

10, октябрь 2015

УДК 004.047

Латентно-семантический анализ теста с применением алгоритма Портера

***Величкевич А.Г.**, студент
Россия, 105005, г. Москва, МГТУ им. Н.Э. Баумана,
кафедра «Компьютерные системы и сети»*

***Черепяхин А.А.**, студент
Россия, 105005, г. Москва, МГТУ им. Н.Э. Баумана,
кафедра «Компьютерные системы и сети»*

*Научный руководитель: **Андреев А.М.**, доцент
Россия, 105005, г. Москва, МГТУ им. Н.Э. Баумана,
кафедра «Компьютерные системы и сети»
k_iu6@bmstu.ru*

Введение

Латентно-семантический анализ — это метод обработки информации на естественном языке, анализирующий взаимосвязь между множеством документов и терминами в них встречающимися, сопоставляющий некоторые факторы (тематики) всем документам и терминам. (В основе метода лежат принципы факторного анализа, в частности выявление латентных связей изучаемых явлений или объектов) [1].

Основные применения данного анализа:

- сравнение двух терминов между собой;
- сравнение двух документов между собой;
- сравнение термина и документа.

В настоящее время многие поисковые службы в интернете активно используют данную методику. Он позволяет находить тексты похожие по смыслу. Латентно-семантический анализ отображает документы и отдельные слова в так называемое «семантическое пространство», в котором и производятся все дальнейшие сравнения. При этом делаются следующие утверждения:

1) Документы - это просто набор слов. Порядок слов в документах игнорируется. Важно только то, сколько раз то или иное слово встречается в документе.

2) Семантическое значение документа определяется набором слов, которые, как правило, идут вместе. Например, в биржевых сводках, часто встречаются слова: «фонд», «акция», «доллар»

3) Каждое слово имеет единственное значение. Это, безусловно, сильное упрощение, но именно оно делает проблему разрешимой.

1. Выбор библиотеки для реализации алгоритма Портера

Перед тем как перейти к примеру анализа, рассмотрим реализацию алгоритма Портера для стемминга. Стеминг — это процесс нахождения основы слова для заданного исходного слова. Основа слова необязательно совпадает с морфологическим корнем слова. Данный алгоритм написан на многих языках программирования, но не реализован на C++. Его реализация была произведена, взяв за основу код на JAVA [2].

Для реализации данного алгоритма понадобились регулярные выражения для работы со строками. В связи с тем, что в стандартной сборке C++ нету регулярных выражений, пришлось выбрать между двумя библиотеками (BOOST и PCRE), в которых они есть.

В ходе анализа библиотек использовались исследования, проведенные компанией boost, которая сравнила время поиска регулярных выражений в английском тексте [3].

Для начала был взят текст объемом 19Мб, в котором производился поиск регулярных выражений, и измерялось время нахождения в различных библиотеках (см. таблицу 1).

Таблица 1

Время нахождения регулярных выражений в тексте, объема 19Мб

Регулярное выражение	Библиотека BOOST, с	Библиотека PCRE, с
Twain	$T_{b1B} = 3,49$	$T_{p1B} = 1$
beman john dave	$T_{b2B} = 1,03$	$T_{p2B} = 1,16$
^[^]*?Twain	$T_{b3B} = 1$	$T_{p3B} = 3,32$
Tom Sawyer Huckleberry Finn	$T_{b4B} = 1,02$	$T_{p4B} = 1,08$

Затем был взят текст объемом 50Кб, в котором был произведен то же подсчет (см. таблицу 2).

Время нахождения регулярных выражений в тексте, объема 50Кб

Регулярное выражение	Библиотека BOOST, с.	Библиотека PCRE, с.
Twain	$T_{b1S} = 1,8$	$T_{p1S} = 1$
beman john dave	$T_{b2S} = 0,5$	$T_{p2S} = 0,8$
^[^]*?Twain	$T_{b3S} = 1$	$T_{p3S} = 2,82$
Tom Sawyer Huckleberry Finn	$T_{b4S} = 1$	$T_{p4S} = 1,28$

В ходе измерений был подсчитан средний результат:

$$T_{\text{ср } BOOST} = \left(\sum_{i=1}^4 T_{biB} + \sum_{i=1}^4 T_{biS} \right) / 2.$$

Откуда $T_{\text{ср } BOOST} = 1,36 \text{ с.}$

$$T_{\text{ср } PCRE} = \left(\sum_{i=1}^4 T_{piB} + \sum_{i=1}^4 T_{piS} \right) / 2.$$

Откуда $T_{\text{ср } PCRE} = 1,56 \text{ с.}$

В результате исследования была выбрана библиотека BOOST, так как она более быстрая в поиске регулярных выражений в тексте и имеет более высокий функционал.

