

Автоматизированная система построения тезауруса AUTHECO

Цель разработки

Автоматизация процесса составления информационно-поискового тезауруса предметной области на базе корпуса текстов предметной области.

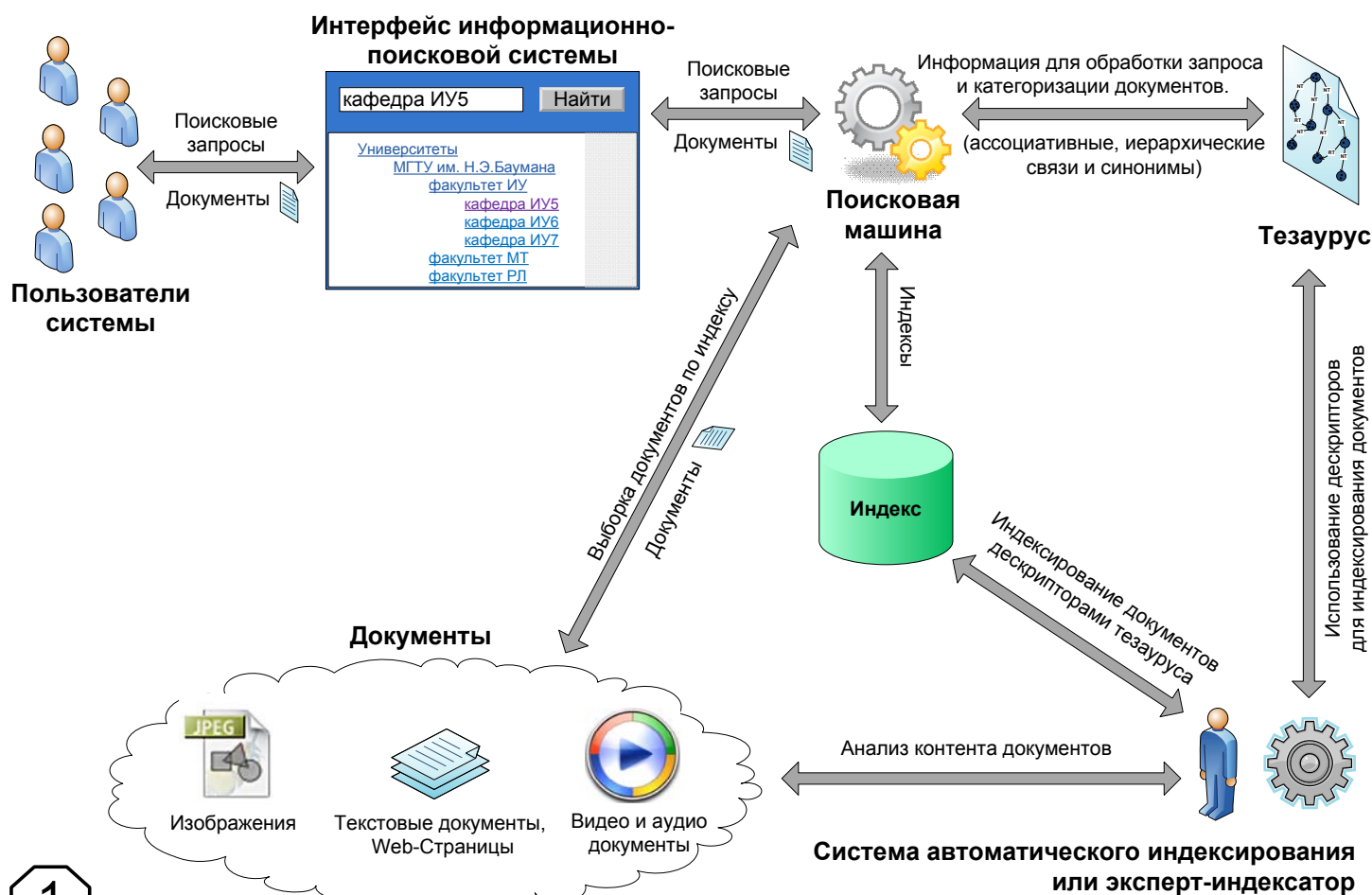
Назначение системы

Система предназначена для использования экспертами в области построения информационно-поисковых тезаурусов. Программный комплекс облегчает процесс составления семантического ресурса производя интеллектуальный анализ текстов и определяя наиболее вероятные кандидатуры слов и связей между лексическими единицами для включения в тезаурус.

Задачи, решаемые в рамках дипломного проекта

1. Анализ предметной области обработки естественных языков (NLP) и информационного поиска (IR) с использованием тезаурусов
2. Выбор методов и алгоритмов, подходящих для автоматизации процесса составления тезауруса.
3. Формализация предметной области и разработка технологии автоматизированного составления тезауруса на базе корпуса текстов данной предметной области.
4. Проектирование архитектуры системы.
5. Конструирование программной системы: кодирование алгоритмов, создание лингвистической базы данных и т.п.

Роль тезауруса в информационно-поисковой системе



Структура и назначение тезауруса

Определение

Информационно-поисковый тезаурус – контролируемый словарь терминов (нормализованной лексики) предметной области, создаваемый для улучшения качества информационного поиска в данной предметной области. Используется для индексирования документов и при поиске информации.

Назначение

Для индексирования документов:

- Предоставление стандартного набора терминов для индексирования документов
- Решение проблемы неоднозначности естественного языка и многозначности слов
- Снижение “шума” в наборе ключевых слов
- Обеспечение последовательности в присваивании индексных терминов

Для поиска информации:

- Оптимизация пользовательских запросов используя семантическую информацию
- Использование иерархических связей для обобщения и уточнения запросов
- Переформирование запроса используя синонимы

Структура

Дескрипторы – основные понятия предметной области удовлетворяющие специальным требованиям:

- Обозначает отдельное понятие
- Может быть однословным или многословным
- Должны быть однозначными
- Должны быть реально использоваться в текстах
- Для уточнений значений – комментарии и другие (точное описание можно найти в стандартах ГОСТ, ISO и ANSI/NISO).

