارهایی وجود دارد: (۱) در نظرگیری فیلدهای جداگانه براساس هر زبان برای هر سند (شکل ۳-۱۱)، (۲) دسته‌بندی اسناد براساس زبان و نمایه‌سازی جداگانه آنها (شکل ۳-۱۲) و (۳) تشخیص خودکار زبان سند (شکل ۳-۱۳).



شکل ۳-۱۱- در نظرگیری فیلدهای جداگانه در تحلیل‌های چند زبانی



شکل ۳-۱۲- دسته‌بندی اسناد براساس زبان در تحلیل‌های چند زبانی



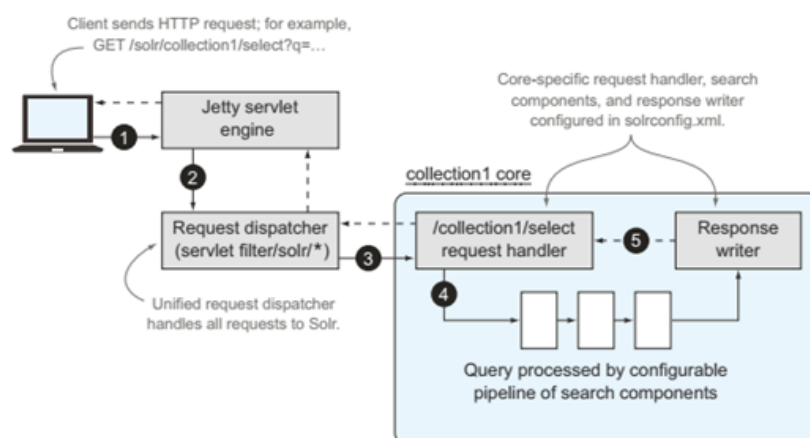
شکل ۳-۱۳- تشخیص خودکار زبان سند در تحلیل‌های چند زبانی

درمورد راهکار ۱ و ۳ نیازی به در نظر گیری هسته‌های پردازش موازی و جداگانه برای سولر نیست اما معماری سیستم بر اساس راهکار ۲ پیچیده بوده و در بسیاری از موارد باید هسته‌های پردازشی برای هر زبان در نظر گرفته شود. همچنین پیاده سازی راهکار ۳ نیازمند استفاده از روش‌های پیچیده در تحلیل متن است.

سومین زیر سیستم طراحی شده در سولر، زیر سیستم پردازش پرس و جو است که وظیفه مدیریت درخواست‌های کاربر را از سیستم جستجو برعهده دارد.

۳-۵ زیرسیستم پردازش پرس و جو

نمای کلی زیرسیستم پردازش پرس و جو در قالب مدیریت درخواست^۱ در شکل ۱۴-۳ شده است. در این شکل درخواستی از سمت کلاینت به صورت HTTP به موتور سرولت Jetty فرستاده می شود (۱). پس از قبول این درخواست، Jetty آن را به توزیع کننده درخواست سولر^۲ می فرستد (۲). در گام بعد توزیع کننده درخواست هسته collection1 را شناسایی کرده و مدیریت درخواست دستور /select را از Schema بازخوانی می کند (۳). این handler درخواست را بر اساس فرآیند ارائه شده در شکل ۱۷-۳ پردازش کرده (۴) و نتایج را برای نمایش در کلاینت تنظیم می کند (۵).



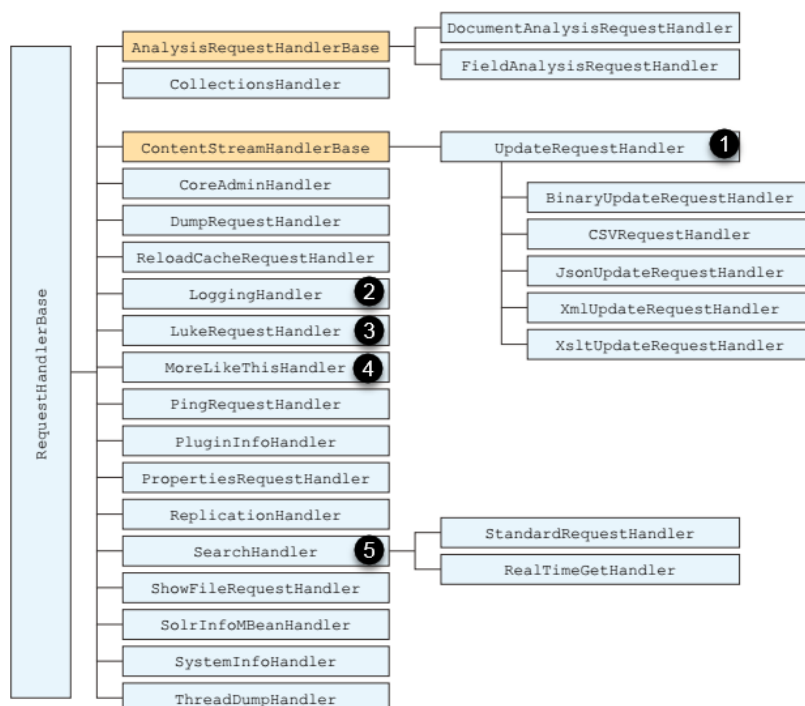
شکل ۱۴-۳- نمای کلی مدیریت درخواست جستجو

بخش مدیریت درخواست از کلاس های متعددی تشکیل شده اند که ساختار درختی آنها در شکل ۱۵-۳ نمایش داده شده و خصوصیات تمامی این کلاس ها باید در ابتدا و در Schema تعیین گردند. به عنوان نمونه، بروزرسانی نمایه ها از طریق کلاس UpdateRequestHandler (۱) و از طریق handler ۵ زیرمجموعه انجام می شود. همچنین ذخیره سازی لاگ از طریق LoggingHandler (۲)، گزارش فراداده مربوط به اسناد از طریق LukeRequestHandler (۳) و امکان مشاهده اسناد مشابه با سند انتخابی کاربر^۳ از طریق MoreLikeThisHandler (۴) ممکن می شود.

^۱ Request handler

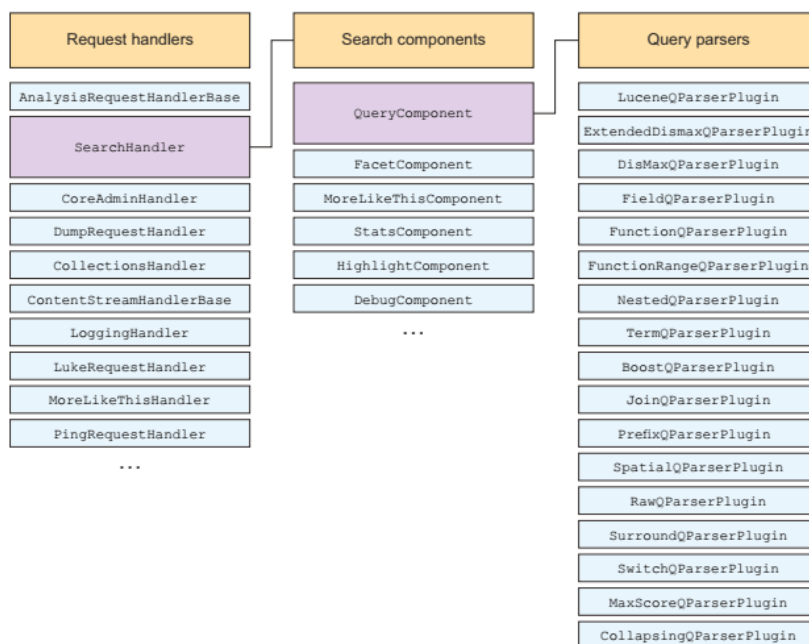
^۲ Solr request dispatcher

^۳ More like this

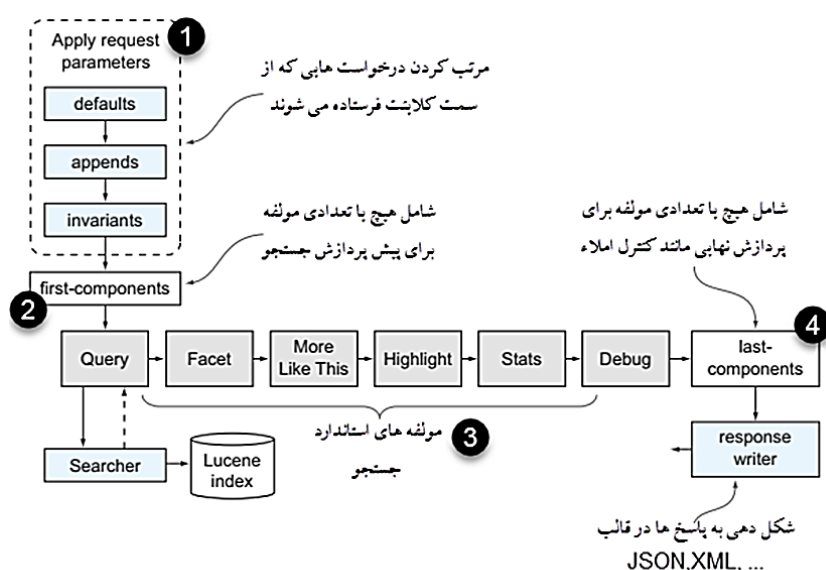


شکل ۳-۱۵- کلاس‌های مدیریت درخواست در سولر

SearchHandler (۵) وظیفه مدیریت جستجوی کاربر را برعهده دارد و به جهت اهمیت این موضوع در این پژوهش در ادامه بر آن تمرکز خواهد شد. نمای درختی کلاس‌های این handler در شکل ۳-۱۶ و فرآیند کلی آن در شکل ۳-۱۷ نمایش داده شده است.



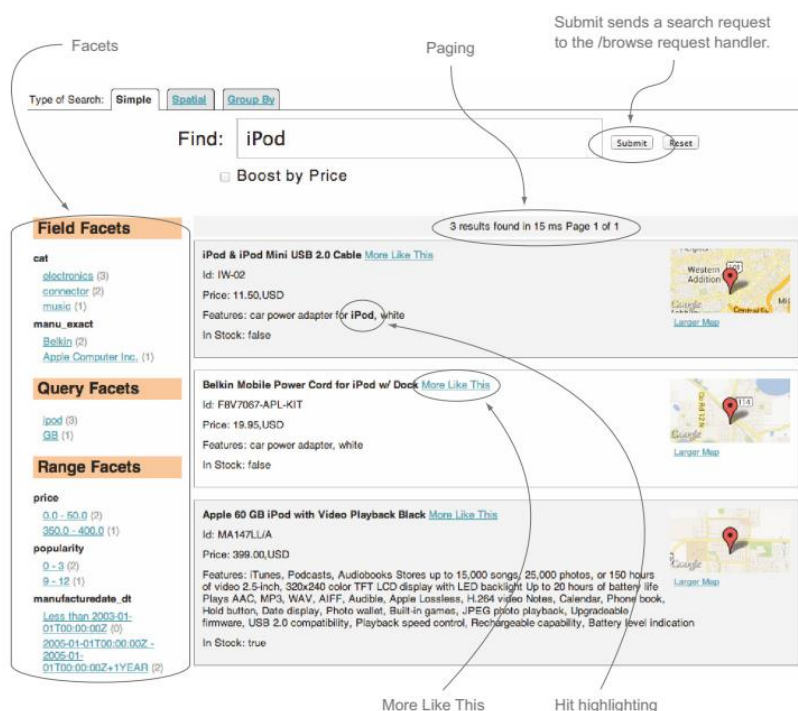
شکل ۳-۱۶- نمای درختی کلاس‌های SearchHandler



شکل ۳-۱۷- فرآیند کلی SearchHandler

بخش ① در شکل ۳-۱۷ وظیفه تنظیم پارامترهای اولیه و مرتب کردن درخواست هایی که از سوی کلاینت فرستاده می شود را برعهده دارد. اعمال تنظیمات و مقادیر پیش فرض، در صورتی که کاربر تغییری را اعمال نکرده باشد، در default و ثابت کردن برخی از پارامترها در invariants انجام شده و در بخش append پارامترهای دیگری به موارد مورد نظر کاربر افزوده می شود. بخش ②، First-components، وظیفه پیش پردازش را برعهده داشته و به عنوان بخشی اختیاری در زنجیره جستجو طراحی شده است. بخش ③ شامل مولفه های جستجو بوده و حداقل یکی از این موارد بایستی در فرآیند جستجو فعال شده باشد (مولفه Query به صورت پیش فرض فعال است). بخش ④، Last-component، مولفه ای اختیاری در زنجیره بوده و برای فعالیت های پس از پردازش طراحی شده است.

نمونه ای از کاربرد مولفه های بخش ③ در شکل ۳-۱۸ نمایش داده شده و در ادامه به صورت تفصیلی به آنها پرداخته خواهد شد.



شکل ۳-۱۸- نمونه‌ای از کاربرد مولفه‌های بخش ۳ شکل ۳-۱۷

۳-۵-۱- مولفه QUERY

این مولفه هسته اصلی فرآیند مدیریت درخواست‌ها بوده و وظیفه یافتن کلیه اسناد مرتبط با عبارت پرس‌وجوی کاربر را با توجه به نمایه‌های موجود در لوسن برعهده دارد. به عبارت دیگر، در این مولفه با کمک تجزیه‌گر عبارت پرس‌وجو^۱، عبارت مورد نظر کاربر به واژگان مورد نیاز تبدیل شده و اسناد مرتبط بازیابی می‌شوند.

در هنگام استفاده از این مولفه، دو پارامتر q (عبارت پرس‌وجوی کاربر (query)) و fq (فیلترهای اعمال شده بر عبارت پرس‌وجو (filter query)) باید تعریف شود. تعریف این دو پارامتر به صورت مجزا مزیت‌هایی را به همراه دارد:

- هدف از پارامتر q محدود کردن نتایج جستجو به اسناد مرتبط و فراهم آوردن لیستی از واژگان مرتبط برای محاسبه ارتباط اسناد با عبارت پرس‌وجو بوده درحالی که هدف از پارامتر fq تنها محدود کردن نتایج جستجو به اسناد مرتبط است. به عبارت دیگر، اگر پرس‌وجوی کاربر به صورت زیر باشد

keywords: "solr in action", category: "technology"

^۱ Query parser

پارامتر $q = solr\ in\ action$ و پارامتر $fq = technology$ خواهد بود و درواقع پارامتر q علاوه بر محدود کردن نتایج جستجو به اسناد مرتبط، لیستی از واژگان مرتبط را نیز فراهم می‌کند.

□ اغلب کاربران از fq های ثابتی در دفعات جستجوی خود استفاده کرده و تنها q را تغییر می‌دهند. لذا جدا بودن این دو پارامتر باعث بهبود سرعت پردازش و کاهش حجم محاسبات در سنجش میزان ارتباط یک سند با عبارت پرس‌وجو خواهد شد.

ترتیب پردازش q و fq نیز اهمیت دارد؛ اول باید q را پردازش کرد و سپس fq را اعمال کرد؟ و یا برعکس؟. پاسخ دادن به این پرسش بستگی به زمینه کاری موتور جستجو و بافت مسئله دارد اما فرآیند کلی در سولر به صورت زیر است:

۱. ابتدا fq ها به صورت مجزا بر مجموعه اسناد موجود در cache سیستم اعمال شده و مجموعه DocSet برای هر fq تشکیل می‌شود.

۲. اگر گام ۱ نتیجه‌ای دربرنداشت، fq ها به صورت مجزا بر مجموعه اسناد موجود در پایگاه داده اعمال شده و مجموعه DocSet برای هر fq تشکیل می‌شود.

۳. نتایج برای هر fq با رابطه AND جمع شده و یک DocSet واحد تشکیل می‌شود.

۴. پارامتر q بر نتایج موجود در DocSet اعمال شده و امتیاز ارتباط محاسبه می‌شود.

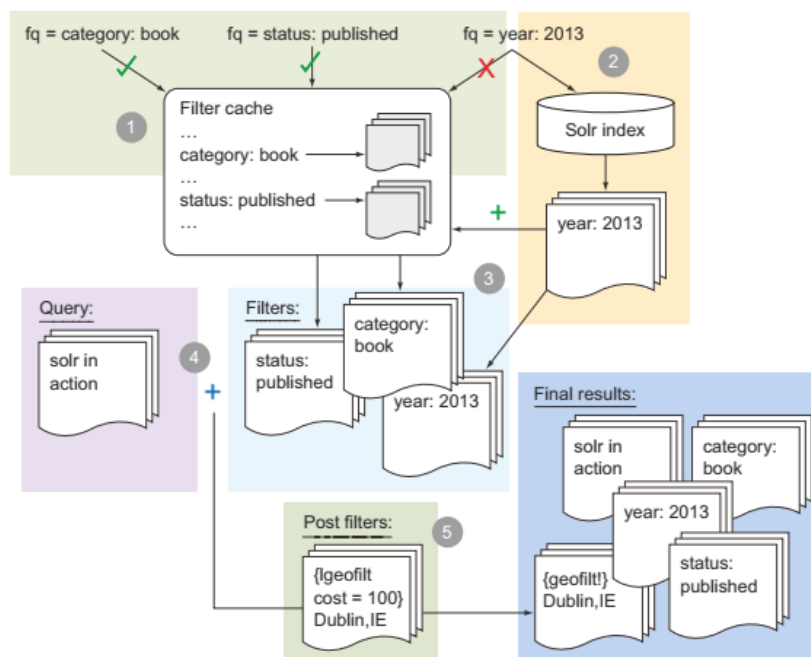
۵. اگر کاربر فیلتر نهایی^۲ تعریف کرده باشد، این فیلتر در بازنمایی نتایج اعمال خواهد شد.

