

① → 10

② → 11

③ → 14

④ → 27

⑤ → 28

⑥ Применение тезаурусов для расширения запросов.
Информационно-поисковый тезаурус (ИПТ) - нормативный словарь терминов предметной области, создаваемый для улучшения качества поиска в данной области.

Расширение запроса:

- для каждого термина t в запросе происходит расширение запроса словами из тезауруса, близкими по смыслу.
- расширение → вставлять в вектор запроса с более низкими весами
- вставлять в бинарное выражение

Используется в предметно-ориентированных системах.

⊕ увеличение полноты поиска

- ⊖ Обратно снижает точность для многозначных слов
- сложность создания и обновления тезауруса

⑦ Метрики качества для систем автоматической классификации/рубрикации. Микро- и макро-усреднение.

P - точность, R - полнота

$$P = \frac{TP}{TP+FP}, R = \frac{TP}{TP+FN}$$

Комбинированная мера: $F1 = \frac{2PR}{P+R}$

Макроусреднение:

- посчитать F1 для каждого класса
- посчитать среднее

Микроусреднение:

- посчитать TP, FP, FN для каждого класса
- суммировать каждый показатель
- посчитать F1 для сумм

8 → 44

9 → 45

10 → 46

11 → 47

12 Ф-ла расчета и особенности применения меры Mutual information

$$MI = \log_2 \frac{f(a,b) \cdot N}{f(a) \cdot f(b)}$$

N - размер корпуса

f - частота встречаемости

- MI:
- завышает значимость редких словосочетаний, но при этом выявляются случайные словосочетания (опечатки)
 - требует порог отсеивания по частоте снизу

• обычно подбирается пороговое значение сверху

13 Что такое термины? (опр, примеры)

Термин - слово или словосочетание, являющееся точным обозначением понятия какой-либо спец. области науки, техники, искусства и т.д.

Примеры: уравнение окружности, разностный метод, формальный параметр.

14 Признаки для автоматического извлечения терминов из текстов.

1) Частотные признаки

• tf-idf

• модификация tf-idf с тестовой и контрастной коллекцией (idf берется по контрастной коллекции)

2) Признаки, использующие контекстную информацию

• C-Value = $TF - \frac{\sum_{p \in P} TF(p)}{|P|}$, P - мн-во словосочетаний с данным словом.

Идея - объединение частотности слов-кандидатов и информации о контекстных словах

3) Лингвистические признаки

• только прил. и сущ.

• сущ. в и.п.

• слова с заглавной буквы не в начале предл.

15) Автоматическое аннотирование. Виды автоматических аннотаций.

Автоматическое аннотирование — автоматическая технология, передающая в краткой форме основное содержание документа

Назначение: быстрое ознакомление с содержанием.

Типы аннотаций:

- 1) абстракты vs. экстракты
↓
порождаются ↓
получаются извлечением
фрагментов изх. текстов
- 2) по содержанию
↓
индикативные (упор на само содержание) информативные (упор на цифры, данные) оценочные (упор на мнения)
- 3) по фокусу → общее содержание
→ в ответ на запрос
- 4) по структуре → связная аннотация
→ ключевые слова
- 5) по кол-ву документов
- 6) по используемому языку

16) Методы и признаки для отбора предложений в экстрактивном методе.

1) Частотные подходы

а) метод *Sim Basic* — наиболее частотные слова с большей вероятностью должны оказаться в аннотации

• на каждой итерации происходит расчет вероятностей слов $\Rightarrow$ предп. с макс. средней вероятностью слов

• пересчет вероятностей для слов из отобранного предп. $\Rightarrow$ снижение повторов

б) использование $Hf - IdF$

↑
в иех. док. ↑
в корпусе, из которого извлечен док-т

2) Подходы, основанные на методах опт.

а) локальная оптимизация — оптимизация выбора следующего предложения

б) глобальная оптимизация — оптимизация нек. ф-ции

• ЛП

в) метод ММР

3) Графовые методы: вершины — предложения, дуги — сходство

4) Методы на основе машинного обучения
Характеристики для извлечения предлж.

