

МИНИСТЕРСТВО НАУКИ И ВЫСШЕГО ОБРАЗОВАНИЯ РОССИЙСКОЙ ФЕДЕРАЦИИ
ФЕДЕРАЛЬНОЕ ГОСУДАРСТВЕННОЕ АВТОНОМНОЕ ОБРАЗОВАТЕЛЬНОЕ
УЧРЕЖДЕНИЕ ВЫСШЕГО ОБРАЗОВАНИЯ
«НОВОСИБИРСКИЙ НАЦИОНАЛЬНЫЙ ИССЛЕДОВАТЕЛЬСКИЙ ГОСУДАРСТВЕННЫЙ
УНИВЕРСИТЕТ» (НОВОСИБИРСКИЙ ГОСУДАРСТВЕННЫЙ УНИВЕРСИТЕТ, НГУ)

На правах рукописи



Пименов Иван Сергеевич

ФОРМАЛЬНОЕ ВЫДЕЛЕНИЕ ПРИЁМОВ И СТРАТЕГИЙ АРГУМЕНТАЦИИ В
ТЕКСТАХ НАУЧНОЙ КОММУНИКАЦИИ

5.9.8 – Теоретическая, прикладная и сравнительно-сопоставительная лингвистика

Диссертация
на соискание ученой степени
кандидата филологических наук

Научный руководитель
доктор филологических наук, доцент
Тимофеева Мария Кирилловна

Консультант
кандидат физико-математических наук
Саломатина Наталья Васильевна

Новосибирск – 2025

ОГЛАВЛЕНИЕ

Введение	5
Глава 1. Теоретические основы представления аргументации	19
1.1 Анализ аргументации как фундаментальная лингвистическая проблема	20
1.2 Лингвистические подходы к анализу аргументации.....	24
1.2.1 Подходы к соотношению аргументации и языка как системы	27
1.2.2 Теория семантических блоков	32
1.2.3 Анализ аргументации в выявлении ключевых слов культуры	34
1.2.4 Дискурсивный подход к аргументации в естественном языке	37
1.2.5. Развитие аргументационной теории в отечественной лингвистике	39
1.3 Формальное моделирование аргументационных структур	44
1.3.1 Аргументационная модель С. Тулмина	45
1.3.2 Теория риторической структуры в анализе аргументации.....	50
1.3.3 Моделирование текстов через Argument Interchange Format	58
1.3.4 Классификация аргументационных схем Д. Уолтоном.....	60
Выводы	63
Глава 2. Методы анализа аргументационных приёмов и стратегий.....	65
2.1 Экспертный подход к разметке аргументации	66
2.2 Методы автоматизации аргументационной разметки.....	72
2.2.1 Обзор подходов к формальному выявлению аргументации	72
2.2.2 Компьютерное распознавание утверждений с аргументацией	78
2.2.3 Методы машинного обучения для распознавания тезисов и схем	80
2.3 Методы оценки качества разметки через согласие аннотаторов	81
2.3.1 Подсчёт коэффициента аргументационного соответствия	82
2.3.2 Оценка согласия через каппу Коэна и альфу Криппендорфа	83
2.4 Методы выявления составных приёмов аргументации	86
2.4.1 Обзор подходов к обработке приёмов аргументации	86
2.4.2 Выявление приёмов аргументации через анализ подграфов	88
2.4.3 Построение частотного спектра для повторяющихся подграфов.....	91
2.5 Методы оценки убедительности аргументации.....	93