گام‌های فوق برای پارامترهای زیر در شکل ۳-۱۹ نمایش داده شده‌اند.

```
q=solr in action &
fq=category:book & fq=status:published & fq=year:2013 &
post filter= {!geofilt tag="Dublin,IE" cost=100 sfield=location}
```

^۱ Search session

^۲ Post filter



شکل ۳-۱۹- گام‌های پردازش q و fq

همانگونه که ذکر شد، لوسن و سولر برای تحلیل اسناد پایگاه داده و تبدیل عبارت پرس‌وجوی کاربر به واژگان مورد نیاز از تجزیه‌گرها استفاده می‌کنند. دو تجزیه‌گر پرکاربرد عبارتند از: لوسن query parser و eDisMax query parser.

LUCENE QUERY PARSER

سولر به صورت پیش‌فرض از Lucene query parser استفاده می‌کند که از طریق کلاس LuceneQParserPlugin بازخوانی می‌شود. خصوصیات اصلی این کلاس به شرح زیر است:

- اگر کاربر فیلد جستجو را مشخص نکرده باشد، جستجو به صورت پیش‌فرض در فیلد محتوا (content) انجام می‌شود. بنابراین دو پرس‌وجوی «solr» و «content:solr» یکسان بوده و نتایج دو پرس‌وجوی «title:apache solr» و «title:apache content:solr» متفاوت خواهند بود.
- عملگرهای +، AND و && برای تعیین کلمات ضروری در نتایج جستجو و عملگرهای OR و || برای کلمات اختیاری در نظر گرفته شده‌اند. همچنین در حالت پیش‌فرض وجود «فاصله» در بین کلمات جستجو به معنای OR است. این امر سبب می‌شود تا بازخوانی افزایش و دقت کاهش یابد (البته این حالت را می‌توان تغییر داد).
- جستجوی دقیق عبارت پرس‌وجو با "..." انجام می‌شود.

مانند آنچه در بخش ۳-۲-۳، تطابق فازی، اشاره شده نزدیکی دو واژه و یا اصلاح واژه براساس فاصله با کمک ~ تعیین می شود.

□ جستجوی Wildcard با عملگرهای * و ؟ انجام می شود.

این کلاس دارای محدودیت هایی نیز هست:

□ به عملگرهای مورد استفاده کاربر در پرس وجو بسیار حساس است و از آنجا که نمی توان انتظار داشت همه کاربران از عملگرهای درست استفاده کنند، در اغلب کاربردها کارایی پایین تر از انتظار دارد.

□ این تجزیه گر قابلیت جستجو در تمامی فیلدها را نداشته و در نتیجه کارایی آن کاهش می یابد (البته می توان از فیلدهای کپی (بخش Schema) استفاده نمود).

به منظور رفع این کاستی ها و بکارگیری تجزیه گر قوی تر، سولر با ترکیب دو تجزیه گر Lucene query parser و disjunction maximum (DisMax) تجزیه گر جدیدی را با نام Extended disjunction maximum (eDisMax) ارائه نموده است که می توان آن را به صورت کلان و کلی در Schema و یا به صورت محلی فراخوانی نمود.

EDISMAX QUERY PARSER

تجزیه گر eDisMax نسخه توسعه یافته DisMax بوده و نسبت به Lucene query parser از حساسیت کمتری نسبت به عملگرهای پرس وجو برخوردار است. همچنین قابلیت جستجو در تمامی فیلدها را دارد و برای روشن تر شدن بحث فرض کنید بخواهیم عبارت «Solr in Action» را در تمامی اسناد پایگاه داده در سه فیلد title، description و author جستجو کنیم. در حالتی که از Lucene query parser استفاده شود عبارت پرس وجو به صورت زیر خواهد بود:

```
((title:solr) OR (description:solr) OR (author:solr)) AND  
((title:in) OR (description:in) OR (author:in)) AND ((title:action)  
OR (description:action) OR (author:action))
```

اما در حالتی که از eDisMax استفاده شود عبارت پرس وجو به صورت زیر خواهد بود و سه فیلد فوق را می توان در پارامتر فیلد عبارت پرس وجو^۱ (qf) تعریف کرد:

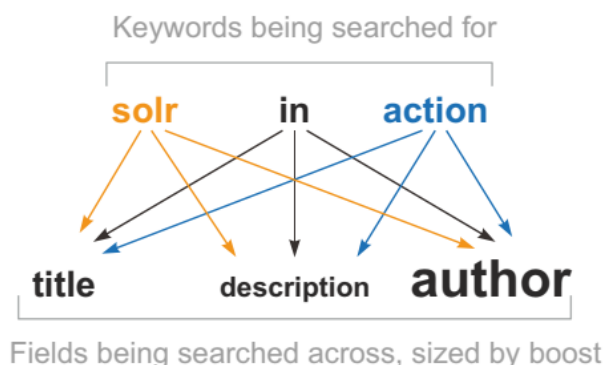
```
q=solr in action&qf=title description author
```

^۱ Query field

همچنین می توان برای هر فیلد ضریب اهمیت (وزن) جداگانه ای نیز تعریف نمود:

`q=solr in action&qf=title^1.5 description author^2`

عبارت پرس وجوی فوق به صورت نمادین در شکل ۳-۲۰ نشان داده شده است.



شکل ۳-۲۰- نمایش نمادین عبارت پرس وجو در eDisMax در حالت تفاوت وزن فیلدها

یکی از قابلیت های کلیدی eDisMax در نظرگیری نزدیکی کلمات پرس وجوی در محاسبه امتیاز سند بازیابی شده است. هرچه کلمات در اسناد بازیابی شده به یکدیگر نزدیک تر باشند امتیاز بیشتری به سند تعلق خواهد گرفت. این امتیازدهی با کمک پارامتر فیلد عبارت^۱ در سه نوع PF ، $PF2$ و $PF3$ انجام می شود. در صورت تنظیم پارامتر PF ، سولر دقیقاً عبارت پرس وجو را جستجو می کند (exact matching)، با تنظیم پارامتر $PF2$ ، سولر عبارت پرس وجو را به ۲-گرم شکسته و عملیات بازیابی را با کمک عبارت (های) جدید انجام می دهد. به طور مشابه با تنظیم پارامتر $PF3$ ، سولر عبارت پرس وجو را به ۳-گرم شکسته و عملیات بازیابی را با کمک عبارت (های) جدید انجام می دهد. به بیانی دیگر، در صورتی که پارامترهای جستجو به صورت زیر تنظیم شده باشند:

```
defType=edismax &
q=big data analytics &
qf=content &
pf=content^10
```

سولر به دنبال یافتن عبارت «big data analytics» در فیلد content بوده و ضریب وزنی آن برابر ۱۰ است. حال اگر:

`Pf2=content^10`

^۱ Phrase field

در این صورت سولر به دنبال یافتن عبارت‌های «big data» و «data analytics» در فیلد content خواهد بود و به طور مشابه اگر:

$$Pf3 = \text{content}^{10}$$

سولر به دنبال یافتن عبارت «big data analytics» در فیلد content است.

پارامترهای دیگری چون فاصله کلمات^۱ (PS, PS2, PS3)، با هدف اصلاح املائی کلمات، گره بازکن^۲ (TIE)، با هدف تعیین نحوه امتیازدهی در صورتی که عبارت پرس‌وجو در بیش از یک فیلد یافت شود) و تقویت عبارت پرس‌وجو^۳ (bq)، با هدف دادن وزن بیشتر به اسنادی که با این پارامتر منطبق هستند) برای تابع eDisMax قابل تنظیم بوده و در صورت اعمال در ست، کارایی تجزیه‌گر افزایش می‌یابد.

پارامتر حداقل انطباق^۴ (mm) پارامتری است که حداقل تعداد قطعی حضور عبارت پرس‌وجو در سند را برای اینکه به عنوان نتیجه برای کاربر نمایش داده شود مشخص می‌کند. با تنظیم این پارامتر می‌توان مقادیر بازخوانی و دقت را براساس نیاز تنظیم نمود.

نکته مهم دیگر، نحوه امتیازدهی به اسناد در این تجزیه‌گر در مقایسه با Lucene query parser است. اگر عبارت پرس‌وجو برای این دو تجزیه‌گر به صورت زیر باشد:

```
q=title:solr OR content:solr
q={!edismax qf=title content}solr
```

در تجزیه‌گر لوسن امتیاز نهایی یک سند برابر است با مجموع امتیاز آن سند برای فیلدهای title و content. در حالی که امتیاز نهایی یک سند برای eDisMax برابر است با بیشینه امتیاز آن سند در فیلدهای title و content. به بیانی دیگر، در تجزیه‌گر لوسن امتیاز یک سند بمعنای میزان ارتباط عبارت پرس‌وجو با آن سند، به عنوان یک موجودیت واحد، است در حالی که در تجزیه‌گر eDisMax، امتیاز یک سند بمعنای بیشینه ارتباط عبارت پرس‌وجو با یک یا چند فیلد مشخص شده سند است. مشکلی که این نوع امتیازدهی در روش eDisMax می‌تواند ایجاد کند آن است که فرض کنید جستجو برای فیلد عنوان و چکیده تنظیم شده باشد، حال اگر در سندی چکیده با متن سند تطابق مناسبی نداشته باشد، در صورت تطابق عبارت پرس‌وجو و چکیده سند امتیاز بالاتری

^۱ Phrase slope

^۲ Tie breaker

^۳ Boost query

^۴ Minimum match

خواهد گرفت نسبت به سندی که چکیده آن انطباق کم تر و متن آن منطبق تر با عبارت پرس وجو است. لذا در انتخاب بین این دو تجزیه گر، بافت و نوع اسناد موجود در پایگاه داده و نوع تحلیل مورد نیاز اهمیت بسیار دارد.

تجزیه گرهای دیگری نیز در سولر طراحی شده اند که می توان از آنها به صورت مجزا و یا ترکیبی استفاده کرد:

□ Field query parser: این تجزیه گر با استفاده از فاصله بین کلمات و بزرگ یا کوچک

بودن کلمات آنها را تجزیه می کند. به عنوان مثال، `{!field f=myfield}Foo Bar` برابر با جستجوی `myfield:"Foo Bar"` در پایگاه داده است.

□ Function Range query parsers: این تجزیه گر براساس تابع داده شده و در بازه ای خاص جستجو را انجام می دهد. به عنوان مثال،

```
fq={!frange l=0 u=2.2} sum(user_ranking,editor_ranking)
```

بدین معنا است که ابتدا جمع دو مقدار `user_ranking` و `editor_ranking` انجام شده و سپس در بازه مشخص شده نتایج بازیابی می شوند.

□ Nested query parser: این تجزیه گر، عبارت پرس جوی تودرتو ایجاد می کند. به عنوان مثال،

```
q:"roses are red" AND _query_:"type:poems"
```

به معنای جستجوی عبارت «`roses are red`» است با این شرط که از نوع «`poems`» باشند.

□ Boost query parser: این تجزیه گر، عبارت پرس وجو را با کمک تابع تعیین شده وزن دار می کند. به عنوان مثال،

```
{!boost b=log(popularity)}foo
```

به معنای آن است که عبارت پرس وجوی «`foo`» وزنی معادل «`log(popularity)`» در جستجو دارد.

□ Prefix query parser: در این تجزیه گر عبارت پرس جوی کاربر به عنوان پیش وند برای جستجو بکار می رود. به عنوان مثال، `{!prefix f=myfield}foo` برابر با جستجوی `myfield:foo*` در پایگاه داده است.

□ Join query parser: این تجزیه گر نتایج بازیابی اطلاعات را براساس ویژگی های تعیین شده با یکدیگر ترکیب می کند. به عنوان نمونه،

`{!join from=manu_id_s to=id}ipod`

بدین معنا است که جستجوگر تمامی اسناد حاوی عبارت ipod را یافته و براساس id تولید کننده آنها را ترکیب کند.

□ Collapsing query parser: این تجزیه گر، نتایج جستجو را براساس ساختار سلسله مراتبی (درختی) با کمک میزان ارتباط بدست آمده نمایش می دهد. به عنوان نمونه، دستور `fq={!collapse field=group_field}` سبب می شود تا نتایج بازیابی شده براساس میزان ارتباط و نیز group_field به صورت سلسله مراتبی نمایش داده شود.

۳-۵-۲ مولفه FACET

دومین مولفه در بخش ③ شکل ۳-۱۷ مولفه Facet است. این مولفه امکان برقراری جستجوی چندوجهی را در سولر مهیا می سازد. با کمک این مولفه کاربر قادر است نتایج جستجو را براساس نیاز خود فیلتر نموده و نمایش کلی از خصوصیات اسناد موجود در پایگاه داده را درخصوص موضوع مورد جستجو مشاهده نماید. به تعبیری دیگر، با کمک این مولفه فراداده مربوط به اسناد مرتبط با عبارت پرس وجو برای کاربر نمایش داده می شود (شکل ۳-۱۸). نکات مربوط به این مولفه به صورت مختصر عبارتند از:

- این مولفه به صورت پیش فرض فعال نیست و باید در Schema تنظیم شود.
- Facet ها به صورت برخط و برای هر جستجو محاسبه می شوند. لذا، زمان بر بوده و در صورت طراحی نامناسب می تواند سبب کاهش کارایی شوند.
- در نسخه های فعلی سولر، این قابلیت در صورتی ممکن است فعال شود که سولر دارای یک هسته پردازشی فعال باشد.

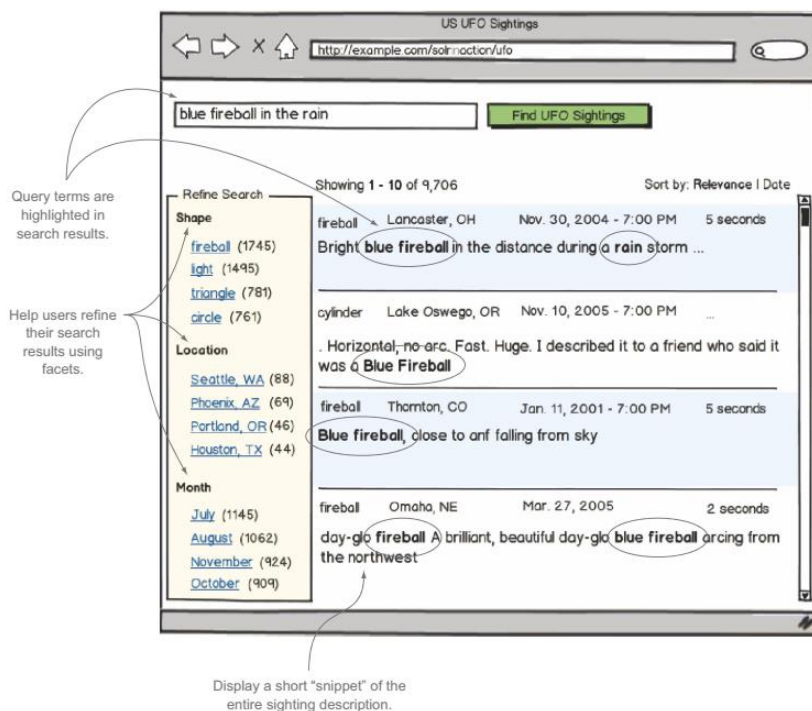
۳-۵-۳ مولفه MORE LIKE THIS

مولفه More like this سومین مولفه در بخش ③ شکل ۳-۱۷ بوده و وظیفه نمایش اسناد مشابه با سند مورد انتخابی کاربر را برعهده دارد (شکل ۳-۱۸). شباهت در این مولفه براساس زاویه بین بردار و TF-IDF دو سند محاسبه می شود؛ این محاسبه می تواند به صورت برخط (در هنگام جستجو) و یا آفلاین باشد. تفاوت بین این دو حالت در میزان استفاده از حافظه و فضای ذخیره سازی و نیز زمان پردازش است. به تعبیری دیگر، در حالت برخط فضای ذخیره سازی کمتر و حافظه بیشتری مورد نیاز است در صورتی که در حالت آفلاین میزان استفاده از حافظه و سرعت پردازش افزایش می یابد.

این مولفه به صورت پیش فرض فعال نبوده و باید در Schema فعال شده و فیلدهایی که برای محاسبه شباهت استفاده می شود مشخص گردند. همچنین می توان لیستی از واژگان مهم برای یافتن شباهت را نیز به همراه وزن آنها تعیین نمود تا در محاسبه شباهت بین دو سند تغییراتی ایجاد شود (این مورد در بخش ۵-۴، شخصی سازی جستجو، بررسی خواهد شد).