2. Реализация алгоритма Портера

Рассматриваемый алгоритм не использует баз основ слов, в отличие от других алгоритмов стемминга. Он применяет последовательно ряд правил, отсекая окончания и суффиксы, основываясь на особенностях языка, в связи с чем работает быстро, но не всегда безошибочно. Данный алгоритм работает по следующим правилам [2]:

1) При поиске окончания из всех возможных выбирается наиболее длинное. Например, в слове «чтение» выбираем окончание «ие» а не «е».

2) **Rv**-область слова после 1 гласной, **r1**- область слова после первой гласной-согласной, **r2**-область **r1** после первой гласной-согласной. Все проверки производятся над областью **rv**. Например, в слове «необъятный» области будут такие: **rv**: «объятный»,

r1: «ъятный», **r2**: «ный». Ниже приведена схема-алгоритма Портера (см. рис. 1).

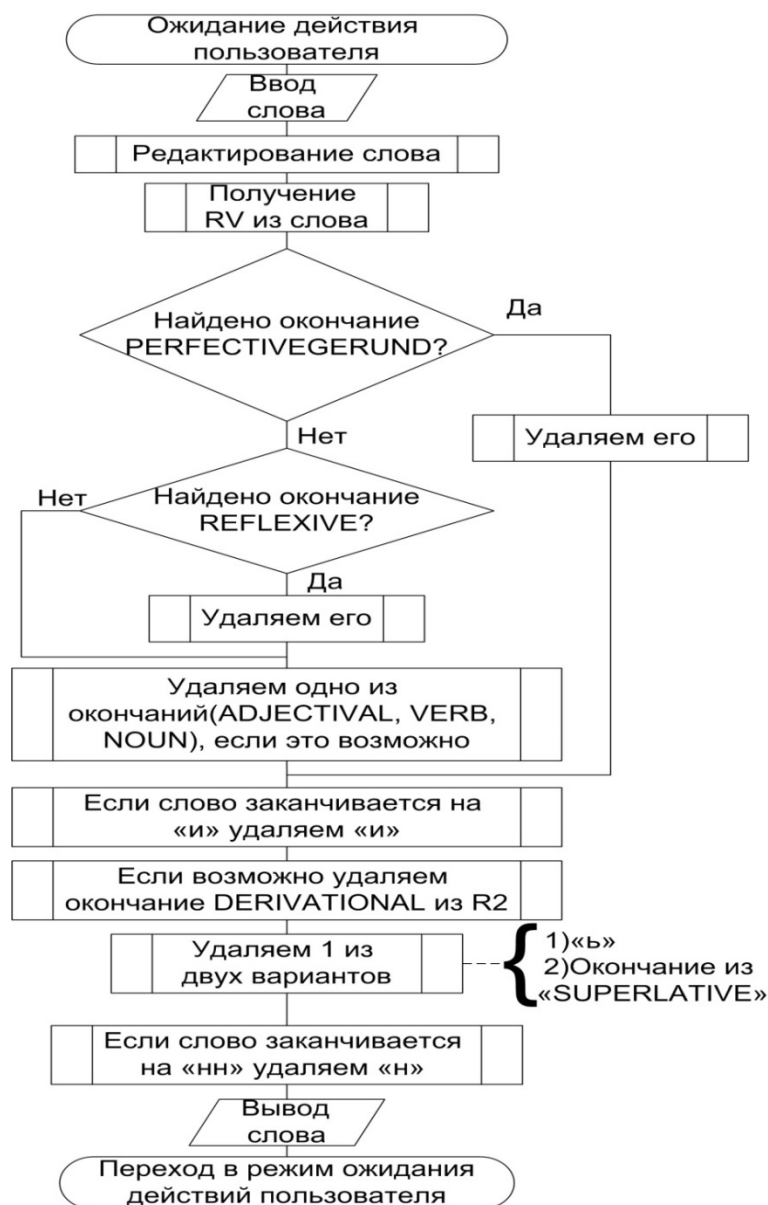


Рис. 1. Схема-алгоритма Портера

Регулярные выражения, которые применяются в алгоритме, представлены ниже. В них описаны наборы букв, на которые может заканчиваться слово, относящиеся к соответствующей части речи:

PERFECTIVEGERUND (деепричастие совершенного вида)

((ив|ивши|ившись|ыв|ывши|ывшись))((?<=[ая]*)(в|вши|вшись)))\$**;

REFLEXIVE (возвратный глагол) (ся|сь)\$;

ADJECTIVE (прилагательное)