Пример:

автомобилестроительный завод, рабочий, сверлильный станок

Отношения между дескрипторами:

NT Родо-видовые (иерархические):

завод NT цементный завод, кирпичный завод

RT Ассоциативные:

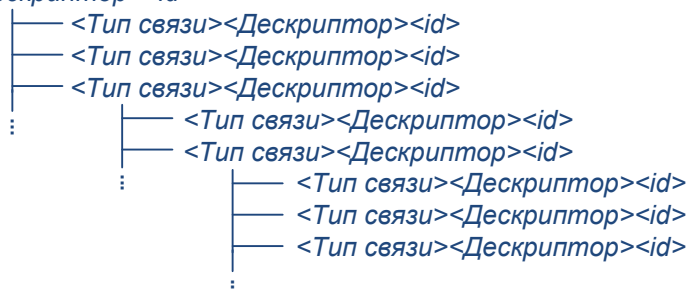
завод RT рабочий, станок, промышленность, Цех

USE Синонимии:

завод, фабрика USE промышленное предприятие.

Мини-тезаурус – подмножество тезауруса состоящее из дескрипторов и отношений относящихся к какой-либо предметной области (экономика, радиотехника, информатика и т.п.).

`<Дескриптор><id>`



`<Тип связи> ::= RT | NT | USE`

`<Дескриптор> ::= <Слово> | <Выражение>`

`<id> – число`

Примеры

Фрагмент технического тезауруса

АНТЕННЫ

NT ВЫСОКОЧАСТОТНЫЕ АНТЕННЫ
NT НИЗКОЧАСТОТНЫЕ АНТЕННЫ
NT ШИРОКОПОЛОСНЫЕ АНТЕННЫ
NT ПАРАБОЛИЧЕСКИЕ АНТЕННЫ
NT ТЕЛЕСКОПИЧЕСКИЕ АНТЕННЫ
NT ШТЫРЕВЫЕ АНТЕННЫ
NT ЩЕЛЕВЫЕ АНТЕННЫ

Фрагмент тезауруса EuroVoc
(мини-тезаурус “Industry”)

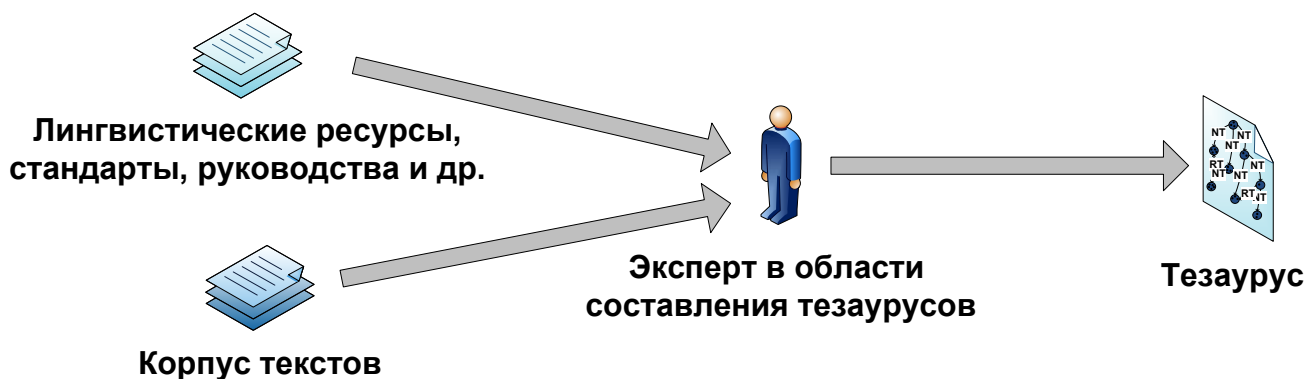
chemical industry
RT biochemistry (3606)
RT chemical pollution (5216)
RT chemistry (3606)
RT industrial chemistry (3606)
NT1 chemical accident
NT1 chemical product
RT chemical compound
RT chemical element
NT1 glass industry
NT2 glass

Фрагмент тезауруса
русского языка “РусТез”

алгоритмические языки
RT языки алгоритмические
программирование
программа
теория алгоритмов
BT программное обеспечение
формальные языки
NT си
паскаль
кобол

Ручная и автоматическая технология построения тезауруса

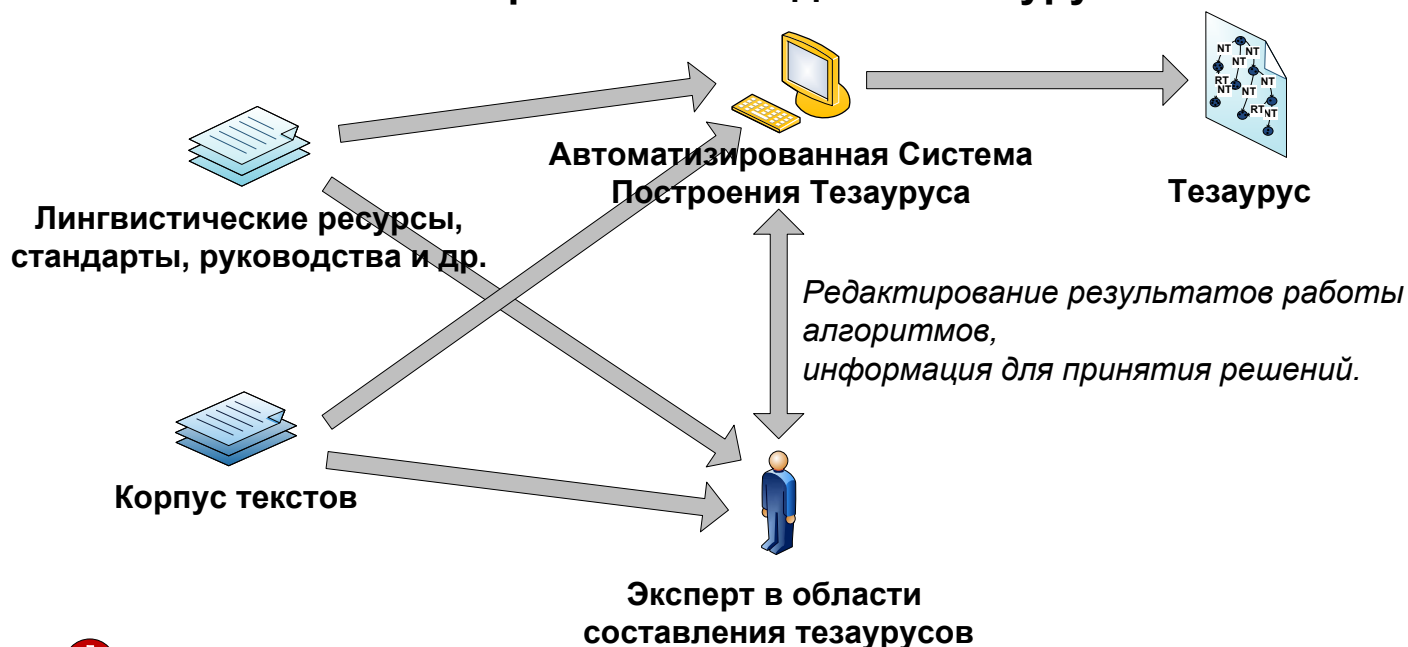
Ручное создание тезауруса



— Проблемы

1. Высокая стоимость создания тезауруса
2. Длительность процесса создания тезауруса (до 1 года)
3. Зависимость результата от квалификации и опыта эксперта
4. Невозможность проанализировать весь корпус текстов
5. Малая гибкость процесса построения тезауруса