- частотные слова текста
- укт заголовка
- присутствие ключевых слов и конструкций
- позиция в тексте
- связность с предыдущим
- новизна информации

17) Метод MMR

Итеративный метод - на каждой итерации произв-ся ранжирование предложений - кандидатов

Q - запрос к системе

S - мн-во предложений-кандидатов

$s \in S$ - рассматриваемое предложение

E - мн-во выбранных предлж.

$$MMR = \arg\max_{s \in S} [\lambda \text{Sim}_1(s, Q) - (1 - \lambda) \max_{s_j \in E} \text{Sim}_2(s, s_j)]$$

18) Метрика Rouge

Идея - сравнение сгенерированного текста с эталонным составленным человеком

$$\text{Rouge-N}(A_i) = \frac{\sum_{M_{ij}} \text{count}(\text{Ngram}(A_i) \cap \text{Ngram}(M_{ij}))}{\sum_{M_{ij}} \text{count}(\text{Ngram}(M_{ij}))}$$

A_i - оцениваемая обзорная аннотация

M_{ij} - ручные аннотации i-го кластера

$\text{Ngram}(D)$ - мн-во всех n-грамм из мн-ва документов D.

19) Метод пираинга для тестирования автоматических аннотаций

• Эксперты выделяют из "эталонных" аннотаций "информационные единицы" - SEM

• Каждой SEM получает вес, равный кол-ву "эталонных" аннотаций, где она встречалась

• Оценка - суммарный вес входящих SEM

• Итоговый рез = $\frac{\text{сумма весов SEM в авторских анно.}}{\text{сумма весов SEM в экспертных анно.}}$

20) Исправление неоварных ошибок на основе применения правила Байеса.

x - неправильное слово

$\hat{w}$ - правильное:

$$\hat{w} = \arg\max_{w \in V} P(w | x) = \arg\max_{w \in V} \frac{P(x | w) P(w)}{P(x)} =$$

$$= \arg\max_{w \in V} P(x | w) P(w).$$

Порождается мн-во кандидатов.

$$P(w) = \frac{C(w)}{T_{\text{корпус}}} \leftarrow \text{кол-во вхождений } w \text{ в корпус}$$

$P(x|w)$ - вероятность перехода/исправления (удаление, вставка, замена, перестановка).

$$P(x|w) = \begin{cases} \frac{\text{del}[w_{i-1}, w_i]}{\text{count}[w_{i-1}w_i]} & \text{- удаление} \\ \frac{\text{ins}[w_{i-1}, x_i]}{\text{count}[w_{i-1}]} & \text{- вставка} \\ \frac{\text{sub}[x_i, w_i]}{\text{count}[w_i]} & \text{- замена} \\ \frac{\text{trans}[w_i, w_{i+1}]}{\text{count}[w_iw_{i+1}]} & \text{- перест.} \end{cases}$$

$\text{del}[x, y]$ - кол-во xy написано как x

$\text{ins}[x, y]$ - кол-во x написано как xy

$\text{sub}[x, y]$ - кол-во y написано как x

$\text{trans}[x, y]$ - кол-во xy написано как yx

21 Исправление ошибок перехода в другое словарное слово на основе применения метода Байеса.

Список кандидатов: самое слово; слова на расстоянии 1-2; слова, близкие по звучанию.

Выбор лучшего кандидата - noisy channel:

• дано предложение $w_1, w_2, \dots, w_n$.

• мин-во кандидатов для каждого w_i :

$$\text{candidate}(w_1) = \{w_1, w_1', w_1'' \dots\}$$

;

$$\text{candidate}(w_n) = \{w_n, w_n', w_n'' \dots\}$$

• выбрать посл-ть $W = \text{argmax } P(W)$

Вероятности подсчитываются по корпусу + channel model как для лексической ошибки + $P(w/w) \approx 0.9 \dots 0.99$

22 Учет контекста при исправлении ошибок написания

• учет условных вероятностей появления одного слова после другого:

$$P(w_1 \dots w_n) = P(w_1) \cdot P(w_2|w_1) \dots P(w_n|w_{n-1})$$

• вероятности насчитываются на корпусе, нужно посчитать $P(w_i)$ - вероятность угадывания и $P(w_n|w_{n-1})$ - вероятность биграмм

• для угадывания $P(w_i) > 0$, но $P(w_i, w_{i-1})$ м.б. = 0

⇒ нужно слаживание

23 Метод тестирования В-О системы. Отличия от тестирования стандартной системы ИТ.

Особенности ВО систем:

• нужно не только выдать документ, но и выдать фрагмент с ответом на вопрос

• длинные формулировки запросов ⇒ меньше

слова, учёт синтаксической структуры, в тексте могут быть синонимы

- тип вопроса и ответа должны совпадать

Методы оценки в TREC:

- засчитываются первые три ответа системы
- Метрика MRR - за правильный ответ на 1-м месте - 1 балл, на втором - 0.5, на третьем - 0.25
- вычисляется среднее по всему мн-ву вопросов

(24) Позиционный индекс в поисковой системе. Зачем нужен, как обрабатывается.

