

различные модели информационного поиска

- векторная модель
- классич. модель ВМ
- ВМ25
- языковая модель

① Основные св-ва естественного языка

- неограниченная семантическая мощность - способность к передаче инф. относительно любой области
- эволютивность - неограниченная способность к бесконечному развитию и модификациям
- проявление языка в виде речи
- энтимность
- двойственность
- словарь и граммат. правила определяются практикой применения и не всегда могут быть формально зафикс.

② Что такое графематический анализ

- начальный этап обработки текста, в ходе которого определяются элементы грамматической структуры (слова, знаки пунктуации, числа, сокращения)

③ Что такое лемматизация

- процесс приведения словоформ к лемме - её нормальной (словарной форме).

④ Как работает словарный морф. анализ

Каждой словоформе сопоставлена основа или лемма
Для служебных слов:

- выделить возможные окончания длиной 0-3 симв.
- для каждого окончания определить код по табл. окончаний и номер флексивного класса по словарю лемм
- проверить их согласованность по морф. таблице
- сохранить вариант, если согласуются

5) Как морф. анализаторы обрабатывают слова, отсутствующие в словаре

Предсказание - выявление аналогий со словоформами, распознаваемыми имеющимся словарем

Алгоритм предсказания:

- 1) предсказание префиксального образования
- 2) предсказание по концовке, взятая из известных словоформ

По префиксу:

- попытка найти словоформу, макс. совпад. справа
- если префикс $\leq M$ символов, а п.г. $\geq N$ симв., то слово разбирается по образцу известной словоформы

По концовке

- опред-ся известные концовки длины K
- сопоставление с ~~известными~~ морф. характеристик.
- заносится в иск. лексикон, если встречается не менее L раз и меньше конкурентов в пределах одной части речи

Всегда предусматривается разбор сущ.

6) Постморф. анализ. Основные методы

представительский

- процедура, предназначенная для упрощения морфологической обработки слов
- выбор правильной леммы
- уточнение морф. характеристик

Основные методы:

- написание правил
- статистические методы (на основе размеченного корпуса)

7) Основные понятия инф. поиска

• IR - процесс поиска в большой коллекции нек. неструктурированного материала, удовлетворяющего информационные потребности.

• Релевантность - св-во документа содержать информацию, которую ищет пользователь

• Оценка качества - экспериментальные процедуры и меры для сравнения результатов работы систем с ожиданиями пользователей

• Потребность пользователя, инф. потребность

8) Виды поисковых систем по охвату и направленности. Особенности разных типов поисковых систем.

По охвату:

1) Интернет-поиск

- собирает рез-ты по общедоступному Интернету
- проблема ранжирования рез-тов
- большие объемы
- Интернет-реклама
- хорошее качество

2) Корпоративный поиск

- собирает инф. разных форматов из совокупности хранилищ
- ранжирование док. разного типа
- относительно малый объем
- хуже кач-во поиска

По направленности:

- мед. поиск
- патентный поиск
- научный
- по законодательству
- ...

9) Особенности научного поиска

- рост публикаций
- устаревание
- проблема литературы до эпохи интернета
- цитаты, которые можно использовать как ссылки

10) Основные этапы обработки текстов в поисковой машине

1) Извлечение текстов

- Краулер идентифицирует и извлекает док-ты для поисковой машины.
- Хранилище документов - тексты, метадаанные
- Быстрый доступ к содержанию

2) Преобразование текстов

- Анализатор обрабатывает послед-ть токенов, распознаёт структурные элем.
- Токенизатор распознаёт слова
- Обработка структуры HTML, XML
- Удаление наиболее частотных слов
- Стеминг (морф. анализ)
- Анализ ссылок
- Извлечение инф. (распозн. имен, мест, дат ...)
- Классификация (по тематике, topicality и т.д.)