Автоматизированное создание тезауруса



Решения проблем

- 1, 2: Уменьшение стоимости и сроков разработки тезауруса за счет снижения трудоемкости процесса составления
- 3: Снижение влияния квалификации эксперта на качество тезауруса за счет использования статистических данных при построении семантического ресурса
- 4: Автоматическая обработка всего корпуса текстов, предоставление статистической информации о словах и словосочетаниях
- 5: Возможность перезапуска процедур автоматического анализа корпуса текстов

Сравнительный анализ аналогов и прототипов

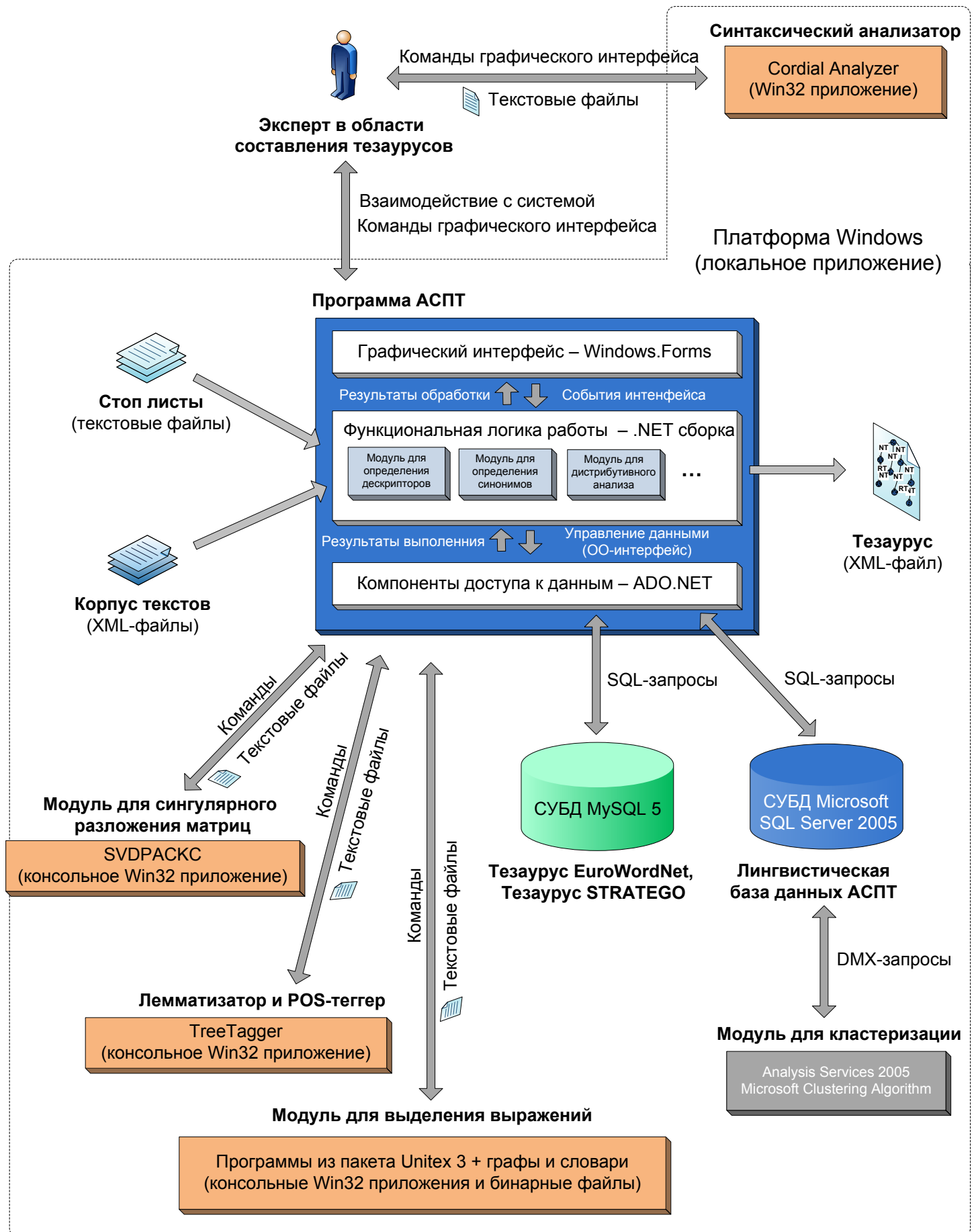
Сравнительный анализ функциональности систем

	K _{зн}	XIP	Oracle Text	GALILEI	a.k.a.	AUTHECO
Готов к использованию	0,1	—	—	+	+	+
Препроцессинг	0,15	+	+	+	—	+
Извлечение выражений из текста	0,15	+	—	+	—	+
Построение множества предпочтительных дескрипторов	0,15	—	—	—	—	+
Нахождение отношений синонимии	0,05	—	—	—	—	+
Определение отношений между дескрипторами	0,2	+	+	—	—	+
Кластеризация документов	0,1	—	+	—	—	+
Средства хранения и редактирования тезауруса	0,1	—	+	—	+	+
Сумма =		0,5	0,55	0,4	0,2	1