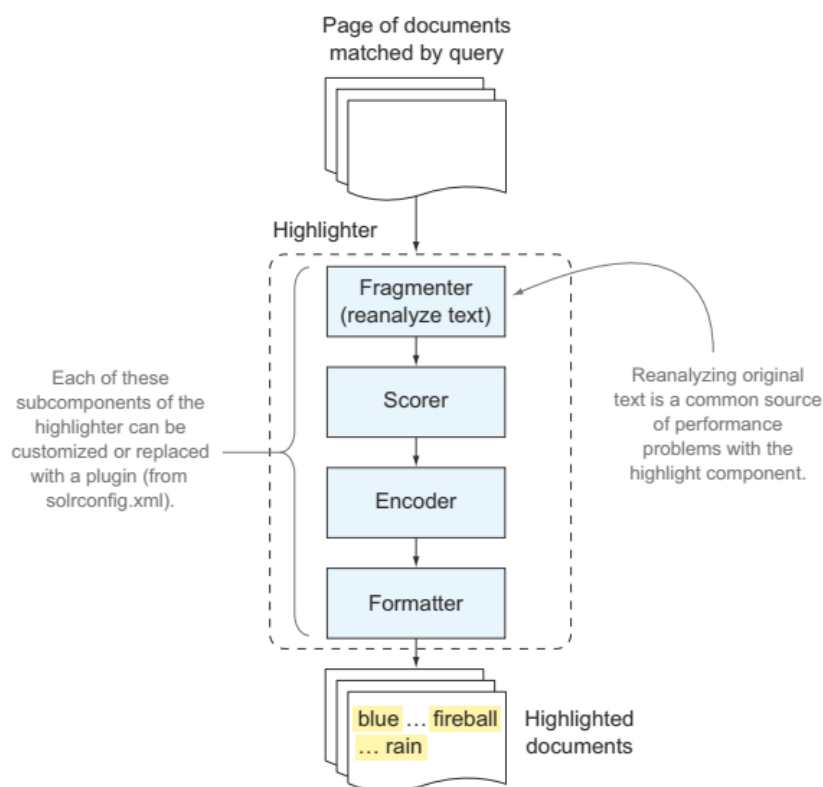
۳-۵-۴ مولفه HIGHLIGHT

پررنگ کردن بخش هایی از سند که با عبارت پرس و جوی کاربر منطبق هستند اتفاقی است که پس از فعال کردن این مولفه در فرآیند مدیریت درخواست جستجو (شکل ۳-۱۷) و در Schema می افتد. نمونه ای از نتایج جستجو پس از فعال کردن این مولفه در شکل ۳-۲۱ نشان داده شده است.



شکل ۳-۲۱- نمونه ای از نتایج جستجو پس از فعال کردن مولفه highlight

یکی از مواردی که در این شکل نشان داده شده نمایش snippet است. در واقع snippet بخشی از متن سند است که با عبارت پرس و جوی کاربر مطابقت دارد و انتخاب آن از طریق فرآیند نشان داده شده در شکل ۳-۲۲ انجام می شود.



شکل ۳-۲۲- فرآیند نمایش snippet و پررنگ کردن بخش‌هایی از سند

ورودی این فرآیند بخش‌هایی از سند است که با عبارت پرس‌وجوی کاربر مطابقت و گام‌های آن عبارتند از:

□ گام ۱: قطعه‌بندی کننده^۱ (تحلیل مجدد متن):

برای پررنگ کردن بخش‌هایی از سند، واژگان سند باید قابل مقایسه با واژگان عبارت پرس‌وجو باشند و همچنین سولر باید از مکان دقیق واژگان در سند اطلاع داشته باشد. لذا باید فرآیند تحلیل متن روی سند اجرا شود (موارد مطرح شده در بخش ۴-۳، زیر سیستم تحلیل متن).

سپس برای تعیین snippet‌ها برای نمایش، متن باید به قطعاتی تقسیم شود. از دو رویکرد برای قطعه‌بندی متن در سولر استفاده می‌شود: GapFragmenter و RegexFragmenter. GapFragmenter رویکرد پیش‌فرض بوده و متن را به تعداد واژگان از پیش تعیین شده^۲ قطعه‌بندی می‌کند (مقدار پیش‌فرض (h1.fragsize) برابر ۱۰۰ کاراکتر است). در

^۱ Fragmenter

^۲ Token boundary

RegexFragmenter متن براساس الگوهای از پیش تعیین شده (مانند رسیدن به نشانه‌هایی چون . ، / و ...) قطعه‌بندی می‌شود.

گام ۲: امتیازدهنده: در این گام برای تمامی قطعات شناسایی شده در گام ۱ امتیاز ارتباط با عبارت پرس‌وجو محاسبه شده و قطعه دارای بالاترین امتیاز برای نمایش انتخاب می‌شود. تابع پیش‌فرض برای محاسبه این امتیاز QueryScorer است و براساس شمارش واژه‌ها در هر قطعه، امتیاز را محاسبه می‌کند.

گام ۳ و ۴: کدگذاری و فرمت‌دهی: در این گام واژه‌های تعیین شده برای پررنگ کردن کدگذاری شده و فرمت آنها به فرمت مشخص شده تغییر می‌یابد.

در این مولفه چند نکته وجود دارد:

□ اگر کاربر در صفحه نمایش نتایج جستجو facetی را انتخاب کند، در این حالت به صورت پیش‌فرض، facet انتخابی کاربر پررنگ نمی‌شود مگر آنکه تغییراتی در تنظیمات مولفه Facet ایجاد شود. به عنوان مثال، اگر کاربر بخواهد نتایج جستجو را، در صفحه نمایش نتایج، به نام دانشگاه خاصی محدود کند، در حالت عادی نام این دانشگاه پررنگ نخواهد شد.

□ فرض کنید عبارت جستجوی کاربر به صورت «blue fireball in the rain» باشد. در این حالت کد پرس‌وجوی لوسن به صورت زیر خواهد بود:

content: blue OR content: fireball OR content: rain

نتایج جستجو در پایگاه داده نیز به صورت زیر است:

1- and orange as well as **blue** light and what appeared to be a **fireball**. After that no more lights and no more **rain**

2- I saw a **blue fireball** in the sky, it was an awesome **fireball**.

حال اگر کاربر به دنبال نتایج دقیق عبارت پرس‌وجو باشد، نظیر "blue fireball" در این حالت کد پرس‌وجوی لوسن به صورت زیر است:

content: blue fireball

و نتایج جستجو در پایگاه داده نیز به صورت زیر خواهد بود:

1- I saw a **blue fireball** in the sky, it was an awesome fireball.

در این حالت واژه fireball دوم پررنگ نمی‌شود.

□ از آنجا که در این مولفه باید متن سند تحلیل شود، عملکرد آن در اسناد حجیم (بیشتر از ۵۰ صفحه) کند است. برای این منظور می‌توان از توابع FastVectorHighlighter و PostingsHighlighter استفاده نمود.

در هر دو این توابع تغییراتی در فرآیند پررنگ کردن بوجود آمده است و هر دو باید به اطلاعات شاخص معکوس دسترسی داشته باشند با این تفاوت که FastVectorHighlighter نیازمند چینه متفوتی از اطلاعات شاخص معکوس بوده و لذا از ابتدا اطلاعات این شاخص را جداگانه ذخیره می‌کند، اما PostingsHighlighter مستقیماً از اطلاعات شاخص معکوس استفاده می‌کند.

در FastVectorHighlighter گام تحلیل متن حذف شده و ضروری است ویژگی‌های termOffsets, termVectors, termPositions برای فیلدهایی که باید پررنگ شوند، فعال شده باشد. اما بقیه تنظیمات مانند حالت پیش فرض است.

بر خلاف FastVectorHighlighter، تابع PostingsHighlighter از رویکرد جمله-آگاه^۱ برای قطعه‌بندی متن استفاده کرده و قطعات با «...» (سه نقطه یا ellipsis) از هم جدا می‌شوند. همچنین PostingsHighlighter برای یافتن میزان ارتباط یک قطعه با عبارت پرس‌وجو از روش BM25 استفاده می‌شود.

برای استفاده از PostingsHighlighter وجود اطلاعات مربوط به offset‌های واژه‌ها و اضافه کردن پارامتر storeOffsetsWithPositions در Schema ضروری است و از آنجا برخی از ابزارهای واژه-واژه‌سازی برای برخی از زبان‌ها نمی‌توانند اطلاعات offset‌های را گزارش کنند استفاده از این تابع محدود خواهد بود.

پارامترهای تنظیم مولفه highlighter متعدد بوده و در بخش ضمیمه (جدول ۷-۷) به برخی از آنها اشاره می‌شود.

۵-۵-۳- مولفه LAST-COMPONENT

مولفه Last-Component مولفه دیگری در فرآیند مدیریت درخواست جستجو (شکل ۳-۱۷، بخش ④) است که وظیفه پردازش نهایی فرآیند جستجوی کاربر را بر عهده دارد. این مولفه

^۱ Sentence-aware approach

اختیاری بوده و شامل ابزارهایی است که زمان زیادی را برای پردازش به خود اختصاص می‌دهند و ممکن است نتایج آنها تاثیر چندانی در نتایج بازیابی شده نداشته باشد.

در این بخش به دو ابزار مهم این مولفه شامل امکانات کنترل املاء^۱ و پیشنهاد خودکار^۲ عبارت پرس‌وجو پرداخته می‌شود که می‌تواند موجب بهبود تجربه کاربر در استفاده از سیستم بازیابی اطلاعات شوند.

کنترل املاء به مسئله کنترل نحوه نگارش عبارت پرس‌وجو پرداخته و در نتیجه سبب می‌شود تا نتایج نامناسب بازیابی نشوند. پیشنهاد خودکار نیز به دنبال فراهم آوردن امکان ارائه پیشنهاد واژگانی جایگزین به کاربر براساس نمایه‌های موجود در پایگاه داده است. به تعبیری دیگر با کمک امکان پیشنهاد خودکار بازخور فوری به کاربر درباره اعتبار عبارت پرس‌وجو داده می‌شود.

کنترل املاء

از منظر کنترل املاء چند سناریو ممکن است روی دهد:

- عبارت پرس‌وجو حاوی یک یا چند واژه نادر ست بوده و تعداد اسناد بازیابی شده صفر است. در این حالت اگر سیستم بتواند واژه جایگزین پیدا کند، بهتر آن است که به صورت خودکار جستجو را با واژه جایگزین انجام داده و پیامی حاوی توضیحات مرتبط به کاربر ارائه نماید، نظیر «... instead of ... Searched for».
- عبارت پرس‌وجو حاوی یک یا چند واژه نادر بوده و تعداد اسناد بازیابی شده بسیار کم است. در این حالت، اگر سیستم بتواند واژه جایگزین بیابد، بهتر است که توضیحی به کاربر ارائه شود نظیر «... Did you mean».
- عبارت پرس‌وجو در ست است و واژه جایگزینی نیز وجود دارد. در این حالت اگر تعداد نتایج بازیابی در دو حالت عبارت کاربر و واژه جایگزین تقریباً برابر باشد، سیستم نباید پیشنهاد(هایی) را به کاربر ارائه کند.
- عبارت پرس‌وجو حاوی واژه‌گانی است که در نمایه وجود نداشته و سیستم قادر نیست واژه جایگزینی را نیز پیشنهاد کند.

^۱ Spell-checking

^۲ Autosuggest

با کمک این سناریوها می توان دو نیازمندی را برای کنترل املاء استخراج کرد: (۱) روشی نیاز است تا واژگان جایگزین برای هر واژه در پرس وجو شناسایی شود و (۲) باید مشخص شود برای هر واژه پرس وجو و جایگزین چه تعداد نتایج بازیابی می شود.

سولر برای ارائه واژگان جایگزین از تابع Spellcheck و نمایه های موجود در پایگاه داده استفاده می کند. پارامترهای اصلی و کلان این تابع برای تنظیم در Schema در جدول ۷-۸ ضمیمه آمده است. تابع پیش فرض DirectSolrSpellChecker بوده و اطلاعات خود را مستقیماً از نمایه اصلی سیستم دریافت می کند. سه پارامتر اصلی این تابع عبارتند از:

□ field: تعیین این که از اطلاعات مربوط به چه فیلدهایی برای ارائه پیشنهاد خود کار استفاده شود؛

□ distanceMeasure: تعیین نحوه محاسبه فاصله بین واژه پرس وجوی کار بر و واژه پیشنهادی. فاصله پیش فرض فاصله دامرا-لون اشتاین است که در بخش تطابق فازی بدان پرداخته شد. البته می توان از فواصل دیگری نظیر موارد مطرح شده در جدول ۷-۲ نیز استفاده کرد؛

□ accuracy: مقداری بین صفر و یک بوده و برای تعیین میزان دقت پیشنهادها تنظیم می شود.

این تابع از ابزارهای تحلیل متن StopFilterFactory برای حذف ایست واژه ها، UAX29URLEmailTokenizer برای حذف فاصله ها و نشانه ها و درعین حال نگهداری آدرس رایانامه ها و آدرس های اینترنتی، از LowerCaseFilter برای کوچک کردن حروف بزرگ، از ASCIIIFoldingFilter برای تبدیل حروف به کدهای ASCII و از EnglishPossessiveFilter برای حذف مضاف الیه های زبان انگلیسی استفاده می کند.

تابع WordBreakSolrSpellchecker دیگر تابع مورد استفاده در کنترل املاء است و در مواردی کاربرد دارد که واژگان عبارت پرس وجو بدون فاصله در کنار یکدیگر قرار گرفته باشند. این ماژول این نوع واژگان را به بخش های مختلف تجزیه کرده و سپس آنها را برای ارائه پیشنهادها جایگزین در نمایه جستجو می کند.

پیشنهاد خودکار

وظیفه اصلی ابزار پیشنهاد خودکار جلوگیری از اشتباهات کاربر در هنگام وارد کردن عبارت پرس و جو بوده و می تواند نقش موثری در کاهش هزینه های پردازشی ایفا نماید (پیاده سازی و بکارگیری آن در سیستم های مبتنی بر تلفن همراه ضروری است). سرعت در ارائه پیشنهادها و در نظرگیری رتبه واژگان (بر اساس TF) برای پیشنهاد از نیازمندی های چنین ابزاری است.

یکی از توابع پر کاربرد در این ابزار (برای داده های حجیم) FSTLookupFactory است که از ساختار اطلاعات خودکار سازی و وضعیت محدود^۱، که نوعی نگاشت بین دو مجموعه نمادهاست، استفاده می کند و لغت نامه ای از پیشنهادها را می سازد.

علی رغم شباهت های بین دو ابزار کنترل املاء و پیشنهاد خودکار، تفاوت های کلیدی نیز وجود دارد: مولفه کنترل املاء نیازمند عبارت کامل پرس و جوی کاربر است در حالی که در پیشنهاد خودکار، در حین وارد کردن عبارت پرس و جو پیشنهادهایی به کاربر ارائه می شود. همچنین، نوع تابع فاصله یکی از پارامترهای تاثیر گذار در کارایی کنترل املاء است ولی مشهور بودن واژگان در پایگاه داده عامل تاثیر گذار در پیشنهاد خودکار است.

۳-۵-۶ مولفه RESPONSE WRITER

پس از بازیابی اطلاعات مرتبط، نتایج بایستی به کاربر نمایش داده شوند. این امر از طریق مولفه Response Writer (شکل ۳-۱۴ و شکل ۳-۱۷) انجام می شود. در این مولفه پاسخ به چند سوال ضروری است: «خروجی در چه فرمتی و از چه نوعی باید باشد؟»، «چه فیلدهایی باید برای کاربر نمایش داده شوند؟»، «هر صفحه باید شامل چه تعداد نتایج باشد؟» و «نحوه مرتب کردن نتایج چگونه باید باشد؟». تنظیمات مربوط به پاسخ های این سوالات در مولفه نمایش نتایج انجام می شود.

فرمت و نوع نتایج خروجی^۲ (wt) می تواند یکی از موارد گزارش شده در جدول ۳-۷ باشد.