(ее|ие|ые|ое|ими|ыми|ей|ий|ый|ой|ем|им|ым|ом|его|ого|ему|ому|их|ых|ую|юю|ая|яя|ою|ею)\$

PARTICIPLE (причастие)

((ивш|ывш|уюш)|((?<=[ая])(ем|нн|вш|юш|щ))))\$;

VERB (глагол)

((ила|ыла|ена|ейте|уйте|ите|или|ыли|ей|уй|ил|ыл|им|ым|ен|ило|ыло|ено|ят|ует|уют|ит|ыт|ены|ить|ыть|ишь|ую|ю)|((?<=[ая])(ла|на|ете|йте|ли|й|л|ем|н|ло|но|ет|ют|ны|ть|ешь|нно))))\$;

NOUN (существительное)

(а|ев|ов|ие|ье|е|иями|ями|ами|еи|ии|и|ией|ей|ой|ий|й|иям|ям|ием|ем|ам|ом|о|у|ах|иях|ях|ы|ь|ию|ью|ю|ия|ья|я)\$;

DERIVATIONAL (Слова с двумя корнями)

(ост|ость)\$;

SUPERLATIVE (прилагательное превосходной степен)

(ейше|ейш)\$;

* - наборам букв, указанным справа от (?<=[ая]), должна предшествовать буква «а» или «я».

** - символ, указывающие на проверку нахождения набора букв в конце слова.

3. Пример латентно-семантического анализа

Далее рассмотрим пример анализа. Было выбрано девять заголовков с различных новостей. Были взяты те заголовки, в которых есть повторяющиеся слова. Далее из заголовков были исключены все стоп-слова (союзы, частицы, предлоги и другие слова, не имеющие смысла). Были убраны слова, встречающиеся только в одной статье. Этот этап упрощает математическое вычисление. Далее была произведена операция стемминга для остальных слов с использованием алгоритма Портера. Если ее не производить, вычисления будут намного сложнее. В итоге у нас остались, так называемые, индекслируемые слова, они выделены жирным шрифтом:

1. Британская **полиция** знает о местонахождении **основателя WikiLeaks**.
2. В **суде США** начинается процесс **против** россиянина, рассылавшего спам.
3. **Церемонию вручения Нобелевской премии** мира бойкотируют 19 **стран**.
4. В **Великобритании арестован основатель** сайта **Wikileaks** Джулиан Ассандж.
5. Украина игнорирует **церемонию вручения Нобелевской премии**.
6. Шведский **суд** отказался рассматривать апелляцию **основателя Wikileaks**.
7. НАТО и **США** разработали планы обороны **стран** Балтии **против** России.

8. **Полиция Великобритании** нашла **основателя WikiLeaks**, но, не арестовала.

9. В Стокгольме и Осло сегодня состоится **вручение Нобелевских премий**.

Далее составляется матрица термин документ, в которой строки это индексируемые слова, а столбцы это номера заголовков. В каждой ячейке пишем, сколько раз встретилось слово в той или иной статье. Матрица-термин документ представлена в таблице 3.

Таблица 3

Матрица термин-документ

	T1	T2	T3	T4	T5	T6	T7	T8	T9
wikileaks	1	0	0	1	0	1	0	1	0
арестова	0	0	0	1	0	0	0	1	0
великобритан	0	0	0	1	0	0	0	1	0
вручен	0	0	1	0	1	0	0	0	1
нобелевск	0	0	1	0	1	0	0	0	1
основател	1	0	0	1	0	1	0	1	0
полиц	1	0	0	0	0	0	0	1	0
прем	0	0	1	0	1	0	0	0	1
прот	0	1	0	0	0	0	1	0	0
стран	0	0	1	0	0	0	1	0	0
суд	0	1	0	0	0	1	0	0	0
сша	0	1	0	0	0	0	1	0	0
церемон	0	0	1	0	1	0	0	0	0

Следующим шагом мы проводим сингулярное разложение полученной матрицы. Сингулярное разложение - это математическая операция, раскладывающая матрицу на три составляющих. Т.е. исходную матрицу M мы представляем в виде [4]:

$$M = U \cdot W \cdot V^T.$$

Где W – диагональная матрица из корней собственных чисел матрицы $M \cdot M^T$, расположенных в порядке убывания. Диагональные элементы матрицы W называются сингулярными числами. Матрица U составляется из собственных векторов для $M \cdot M^T$. Матрица V составляется из собственных векторов для $M^T \cdot M$, где U и V^T – ортогональные матрицы. В итоге получаем (см. рис. 2):