Сравнительный анализ по методам обработки данных

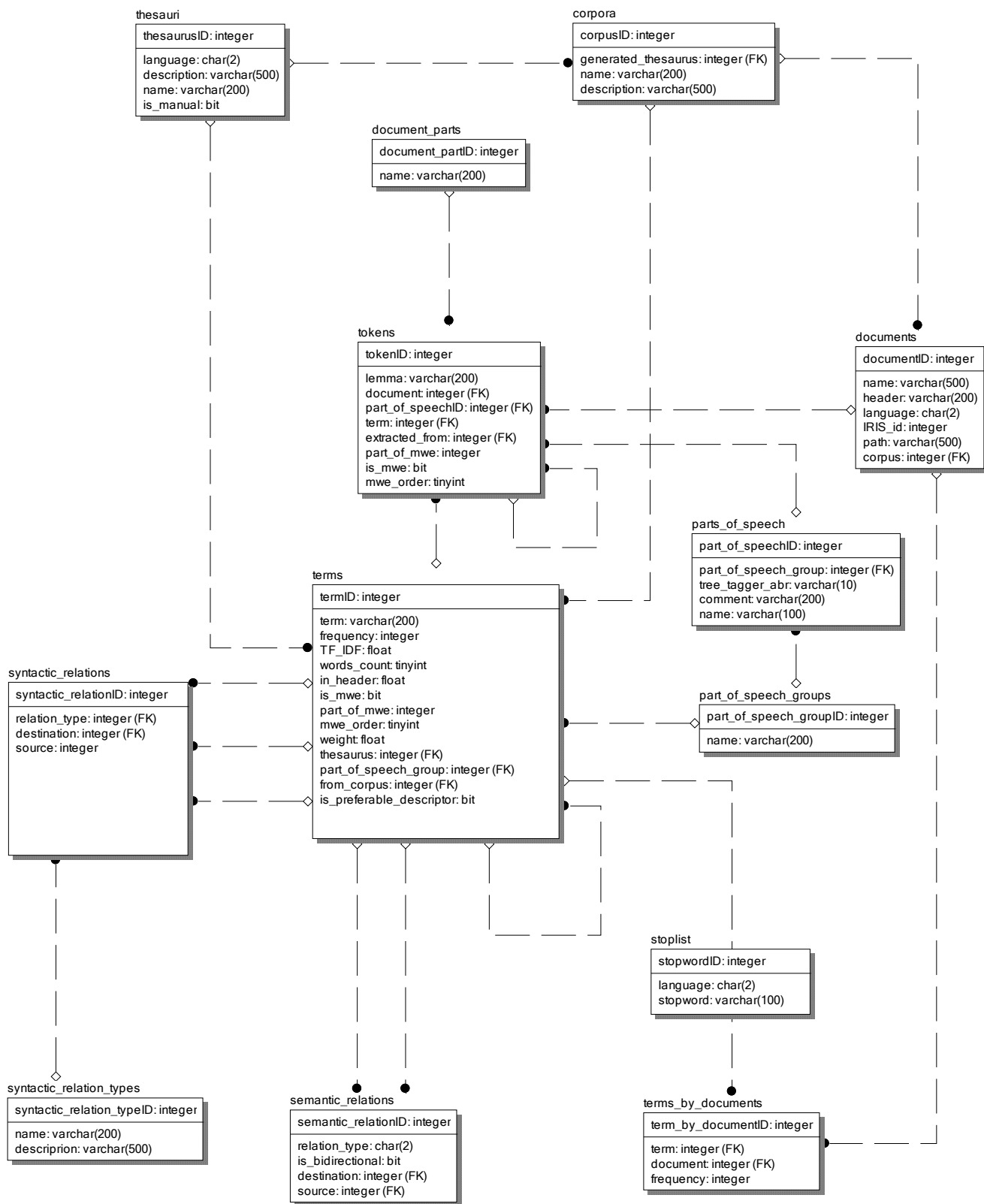
	Препроцессинг	Извлечение выражений из текста	Построение множества предпочтительных дескрипторов	Нахождение отношений синонимии	Определение отношений между дескрипторами	Кластеризация документов
XIP	Лемматизация, Части речи, Синтаксические зависимости	N-граммы, специальные грамматики и др. лингвистические ресурсы	—	—	Дистрибутивно-статистический анализ, Методы Data Mining, Алгоритм Apriori	—
Oracle Text	Лемматизация, Исправление опечаток	—	—	—	Методы Data Mining: алгоритм Apriori, Association Rules и др.	Enhanced K-Means, Orthogonal Portioning Clustering
GALILEI	Лемматизация	—	—	—	—	—
a.k.a.	—	—	—	—	—	—
AUTHECO	Лемматизация, Части речи, Синтаксические зависимости	Лингвистические графы и словари	Функция отбора, основанная на статист. и лингвист. данных	Использование базы синонимов	Дистрибутивно-статистический анализ	Латентно-семантический анализ + K-Means

Архитектура системы AUTHESO (АСПТ)