2.5.1 Анализ рассуждений по выраженности логоса, пафоса, этоса	93
2.5.2 Анализ методов убеждения на основе маркеров и схем.....	96
2.5.3 Оценка выраженности методов убеждения в целостных текстах	101
2.6 Кластеризация текстов по сложности аргументации	104
2.6.1 Традиционные показатели сложности текстов	104
2.6.2 Кластеризация аргументационных графов.....	106
2.6.3 Оценка качества кластеризации	107
Выводы	108
Глава 3. Выявление элементов аргументационных стратегий в графах аргументации	111
3.1 Выявление элементарных приёмов аргументации	111
3.1.1 Характеризация аргументационного корпуса научных статей.....	112
3.1.2 Оценки согласия аннотаторов на размеченном корпусе	113
3.1.3 Анализ встречаемости отдельных аргументационных схем	116
3.2 Компьютерное распознавание предложений, содержащих аргументацию..	123
3.2.1 Обучение классификатора с учётом специфики корпуса текстов	123
3.2.2 Статистическая фильтрация лемм на основе их частоты	125
3.2.3 Результаты автоматической классификации предложений.....	127
Выводы	130
Глава 4. Стратегии аргументации как совокупность общих приёмов.....	133
4.1 Выявление составных приёмов аргументации.....	134
4.1.1 Составные приёмы аргументации в научных текстах	135
4.1.2 Приёмы аргументации в текстах разных жанров	140
4.2 Выявление методов убеждения в аргументационных структурах	154
4.2.1 Оценка этоса и пафоса по аргументационным графам	154
4.2.2 Оценка авторитетности и эмоциональности текстов	156
4.3 Компьютерное распознавание этоса и пафоса в текстах	158
4.3.1 Построение бинарных классификаторов для пафоса и этоса.....	159
4.3.2 Результаты распознавания аргументов с этосом и пафосом	160
4.4 Классификация аргументов по функциональной близости.....	163

4.4.1 Функциональное сопоставление моделей рассуждения	164
4.4.2 Представление цепочек аргументов через функциональные блоки.....	167
4.4.3 Сочетаемость функциональных групп в цепочках рассуждений	169
4.5 Кластеризация аннотаций по структурной сложности	175
4.5.1 Оценка аргументационной сложности отдельного текста	176
4.5.2 Сопоставление текстов корпуса по аргументационной сложности	179
4.5.3 Выявление отличительных признаков между группами текстов	180
4.6 Использование приёмов аргументации в жанровой атрибуции	183
4.6.1 Высокоуровневые признаки в распознавании жанров.....	184
4.6.2 Представление текстов через совокупность приёмов.....	185
4.6.3 Идентификация жанров посредством приёмов разных видов	188
Выводы	199
Заключение.....	202
Список литературы.....	205
Приложение А Маркеры главного тезиса.....	222
Приложение Б Маркеры аспектов содержания.....	224
Приложение В Актуальные схемы аргументации в научных статьях.....	226
Приложение Г Пример аргументационной разметки короткой научной статьи...	234
Приложение Д Наиболее сложный приём аргументации в научных статьях.....	247
Приложение Е Лексико-грамматические шаблоны с маркерами этоса и пафоса .	251

Введение

Диссертационное исследование посвящено проблеме формального анализа структур аргументации в текстах научной коммуникации (преимущественно коротких научных статьях) на русском языке методами компьютерной обработки. Такие структуры рассматриваются на уровне приёмов и стратегий аргументации. Приёмы соответствуют организованным, повторяющимся в разных текстах структурам из одной или более типовых моделей рассуждения (абстрактных логических схем, регулирующих построение аргументационных связей между высказываемыми утверждениями в процессе доказательства некоторой идеи). Совокупность приёмов, применяемых на уровне целостного полного текста для доказательства его ключевой идеи, образует стратегию аргументации, реализуемую автором этого текста. Полезно сразу отметить, что термин «доказательство» понимается в работе в широком смысле как поддержка одного тезиса другим в аргументационной структуре текста. Расширенное толкование связано с моделированием аргументации в терминах аргументационных схем Д. Уолтона, где разные схемы (модели рассуждения) соответствуют разным видам обоснования в классическом понимании, от доказательств в строгом определении до логических ошибок (таких как *ad hominem*).

Формальный анализ прагматических языковых явлений, к числу которых относится аргументация, является фундаментальной лингвистической проблемой. Сложность аргументации как явления и теоретическая значимость её изучения иллюстрируются многоаспектностью аргументационных исследований в рамках различных дисциплин: как собственно лингвистики, теоретической и прикладной, так и риторики [Perelman, Olbrechts-Tyteca, 1969], философии [Van Eemeren, Grootendorst, 2004], формальной логики [Hintikka, 1989]. Задача формальной обработки аргументационных структур предполагает использование методов компьютерной лингвистики, в частности, направления Argument Mining, извлечения аргументов из текстов на естественном языке [Lawrence, Reed, 2019]. Однако ввиду сложности аргументации как языкового явления исследования по её

компьютерной обработке активизировались сравнительно недавно. При этом такие работы обращены преимущественно к текстам на английском и других зарубежных языках [Lawrence, Reed, 2014].