3) Создание индекса

- Статистика по документам (частота слов и т.д.)
- Определение весов для индексных термов
- Инвертирование табл. документ-терм
- Распред. хранение

11) Основные этапы обработки запроса - " -

1) Взаимодействие с пользователем

- Ввод запроса (интерфейс и парсер)
- Трансформация запроса (спеллинг, подсказки, расширение)
- Вывод рез-тов

2) Ранжирование

- приращивание веса соответствии документа запросу

- Оптимизация выполнения запроса
- Распред. выполнение

3) Оценка качества

12) Булевская модель IR. + и -

Два возможных рез-та TRUE и FALSE

Поиск по полному совпадению

Запрос составляется с использованием булевских операторов и операторов близости.

Для каждого термина t хранится список документов, содержащих t . Документ $\rightarrow docID$

Создание инверт. индекса:

- 1) преобразование док. в поток токенов
- 2) модификация токенов
- 3) пары (token, docID)
- 4) сортировка по токенам и затем docID

Для булевских выр. пересечение, объединение, дополнение
списка ~~текст~~ док. где входящих в выр. токенов

- ⊕ • рез-ты предсказуемы, их легко объяснить
- могут быть встроены различные признаки
- эффективная обработка

- ⊖ • качество выдачи зависит от пользователя
- простые запросы дают слишком много рез-тов
- длинные запросы сложно составлять

13) Как измеряется кач-во булевского поиска

Рез-ты не ранжированы

Поисковая система раздает коллекцию на
ли-ва: выдано - не выдано

Эксперт раздает на: релевантный - нерелевантный

Меры качества:

- точность P - доля релевантных в выданных
- полнота R - доля выданных среди релевантных
- F-мера $F1 = \frac{2PR}{R+P}$

14) Алгоритм сопоставления запроса с док-тами (Merge)

Пересечение списков $doc \neq 0$

answer $\leftarrow \{ \}$

while $(p1 \neq ML) \& \& (p2 \neq Nil)$ do

if $doc \neq 0(p1) = doc \neq 0(p2)$ then

ADD (answer, $doc \neq 0(p1)$)

$p1 \leftarrow next(p1)$

$p2 \leftarrow next(p2)$

else if $doc \neq 0(p1) < doc \neq 0(p2)$ then

$p1 \leftarrow next(p1)$

else

$p2 \leftarrow next(p2)$

Если $doc \neq 0$ не равны, двигаться по списку, где $doc \neq 0$ меньше

Необходимо, чтобы списки были отсортированы по $doc \neq 0$.

Тогда будет $O(len(p1) + len(p2))$ операций.

15) Векторная модель инф. поиска

Документы и запрос представляются в виде векторов N -мерного пространства.

N - общее кол-во термов.

$v(i)$ - вес i -го термина в документе.

Для каждого док-та считается вектор:
для каждого термина - tf в этом документе.

Для каждого ~~слова~~ термина считается idf :

$$idf = \log_{10} \frac{N}{df_t}, \text{ где } df_t - \text{кол-во док-тов, содержащих}$$

Формализация близости векторов - косинусная мера (вес)
Документы ранжируются по мере сходства веса
Выдаются первые K док-тов.

16) Смысл показателей idf и $tf \cdot idf$.

tf - сколько раз терм встречается в документе

df - в скольких документах встречается терм

$idf = \log_{10} \frac{N}{df}$ - увеличивается для редких термов

$tf \cdot idf$ - увеличивается при росте количества упоминаний термина в документе и для редких термов.

17) Классическая процедура оценки качества инф. поиска

Мера кач-ва инф. поиска - удовлетворенность пользователя
Приближение такой оценки - релевантность
Классич. процедура оценки.

1) Составили список запросов и ограничи кол-во документов

2) Для каждой пары запрос-документ выставили

экспертную оценку релевантности

3) Ответ системы рассматривается как множ-во или последовательность оценок релевантности

н) Построим метрики на полученных рез-тах

18) Что такое РОМИП, какие задачи решаются?