Процесс индексирования в ПС:

- идентифицирует и сохраняет тексты для индексирования
- трансформирует документы в индексные термины
- берёт индексные термины и создаёт индексы для быстрого поиска

Создание индекса:

- собирается статистика по документам (частоты и позиции слов и других признаков)
- Эта инф. используется в ранжирующем алгоритме
- вычисление веса для индексного термина (ед. $\text{tf} \cdot \text{idf}$)

~~инвертирование индекса - преобразование матрицы документ-терм в данные терм-документ ($t \rightarrow \{docid\}$)~~

- ~~• сжатие данных, распределённое хранение индекса - для быстрой обработки запросов~~

(25) PageRank (зачем нужен, как вычисляется)

Пользователь кликает со случайной страницы, на каждом шаге переходит на следующую с равной вероятностью $\Rightarrow$ в пределе каждая страница получает рейтинг посещения $\sim$ вес страницы. С тупиковых страниц - переход на случайную.

$$PR(A) = d + (d-1) \left(\frac{PR(T_1)}{C(T_1)} + \dots + \frac{PR(T_n)}{C(T_n)} \right)$$

$PR(T_n)$ - исходная значимость

$C(T_n)$ - кол-во исходящих ссылок со стр.

d - значимость страницы без входящих ссылок (коэф. телепортации)

Матем. основа - марковские цепи из n состояний.

P - матрица переходов $n \times n$. PageRank - вектор стационарного состояния ($a = aP$ или считать итеративно $xP, xP^2, \dots$, x - начальный вектор).

24) Позиционный индекс

Используется для обработки фразовых запросов

Хранятся позиции терма в каждом документе

Для сопоставления используется алгоритм merge.

Позиционный индекс стандартно используется для обработки запросов на близость и фразовых запросов.

26) Методы приближенного вычисления сходства документов в реальных поисковых системах.

1) Неточная векторная модель:

- найти A кандидатов $K < |A| \ll N$; A не обязательно содержит все K лучших документов, но содержит многие из них

- всегда K лучших из A

- Методы: рассмотреть только слова с высоким idf • рассмотреть только документы с большим кол-вом терминов запроса

2) Shapiron lists

- заранее для каждого терма t вычисляются

r документов с max весом t .

- r выбирается во время индексирования и может зависеть от t .

- во время запроса считаются только веса документов, входящих хотя бы в один топ- r одного из термов запроса

- выбирается K лучших документов из топ-листов

3) Учет качества документа

- авторитетность - кач-во документа, не зависящее от запроса

- пример факторов: стр. Wiki, статьи из СМН, научные статьи с большим кол-вом цитиров.

- $g(d)$ - вес кач-ва

- $net\text{-}score = \alpha * g(d) + (1 - \alpha) * \cos(q, d)$

- хранить список лучших документов для слова по $net\text{-}score$

- ищем лучшие K результатов только из этого списка

27) Обработка фразовых запросов и вопросов с указанием близости слов

Для фразовых запросов недостаточно хранить только индекс терм-документ

- 1) Биграммный индекс: строится на парах слов из запроса

- длинные запросы нарезаются на пары

→ могут встречаться док, которые содержат пары слов, но не фразу;
сверхбольшой индекс;
проблемы со случайным вхождением слов в биграмму

2) Позиционный индекс - см. Л24

Обработка запросов с операторами близости:

- может использоваться только позиционный индекс
- возможно комбинирование позиционного и биграммного

Разборщик запросов - запрос выполняется как несколько запросов к системе, пока не будет набрано K лучших документов.

28 Как учитываются при поиске различные зоны документа

Документ имеет много частей, имеющих свой смысл: автор, заголовок, дата публикации, язык, формат - метадаанные

Создаются отдельные индексы по зонам метадаанных - параметрический индекс

Многоуровневые списки:

- разбиение индексов по уровням

- инвертированный список разбивается на уровни снижающиеся важности
- при поиске используются индексы большей важности

29 Какие факторы помимо веса $tf \cdot idf$ учитываются в поисковых моделях, как создаются многофакторные модели и т.