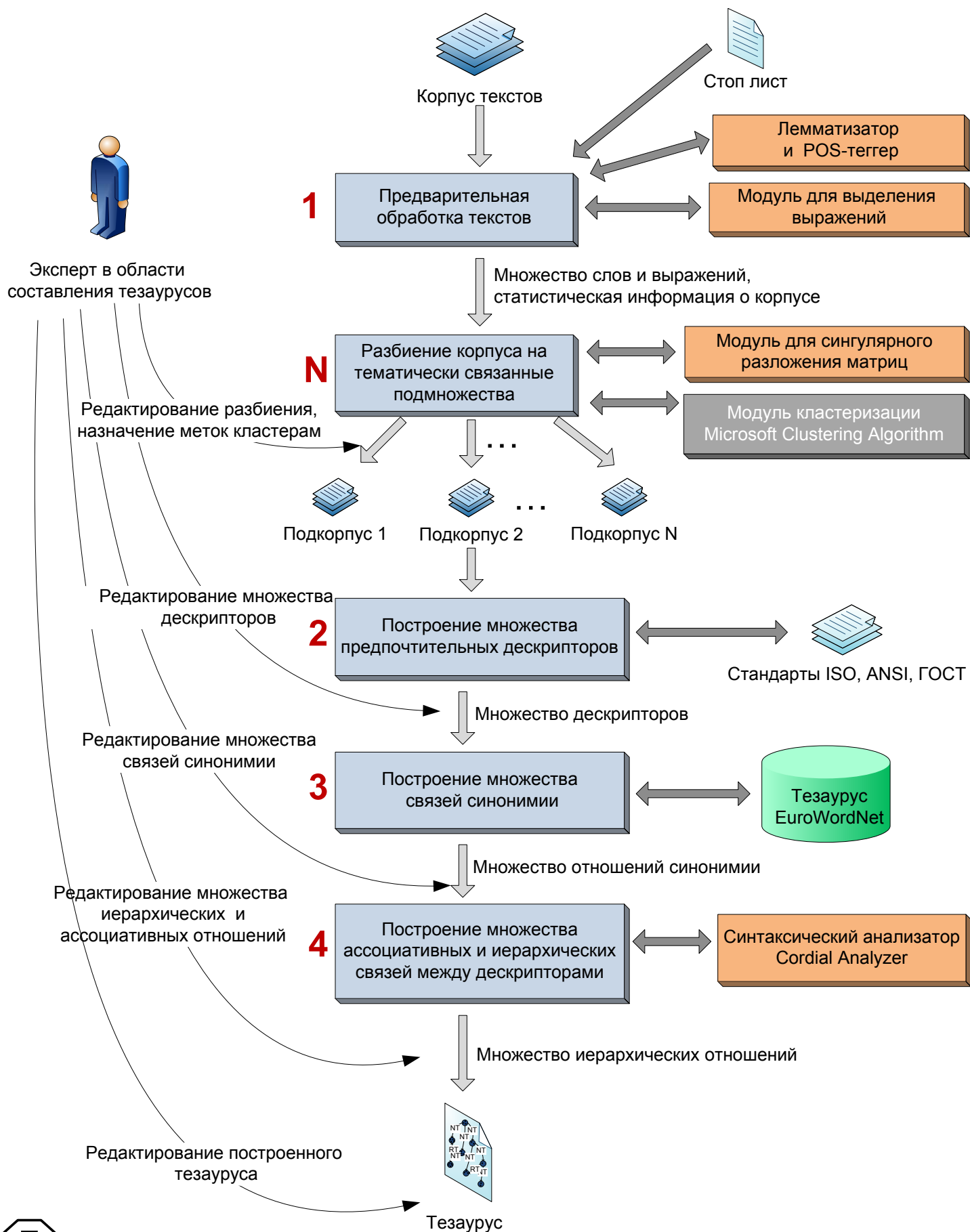


Компоненты разработанные в процессе дипломного проектирования

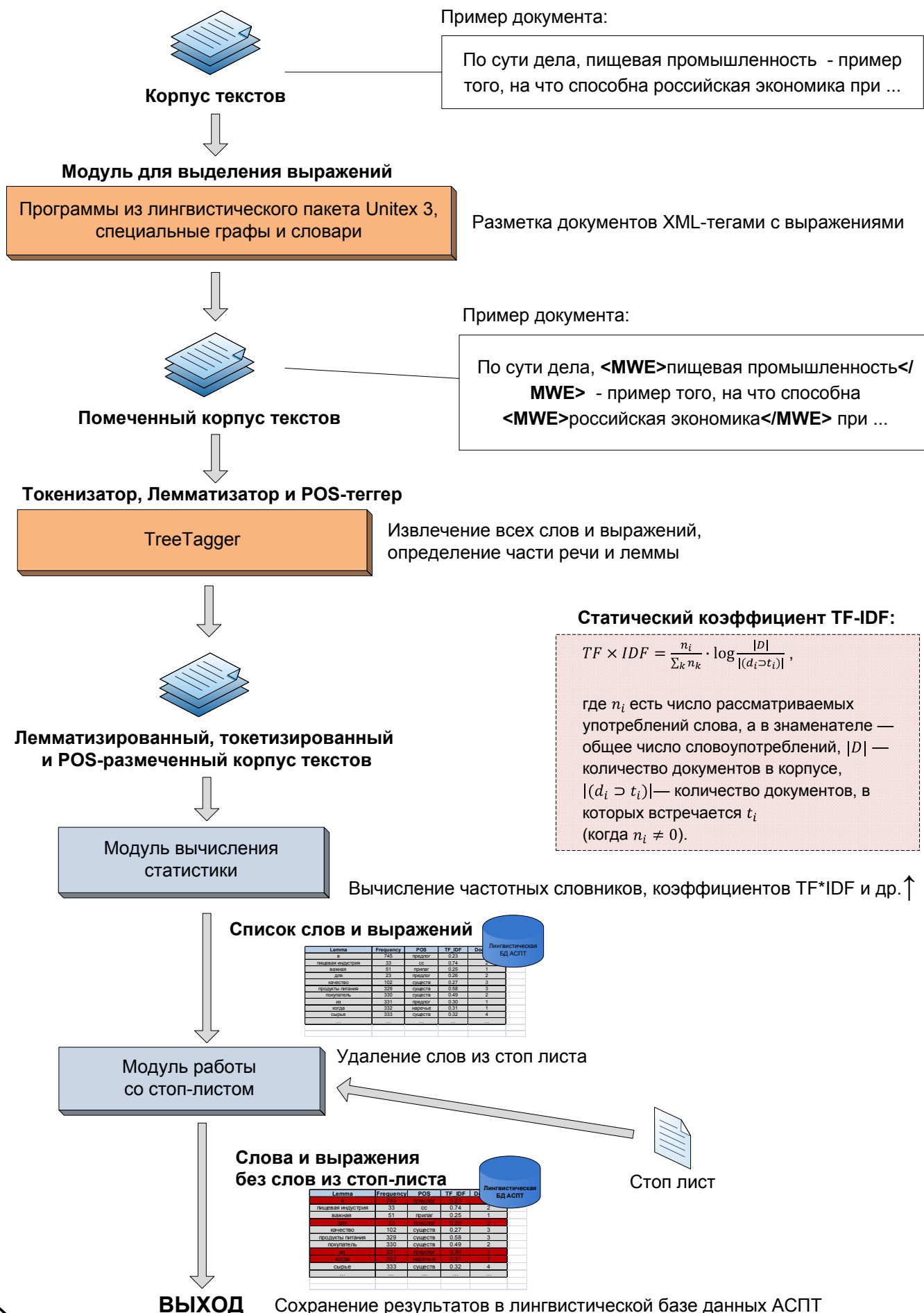
Даталогическая схема лингвистической базы данных АСПТ