РОМИП - российский семинар по методам инф. поиска
Проведение независимой оценки методов инф. поиска

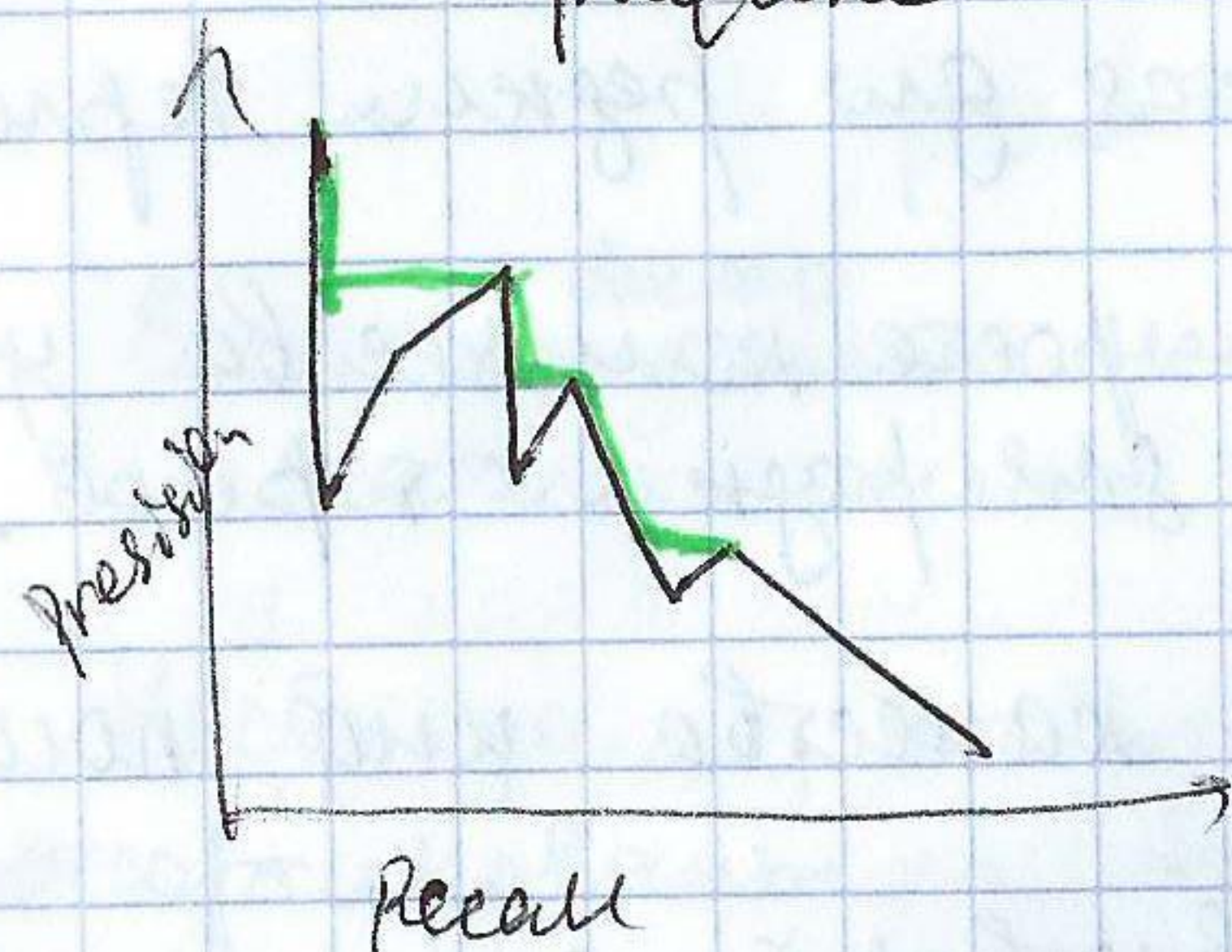
Набор - дорожек - сесий, посвященных конкретным проектам с опред. задачами и правилами оценки (например: классиф-я, кластеризация, аннотирование текстов и т.д.)