جدول ۳-۷- پارامتر نوع نتایج خروجی

نوع نتایج خروجی (wt)	کلاس مورد استفاده در سولر
CSV	CSVResponseWriter
json	JSONResponseWriter
php	PHPResponseWriter

^۱ Finite state automaton (FST)

^۲ Writer type

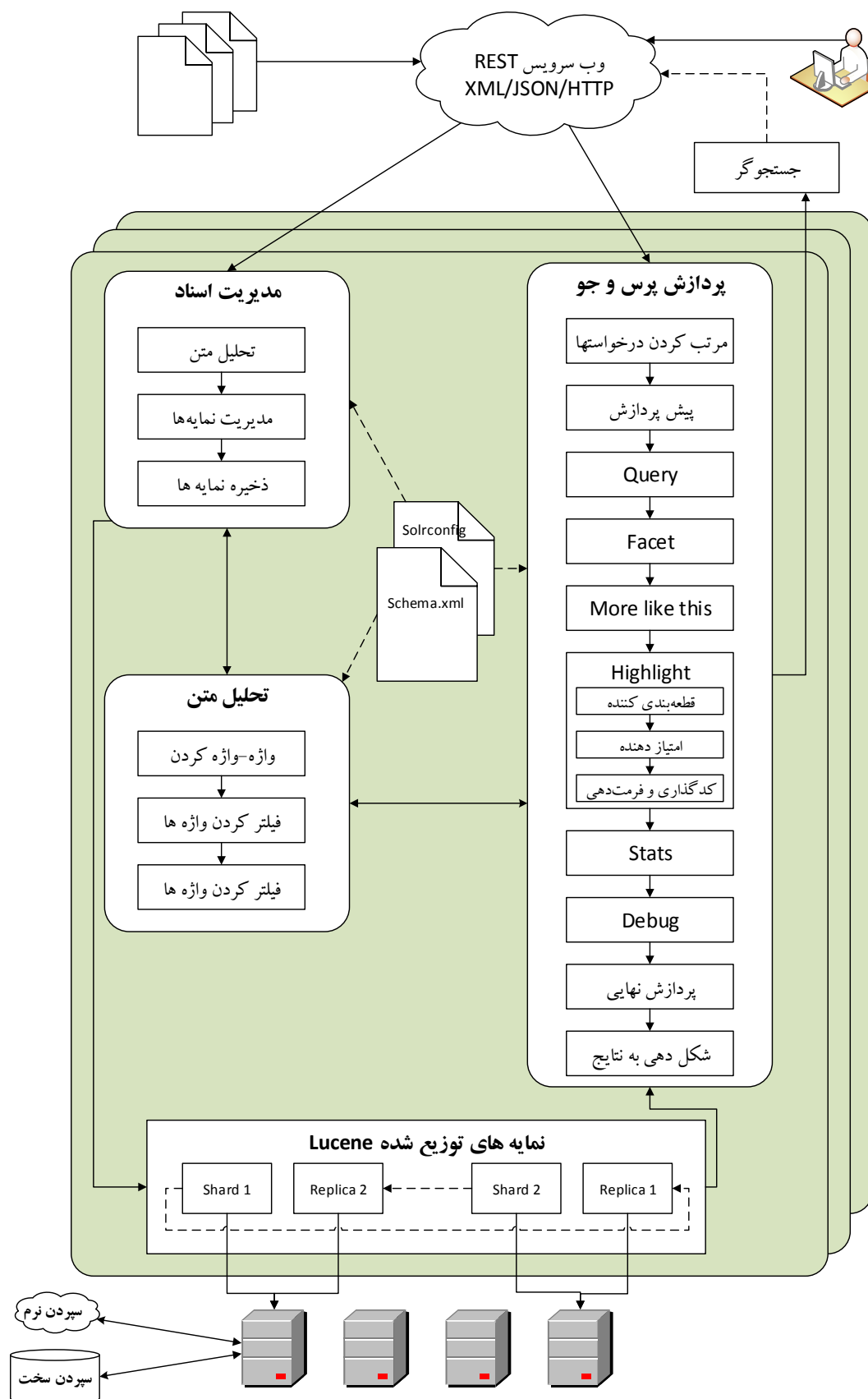
PHPSerializedResponseWriter	phps
PythonResponseWriter	python
RubyResponseWriter	ruby
XMLResponseWriter	xml
XSLTResponseWriter	xslt
BinaryResponseWriter	javabin

تعداد فیلدهایی که برای کاربر نمایش داده می شود بر سرعت نمایش نتایج و حجم اطلاعات تبادل شده موثر است. یکی از راه های بهبود سرعت نمایش نتایج، استفاده از صفحات متعدد است. به عبارت دیگر، در صفحه نخست نمایش نتایج معمولاً نتایج با بالاترین امتیاز نمایش داده شده و کاربر با جابجا شدن در صفحات می تواند سایر نتایج را ببیند. سولر برای مرتب کردن نتایج، از میزان انطباق هر سند با عبارت پرس وجو استفاده می کند و نتایج را در حالت پیش فرض به صورت نزولی نمایش می دهد. اگر دو سند دارای امتیازهای یکسان باشند، براساس ID سند در لو سن مرتب می شوند. البته برای مرتب کردن نتایج باید حافظه زیادی به جستجو تخصیص داد و این امر در صورتی که کاربران زیادی به صورت همزمان از سیستم استفاده کنند می تواند مشکل زا باشد.

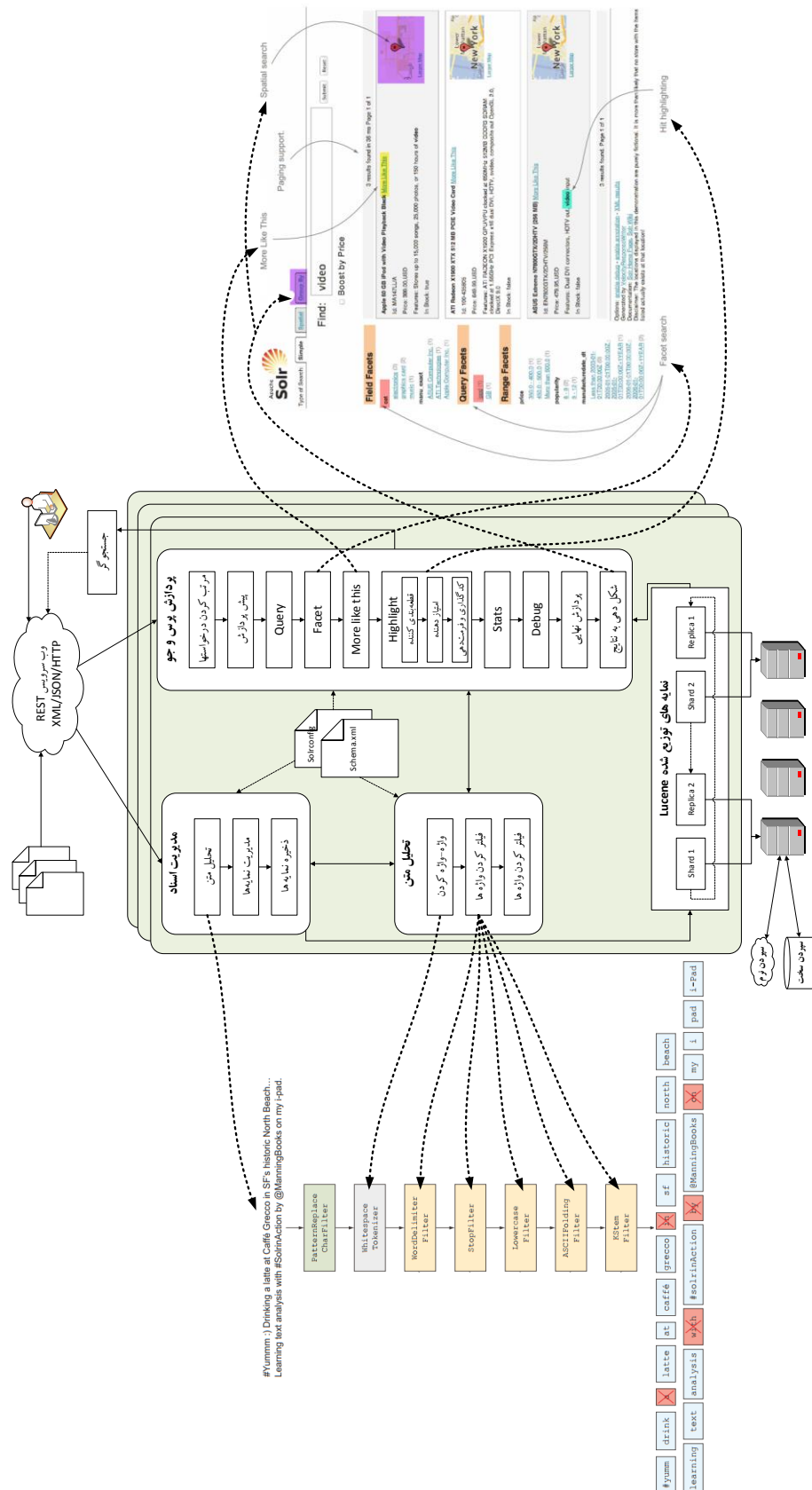
۳-۶ جمع بندی

فرآیند آنچه در سولر رخ می دهد را می توان در شکل ۳-۲۳ خلاصه نمود. در طی این فرآیند اسناد ورودی به پایگاه داده توسط لو سن نمایه شده و ذخیره می شوند. درخواست های جستجوی کاربر نیز تحلیل شده و نتایج براساس تنظیمات مشخص شده در Solrconfig و Schema بازیابی شده و به کاربر نمایش داده می شوند. فرآیند تحلیل متن جزء کلیدی سولر و لو سن بوده و به عنوان ابزاری در زیرسیستم های مدیریت اسناد و پردازش پرس وجو استفاده می شود. قابلیت تحلیل های چند زبانی نیز از دیگر قابلیت های سولر و لو سن است که برای اجرای آن سیستم نیازمند معماری خاصی است. شکل ۳-۲۴ نیز این فرآیند و عملکرد آن را در قالب یک مثال نمایش می دهد.

علاوه بر این، به غیر از موارد مطرح شده در این فصل، ملاحظات و تنظیمات دیگری نیز در سولر وجود دارد که در این گزارش در قالب مباحث پیشرفته و در فصل بعد مورد بررسی قرار خواهند گرفت.



شکل ۳-۲۳- فرآیند پردازش اسناد و پرس و جو در سولر



شکل ۳-۲۴- فرآیند پردازش اسناد و پرس و جو در سولر در قالب مثال

فصل چهارم

مباحث پیشرفته در موتور جستجوی

سولر

۴-۱ مقدمه

در فصل پیش توضیحاتی در مورد موارد و نکات پایه‌ای سولر ارائه شد. در این فصل به بررسی برخی نکات پیشرفته‌تر شامل نحوه فراخوانی و استفاده از داده‌های منابع بیرونی، عملیات جستجوی پیچیده، نحوه بهبود ارتباط سند با عبارت پرس‌وجو و شخصی‌سازی جستجو اختصاص دارد.

۴-۲ استفاده از داده‌های منابع بیرونی

داده جزء بنیادین اغلب سیستم‌ها، خصوصاً سیستم‌های جستجو و بازیابی اطلاعات، است. این داده‌ها می‌توانند از منابع درونی سولر (ذخیره و نمایه شده در لوسن) و یا از منابع بیرونی تامین شوند.

سه مکانیزم اصلی برای استفاده از داده‌های منابع بیرونی (بیرون از نمایه سولر) وجود دارد: (۱) استفاده از دستوراتی در توابع پردازش پرس‌وجو برای دسترسی به منابع بیرونی؛ (۲) استفاده از ارتباطات بین اسناد در هسته‌های مجزای سولر فعال در یک سرور و یا (۳) استفاده از فیلد خاصی بانام ExternalFileField که می‌تواند مقادیر را از فایل بیرونی دریافت کند.

فیلد ExternalFileField امکان بروزرسانی یک فیلد خاص را مستقل از فرآیند نمایه‌سازی در سولر فراهم می‌کند. این امر با تعیین فایل متنی حاوی نگاشت ID هر سند در سرور سولر با مقدار آن در فایل بیرونی محقق می‌شود. به تعبیری دیگر سولر با منبع بیرونی از طریق مشخصه سند

ارتباط دارد. تعریف این فیلد زمانی مفید است که بخواهیم فیلد خاصی را برای تعداد زیادی از اسناد بروزرسانی کنیم بدون آنکه بقیه فیلدها را تغییر دهیم.

ExternalFileField مانند سایر فیلدها (البته با پارمترهایی بیشتر) باید در Schema (به عنوان مثال به صورت زیر) تعریف شود:

```
<fieldType name="popularity" keyField="id" defVal="0"
stored="false" indexed="false" class="solr.ExternalFileField"
valType="pfloat"/>
```

در این مثال، فیلد popularity مقدار خود را برای هر سند از فایل بیرونی می‌خواند که در دایرکتوری نمایه سولر در سرور قرار داده شده است. اسم این فایل باید به صورت external_{field} یا external_{field}.* باشد که در آن {field} نام ExternalFileField است که در Schema تعریف شده است. بنابراین برای مثال فوق نام این فایل خارجی می‌تواند external_polularity.2013_11_03_09_00 و یا external_popularity.txt باشد. این فایل می‌تواند دارای اطلاعاتی نظیر زیر باشد که به عنوان مثال در آن مقدار معرفیت doc123 برابر ۲۲ است:

```
doc123=22
doc456=15.7
doc789=19.4
```

فیلد ExternalFileField خصوصیات مهمی دارد:

- جستجو براساس این فیلد امکان پذیر نیست. به بیانی دیگر، سولر قادر نیست براساس مقادیر این فیلد جستجو کند.
- لزومی ندارد همه اسناد موجود در پایگاه داده دارای اطلاعاتی مخصوص به خود در فایل بیرونی باشند و برعکس. به عنوان مثال، سند doc709 ممکن است مقداری برای popularity نداشته باشد و یا سند doc456 در پایگاه اسناد وجود نداشته باشد.
- در حال حاضر فیلد ExternalFileField تنها می‌تواند از نوع float باشد و این امر می‌تواند محدودیت‌هایی را ایجاد نماید. برای رفع این مشکل می‌توان پس از بازخوانی مقادیر خارجی توابعی را تعریف کرد تا نوع اطلاعات را تغییر دهد.
- مقادیر موجود در فایل بیرونی تنها یکبار و در هنگام اولین بار استفاده در حافظه بارگذاری می‌شوند. در نتیجه اگر تغییری در این مقادیر ایجاد شود سولر متوجه آن نخواهد شد. برای

رفع این مشکل می‌توان event listener تعریف کرد تا به محض رخداد تغییر، فایل بیرونی مجدداً بارگذاری شود.

۴-۳ عملیات جستجوی پیچیده

سولر این امکان را برای کاربران خود فراهم می‌آورد تا علاوه بر انجام جستجو براساس عبارت پرس‌وجو، توابعی را (به صورت از پیش طراحی شده، تودرتو^۱ و یا جدید) اجرا کرده و از نتایج آن به صورت عبارت پرس‌وجوی جدید و یا پارامتری برای تعیین میزان ارتباط نتایج بازیابی شده استفاده کند. برخی از توابع ریاضی در جدول ۷-۱ ضmann، توابع فاصله در جدول ۷-۲ ضmann و برخی از توابع منطق بولی در جدول ۷-۳ ضmann گزارش شده است. اما به جهت اهمیت توابع محاسبه میزان ارتباط سند و عبارت پرس‌وجو، برخی از مهمترین این توابع در جدول ۴-۱ آورده شده‌اند.

جدول ۴-۱- توابع محاسبه میزان ارتباط سند و عبارت پرس‌وجو

توضیحات	دستور (syntax) تابع
تعداد اسناد حاوی واژه term در فیلد fieldName	<code>docfreq(fieldName, term)</code>
محاسبه idf برای value در فیلد fieldName	<code>idf(fieldName, value)</code>
تعداد اسناد نمایه شده شامل اسناد حذف شده‌ای که از حافظه اصلی (durable storage) پاک نشده‌اند	<code>maxdoc()</code>
نرم محاسبه و ذخیره شده در نمایه برای فیلد fieldName	<code>norm(fieldName)</code>
تعداد اسناد نمایه شده بدون در نظرگیری اسناد حذف شده‌ای که از حافظه اصلی پاک نشده‌اند	<code>numdocs()</code>
امتیاز subquery برای همه اسناد مرتبط با subquery و امتیاز defaultScore برای اسناد نامرتبط	<code>query(subquery, defaultScore)</code>
تعداد واژه‌های (token) نمایه شده در field (این دو دستور عملکردی مشابه دارند)	<code>sumtotaltermfreq(field)</code> <code>sttf(field)</code>
تعداد دفعاتی که واژه term در فیلد fieldname تکرار شده است	<code>termfreq(fieldName, term)</code>

^۱ Nested

فراوانی کلمه (tf) برای واژه term در فیلد fieldname	tf(fieldName, term)
براساس شباهت آن فیلد	
تعداد دفعاتی که واژه term در کل نمایه تکرار می‌شود	totaltermfreq(fieldName, term)
(این دو دستور عملکردی مشابه دارند)	ttf(fieldName, term)

نمونه‌ای از نحوه بکارگیری توابع فوق در برای یافتن عبارت «microsoft office» در فیلد content در ادامه آمده است. در بخش ابتدایی این دستور ابتدا فراوانی کلمه (TF) برای واژه «microsoft» در فیلد content و سپس فراوانی معکوس سند (IDF) برای همین واژه در فیلد content محاسبه شده و نتایج این دو محاسبه در هم ضرب می‌شوند. این عمل برای واژه «office» در فیلد content نیز تکرار شده و نتیجه کلی ضرب TF و IDF برای این دو واژه با هم جمع خواهد شد.

```
fq={!cache=false}content:"microsoft office"&
q={!func}sum(
    product(
        tf(content, "microsoft"),
        idf(content, "microsoft")
    ),
    product(
        tf(content, "office"),
        idf(content, "office")
    )
)
```

۴-۴ ارتقاء ارتباط سند با عبارت پرس‌وجو

همانطور که در بخش ۴-۲-۳، ارتباط، ذکر شد، محاسبه ارتباط سند و عبارت پرس‌وجو و رتبه‌دهی به اسناد بازیابی شده، از مزیت‌های موتورهای جستجو نسبت به پایگاه‌های داده است.^۱ در این بخش، روش معمول سولر برای این منظور بیان شد اما برای ارتقاء امتیاز محاسبه شده و بهبود نتایج بازیابی شده راهکارهایی وجود دارد:

□ راهکار ۱- وزن‌دهی به فیلدی معین: اگر فیلدی حاوی اطلاعاتی باشد که کاربران اغلب به دنبال آن هستند می‌توان در محاسبه ارتباط به این فیلد وزن بیشتری اختصاص داد. برای این پیاده‌سازی این راهکار چند روش وجود دارد:

^۱ پایگاه‌های داده قادرند در بین داده‌های خود جستجو کنند اما نمی‌توانند نتایج را رتبه‌بندی نمایند.