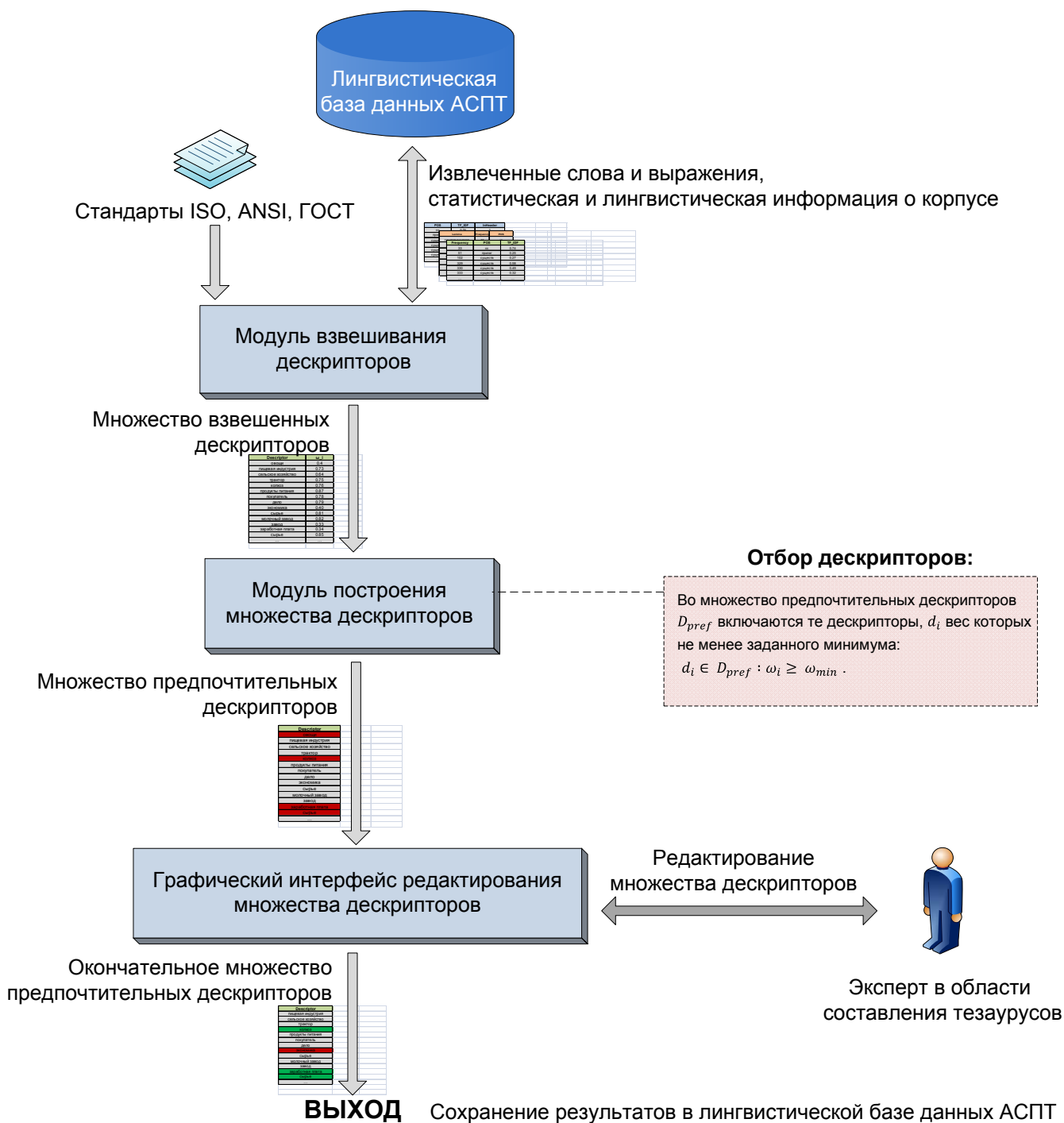
Технология автоматизированного построения тезауруса



Предварительная обработка текстов (1)



Построение множества предпочтительных дескрипторов (2)

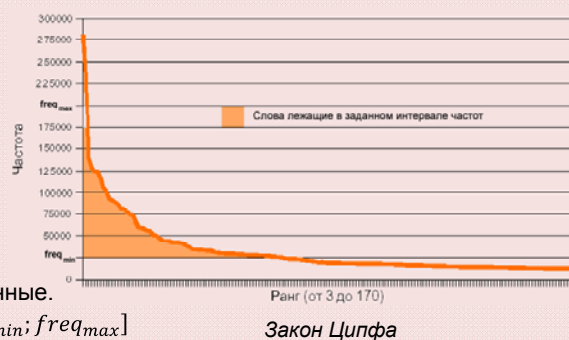


Взвешивание дескрипторов функцией отбора:

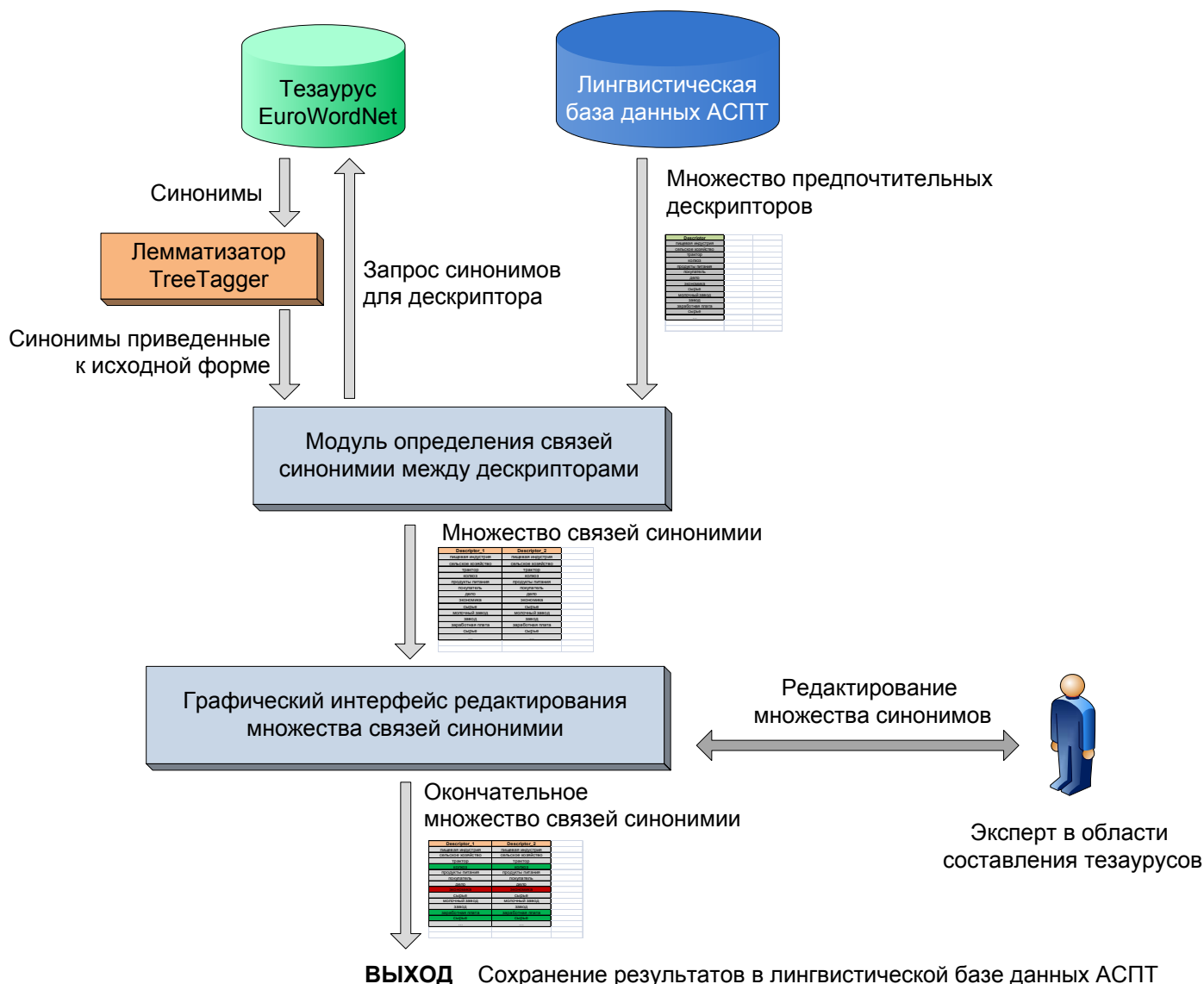


Функция представляет собой комбинацию ограничений на данные переменные.

К примеру, $freq_i$ должен лежать в определенном интервале $freq_i \in [freq_{min}; freq_{max}]$ в соответствии с законом Ципфа (см. рисунок).



Построение множества связей синонимии (3)



Алгоритм определения связей синонимии между дескрипторами

Множество дескрипторов

Пусть множество *предпочтительных дескрипторов* – $D_{pref} = \{d_1, d_2, \dots, d_n\}$,

где n – мощность множества дескрипторов;

причем d_i – дескриптор, являющийся упорядоченным кортежем, таким что $d_i = (w_1, w_2, \dots, w_k)$, $k > 0$, где w_i – слово.

Множество связей синонимии

Пусть S_{ewn} – непустое множество всех связей синонимии содержащихся в тезаурусе EuroWordNet, такое что $S_{ewn} = \{s_1, s_2, \dots, s_m\}$,

где s_i – отношение синонимии которое можно представить в виде неупорядоченной пары вида $(d_i; d_j)$ или записывая по-другому $d_i \rightleftharpoons d_j$.

Алгоритм нахождения синонимов

Алгоритм нахождения синонимов в заданном множестве дескрипторов заключается в построении множества отношений

синонимии $S_D = \{s_1, s_2, \dots, s_l\}$; $S_D \subseteq S_{ewn}$.

1. Для каждого дескриптора d_i из множества D_{pref} находим такое подмножество

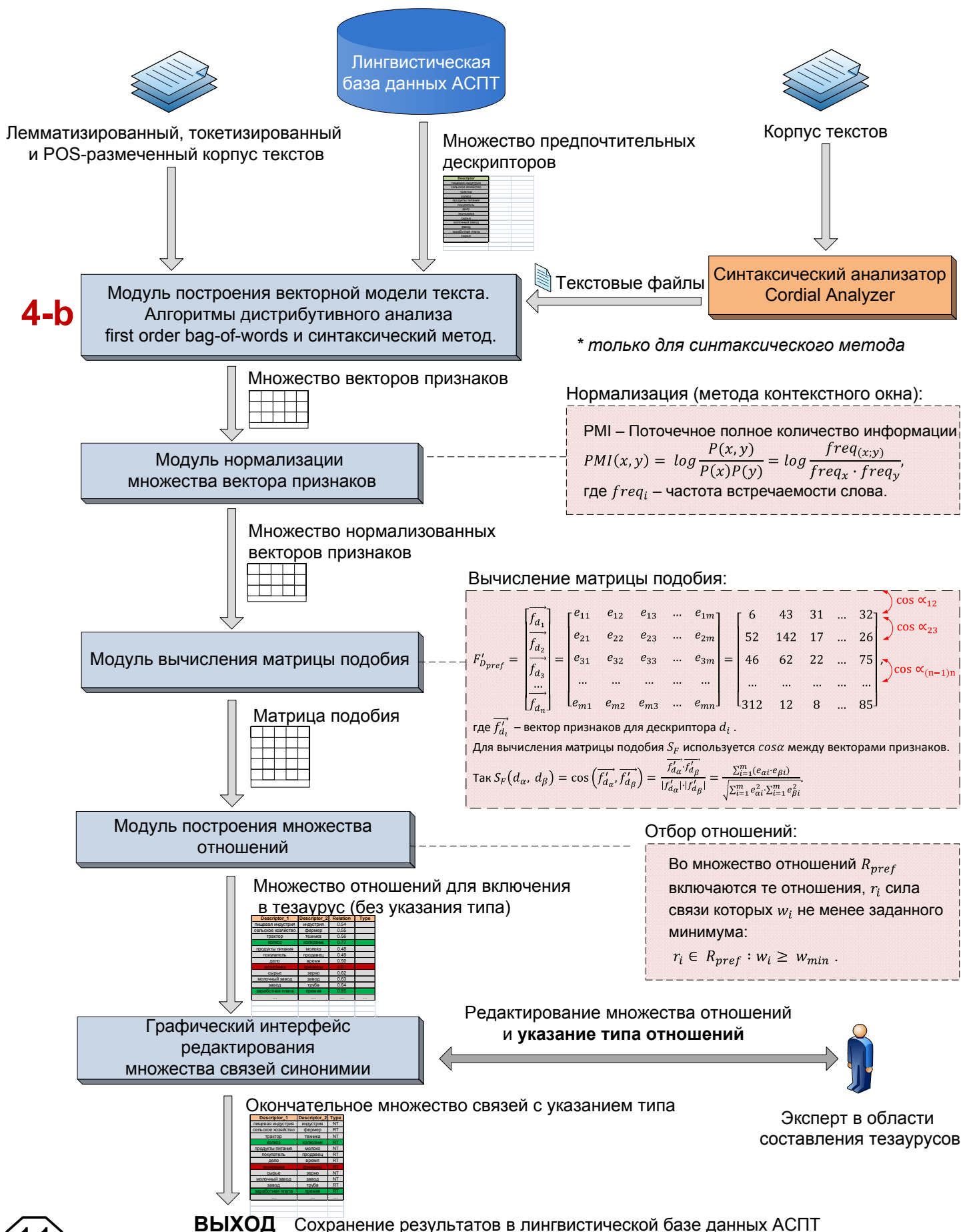
$S_i \subseteq S_{ewn}$ что $S_i = \{s_1^i, s_2^i, \dots, s_t^i\}$; $\forall s_j^i = (d_i; d_k)$; $j \in [1; t]$, $d_k \in D_{pref}$.