19) Что такое кривая полнота-точность?

Кривая строится для совокупности запросов.

Изначально - набор точек, нужно интерполировать.

Идея интерполяции - если локально точность возросла с увеличением полноты, то засчитываем её так, т.е. берём макс. точность на графике



20) Что такое 11-точечный график P-RCS?

Рассм-ся значения полноты 0, 0.1, ..., 1.0 - 11 шт.

$r_i \in \{0, 0.1, \dots, 1\}$, q_j - запрос. Для каждого запроса ~~выбирается~~ и значения точности вычисляются

значение интерп. точности $p(r_i, q_j)$:

- если $r_i > \text{recall}(q_j)$, то $p(r_i, q_j) = 0$

- если $r_i \leq \text{recall}(q_j)$, то $p(r_i, q_j) = \max_{n: \text{pos}(r_i, q_j) \geq \text{pos}(r_n, q_j)} (\text{precision}(n))$

где $\text{pos}(r_i, q_j)$ - мин длина списка рез-тов для запроса q_j , на которой достигается r_i .

21) Что такое пуллы в ИР? Сложности

Для каждого запроса:

- собрать рез-ты систем участников глубины A
- выбрать B первых рез-тов
- удалить дубликаты
- проставить оценки релевантности
- не оцененные зап. считаются нерелевантными
- оценить весь ответ системы (глубины A)

Сложности:

- взаимное усиление систем
- переоценка систем, не участвовавших в оценке
- полученная оценка - оценка снизу

22) Оценка кач-ва в поисковых сист. Интернет

Особенности:

- невозможно измерить полноту
- для оценки выбираются k первых документов
- нужно учитывать при оценке, что релевантные документы должны ~~быть~~ показываться раньше

23) Шкалы оценок. Мера NDCG

TRC:

раньше - бинарная

сейчас - высоко релевантный, релевантный, нерелевантный

РОМИП: соответствует, скорее соответствует, возможно соответствует, не соответствует, не м.б. оценен

NDCG: g_i - оценка i -го документа

$$DCG_{\lambda} = g_{\lambda} + \sum_{i=2}^{\lambda} \frac{g_i}{\log i}$$

IDCG - наилучшее упорядочение по запросу

$$NDCG = \frac{DCG}{IDCG}$$

24) Что такое информационно-поисковые тезаурусы? Зачем нужны и где применяются?

Инф.-поисковый тезаурус - нормативный словарь терминов предметной области, создаваемый для улучшения качества инф. поиска в данной предметной области.

Цели:

- перевод языка автора на нормативный язык, используемый для индексации и поиска
- обеспечение последовательности в присвоении индексных терминов
- обозначение отношений между терминами
- облегчение инф. поиска

Использование:

- в настоящее время - рост интереса к предметно-ориентированному и корпоративному поиску
- индексирование и расшир. запросов

25) Что такое рубрикаторы. Чем отличаются от инф.-поисковых тезаурусов?

Рубрикатор - таблица иерархической классификации, содержащая полный перечень включенных в систему классов и предназначенная для систематизации инф. фондов, массивов и изданий и поиска в них.

Термины тезауруса явл-ся фундаментально языковыми, в то время как рубрики соответствуют концептуальным категориям.

26) Методы расширения запросов пользователей при ИР.