Таким образом, актуальность диссертационного исследования проявляется в двух аспектах: теоретическом и методологическом. **Теоретическая актуальность** обуславливается важностью изучения формальной организации аргументации как явления прагматического языкового уровня в текстах на русском языке, в частности, в научных статьях (направленных на доказательство исследовательских выводов). **Методологическая актуальность** работы основывается на необходимости интеграции традиционных лингвистических подходов к изучению аргументации и методов компьютерной обработки текстов для формального анализа структур доказательства и рассуждения с учётом языковой и жанровой специфики текстов научной коммуникации на русском языке. При этом лингвистическая сложность аргументации как прагматического явления обуславливает потребность в развитии и адаптации методов компьютерной обработки текстов, преимущественно применяемых для анализа более простых единиц иных языковых уровней (морфологического, лексического, синтаксического). Соответственно, моделирование и формальный анализ аргументационных структур обеспечивают возможность для расширения методологического аппарата лингвистики в целом.

Степень разработанности проблемы исследования. Ввиду сложности аргументации как языкового явления исследования по её компьютерной обработке (Argument Mining) образовали отдельное целостное направление с устоявшимся методологическим аппаратом сравнительно недавно (первая специализированная конференция по Argument Mining состоялась в 2014 г. [Lawrence, Reed, 2014]). Исследования в этом направлении обращены в основном к текстам на английском языке и таким жанрам, как юридические диспуты [Walker, Vazirova, Sanford, 2014], новостные статьи [Wachsmuth et al., 2018], политические дебаты [Lindahl, Borin, Rouces, 2019], онлайн-дискуссии [Hidey,

McKeown, 2018]. Тем не менее, отдельные работы посвящены анализу аргументации в научных статьях, примером чего является статья [Green, 2015].

Исследования по компьютерной обработке аргументов в текстах на русском языке представлены в ограниченном количестве из-за недостаточного объёма доступных корпусов, как указано в [Котельников, 2018]. Впрочем, значимость темы обуславливает активную разработку аргументационных корпусов русского языка в последнее время: посредством перевода иноязычных корпусов [Fishcheva, Kotelnikov, 2019], через разметку научно-популярных статей [Сидорова и др., 2020] либо сообщений в социальных сетях [Kotelnikov et al., 2022]. Хотя научные статьи на русском языке ранее не становились объектом компьютерного анализа аргументации, известны исследования научно-популярных текстов, например, [Ким, Ильина, 2020; Саломатина и др., 2020; Zagorulko et al., 2020; Ильина, 2023].

Объект диссертационного исследования – аргументационные структуры в текстах научной коммуникации на русском языке (преимущественно в научных статьях). **Предметом** исследования являются средства формальной организации элементов аргументации (утверждений, объединяющих их связей, типовых моделей рассуждения в основе этих связей) при доказательстве основной идеи целостного текста, обеспечивающей его смысловую связность.

Основная **гипотеза** исследования: формальный анализ аргументационных структур посредством интеграции сведений об их лингвистическом выражении (языковом оформлении утверждений) и их абстрактной логической организации (на уровне типовых моделей рассуждения в основе связей между утверждениями) позволяет выявлять составные приёмы и стратегии аргументации, характеризующие специфику доказательств в отдельных целостных текстах. На Рис. 1 показан пример искомого составного приёма: прямоугольные блоки содержат утверждения из конкретных текстов, овальные указывают на типовые абстрактные модели (схемы аргументации, элементарные приёмы), реализуемые при обосновании одних утверждений другими. Содержание утверждений различается между разными текстами (слева и справа), тогда как модели в связях между утверждениями совпадают. Одинаковы и их конфигурации в образуемой

структуре (составном приёме) из трёх элементарных приёмов, и эта структура соответствует составному приёму аргументации.