Иначе говоря, отбираем те связи синонимии, которые содержат дескрипторы только из D_{pref} .

2. Добавляем множество S_i к искомому: $S_D = S_D \cup S_i$.

3. Переходим к следующему дескриптору. Если все дескрипторы перебраны тогда выход.

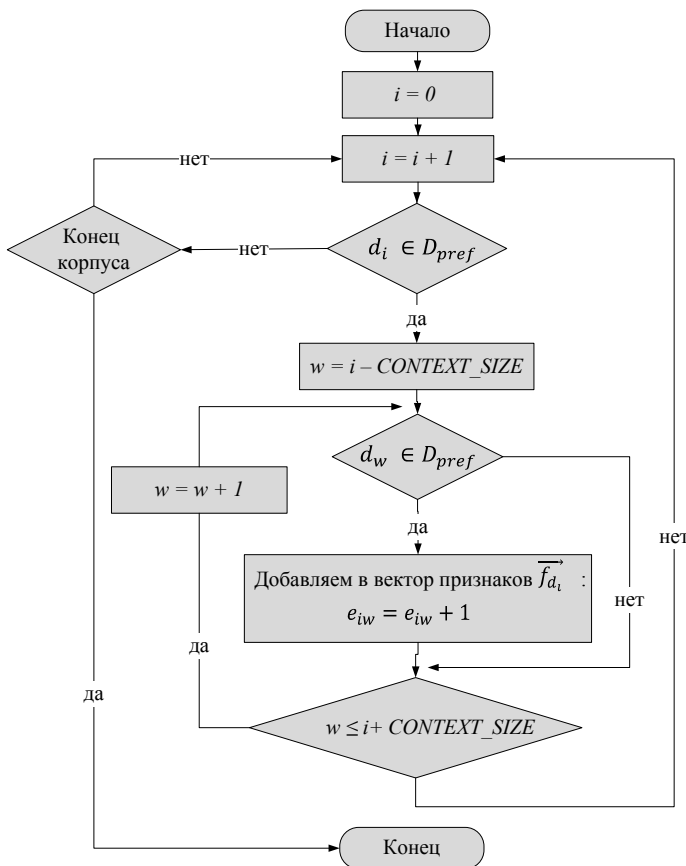
Построение множества ассоциативных и иерархических связей между дескрипторами (4-а)



Построение векторной модели текста в алгоритмах дистрибутивно-статистического анализа (4-b)

Метод контекстного окна (first order bag-of-words)

Дистрибутивная гипотеза – слова появляющиеся в одинаковых контекстах связаны.



	D_{pref}				
	$d_1 = \text{промышленность}$	$d_2 = \text{сельское хозяйство}$	$d_3 = \text{удобрения}$	...	$d_n = \text{рабочий}$
$d_1 = \text{промышленность}$	1	23	13	...	43
$d_2 = \text{сельское хозяйство}$	23	1	98	...	32
$d_3 = \text{удобрения}$	13	98	1	...	45
...	...	...	...	...	...
$d_n = \text{рабочий}$	43	32	45	...	1

Формат матрицы признаков, где $\vec{f}_{d_i}$ – вектор признаков дескриптора d_i

Параметр алгоритма – размер контекстного окна (CONTEXT_SIZE).

1 слово ≤ CONTEXT_SIZE ≤ 1 параграф

i – указатель на обрабатываемый дескриптор

w – указатель на слово обрабатываемое внутри контекстного окна

Пример для контекстного окна = 1 слово:

[Снятие омонимии] полезно во многих приложениях...
 [Снятие омонимии полезно] во многих приложениях...
 Снятие [омонимии полезно во] многих приложениях...
 Снятие омонимии [полезно во многих] приложениях...
 Снятие омонимии полезно [во многих приложениях] ...

GSR

	(*, подлежа, работать)	(*, сказ, фермер)	(*, пр_дополн, трактор)	...	(*, type_k, d_b) = gsr_j	...	(*, подлежа, пахать)
$d_1 = \text{промышленность}$	1	0	13	...	...	...	43
$d_2 = \text{сельское хозяйство}$	23	1	98	...	...	...	32
...	...	...	...	...	...	...	...
d_i	...	...	...	...	e_{ij}	...	...
...	...	...	...	...	...	...	...
$d_n = \text{рабочий}$	43	32	45	...	...	...	1

Формат матрицы признаков,

где $\vec{f}_{d_i}$ – вектор признаков дескриптора d_i

Синтаксические зависимости в предложении:

«Лемматизация играет важную роль в задачах компьютерной лингвистики.»

(лемматизация, **подлежа_гл**, играть); (играть, **сказ**, лемматизация);

(играть, **прям_дополн_гл**, важная роль)...($d_i, type_k, d_j$).

d_i, d_j – дескрипторы, причем $d_i \in D_{pref}$

$type_k$ – тип синтаксической связи $type_k \in T$

Синтаксический метод

$SR_{all} = \{sr_1, sr_2, \dots, sr_n\}$ – множество всех извлеченных синтаксических зависимостей, причем sr_i – упорядоченная тройка вида $(d_a, type, d_b)$.

$SR_{pref} \subseteq SR_{all}$; $SR_{pref} = \{sr \mid (sr \in SR) \wedge (type \in T_{pref}) \wedge (d_a \in D_{pref})\}$, где

$T_{pref} = \{type_1, type_2, \dots, type_k\} = \{\text{подлежа_гл_v}, \text{прям_дополн_гл_v}, \text{обстоят_гл_v}, \dots, \text{имеет_предл_p}\}$ – множество рассматриваемых в методе синтаксических зависимостей.

$T_{pref} \subseteq T$ – множество всех извлеченных типов отношений с использованием синтаксического анализатора.

$GSR = \{gsr_1, gsr_2, \dots, gsr_m\} = \{gsr_i \mid (\exists sr = (d_a, type, d_b) \in SR_{pref})\}$, таким образом GSR содержит множество элементов типа $(*, \text{подлежа_гл}, \text{работать})$