глобальные

- ручные тезаурусы
- автоматич. эксперт. тезаурус

локальные

- обратная связь по релевантности
- псевдорелевантности

27) Что означает термин relevance feedback? Сок. принципы

Пользователь оценивает документы в поисковой выдаче:

- задает простой запрос
- размечает часть результатов как релевантные/нерелевантные
- система улучшает соответствие документов запросу на основе разметки пользователя
- процедура м.б. итеративной

Венковская идея: сформулировать хороший запрос трудно, если пользователь не знаком с коллекцией, поэтому запрос строится итеративно

28) Алгоритм Фоккио для relevance feedback

Документы - точки в многомерном пр-ве

Центроид - ЦМ этих точек: $\mu(C) = \frac{1}{|C|} \sum_{d_j \in C} d_j$, C - мн-во док.

C_r - релевантные, C_{nr} - нерелевантные

Вектор ищет запрос $q_{opt} = \arg \max (\cos(q, \mu(C_r)) - \cos(q, \mu(C_{nr})))$

$$q_{opt} = \frac{1}{|C_r|} \sum_{d_j \in C_r} d_j - \frac{1}{|C_{nr}|} \sum_{d_j \in C_{nr}} d_j$$

$$\text{На практике: } q_m = \alpha q_0 + \beta \cdot \frac{1}{|C_r|} \sum_{d_j \in C_r} d_j - \gamma \frac{1}{|C_{nr}|} \sum_{d_j \in C_{nr}} d_j$$

C_r и C_{nr} - мн-во известных рел. и нерел. векторов.

α, β, γ - веса

29) Проблемы расширения запроса при помощи обратной связи по релевантности

- длинные ~~запросы~~ запросы неэффективны:
 - большое ожидание для полнот.
 - высокая стоимость для системы
- пользователи часто не хотят размечать документы
- трудно понять, почему данный документ был выведен после relevance feedback

30) Основные методы автоматической рубрикации текстов

- ручное рубрицирование
- автоматическое
 - инженерный подход
 - метод машинного обучения
- полуматематическое

31) По какому инженерному методу классиф. текстов? $\oplus$ и $\ominus$

Рубрикатор создается вручную, содержание рубрики можно выразить ограниченным кол-вом понятий в виде формулы

Эксперты описывают смысл рубрики в виде булевских выражений, правил продукции

- $\oplus$ • правила обладают хорошими возможностями для масштабирования

- эксперты могут создать набор правил, превосходящий любой автоматический генератор классификаторов

- $\ominus$ • сложное создание и сопровождение правил

- трудно найти хороших экспертов

32) $\oplus$ и $\ominus$ — ручного рубрицирования

- $\oplus$ высокая точность

- $\ominus$ низкая скорость

33) Метод Байеса для автоматической классиф. текстов

C — набор классов, d — документ

$$P(c|d) = \frac{P(d|c)P(c)}{P(d)}, \quad c \in C$$

$$d = (x_1, \dots, x_n)$$

$$c = \underset{c_j \in C}{\operatorname{argmax}} P(c_j | x_1, \dots, x_n) = \underset{c_j \in C}{\operatorname{argmax}} \frac{P(x_1, \dots, x_n | c_j) \cdot P(c_j)}{\underbrace{P(x_1, \dots, x_n)}_{\text{один и тот же для всех классов}}}$$

$$= \underset{c_j \in C}{\operatorname{argmax}} P(x_1, \dots, x_n | c_j) \cdot P(c_j)$$

34) Метод Роккио — " —

D_c — мн-во документов $\in$ классу c .

$v(d)$ — вектор $\{f, idf\}$ для документа d .

Централизация: $\mu(c) = \frac{1}{|D_c|} \sum_{d \in D_c} v(d)$ вычисляется при обучении

Классификация основывается на близости по кос. мере

35) Метод Knn — " —

Чтобы классифицировать документ d в класс c нужно:

- определить k -окружение из k ближайших соседей d
- подсчитать число документов в N , кот. относятся к c .
- считать $P(c|d)$ как i/k
- выбрать класс $\underset{c_j}{\operatorname{argmax}} P(c_j|d)$

36) Основное ~~принцип~~ принципов метода SVM — «—».