Факторы для агрегирования весов:

- скалярное pr -е
- близость слов запроса
- ссылки
- связь-е тематике запроса и док-та

Комплексная ф-я релевантности:

- частота термов
- совместно встречаемость термов
- зоны док-та
- полное вхождение слов в запрос
- особая обработка длинных документов

$$W = k_f W_f + k_p W_p + k_{ps} W_{ps}$$

↑ частотность термов по $tf \cdot idf$ ↑ встречаемость пар соседних слов запроса ↑ вес наилучшего пассажа в док-те

30 Методы поискового спама

- повторы терминов
- "мета-теги" с избыточным повторением
- клоаксинг - информация, выдаваемая краулеру и пользователю, различается
- дорбей - страницы оптимизируются под заданное ключевое слово, которое перенаправляет на нужную страницу
- ссылокный спам, линкфермы (покупка ссылок)
- порождение искусственных текстов

31 Методы противодействия поисковому спаму

- сигналы качества - голоса авторов, голоса пользователей
- ограничения на мета-слова
- анализ ссылок
- использование машинного обучения
- редакторская политика: чёрные листы, просмотр частотных запросов, жалобы пользователей, шаблоны для обнаружения подозрительных сайтов

? 32 Байесовский классификатор в задаче обнаружения поискового спама

Необходима коллекция известного спама - используется для обучения.

Классификатор определяет вероятность того, что данная страница/сайт содержит спам.

33 Методы порождения искусственных текстов и зачем это нужно в поисковой оптимизации?

- 1) Цепи Маркова (ищется в исходном тексте каждое следующее слово)
- 2) Метод фокуса внимания
- 3) Метод с использованием словарей

При поисковой оптимизации сайтов исп-ся для генерации названий, описаний и содержания целых сайтов с большим кол-вом страниц и "уникальным контентом"

34 Алгоритм HITS

В ответ на ~~запрос~~ запрос вместо упор. ли-ва страниц выдаётся 2 ли-ва взаимосвязанных страниц:

- хабы - хорошие списки ссылок по теме
- авторитеты - часто упоминаются в хабах

Хорошая хаб-страница указывает на многие авторитетные страницы

Хорошая авторитетная страница указывается большим кол-вом хороших хабов

Для каждой страницы рекурсивно вычисляются значения как пофудника и авторитета

Построение базового набора страниц:

- 1) по запросу получаем страницы, содержащие слова запроса - это корневой набор

- 2) добавляем $\forall$ страницу, которая указывает на стр. из корневого набора или на которую есть из него ссылка $\Rightarrow$ получим базовый набор

Разделение хабов и авторитетов:

- 1) Минимизация hub score $h(x)$ и authority score $a(x)$

- 2) Итеративно пересчитываем: $\forall x$
$$h(x) \leftarrow \sum_{x \rightarrow y} a(y)$$

$$a(x) \leftarrow \sum_{y \rightarrow x} h(y)$$

$a(x)$ вычисляется как сумма значений оценок посреднических страниц, которые на неё указывают

$h(x)$ вычис-ся как $\sum$ значений оценок авторитетности страниц, на кот. она указывает

Значение в итоге сходится.

35) Метод шиплов для определения дубликатов документов

Шиплы - пословные n-граммы

- 1) Канонизация текстов - очищение текстов от слов и символов, которые не должны участвовать в сравнении; нормализация
- 2) Разбиение на шиплы внахлест
- 3) Вычисление контрольных сумм шиплов через 24 независимые статические хэш-функции.

4) Случайная выборка из каждого набора хэшей

5) Сравнение полученных массивов на соответствие

36) Особенности использования кликов пользователя в качестве фидбека

Отсутствие кликов - не релевантно/не верно

Неравномерность позиций относ-но кликов:

- более высокие позиции получают больше кликов
- справедливо, если переставить выдачу $\Rightarrow$ клики информативны, но смещены

Гипотеза о наблюдении:

- документ должен быть прочитан через кликом; условная вероятность клика после прочтения зависит от релевантности док-та
- вероятность клика - ~~не~~ глобальный коэффициент (вер-ть увидеть, зависит от позиции), локальный коэффициент (зависит от пары <запрос, документ>)

37) Классификация запросов по цели. Особенности обработки разных типов запросов.

- 1) Информационные запросы - пользователю важен конкретный результат
- 2) Навигационные запросы - для поиска конкретного адреса веб-ресурса

3) Транзакционные Запросы - для
приобретения конкретных товаров и услуг

Иная представление о классах запросов, можно
лучше анализировать выдачу, строить гипотезы
для изучения атрибутов ранжирования,
совершенствовать и оптимизировать сайт.