- وزن‌دهی به فیلد خاص در هنگام نمایه‌سازی. مثال زیر را در نظر بگیرید:

```
<add>  
<doc>  
<field name="id">1</field>  
<field name="restaurant_name" boost="10.0">Red Lobster</field>
```

مشکل این روش آن است که اگر بخواهیم مقدار وزن را با سعی و خطا تعیین کنیم نیاز است تمامی اسناد مجدداً نمایه شوند و لذا این روش در اغلب موارد مناسب نخواهد بود.

- وزن‌دهی به فیلد خاص در زمان پردازش جستجو: این راهکار انعطاف‌پذیری بیشتری نسبت به روش نخست دارد. مثال زیر را در نظر بگیرید:

```
q=red lobster&  
qf=restaurant_name^10 description
```

در این مثال عبارت پرس‌وجو «red lobster» بوده و وجود این عبارت در فیلد نام رستوران وزنی ۱۰ برابر وجود آن در توضیحات خواهد داشت.

- راهکار ۲- وزن‌دهی به واژگان خاص: در این راهکار به واژگان خاص در عبارت پرس‌وجوی کاربر وزن بیشتری داده می‌شود بجای آنکه به کل عبارت و یا فیلد مورد جستجو وزن داده شود. البته این روش را می‌توان با راهکار ۱ نیز ادغام نمود. به عنوان مثال:

```
q=red^2 lobster^8&  
qf=restaurant_name^10 description
```

در مثال فوق وزن واژه «red» اگر در فیلد نام رستوران قرار داشته باشد برابر خواهد بود با $2 \times 10 = 20$ و وزن واژه «lobster» اگر در فیلد نام رستوران قرار داشته باشد برابر خواهد بود با $8 \times 10 = 80$ (ضرب وزن‌ها بدین صورت برای تمامی راهکارهای وزن‌دهی در سولر برقرار است).

- راهکار ۳- وزن‌دهی به واژگان سند: در این راهکار، واژگان سند و نقش آنها در جمله شناسایی شده و به نقش‌های خاص (مثلاً اسم) وزن بیشتری در هنگام تحلیل متن داده می‌شود. این امر از طریق Payloads در سولر محقق می‌شود که به صورت پیش‌فرض

فعال نیست. Payloads آرایه‌ای است که در کنار واژگان نمایه شده در سند قرار گرفته و می‌توان اطلاعاتی را درباره آن واژه در آن ذخیره نمود^۱.

□ راهکار ۴- وزن‌دهی با کمک تابع: استفاده از توابع در وزن‌دهی می‌تواند سبب تغییر در امتیازهای اسناد شود. به عنوان مثال:

```
q=restaurant_name:(Burger King) AND  
_query_:"{!func}recip(geodist(location,37.765,-122.43),1,10,1)"
```

در این تابع رستوران‌هایی که به مکان مشخص شده توسط geodist نزدیک‌تر باشند وزنی ۱۰ داده می‌شود و بنابراین رستوران‌های نزدیک‌تر امتیاز بالاتری دریافت خواهند کرد. بنابراین تابع امتیازدهی به صورت زیر محاسبه می‌شود:

```
score("Burger") + score("King") + LocationProximityBoost
```

□ راهکار ۵- وزن‌دهی به نزدیکی واژگان: اگر در عبارت پرس‌وجوی کاربر دو واژه در کنار هم قرار گرفته باشند، اسنادی باید امتیاز بالاتری بگیرند که در آنها این دو واژه نزدیک‌تر به هم هستند. برای اعمال این روش دو راهکار وجود دارد:

○ بکارگیری eDisMax query parser و استفاده از پارامتر pf نظیر:

```
defType=edismax&  
q=big data analytics&  
qf=content&  
pf=content~3^10
```

از پارامترهای دیگر این تجزیه‌گر، pf2 و pf3 نیز می‌توان استفاده کرد. اگر pf2=content، اسناد دارای عبارت‌های «big data» و یا «data analytics» وزن بیشتری می‌گیرند و اگر pf3=content، اسناد دارای عبارت «big data analytics» وزن بیشتری می‌گیرند.

○ بکارگیری Lucene query parser و استفاده از دستور ... term1 term 2

customer"~2^10 service"~2^10. به عنوان مثال، عبارت پرس‌وجوی "customer service" را بازیابی می‌کند که در آنها دو واژه customer و service با ۲ فاصله از هم قرار گرفته باشد و وزن این اسناد برابر ۱۰ خواهد بود. ممکن است کاربر بخواهد اسناد حاوی این دو واژه را بیابد اما اسنادی اهمیت

^۱ برای اطلاعات بیشتر به <https://issues.apache.org/jira/browse/SOLR-1485> مراجعه نمایید.

بیشتر داشته باشند که در آنها این دو واژه در فاصله نزدیک به هم قرار دارند. در این حالت، نمونه‌های پرس‌وجوی زیر می‌توانند نتایج مورد نظر کاربر را حاصل کنند:

1. `q=+customer +service OR "customer service"`
2. `q=+customer +service OR "customer service"~2^10`
3. `q=customer OR service +{some other terms}+(*:* OR "customer service"^100)`
4. `q=+customer +service +representative OR "customer service representative"~10000^10`

در حالت ۱، اسنادی که حاوی این دو واژه باشند بازیابی می‌شوند و اسنادی که در آنها این دو واژه در کنار یکدیگر قرار داشته باشند وزنی دو برابر دریافت می‌کنند. در حالت ۲، اسنادی که حاوی این دو واژه باشند بازیابی می‌شوند و اسنادی که در آنها این دو واژه با ۲ فاصله از هم قرار داشته باشند وزنی ۱۰ برابر دریافت می‌کنند. در حالت ۳، اسنادی که حاوی یکی از این دو واژه باشند بازیابی می‌شوند و اسنادی که در آنها این دو واژه در کنار یکدیگر قرار داشته باشند وزن ۱۰۰ برابر دریافت می‌کنند. در حالت ۴ نیز اسنادی بازیابی می‌شوند که حاوی سه واژه `customer`، `service` و `representative` بوده و مکان قرارگیری و فاصله آنها از یکدیگر اهمیت چندانی ندارد. اما هر چه به هم نزدیک‌تر باشند امتیاز بیشتری دریافت خواهند کرد.

□ راهکار ۶- ارتقاء ارتباط اسناد مهم: یکی از راه‌های بهبود امتیاز ارتباط سند و عبارت پرس‌وجوی کاربر، در نظر گرفتن اهمیت اسناد در بازیابی اطلاعات است. یکی از راهکارها برای اجرای این روش تخصیص وزن خاص به سند در زمان نمایه‌سازی است (این وزن برای تمامی فیلدهای آن سند اعمال می‌شود). اشکال این راهکار، مانند راهکار نخست روش ۱، آن است که اگر بخواهیم مقدار وزن را با آزمایش تعیین کنیم نیاز است تا تمامی اسناد مجدداً نمایه شوند. راهکار دیگر تخصیص وزن به اسناد در زمان جستجو است که این امر می‌تواند از طریق در نظر گیری محبوبیت سند اجرایی شود.

شخصی سازی سیستم های بازیابی اطلاعات با ارائه پیشنهادهای مفید به کاربران کمک می کند تا تعامل بهتر و کاراتری با سیستم داشته باشند. سیستم های بازیابی اطلاعات را می توان نوعی از سیستم های پیشنهاددهنده دانست که در آنها براساس عبارت پرس و جوی کاربر و یا مشخصات سند انتخاب شده، اسناد مشابه به کاربران ارائه می شود. سولر این قابلیت را دارد تا از طریق راهکارهای زیر به عنوان یک سیستم جستجوی پیشنهاددهنده به کاربران خود کمک کند:

□ راهکار ۱- تطابق براساس ویژگی ها^۱: برای پیاده سازی این راهکار، اسناد در پایگاه داده باید دارای برچسب هایی بوده و سولر بر اساس پروفایل رفتاری کاربر و برچسب های اسناد پیشنهادهای لازم را به وی ارائه می کند.

□ راهکار ۲- تطابق سلسله مراتبی^۲: در این روش ویژگی های اسناد و پروفایل رفتاری کاربر به صورت سلسله مراتبی و براساس اهمیت مرتب شده و پیشنهادها براساس این ترتیب ارائه می شود.

□ راهکار ۳- مشابه این...^۳: در این راهکار نیازی به دسترسی به پروفایل کاربر و یا برچسب های از قبل تعیین شده نیست. بلکه در این روش تشابه دو سند به یکدیگر محاسبه شده و مشابه ترین سند با سند انتخابی کاربر به وی نمایش داده می شود. More Like This Handler در سولر وظیفه استخراج واژگان مهم سند را بر اساس روش TF-IDF بر عهده داشته و به صورت خود کار از این واژگان به عنوان کلیدواژه برای یافتن اسناد مشابه استفاده می کند. برای استفاده از این handler گام های زیر باید طی شوند:

گام ۱: فعال سازی در Schema با دستور:

```
<requestHandler name="/mlt" class="solr.MoreLikeThisHandler" />
```

گام ۲: تعیین فیلدهایی که این handler برای مشابهت یابی در آنها باید واژگان مهم را استخراج کند توسط پارامتر mlt.fl (اگر جزئیات و اطلاعات آماری این واژگان نیز نیاز باشد با کمک دستور mlt.interestingTerms=details اطلاعات مربوطه بازگردانده می شوند).

^۱ Attribute-based matching

^۲ Hierarchical matching

^۳ Like this ...

گام ۳: تنظیم کردن مقدار true برای پارامترهای termVectors و stored در Schema برای فیلدهایی که برای مشابهت یابی استفاده می‌شوند.

گام ۴: تعیین وزن برای واژگان استخراج شده. این وزن‌ها به صورت پیش فرض برابر ۱ هستند اما می‌توان این وزن را از طریق پارامتر `mlt.boost=true` به مقدار TF-IDF هر واژه تغییر داد.

ممکن است شرایطی وجود داشته باشد که در طی آن لیست واژگان مهم دارای نویز باشد (مانند ایست‌واژه). برای جلوگیری از این اتفاق می‌توان از تحلیل ادات متن^۱ و استخراج اسم‌ها از سند استفاده کرد. همچنین به منظور افزایش سرعت می‌توان میزان شباهت اسناد به یکدیگر را به صورت آفلاین محاسبه کرد، هرچند که در این حالت فضای زیادی برای ذخیره‌سازی مورد نیاز خواهد بود.

□ راهکار ۴- تطابق براساس مفهوم^۲: روش‌های قبلی بر اساس ویژگی‌های یک سند و یا واژگان مهم آن سند، پیشنهادهایی را به کاربر ارائه می‌کردند. اما سولر این قابلیت را دارد که اسناد را خوشه‌بندی نماید (این خوشه‌بندی در سولر به صورت offline انجام شده و در بازه‌های زمانی مشخصی بروزرسانی می‌شود). این امر با کمک مولفه clustering محقق می‌شود. پس از بازیابی نتایج بر اساس عبارت پرس‌وجوی کاربر، سولر خوشه‌هایی که اسناد با بالاترین امتیاز بدان‌ها تعلق دارند را یافته و سایر کلید واژگان را شناسایی می‌کند. به عنوان مثال ممکن است کاربر عبارت «Net Jobs» را جستجو کند اما در نتایج، علاوه بر Net Jobs، نتایج مربوط به واژگان developer, software engineer, C# و net developer نیز دیده می‌شود (البته باید دقت داشت اسنادی که حاوی عبارت پرس‌وجوی کاربر هستند وزن بالاتری داشته باشند). در این روش تعداد نتایج بازیابی شده تغییر نخواهد کرد بلکه ترتیب نمایش آنها متفاوت است. استفاده از این روش باعث بهبود دقت شده و بازخوانی نیز بهبود یافته و یا ثابت می‌ماند.

برای اجرای این روش باید ابتدا ClusteringComponent و روش خوشه‌بندی مورد استفاده در Schema تعریف شده و سپس در Search Handler فعال شود (سولر به

^۱ Part of speech (POS)

^۲ Concept-based matching

صورت پیش فرض از موتور خوشه‌بندی Carrot² برای خوشه‌بندی اسناد بازیابی شده استفاده می‌کند^۱:

```
<searchComponent name="clustering" enable="true"
  class="solr.clustering.ClusteringComponent">
  <lst name="engine">
    <str name="name">default</str>
    <str name="carrot.algorithm">
      org.carrot2.clustering.lingo.LingoClusteringAlgorithm
    </str>
    <str name="carrot.resourcesDir">clustering/carrot2</str>
  </str>
</searchComponent>
<requestHandler name="/clustering" enable="true"
  class="solr.SearchHandler">
  <lst name="defaults">
    <bool name="clustering">true</bool>
    <str name="clustering.engine">default</str>
    <bool name="clustering.results">true</bool>
    <str name="fl">*,score</str>
  </lst>
  <arr name="last-components">
    <str>clustering</str>
  </arr>
</requestHandler>
```

این راهکار زمانی مناسب است که عبارت پرس و جوی کاربر محدود بوده و موتور جستجو قصد داشته باشد نتایج مناسب‌تری را به کاربر نمایش دهد.

□ راهکار ۵- انطباق جغرافیایی^۲: از این روش برای مواقعی استفاده می‌شود که موقعیت مکانی فعلی کاربر مهم است و یا نتایج باید بر اساس موقعیت مکانی مورد نظر کاربر مرتب شوند.

□ راهکار ۶- پالایش همکارانه^۳: این روش یکی از پرکاربردترین روش‌ها در حوزه یادگیری ماشین برای ارائه پیشنهاد(ها) است. یکی از فرض‌های اصلی در این روش آن است که کاربران مشابه، علائق مشابهی نیز دارند و بنابراین می‌توان اسناد مشابه و یا مکمل را به کاربر فعلی، بر اساس انتخاب‌ها و یا ترجیحات سایر کاربران، پیشنهاد داد. درواقع، این روش نیازی به اطلاع از موضوع و محتوای سند ندارد و صرفاً براساس رفتار کاربران

^۱ <https://project.carrot2.org/>

^۲ Geographical matching

^۳ Collaborative filtering

بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی □ ۸۰

پیشنهادهایی را ارائه می‌کند. شکل ۴-۱ و شکل ۴-۲ نمایی از سیستم پیشنهاد دهنده وبسایت آمازون را برای کالاهای مشابه و مکمل نشان می‌دهند.



شکل ۴-۱- سیستم پیشنهاد دهنده آمازون برای کالاهای مشابه



شکل ۴-۲- سیستم پیشنهاد دهنده آمازون برای کالاهای مکمل

در شکل ۴-۱، کاربری که تمایل به خرید فصل نخست سریال The Big Bang Theory را دارد با احتمال زیاد فصل دوم و سوم این سریال را نیز خواهد خرید. به عبارت دیگر مشاهده رفتار کاربران ارتباط یا همبستگی بین اسناد را مشخص می‌کند. در شکل ۴-۲ نیز پیشنهاد کالاهای مکمل بر اساس کالاهای انتخابی کاربر به وی ارائه می‌شود (کالاها با کالای انتخابی کاربر ارتباط دارند اما متفاوتند).

به تعبیری دیگر، این روش از دانش جمعی و اطلاعات سایر کاربران برای تنظیم پارامترهای جستجو در الگوریتم‌های پیشنهاددهنده استفاده کرده و با کمک امتیاز ارتباط کسب شده توسط روش TF-IDF، رتبه اسناد بازیابی شده را تغییر می‌دهد. نکته دیگر در این روش آن است که با کمک آن می‌توان با استفاده از رفتار سایر کاربران پیشنهاداتی را درباره اسناد مشابه به کاربر فعلی ارائه داد و یا آنکه بر اساس اسناد و نحوه مشاهده آنها توسط کاربران، کاربران مشابه را به هم معرفی نمود. برای این منظور باید دو جدول مانند شکل

۳-۴ برای ارتباط بین کاربران و اسناد محاسبه شده و با کمک آنها پیشنهادهاى لازم ارائه شود.