Метод опорных векторов:

- точки имеют вид $\{(x_1, c_1), \dots, (x_n, c_n)\}$, где $c_i = \pm 1$ или -1 в зависимости от принадлежности к классу; x_i — p -мерный вектор.
- разделяющая гиперплоскость имеет вид $w \cdot x - b = 0$, w — $1-p$ к разделяющей гиперплоскости
- нужны 2 || гиперплоскости, они м.б. описаны равенствами
$$\begin{cases} w \cdot x - b = 1 \\ w \cdot x - b = -1 \end{cases}$$
- нужно максимизировать расстояние между плоскостями, $d = \frac{2}{|w|} \Rightarrow |w| \rightarrow \min$, при условиях:
$$\begin{cases} w \cdot x - b \geq 1, c_i = 1 \\ w \cdot x - b \leq -1, c_i = -1 \end{cases}$$

Результатирующая Ф-ла:

$$f(x) = \sum d_i c_i \langle x_i, x \rangle + b$$

x_i — опорные вектора

x — вектор текущего документа

d_i — вес опорных векторов

$$c_i \in \{-1, 1\}$$

37) + и — методов машинного обучения для рубрикации текстов

⊕. эффективно при наличии качественно размеченной обучающей коллекции

⊖. низкая эффективность при большом числе рубрик

- трудно интерпретируемые результаты

38) Возможности применения методов машинного обучения для рубрикации текстов в зависимости от размера обучающей коллекции

Кол-во данных.

- нет данных для обучения — правила пишутся вручную

- мало данных — классификация "с учителем" + нужно получить размеченные данные по возможности быстрее

- достаточное количество данных — подходит для применения SVM, можно использовать булевские модели

- большое кол-во данных — хорошо подходит в теории, но SVM и KNN непрактичны; можно использовать Naïve Bayes

~~Классификация~~

39) Что такое кластеризация текстов? Чем отличается от классификации?

Кластеризация — разделение совокупности документов на однородные группы, но четко отличающиеся друг от друга подмножества (кластеры)

Цель классификации (разновидность обучения с учителем) — разделить данные по категориям, установленным экспертом.

Кластеризация — без учителя, разделение на основе метрики.

40) Метод k-медиан для кластеризации текстов.

Документы - вектора над $\mathbb{R}$

Кластеры базируются на понятии центра точки в кластере c : $\mu(c) = \frac{1}{|c|} \sum_{x \in c} x$

Присваивание документов к кластеру базируется на сх-ве с текущими центрами

Алгоритм:

1) выберем K случайных векторов $\{s_1, \dots, s_K\}$ как исходные м-во (seeds) - центры

2) до выполнения критерия остановки:
для каждого c_i

$$c_i = \max_{s_j \in \text{seeds}} \text{similarity}(x_i, s_j)$$

3) обновляем м-во $\{s_1, \dots, s_K\}$ заменив на центры текущих кластеров:
 $\forall c_j$
 $s_j = \mu(c_j)$

Условия остановки:

- опред. число итераций
- не меняется разделение документов
- не меняются позиции центров

41) Агломеративная кластеризация - основной принцип и процесс

- Начинает рассмотрение документов как отдельных кластеров
- Итеративно объединяет ближайшую пару кластеров, пока не останется один
- История объединения образует иерархию.

Способы определения, что такое наиболее сходная пара:

- single-link - сходство по наиболее похожим документам
- complete-link - по наиболее непохожим
- centroid - по наиболее похожим центрам
- average-link - среднее ссз между парами элем. двух кластеров

? 42) Метод тестирования автоматической кластеризации

Критерии

внутренний:

- внутри кластера
высокое сх-во
- между кластерами
- низкое

внешний:

- измерение способности кластеризации обнаруживать скрытые классы объектов в эталонных данных

Внешняя оценка кач-ва - чистота. Нужны размеченные данные: документы с правильными кластерами. Тестируемая система порождает $w_1, \dots, w_k$ из n_i элем.