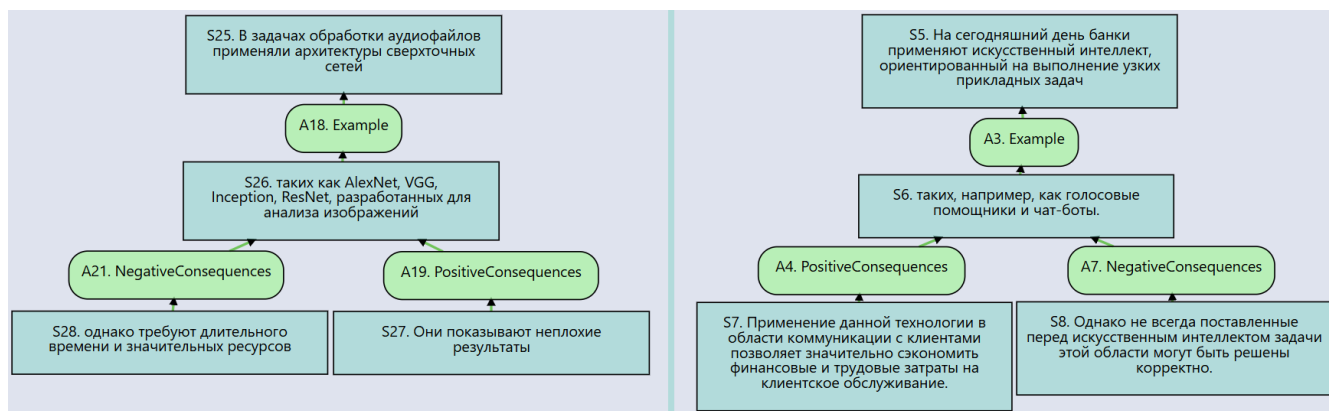


Рисунок 1 – Пример реализации составного приёма аргументации

В соответствии с основной гипотезой **цель работы** может быть сформулирована следующим образом: разработать и экспериментально проверить систему методов формального анализа аргументации в текстах научной коммуникации для распознавания и характеристики специфики доказательств в отдельных текстах на уровне как их языкового оформления, так и структурной организации. Под экспериментом в работе понимается практическое использование разработанных методов для компьютерной обработки текстовых данных из используемого корпуса с последующей оценкой эффективности методов на уровне как количественных показателей, так и содержательной интерпретации результатов их применения.

Достижение поставленной цели предполагает решение следующих задач:

- 1) определить теоретические основы для формального представления аргументации с учётом как традиционных лингвистических подходов к её анализу, так и специфики методов компьютерной обработки текстов;
- 2) подготовить корпус текстов (научных статей на русском языке) с многоаспектной аргументационной разметкой (согласно обозначенным теоретическим принципам, причём под многоаспектностью понимается характеристика аргументации на уровнях утверждений, связей между ними,

типовых семантических моделей в основе этих связей) для подготовки и проверки автоматических методов обработки аргументации, выраженной на естественном языке. Создание корпуса предполагает разработку методологии многоаспектной разметки аргументации в научных статьях, которая повысит согласованность разметки текстов разными аннотаторами;

3) реализовать набор методов компьютерного анализа аргументационных структур, где отдельные методы предназначены для решения следующих подзадач:

- распознавания предложений, содержащих аргументацию, связей между этими предложениями и отдельных частотных элементарных приёмов (экспертная разметка используется для машинного обучения классификаторов, применяемых для распознавания, и автоматической проверки качества их работы);

- выявления элементарных и составных приёмов аргументации из размеченных текстов с учётом их жанровой специфики в многожанровой коллекции;

- обобщения цепочек рассуждения согласно функциональной близости образующих их моделей рассуждения;

- оценки аргументов по выраженности методов убеждения;

- распознавания элементов методов убеждения (пафос, этос) в аргументах из неразмеченных текстов;

- выявления текстов со сходной организацией аргументации и ранжирования их по сложности рассуждений путём интеграции сведений об их аргументационной специфике во всех отмеченных выше аспектах;

- жанровой классификации текстов на основе сведений о реализованных в них приёмах аргументации;

4) экспериментально проверить реализованные методы на подготовленном корпусе текстов, интерпретировать результаты и обозначить особенности построения аргументации в текстах научной коммуникации на русском языке.