Document	"Users who bought this product" field	Term	Documents
doc1	user1, user4, user5	user1	doc1, doc5
doc2	user2, user3	user2	doc2
doc3	user4	user3	doc2
doc4	user4, user5	user4	doc1, doc3, doc4, doc5
doc5	user4, user1	user5	doc1, doc4
...	...	...	...

شکل ۴-۳- جدول ارتباط بین کاربران و اسناد

□ راهکار ۷- راهکار ترکیبی: ترکیب راهکارهای موجود می تواند سبب بهبود ارتباط بین سند و عبارت پرس جوی کاربر شود. به عنوان نمونه:

با ترکیب راهکار «مشابه این...» و راهکار «تطابق براساس مفهوم» می توان شباهت های بین اسناد را براساس کلمات کلیدی خوشه ها محاسبه نمود. این ترکیب می تواند نتایج false positive زیادی را به همراه داشته باشد اما اگر از فیلتر مبتنی بر دسته بندی (وزن دهی) استفاده شود نتایج بهتر خواهد شد.

همچنین ترکیب راهکار «پالایش همکارانه» با راهکار «تطابق براساس ویژگی ها» یا به عبارتی ترکیب اطلاعات مربوط به رفتار کاربران با ویژگی های اسناد، سیستم پیشنهاددهنده قدرتمندتری را می توان ایجاد نمود. علاوه براین، اضافه کردن ویژگی های جغرافیایی به هریک از روش ها نتایج بهتری را فراهم آورده و می تواند سبب افزایش کارایی شود.

۴-۶ جمع بندی

هدف از نگارش این فصل، بررسی مباحث پیشرفته تر در موتور جستجوی سولر بود. بدین منظور مباحثی در خصوص منابع داده بیرونی و نحوه فراخوانی آن ارائه شد. همچنین عملیات جستجوی پیچیده و نحوه بهبود ارتباط سند با عبارت پرس وجو بررسی گردید و از آنجا که سیستم های جستجو نوعی از سیستم های پیشنهاددهنده هستند، درخصوص شخصی سازی آنها نیز مباحثی مطرح شد.

فصل پنجم

موتور جستجوی گنج

۵-۱ مقدمه

در فصل‌های پیشین به نکات پایه‌ای سولر و نیز برخی نکات پیشرفته‌تر شامل نحوه فراخوانی و استفاده از داده‌های منابع بیرونی، عملیات جستجوی پیچیده، نحوه بهبود ارتباط سند با عبارت پرس وجو و شخصی سازی جستجو اختصاص داشت. در این فصل نگاهی خواهیم داشت به آنچه در موتور جستجوی گنج رخ می‌دهد. برای این منظور از دو سند اصلی این سامانه، Schema و Solrconfig استفاده خواهد شد و توابع مورد استفاده در هریک تشریح می‌گردد.

۵-۲ سند SOLRCONFIG سامانه گنج

یکی از اسناد مهم برای تنظیم سیستم بازیابی اطلاعات، سند Solrconfig است. این سند برای سامانه گنج در ادامه تشریح می‌شوند.

متن سند Solrconfig سامانه گنج به شرح زیر است. به جهت تسهیل در تفسیر سند، توضیحات مربوط به هر بخش در ادامه آن آمده است. همچنین اصل سند در ضمیمه آورده خواهد شد.^۱

۵-۳ سند SCHEMA سامانه گنج

^۱ به جهت اهمیت این سند، متن آن در این گزارش آورده نشده است.

همانگونه که ذکر شد، ساختار نمایه سازی و نوع-فیلد و فیلدهای مورد استفاده در سند Schema تعیین و ذخیره می شود. متن این سند برای سامانه گنج و تشریح آن در ادامه آورده شده است^۱.

۵-۴ جمع بندی

این فصل به بررسی دو سند اصلی سامانه گنج، Schema و Solrconfig اختصاص داشت و تلاش شد تا آنچه در موتور جستجوی گنج رخ می دهد بررسی شود. براساس مطالب عنوان شده، کاستی هایی وجود دارد که برای رفع آنها و بهبود عملکرد سامانه فعالیت ها و یا پژوهش هایی باید برنامه ریزی شوند. به این موارد در فصل بعد پرداخته خواهد شد.

^۱ به جهت اهمیت این سند، متن آن در این گزارش آورده نشده است.

فصل ششم

راهکارها و روش‌های بهبود

۶-۱ مقدمه

با بهره‌گیری از مباحث مطرح شده در فصل قبل (وضعیت فعلی) می‌توان پیشنهادها و راهکارهایی را به منظور بهبود عملکرد سامانه گنج و بازیابی اطلاعات در زبان فارسی و برای این سامانه ارائه داد. این فصل به این موارد اختصاص دارد.

۶-۲ راهکارهای بهبود بازیابی اطلاعات در سامانه گنج

طبق نیازمندی‌های عمومی سامانه‌های بازیابی اطلاعات اشاره شده در (Manning and Schütze ۱۹۹۹) یک سیستم بازیابی اطلاعات برای پردازش اسناد موجود در پایگاه داده خود بایستی دارای اجزا و قابلیت‌های زیر باشد:

- شاخص معکوس به همراه اطلاعات مکانی نمایه‌ها و فراوانی واژگان در هر سند
- فهرستی از ایست‌واژه‌ها
- قابلیت تحلیل متن (شامل واژه-واژه سازی و ریشه‌یابی و یا بُن‌واژه‌سازی)
- قابلیت محاسبه ارتباط بین عبارت پرس‌وجوی کاربر و اسناد بازیابی شده

همچنین با بررسی وضعیت فعلی دو فایل Solrconfig و Schema سامانه گنج نکات زیر را می‌توان از جمله کاستی‌های بازیابی اطلاعات سامانه گنج دانست:

□ در حال حاضر برای برای ریشه‌یابی واژگان در د اسناد در زمان نمایه سازی (بخش ۴ سند Shema) و تحلیل عبارت پرس و جو (بخش ۵ سند Schema) از دستور `solr.PorterStemFilterFactory` (روش پورتر) استفاده می‌شود که روشی برای ریشه‌یابی در زبان انگلیسی است.

□ در حال حاضر در سامانه گنج از روش `solr.StandardTokenizerFactory` برای واژه-واژه سازی متن و عبارت پرس و جو استفاده می‌شود. این روش مانند روش `solr.KeywordTokenizerFactory` (بخش ۲ سند Schema) برای زبان انگلیسی طراحی شده و بکارگیری آن برای اسناد موجود در سامانه گنج و یا عبارت‌های پرس و جو که اغلب به زبان فارسی هستند نتیجه‌ای را دربر نخواهد داشت.

□ علاوه بر موارد فوق، در بخش ۲ سند Schema به نوع-فیلد `autosuggest` با وظیفه ارائه پیشنهاد‌های عبارت پرس و جو به کاربر در زمان ورود عبارت پرس و جو اشاره شده که از کلاس `solr.LowerCaseFilterFactory` برای کوچک‌سازی واژگان استفاده می‌شود. این روش نیز برای زبان انگلیسی طراحی و توسعه داده شده و استفاده از آن برای واژگان فارسی کارایی ندارد. همچنین در حال حاضر برای سامانه گنج پیشنهاد‌های عبارت پرس و جو با کمک نمایه‌های موجود در Lucene و با انطباق ۳ کاراکتر وارد شده توسط کاربر با نمایه‌های ذخیره شده انجام می‌شود و در عمل از واژه-واژه کردن و کوچک‌سازی واژگان استفاده‌ای نمی‌شود.

□ در سامانه گنج از دستور `solr.PatternReplaceCharFilterFactory` برای جایگزین کردن الگوهای نوشتاری استفاده می‌شود. الگوی مورد استفاده در سامانه گنج چه در زمان نمایه سازی و چه در زمان تحلیل عبارت پرس و جو یکسان بوده و به صورت تجربی استخراج شده‌اند.

□ در سامانه گنج برای جایگزین کردن کاراکترهای متن با کاراکترهای از قبل تعیین شده از دستور `solr.MappingCharFilterFactory` استفاده می‌شود. فایل مورد استفاده

برای این منظور همان فایل اصلی طراحی شده برای سولر برای زبان انگلیسی بوده و تغییری در آن اعمال نشده است.

□ روش محاسبه ارتباط بین عبارت پرس وجوی کاربر و اسناد بازیابی شده در سولر به صورت پیش فرض و در سامانه گنج روش TF-IDF است که روشی پایه‌ای برای بازیابی اطلاعات است (هرچند که تغییرات و بهبودهایی در فرمول آن انجام شده است) (رابطه (3-1)). همچنین، از آنجا که فرآیند تحلیل متن چندانی بر روی اسناد و عبارت‌های پرس وجوی کاربر در سامانه گنج انجام نمی‌شود، نمی‌توان انتظار کارایی بالایی از روش TF-IDF داشت.

□ در سامانه گنج عبارت پرس وجوی کاربر تجزیه نمی‌شود چرا که در بخش ۸ سند Solrconfig، مقدار پارامتر `handleSelect` برابر `false` قرار داده شده است و این سبب می‌شود تا عبارت پرس وجو به تجزیه‌گر فرستاده نشود. به بیانی دیگر، در سامانه گنج عبارت پرس وجوی کاربر تجزیه نشده و فرآیند تحلیل متن بر آن انجام نمی‌شود.

□ طبق بخش ۹ و ۱۸ سند Solrconfig، سامانه گنج برای کنترل املاء و پیشنهاد خود کار از لغت‌نامه پیش فرض سولر استفاده می‌کند که در عمل برای زبان فارسی غیرقابل استفاده است. به همین دلیل است که در صورتی که کاربر اشتباه املائی در عبارت پرس وجو داشته باشد، هیچ پیشنهادی به وی برای اصلاح ارائه نمی‌شود. چنین کاستی در زمان نمایه‌سازی و سپردن اسناد نیز وجود دارد. بدین معنا که فرآیند کنترل املاء برای اسنادی که وارد فرآیند نمایه‌سازی می‌شوند انجام نمی‌شود. لذا، بسیار محتمل است در فیلدهای نمایه شده اسناد پایگاه گنج اشتباهات املائی وجود داشته و این امر سبب می‌شود تا این سامانه عملکرد نامناسبی در بازیابی اطلاعات داشته باشد.

□ از قابلیت‌های `Facet`، `More like this` و `Highlight` در سامانه گنج استفاده نمی‌شود. با کمک قابلیت `More like this` سولر می‌تواند اسناد مشابه با سند مورد انتخابی کاربر را به او نمایش دهد. با این حال، این قابلیت برای سامانه گنج غیرفعال است. با کمک قابلیت `Facet` کاربر قادر است نتایج جستجو را براساس نیاز خود فیلتر نموده و نمایی کلی از خصوصیات اسناد موجود در پایگاه داده را درخصوص موضوع مورد

جستجو مشاهده نماید (شکل ۳-۱۸). در حال حاضر این قابلیت برای سامانه گنج غیرفعال است.

پررنگ کردن بخش‌هایی از سند که با عبارت پرس‌وجوی کاربر منطبق هستند اتفاقی است که پس از فعال کردن مولفه Highlight می‌افتد (شکل ۳-۲۱). در حال حاضر این قابلیت برای سامانه گنج غیرفعال است و آنچه با عنوان پررنگ کردن در صفحه نتایج نمایش داده می‌شود قابلیت است که طراحان سامانه گنج آن را طراحی و پیاده‌سازی نموده‌اند و مجزا از این قابلیت سولر است.

بنابراین براساس نیازمندی‌های عمومی سامانه‌های بازیابی اطلاعات و کاستی‌های فوق راهکارهای بهبود بازیابی اطلاعات در زبان فارسی و برای سامانه گنج را می‌توان به صورت زیر بیان نمود. این راهکارها در دو دسته راهکارهای نیازمند پژوهش بیشتر و عمیق‌تر و راهکارهای فنی و عملیاتی دسته‌بندی شده‌اند.

۱-۲-۶- راهکارهای نیازمند پژوهش

□ استفاده از روش‌های تحلیل متن برای زبان فارسی؛

○ در حال حاضر در سامانه گنج برای واژه-واژه سازی متن و عبارت پرس‌وجو از روش استاندارد دی که برای زبان انگلیسی طراحی شده استفاده می‌شود و لذا بکارگیری آن برای اسناد موجود در سامانه گنج و یا عبارت‌های پرس‌وجو که اغلب به زبان فارسی هستند نتیجه‌ای را دربر نخواهد داشت. بنابراین بکارگیری و یا توسعه روش (هایی) که بتوانند فعالیت واژه-واژه سازی را در سامانه گنج انجام دهند ضروری است.

○ در حال حاضر در سامانه گنج برای ریشه‌یابی از روش پورتر استفاده می‌شود و از آنجا که روش پورتر روشی برای ریشه‌یابی در زبان انگلیسی است، توسعه و پیاده‌سازی روش‌های دیگر برای زبان فارسی از ضروریات است.

○ در حال حاضر در سامانه گنج از تجزیه‌گر استفاده نشده و عبارت پرس‌وجوی کاربر به تجزیه‌گر فرستاده نمی‌شود. به بیانی دیگر، در سامانه گنج عبارت پرس‌وجوی کاربر تجزیه نشده و فرآیند تحلیل متن بر آن انجام نمی‌شود. این

موضوع چالشی جدی برای بازیابی اطلاعات بوده و ضروری است برای بهبود نتایج بازیابی اطلاعات در سامانه گنج بدان پرداخته شود.

○ در سامانه گنج برای کنترل املاء و پیشنهاد خود کار از لغت نامه پیش فرض سولر استفاده می شود که در عمل برای زبان فارسی غیر قابل استفاده است. به همین دلیل است که در صورتی که کاربر اشتباه املائی در عبارت پرس وجو داشته باشد، هیچ پیشنهادی به وی برای اصلاح ارائه نمی شود. چنین کاری در زمان نمایه سازی و سپردن اسناد نیز وجود دارد. بدین معنا که فرآیند کنترل املاء برای اسنادی که وارد فرآیند نمایه سازی می شوند انجام نمی شود. بنابراین ضروری است با بررسی تجارب گذشته و نیز نمایه های موجود، لغت نامه ای برای کنترل املائی کاربران و اسناد نمایه شده تدوین گردد.

○ در سامانه گنج از دستور `solr.PatternReplaceCharFilterFactory` برای جایگزین کردن الگوهای نوشتاری استفاده می شود. الگوی مورد استفاده در سامانه گنج چه در زمان نمایه سازی و چه در زمان تحلیل عبارت پرس وجو یکسان بوده و به صورت تجربی استخراج شده اند. به نظر می رسد تعیین دقیق تر و علمی تر این الگوها می تواند سبب شود تا داده های ورودی بهتری برای فرآیند تحلیل متن حاصل شوند.

○ در سامانه گنج برای جایگزین کردن کاراکترهای متن با کاراکترهای از قبل تعیین شده از دستور `solr.MappingCharFilterFactory` استفاده می شود. فایل مورد استفاده برای این منظور همان فایل اصلی طراحی شده برای سولر بوده و تغییری در آن اعمال نشده است و از آنجا که سولر برای زبان انگلیسی طراحی شده است، تدوین آن برای زبان فارسی ضروری است.

□ ارتقاء روش محاسبه ارتباط بین عبارت پرس وجوی کاربر و اسناد بازیابی شده

روش محاسبه ارتباط بین عبارت پرس وجوی کاربر و اسناد بازیابی شده در سامانه گنج روش TF-IDF است که روشی پایه ای برای بازیابی اطلاعات بوده و روش های جایگزین دیگری با کارایی بالاتر در ادبیات برای این منظور ارائه شده اند. همچنین، از آنجا که

فرآیند تحلیل متن چندانی بر روی اسناد و عبارت‌های پرس‌وجوی کاربر در سامانه گنج انجام نمی‌شود، نمی‌توان انتظار کارایی بالایی از روش TF-IDF داشت. بنابراین پیشنهاد می‌شود در پژوهش‌های آتی علاوه بر توسعه روش‌های تحلیل متن برای زبان فارسی، روش‌های جدیدتر بازیابی اطلاعات نیز توسعه یابند.

علاوه بر دو راهکار فوق ضروری است برخی از قابلیت‌های سولر برای زبان فارسی و در سامانه گنج فعال شوند. این قابلیت‌ها عبارتند از:

- فعال‌سازی قابلیت More like this: با کمک قابلیت More like this سولر می‌تواند اسناد مشابه با سند مورد انتخابی کاربر را به او نمایش دهد. در حال حاضر این قابلیت برای سامانه گنج غیرفعال است و از آنجا که این قابلیت برای کاربران جذاب است پیشنهاد می‌شود پژوهش‌های جدی در این زمینه انجام شود.
- فعال‌سازی قابلیت Facet: با کمک این قابلیت کاربر قادر است نتایج جستجو را براساس نیاز خود فیلتر نموده و نمایی کلی از خصوصیات اسناد موجود در پایگاه داده را درخصوص موضوع مورد جستجو مشاهده نماید. به تعبیری دیگر، با کمک این مولفه فراداده مربوط به اسناد مرتبط با عبارت پرس‌وجو برای کاربر نمایش داده می‌شود. در حال حاضر این قابلیت برای سامانه گنج غیرفعال است و از آنجا که این قابلیت برای کاربران بسیار کاراست پیشنهاد می‌شود بدان پرداخته شود.
- فعال‌سازی قابلیت Highlight: با استفاده از این قابلیت بخش‌هایی از سند که با عبارت پرس‌وجوی کاربر منطبق هستند پررنگ می‌شود. در حال حاضر این قابلیت برای سامانه گنج غیرفعال است و از آنجا که این قابلیت نیازمند فرآیند تحلیل متن است (شکل ۳-۲۲) پیشنهاد می‌شود برای فعال‌سازی آن پژوهش‌های بیشتری صورت پذیرد.