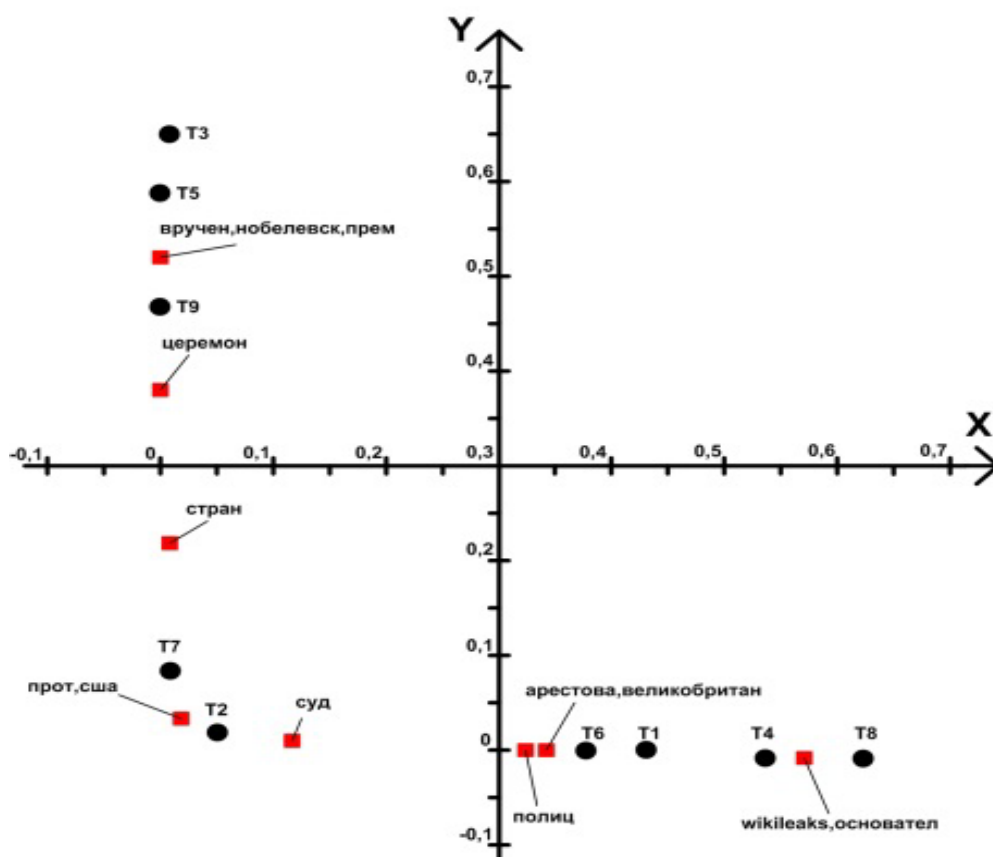


Рис. 4. Семантическое пространство

Из данного графика видно, что статьи образуют три независимые группы. Первая группа статей располагается рядом со словом «wikileaks» и «основател», и действительно, если посмотреть названия этих статей становится понятно, что они имеют отношение к wikileaks и к основателю. Другая группа статей образуется вокруг слова «США», и действительно, в них идет обсуждение США. В связи с этим можно сказать, что произведен кластерный анализ текста.

Заключение

На практике, конечно, количество групп будет намного больше, пространство будет не двумерным а многомерным, но принцип остается тот же. Можно определять местоположения слов и статей в пространстве и использовать эту информацию для определения тематики статьи.

Этот подход используется в поисковой машине Google. Например, запрос «apple computer» будет интерпретирован как относящийся к компьютерной компании Apple, поэтому среди результатов будут документы, описывающие технологии этой компании (даже если эти документы не содержат слов «apple» или «computer»).

В ходе исследования был реализован алгоритм Портера для стемминга на языке C++ и освоена математическая модель, лежащая в основе метода латентного семантического анализа, а также изучен используемый в нём математический аппарат. Были проанализированы статьи с использованием рассматриваемого алгоритма.

Список литературы

1. Латентно-семантический анализ. Режим доступа: <https://ru.wikipedia.org/wiki/Латентно-семантический-анализ> (дата обращения 10.05.2015).
2. Стеммер Портера для русского языка. Режим доступа: <http://www.algorithmist.ru/2010/12/porter-stemmer-russian.html> (дата обращения 15.05.2015).
3. Сравнение библиотек. Режим доступа: http://www.boost.org/doc/libs/1_41_0/libs/regex/doc/gcc-performance.html (дата обращения 12.05.2015).
4. Кристофер Д. Маннинг, Прабхакар Рагхаван, Хайнрих Шютце. Введение в информационный поиск. М.:Издательство Вильямс, 2011. 506 с.