$$\text{Purity}(w_i) = \frac{1}{n_i} \max_j (n_{ij}), j \in C$$

$$\text{Rand-index} = \frac{A + D}{A + B + C + D}$$

43) Особенности кластеризации потока новостей в реальном времени

- корпус документов постоянно пополняется
- временное окно
- разные размеры

- наличие дубликатов $\Rightarrow$ определение первоисточника
- ошибки при сборе новостных сообщений
- спамерские технологии источников

44 Анализ тональности как задача классификации: методы, признаки, меры качества

Методы:

- операция оценочных выражений + правила их комбинирования
- машинное обучение

Признаки для классификации машинным обучением:

- слова с весами (упоминания)
- произвольные биграммы
- знаки препинания
- смайлики

Оценка кач-ва:

- точность-полнота-F-мера
- асимметрия (правильность) - отношение правильно классиф. ко всем остальным

45 Вероятностная модель $\pm R$: осн. идея и отличия от векторной модели

Дан запрос q и идеальное мн-во R релевантных документов. Задача - определить св-ва этого мн-ва.

Вероятностная модель пытается оценить вероятность, что пользователь оценит документ d_j как релев. посредством отношения

$$\frac{P(d_j \text{ relevant to } q)}{P(d_j \text{ nonrelevant to } q)}$$

Различие от векторной модели: сходство между запросом и документом считается по ф-ле из теории вероятности, а не как кос. мера.

$$\text{sim}(d_j, q) = \frac{P(R|d_j)}{P(\bar{R}|d_j)} = \frac{P(d_j|R) \cdot P(R)}{P(d_j|\bar{R}) \cdot P(\bar{R})} \approx$$

$$\approx \log \frac{P(d_j|R)}{P(d_j|\bar{R})} + \log \frac{P(R)}{P(\bar{R})}$$

$\exists k_i$ -слово, $q = (d_1, \dots, d_t)$

$$P(d_j|R) = \prod_{i=1}^t P(k_i = d_i | R)$$

$$P(d_j|\bar{R}) = \prod_{i=1}^t P(k_i = d_i | \bar{R})$$

$$\Rightarrow \text{sim}(d_j, q) \approx \frac{\prod_{g_i(d_j)=1} P(k_i|R) \cdot \prod_{g_i(d_j)=0} P(\bar{k}_i|R)}{\prod_{g_i(d_j)=1} P(k_i|\bar{R}) \cdot \prod_{g_i(d_j)=0} P(\bar{k}_i|\bar{R})} \approx$$

$$\approx \sum_{i=1}^t \left(\log \frac{P(k_i|R)}{1 - P(k_i|R)} + \log \frac{1 - P(k_i|\bar{R})}{P(k_i|\bar{R})} \right)$$

46 Что такое языковые статистические модели?

Языковая модель - матем. модель которая вычисляет вероятность последовательности слов или усл. вероятность следования слова в контексте.

Стат. модель порождения текста определяет вер-ть слова для данного языка

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

47 Языковая модель IR

Каждый запрос трактуется как случайный процесс.

Подход:

- настраиваем языковые модели для каждого документа
- выводим вероятность порождения запроса на основе модели каждого док.
- ранжируем док. по вероятностям

Ранжирующая ф-ла: $P(Q, d) = p(d) p(Q|d) = p(d) p(Q|M_d)$

$$p(Q|M_d) = \prod_{t \in Q} \frac{t f_{t,d}}{d l_d}$$

M_d - яз. модель док. d

$d l_d$ - число слов в d

48 Вопросно-ответные системы: осн. компоненты, тестирование

Этапы работы:

- 1) анализ вопроса (опред. типа вопроса и формирование запроса к поисковой системе)
- 2) поиск релев. док. или абзацев (упор. список, n первых)
- 3) анализ полученных док. (содержит ли тип ответа)
- 4) извлечение ответа заданного типа