۶-۲-۲- راهکارهای فنی و عملیاتی

- فعال کردن JMX برای کنترل و ردیابی عملکرد سولر؛
این می‌تواند ابزاری مناسب برای پشتیبانی از سامانه گنج باشد.
- جایگزینی سیاست‌های مدیریت Cache؛

در حال حاضر، روش مورد استفاده برای مدیریت Cache، نگهداری نتایج جستجوهای قبلی، نگهداری اسناد برای جستجوهای بعدی و مدیریت حافظه کوتاه مدت در سامانه گنج روش پیش فرض LRU است. سیاست Least Frequently Used (LFU) نیز در سولر تعریف شده است که در طی آن اسناد یا نتایج بازیابی شده با کمترین استفاده (غیر جذاب ترین) از حافظه پاک می شود. بنابراین می توان روش LFU را در مواردی امتحان کرده، نتایج و رضایت کاربران را با حالت پیش فرض سنجیده و در صورت نیاز تنظیمات را تغییر داد.

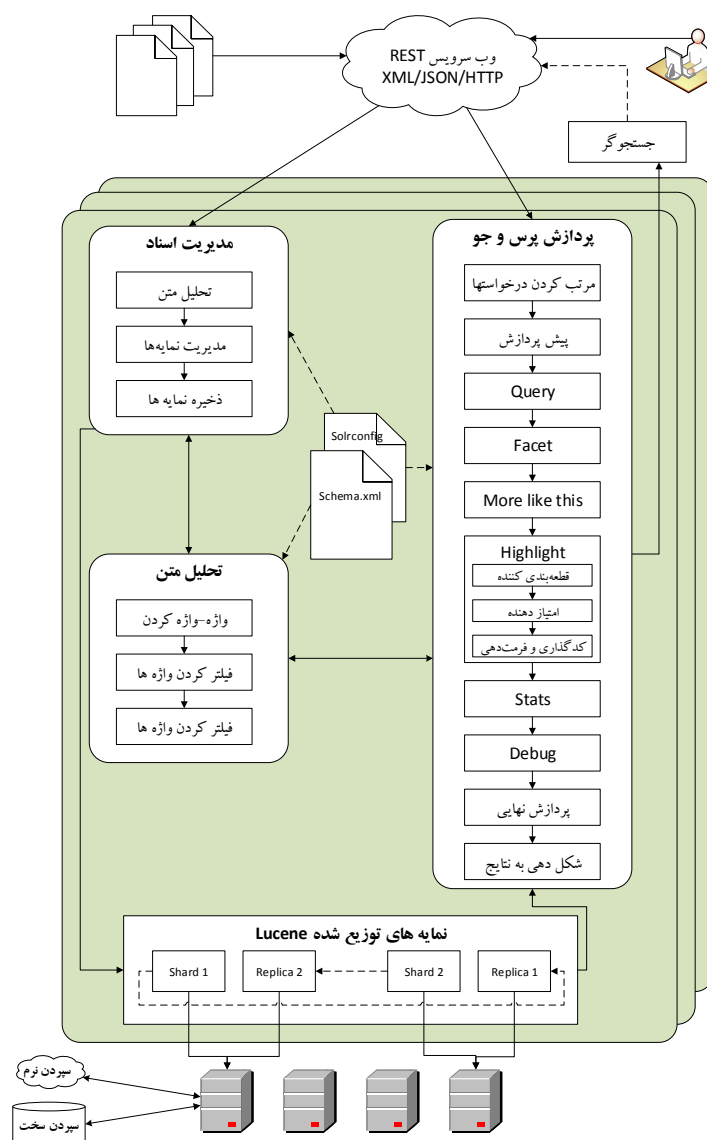
□ فعال کردن قابلیت جستجوی جغرافیایی.

یکی از امکانات و قابلیت های گنج پشتیبانی از جستجوی جغرافیایی است که در حال حاضر برای سامانه گنج فعال نیست.

البته موارد دیگری نیز در این رابطه وجود دارند که به علت عدم تمرکز طرح حاضر بر کدهای اجرایی سولر بدان پرداخته نشده است.

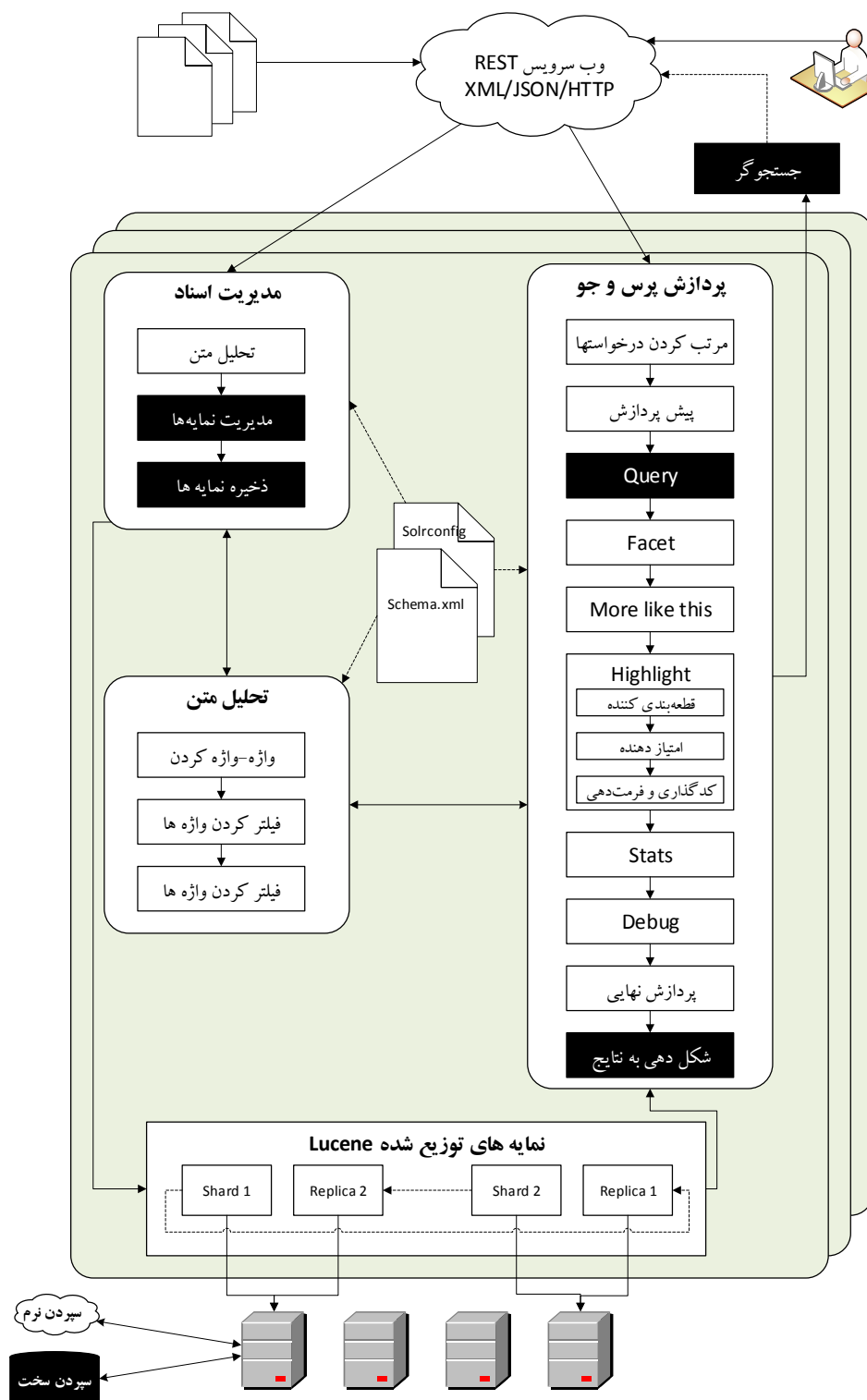
جمع‌بندی و نتیجه‌گیری

فرآیند پردازش اسناد و پرس و جو در سولر به صورت خلاصه در شکل زیر نمایش داده شده است. در طی این فرآیند اسناد ورودی به پایگاه داده توسط لوسن نمایه شده و ذخیره می‌شوند. درخواست‌های جستجوی کاربر نیز تحلیل شده و نتایج براساس تنظیمات مشخص شده در Solrconfig و Schema بازیابی شده و به کاربر نمایش داده می‌شوند. فرآیند تحلیل متن جزء کلیدی سولر و لوسن بوده و به عنوان ابزاری در زیرسیستم‌های مدیریت اسناد و پردازش پرس و جو استفاده می‌شود.



فرآیند پردازش اسناد و پرس و جو در سولر

سامانه گنج و موتور جستجوی آن نیز از چنین ساختاری برخوردار است. با بررسی های انجام شده می توان عنوان داشت که تنها مراحل پررنگ شده در شکل زیر در سامانه گنج و برای زبان فارسی انجام شده و بقیه مراحل همچنان غیرفعال است (برخی نیازمند پژوهش های عمیق تر و توسعه و پیاده سازی روش های مبتنی بر زبان فارسی و برخی نیز نیازمند فعال سازی فنی و عملیاتی است).



از سویی دیگر، با توجه به ضرورت وجود قابلیت تحلیل متن (شامل واژه-واژه سازی و ریشه یابی و یا بُن واژه سازی) و قابلیت محاسبه ارتباط بین عبارت پرس و جوی کاربر و اسناد بازیابی شده در سامانه های بازیابی اطلاعات در (Manning and Schütze ۱۹۹۹)، نکات زیر را می توان از جمله راهکارهای ضروری و نیازمند پژوهش بیشتر و عمیق تر برای بهبود بازیابی اطلاعات در زبان فارسی در سامانه گنج دانست.

□ استفاده از روش های جایگزین برای ریشه یابی در زبان فارسی و پیاده سازی آن در سامانه گنج (در حال حاضر از روش پیش فرض پورتر استفاده می شود که در زبان فارسی کارایی ندارد).

□ بکارگیری و پیاده سازی روش های برای واژه-واژه سازی متن و عبارت پرس و جو در سامانه گنج (در حال حاضر از روش های StandardTokenizer و KeywordTokenizer برای واژه-واژه سازی متن و عبارت پرس و جو استفاده می شود که برای زبان انگلیسی طراحی شده و بکارگیری آن برای اسناد موجود در سامانه گنج و یا عبارت های پرس و جو که اغلب به زبان فارسی هستند نتیجه ای را دربر نخواهد داشت).

□ استفاده از تجزیه گر در سامانه گنج (در حال حاضر عبارت پرس و جوی کاربر تجزیه نشده و فرآیند تحلیل متن بر آن انجام نمی شود).

□ توسعه و ارتقاء روش بازیابی اطلاعات در سامانه گنج (روش جستجوی مورد استفاده، روش TF-IDF است و از آنجا که فرآیند تحلیل متن چندانی بر روی اسناد و عبارت های پرس و جو در سامانه گنج انجام نمی شود، نمی توان انتظار کارایی بالایی از روش TF-IDF داشت).

راهکارهای نیازمند پژوهش دیگری نیز وجود دارند که توجه به آنها می تواند بهبود مناسبی را در بازیابی اطلاعات در سامانه گنج ایجاد نماید.

□ تعیین دقیق تر و علمی تر جایگزینی الگوهای نوشتاری برای سامانه گنج (در حال حاضر الگوی مورد استفاده چه در زمان نمایه سازی و چه در زمان تحلیل عبارت پرس و جو یکسان بوده و به صورت تجربی استخراج شده اند. به نظر می رسد تعیین دقیق تر و علمی تر این الگوها می تواند سبب شود تا داده های ورودی بهتری برای فرآیند تحلیل متن حاصل شوند).

□ تدوین فایل مورد نیاز برای برای جایگزین کردن کاراکترهای متن با کاراکترهای از قبل تعیین شده در سامانه گنج (در حال حاضر فایل مورد استفاده همان فایل اصلی طراحی شده برای سولر بوده و تغییری در آن اعمال نشده است و از آنجا که سولر برای زبان انگلیسی طراحی شده است، تدوین آن برای زبان فارسی ضروری است).

□ فعال سازی و پیاده سازی کنترل املاء در سامانه گنج (در حال حاضر از لغت نامه پیش فرض سولر برای کنترل املاء و پیشنهاد خود کار استفاده می شود که در عمل برای زبان فارسی غیر قابل استفاده است. به همین دلیل است که در صورتی که کاربر اشتباه املائی در عبارت پرس وجو داشته باشد، هیچ پیشنهادی به وی برای اصلاح ارائه نمی شود. چنین کاستی در زمان نمایه سازی و سپردن اسناد نیز وجود دارد. بدین معنا که فرآیند کنترل املاء برای اسنادی که وارد فرآیند نمایه سازی می شوند انجام نمی شود. لذا، بسیار محتمل است در فیلدهای نمایه شده اسناد پایگاه گنج اشتباهات املائی وجود داشته و این امر سبب می شود تا این سامانه عملکرد نامناسبی در بازیابی اطلاعات داشته باشد).

علاوه بر موارد فوق، موتور جستجوی سولر دارای قابلیت هایی است که در حال حاضر در از آنها در سامانه گنج استفاده نمی شود و فعال سازی و پیاده سازی آنها می تواند در جلب نظر کاربران این سامانه موثر باشد. این قابلیت ها عبارتند از:

□ قابلیت More like this

□ قابلیت Facet

□ قابلیت Highlight

□ قابلیت جستجوی جغرافیایی

همچنین، در حال حاضر، روش مورد استفاده برای مدیریت Cache سامانه گنج، روش پیش فرض LRU است که می توان آن را با روش LFU جایگزین کرده، نتایج و رضایت کاربران را با حالت پیش فرض سنجیده و در صورت نیاز تنظیمات را تغییر داد.

ضمائم

۷-۱ توابع

سولر این امکان را برای کاربران خود فراهم می آورد تا علاوه بر انجام جستجو براساس عبارت پرس وجو، توابعی را اجرا کرده و از نتایج آن به صورت عبارت پرس وجوی جدید و یا پارامتری برای تعیین میزان ارتباط نتایج بازیابی شده استفاده کند. برخی از این توابع در ادامه آورده شده اند.

۷-۱-۱- توابع ریاضی

برخی از توابع ریاضی قابل استفاده در سولر به شرح زیر هستند:

جدول ۱-۲- توابع ریاضی در سولر

توضیحات	دستور (syntax) تابع
قدر مطلق مقدار x	$\text{abs}(x)$
سینوس، آرک سینوس و سینوس هیپربولیک x	$\sin(x)$, $\text{asin}(x)$, $\sinh(x)$
کسینوس، آرک کسینوس و کسینوس هیپربولیک x	$\cos(x)$, $\text{acos}(x)$, $\cosh(x)$
تانژانت، آرک تانژانت و تانژانت هیپربولیک x	$\tan(x)$, $\text{atan}(x)$, $\tanh(x)$
تابع نمایی x	$\exp(x)$
لگاریتم طبیعی x	$\ln(x)$
لگاریتم x	$\log(x)$
تبدیل x به رادیان	$\text{rad}(x)$
ریشه دوم x	$\text{sqrt}(x)$

$\frac{x}{y}$	div(x,y)
$x - y$	sub(x,y)
x^y	pow(x,y)
مقدار تابع $f(x) = \frac{a}{mx+b}$	recip(x,m,a,b)
محاسبه عدد پی	pi()
ضرب تمامی ورودی ها	product(x,...n)
	mul(x,...n)
جمع تمامی ورودی ها	add(x,...n)

۷-۱-۲- توابع فاصله

تابع فاصله پیش فرض در سولر تابع فاصله دامرا-لون اشتاین است اما توابع دیگری نیز وجود دارد که لیستی از آنها در جدول ۷-۲ آمده است.

جدول ۷-۲- توابع فاصله در سولر

توضیحات	دستور (syntax) تابع
محاسبه فاصله دو بردار/نقطه در فضای n بعدی براساس معیار فاصله تعیین شده با power: $power = 0 \rightarrow$ Sparseness calculation $power = 1 \rightarrow$ Manhattan distance $power = 2 \rightarrow$ Euclidean distance $power = \text{infinity} \rightarrow$ Infinit norm	dist(power,x1...n1,x2...n2)
محاسبه فاصله اقلیدسی	sqdist(x1...n1,x2...n2)
محاسبه فاصله Haversine	hsin(radius INKM, isDegrees,x1,y1,x2,y2)
محاسبه فاصله جغرافیایی محاسبه شباهت یا فاصله کاراکترها در یک رشته. نتیجه محاسبه بین 0 (بدون شباهت) تا 1 (کاملاً مشابه) تغییر می کند. disType می تواند مقادیر زیر را داشته باشد:	geohash(lat,lon) geodis(sfield,lat,lon) strdis(s1,s2,disType) strdis(s1,s2,disType,ngramSize)

<i>jw</i> → Jaro-Winkler <i>edit</i> → Levenshtein distance <i>ngram</i> → NGram distance سایر توابع مورد نیاز برای فاصله را می توان در رابط کاربری StringDistance تعیین و یا تعریف کرد.	
--	--

۷-۱-۳- توابع منطق بولی

جدول ۷-۳- توابع منطق بولی در سولر

توضیحات	دستور (syntax) تابع
true اگر هر دو <i>x</i> و <i>y</i> مقدار true داشته باشند	<code>and(x,y)</code>
true اگر مقداری برای <i>x</i> وجود داشته باشد	<code>exists(x)</code>
مقدار trueValue را بر می گرداند اگر <i>x</i> درست باشد، در غیر این صورت مقدار falseValue را بر می گرداند	<code>if(x, trueValue, falseValue)</code>
false اگر <i>x</i> true باشد و <i>x</i> false اگر <i>x</i> false باشد	<code>not(x)</code>
true اگر <i>x</i> یا <i>y</i> یا هر دو true داشته باشند	<code>or(x,y)</code>
true اگر <i>x</i> یا <i>y</i> ، true داشته باشند و false اگر هر دو true داشته باشند	<code>xor(x,y)</code>

۷-۲ عوامل محاسبه ارتباط

جدول ۷-۴- عوامل محاسبه ارتباط در سولر

توضیحات	دستور (syntax) عامل
تعداد اسنادی حاوی واژه مورد نیاز در فیلد مشخص شده	<code>docfreq(fieldName,term)</code>
محاسبه idf برای ارزش خاص در فیلد مشخص شده	<code>idf(fieldName,value)</code>
تعداد کل اسناد (شامل تعداد اسنادی که از حافظه کوتاه مدت حذف شده اند اما همچنان در durable storage قرار دارند)	<code>maxdoc()</code>
تعداد واژگان (token) نمایه شده	<code>sumtotaltermfreq(field)</code> <code>sttf(field)</code>

termfreq(fieldname,term)	تعداد تکرار واژه خاص در فیلد مشخص شده در یک سند
tf(fieldname,term)	تعداد تکرار واژه خاص در فیلد مشخص شده
totaltermfreq(fieldname,term)	تعداد دفعات تکرار واژه در کل اسناد نمایه شده
ttf(fieldname,term)	

۷-۳ کلاس‌های محاسبه ارتباط

سولر به صورت پیش فرض از DefaultSimilarity برای محاسبه شباهت بین عبارت پرس و جوی کاربر و اسناد موجود استفاده می کند اما کلاس ها و روش های دیگری نیز وجود دارند که می توان از آنها بجای کلاس پیش فرض استفاده کرد (جدول ۷-۵). این کلاس ها را می توان به صورت کلی برای کل سیستم جستجو در Schema فعال کرد و یا به صورت محلی و برای فیلدی خاص (خصوصیت similarity برای هر فیلد به صورت مجزا قابل تعریف است).

جدول ۷-۵- کلاس های محاسبه ارتباط در سولر

دستور (syntax) عامل	توضیحات
DefaultSimilarity	محاسبه براساس رابطه (1-3)
BM25Similarity	محاسبه براساس مدل BM25
DFRSimilarity	محاسبه براساس مدل Divergence from randomness (DFR)
IBSimilarity	محاسبه براساس مدل Information based (IB) که مشابه DFR است اما ورودی های ساده تری دارد
LMDirichletSimilarity	استفاده از مدل های زبانی برای اصلاح وزن واژه در پیکره واژگانی
LMJelinekMercerSimilarity	استفاده از مدل های زبانی مانند LMDirichletSimilarity اما با پیاده سازی آسان تر
SweetSpotSimilarityFactory	توسعه یافته DefaultSimilarity با قابلیت تنظیم پارامترهای lenghtNorm و tf

۷-۴ خصوصیات UPDATE HANDLER

درخواست اضافه کردن اسناد از طریق واسطه‌هایی به سولر ارائه شده و از طریق Update Handler اجرا می‌شود. خصوصیات این Handler به شرح جدول ۶-۷ است.

جدول ۶-۷- خصوصیات Update Handler

نوع درخواست‌ها	توضیحات
Add	اضافه کردن یک یا چند سند
Delete	حذف یک سند با استفاده از ID
Delete by query	حذف سندی که با پرس و جوی لوسن مطابقت دارد
Atomic update	بروزرسانی یک یا چند فیلد در یک سند
Commit	سپردن اسناد به نمایه‌سازی
Optimize	بهینه کردن نمایه‌ها با ادغام بخش‌ها و حذف موارد حذف شده

۷-۵ پارامترهای مولفه HIGHLIGHT

برخی از پارامترهای مولفه highlight برای تنظیم در Schema به شرح است.

جدول ۷-۲- برخی از پارامترهای مولفه highlight

پارامتر	توضیحات	مقدار پیش فرض
h1	فعالسازی مولفه highlight	False
h1.snippets	پیشینه تعداد snippet‌ها برای نمایش	1
h1.f1	فیلدهایی که snippet‌ها براساس آن ساخته شود	--
h1.fragmenter	روش قطعه‌بندی کردن متن	Gap
h1.fragmentsize	اندازه قطعات	100
h1.q	تعیین عبارت برای پررنگ کردن	--
h1.alternateField	اگر سیستم قادر نباشد snippet بسازد، چه مواردی را در صفحه نتایج نمایش دهد	--
h1.formatter	تعیین فرمت عبارت پررنگ شده	simple
h1.maxAnalyzedChars	پیشینه تعداد کاراکترهای مورد بررسی در مواردی که اسناد حجیم باشند	51200

۷-۶ پارامترهای مولفه SPELLCHECK

برخی از مهمترین پارامترهای مولفه Spellcheck برای تنظیم در Schema به شرح است.

جدول ۷-۸- برخی از پارامترهای مولفه Spellcheck

پارامتر	توضیحات
spellcheck	فعالسازی مولفه Spellcheck
spellcheck.build	اگر true، دستور به سولر برای ساخت دایرکتوری spell-check (درهنگام بکارگیری تابع DirectSolrSpellChecker فعالسازی این پارامتر ضروری است)
spellcheck.count	تعیین بیشینه تعداد پیشنهادها برای اصلاح
spellcheck.dictionary	معرفی نام و آدرس لغت نامه ای که پیشنهادها براساس آن ارائه می شوند
spellcheck.onlyMorePopular	اگر true، محدود کردن پیشنهادها به مواردی که بازخوانی بالاتری دارند
spellcheck.maxResultsForSuggest	تعیین حداقل تعداد اسناد بازیابی شده برای فعال سازی Spellcheck
spellcheck.collate	اگر true، جستجو براساس عبارت پرس و جوی کاربر انجام می شود و پیام ...? Did you mean برای کاربر نمایش داده می شود.
spellcheck.maxCollations	تعیین بیشینه تعداد پرس و جوی تولید شده توسط سولر
spellcheck.maxCollationTries	تعیین بیشینه تعداد دفعات جستجوی جایگزین

منابع

ارشادی م.، رجیبی ت.، شیرانی ف.، رضایی ن. (۱۳۹۴). کاربرد تکنیک تحلیل ریشه در حل مشکلات کیفی سامانه‌های اطلاعاتی تحقیقاتی: مطالعه موردی سامانه اشاعه اطلاعات پایان‌نامه‌ها/ رساله‌های دانش‌آموختگان داخل کشور (گنج). مدیریت اطلاعات، ۱(۲)، ۷۵-۸۹.

زرین بال م. (۱۳۹۷). طراحی مدل مفهومی برای معنابخشی به خدمات ارائه شده در سامانه گنج. تهران.

فتاحی س. (۱۳۹۶). تحلیل رفتار جست‌وجوی اطلاعات پژوهشگران در موتور جست‌وجوی سامانه گنج. تهران.

Ai, Qingyao, Liu Yang, Jiafeng Guo, and W Bruce Croft. 2016. "Analysis of the Paragraph Vector Model for Information Retrieval." In *Proceedings of the 2016 ACM International Conference on the Theory of Information Retrieval*, 133–42. ACM.

Aly, Robin, and Thomas Demeester. 2011. "Towards a Better Understanding of the Relationship between Probabilistic Models in IR." In *Conference on the Theory of Information Retrieval*, 164–75. Springer.

Amati, Gianni, and Cornelis Joost Van Rijsbergen. 2002. "Probabilistic Models of Information Retrieval Based on Measuring the Divergence from Randomness." *ACM Transactions on Information Systems* 20 (4): 357–89. <https://doi.org/10.1145/582415.582416>.

Bookstein, Abraham. 1980. "Fuzzy Requests: An Approach to Weighted Boolean Searches." *Journal of the American Society for Information Science* 31 (4): 240–47. <https://doi.org/10.1002/asi.4630310403>.

Brown, Peter J, and Gareth J F Jones. 2001. "Context-Aware Retrieval: Exploring a New Environment for Information Retrieval and Information Filtering." *Personal and Ubiquitous Computing* 5 (4): 253–63.

Bruza, Peter, and Dawei Song. 2003. "A Comparison of Various Approaches for Using Probabilistic Dependencies in Language Modeling." In *Proceedings of the 26th Annual International ACM SIGIR Conference on Research and Development in Informaion Retrieval*, 419–20. ACM.

Bush, Vannevar. 1945. "As We May Think." *The Atlantic*, 1–10.

Chen, Hsinchun. 1995. "Machine Learning for Information Retrieval: Neural Networks, Symbolic Learning, and Genetic Algorithms." *Journal of the American Society for Information Science* 46 (3): 194–216. [https://doi.org/10.1002/\(SICI\)1097-4571\(199504\)46:3<194::AID-ASI4>3.0.CO;2-S](https://doi.org/10.1002/(SICI)1097-4571(199504)46:3<194::AID-ASI4>3.0.CO;2-S).

Damerau, Fred J. 1964. "A Technique for Computer Detection and Correction of Spelling Errors." *Communications of the ACM* 7 (3): 171–76.

Dumais, S T, G W Furnas, T K Landauer, S Deerwester, and R Harshman. 1988.

- “Using Latent Semantic Analysis to Improve Access to Textual Information.” In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, 281–85. CHI '88. New York, NY, USA: ACM. <https://doi.org/10.1145/57167.57214>.
- Fang, Hui, and ChengXiang Zhai. 2005. “An Exploration of Axiomatic Approaches to Information Retrieval.” In *Proceedings of the 28th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, 480–87. ACM.
- Grainger, Trey, and Timothy Potter. 2014. *Solr in Action*. New York: Manning Publications Co. <https://doi.org/10.1017/CBO9781107415324.004>.
- Grossman, David A, and Ophir Frieder. 2012. *Information Retrieval: Algorithms and Heuristics*. Vol. 15. Springer Science & Business Media.
- Hiemstra, Djoerd. 2000. “A Probabilistic Justification for Using Tf× Idf Term Weighting in Information Retrieval.” *International Journal on Digital Libraries* 3 (2): 131–39.
- Hofmann, Katja, Shimon Whiteson, and Maarten de Rijke. 2013. “Balancing Exploration and Exploitation in Listwise and Pairwise Online Learning to Rank for Information Retrieval.” *Information Retrieval* 16 (1): 63–90.
- Klouche, Khalil, Tuukka Ruotsalo, Luana Micalef, Salvatore Andolina, and Giulio Jacucci. 2017. “Visual Re-Ranking for Multi-Aspect Information Retrieval.” In *Proceedings of the 2017 Conference on Conference Human Information Interaction and Retrieval*, 57–66. ACM.
- Kwok, K L. 1989. “A Neural Network for Probabilistic Information Retrieval.” In *12th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, 23:21–30. New York, NY, USA: ACM. <https://doi.org/10.1145/75335.75338>.
- Liu, Xiaodong, Jianfeng Gao, Xiaodong He, Li Deng, Kevin Duh, and Ye-Yi Wang. 2015. “Representation Learning Using Multi-Task Deep Neural Networks for Semantic Classification and Information Retrieval.”
- Manning, Christopher D., Prabhakar Ragahvan, and Hinrich Schutze. 2009. *An Introduction to Information Retrieval*. Cambridge, England: Cambridge University Press. <https://doi.org/10.1109/LPT.2009.2020494>.
- Manning, Christopher D, and Hinrich Schütze. 1999. *Foundations of Statistical Natural Language Processing*. MIT press.
- Maron, M E, and J L Kuhns. 1960. “On Relevance, Probabilistic Indexing and Information Retrieval.” *Journal of the ACM* 7 (3): 216–44. <https://doi.org/10.1145/321033.321035>.
- McCandless, Michael, Erik Hatcher, and Otis Gospodnetic. 2010. *Lucene in Action*. 2nd ed. Stamford: Manning Publications Co.

- Metzler, Donald, and W Bruce Croft. 2004. "Combining the Language Model and Inference Network Approaches to Retrieval." *Information Processing & Management* 40 (5): 735–50.
- Oard, Douglas W, and Anne R Diekema. 1998. "Cross-Language Information Retrieval." *Annual Review of Information Science and Technology (ARIST)* 33: 223–56.
- Ponte, Jay M, and W Bruce Croft. 1998. "A Language Modeling Approach to Information Retrieval." In *Proceedings of the 21st Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, 275–81. ACM.
- Ribeiro, Berthier A N, and Richard Muntz. 1996. "A Belief Network Model for IR." In *Proceedings of the 19th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, 253–60. ACM.
- Rijsbergen, C J Van. 1986. "A Non-Classical Logic for Information Retrieval." *The Computer Journal* 29 (6): 481–85. <http://dx.doi.org/10.1093/comjnl/29.6.481>.
- Robertson, S., and K. Sparck-Jones. 1976. "Relevance Weighting of Search Terms." *Journal of the American Society for Information Science* 27 (3): 129–46. <https://doi.org/10.1002/asi.4630270302>.
- Robertson, S., S. Walker, S. Jones, Micheline Hancock-Beaulieu, and Mike Gatford. 1995. "Okapi at TREC-3." *Nist Special Publication Sp* 109: 109.
- Robertson, S., S Walker, and M Hancock-Beaulieu. 1995. "Large Test Collection Experiments on an Operational, Interactive System: Okapi at TREC." *Information Processing & Management* 31 (3): 345–60. [https://doi.org/https://doi.org/10.1016/0306-4573\(94\)00051-4](https://doi.org/https://doi.org/10.1016/0306-4573(94)00051-4).
- Rocchio, Joseph John. 1971. "Relevance Feedback in Information Retrieval." *The SMART Retrieval System: Experiments in Automatic Document Processing*, 313–23.
- Salton, G. 1971. *The SMART Retrieval System-Experiments in Automatic Document Processing*. Upper Saddle River, NJ, USA: Prentice-Hall, Inc.
- Salton, G, A Wong, and C S Yang. 1975. "A Vector Space Model for Automatic Indexing." *Communications of the ACM* 18 (11): 613–20. <https://doi.org/10.1145/361219.361220>.
- Shen, Yelong, Xiaodong He, Jianfeng Gao, Li Deng, and Grégoire Mesnil. 2014. "A Latent Semantic Model with Convolutional-Pooling Structure for Information Retrieval." In *Proceedings of the 23rd ACM International Conference on Conference on Information and Knowledge Management*, 101–10. ACM.
- Zaragoza, Hugo, Djoerd Hiemstra, and Michael Tipping. 2003. "Bayesian Extension to the Language Model for Ad Hoc Information Retrieval." In

بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی □ ۱۰۸

*Proceedings of the 26th Annual International ACM SIGIR Conference on
Research and Development in Informaion Retrieval*, 4–9. ACM.

Zhai, ChengXiang. 2008. “Statistical Language Models for Information Retrieval.”
Synthesis Lectures on Human Language Technologies 1 (1): 1–141.

ABSATRACT

Information Retrieval (IR) is generally defined as finding document(s) (unstructured data) from dataset(s) to satisfy users' informational needs. Open-source search engine platforms are designed to enable the websites to find the related documents or data to users' query. Apache Lucene and Apache Solr are among the most well known open-source search platforms. Ganj search engine uses these two open-source platforms to search among the documents in its dataset. However, no research has ever been conducted in IranDoc to thoroughly investigate Apache Lucene and Solr. Thus, in order to improve Ganj search engine, especially its capabilities in Persian, the main structure of Apache Lucene and Solr should be investigated. This is the aim of the following research. Here, the main structure of Apache Lucene and Solr and Ganj modifications has been investigated and some recommendations for improving the results are proposed. These recommendations are divided in two categories: recommendations needing further research and functional recommendations.

Keywords: Information retrieval; Apache Lucene; Apache Solr; Ganj search engine



**Iranian Research Institute
for Information Science and Technology
(I r a n D o c)**

Research Project

Investigating Ganj search engine to improve Persian information retrieval

By:

Marzieh Zarinbal Masouleh

September 2019