Материалом исследования выступает корпус из 11495 тезисов и 10123 аргументов, образующих структуры аргументации для 100 научных статей с ручной разметкой аргументации тремя экспертами (для каждого текста подготовлено по две версии разметки независимо двумя экспертами). Размеченные статьи принадлежат направлениям лингвистики и информационных технологий (по 50 статей обеих тематик), характеризуются объёмом от 800 до 2000 слов (без учёта аннотаций и библиографических списков). Публикация анализируемых статей в научных изданиях указывает на достижение ими академического уровня, достаточного для прохождения рецензирования. Разметка корпуса проведена по ходу исследования, в связи с чем на отдельных этапах задействованы ранние версии корпуса меньшего объёма. На отдельных этапах исследования (при анализе аргументации в текстах разных жанров) привлекаются данные из параллельно размечавшихся корпусов научных новостей и научно-популярных текстов, количественные характеристики которых даны в таблице 1.

Таблица 1 – количественные характеристики корпусов разных жанров

	Научные статьи (из библиотеки «КиберЛенинка»)	Научные новости (портал «Поиск»)	Короткие научно- популярные тексты (из различных журналов)	Научно- популярные тексты (Habr)
Число текстов	100 (с двойной разметкой)	28	95	20
Число утверждений	11495	859	2005	1681
Число аргументов	10123	758	1386	1415

Разметка аргументации проведена в соответствии со стандартом AIF [Rahwan, Reed, 2009] и классификацией типовых моделей рассуждения (схем аргументации) Д. Уолтона [Walton, Reed, Macagno, 2008]. Разметка проведена на трёх уровнях аргументации: выявление утверждений (высказываний с аргументацией, принадлежащих целостной структуре рассуждений в тексте), распознавание связей между утверждениями (что включает определение их ролей как посылок либо заключений в составе аргументов), указание типовых моделей рассуждения из заданного набора для связей (с помощью моделей, также называемых схемами, уточняется семантика перехода от посылок к заключению). Корпусы с разметкой доступны на платформе ArgNetBank Studio [Сидорова и др., 2020], аннотирование текстов осуществлено с помощью инструментов этой платформы.

Для оценки надёжности разметки проведён подсчёт количественных коэффициентов согласия аннотаторов на разных этапах создания корпуса. Использован разработанный в ходе исследования специализированный коэффициент аргументационного согласия (КАС), позволяющий анализировать расхождения в разметке аргументации на всех её уровнях, а также два универсальных коэффициента с учётом случайного согласия между разметчиками, каппа Коэна и альфа Криппендорфа. Значения коэффициентов подсчитаны для каждого уровня разметки (выявление утверждений, связей, схем, в том числе на уровне групп функционально близких схем). В таблице 2 приведены значения двух коэффициентов согласия для рабочей версии корпуса научных статей (каппа Коэна не указана, так как подсчитывается по парам аннотаторов).

Приведённые показатели соответствуют нижним границам оценки согласия разметчиков: случаи формальных расхождений аннотаторов при их фактическом согласии учитываются как несоответствия в разметке (например, случаи разной конфигурации связей, когда оба разметчика указывают связи по одинаковым моделям рассуждения между одними и теми же тезисами, однако один аннотатор строит последовательные связи, а другой – параллельные, оцениваются

компьютером как расхождения в разметке, как и сложные случаи в сегментации утверждений).

Таблица 2 – коэффициенты согласия разметчиков в рабочей версии корпуса

Коэффициент	Утверждения	Связи	Схемы	Функциональные группы схем
КАС	0.87	0.54	0.52	0.71
Альфа Кrippendorfa	0.51	0.66	0.43	0.55

Для решения поставленных задач на базе указанных материалов исследования применены **методы** разных групп:

- 1) анализа аргументации на основе дискурсивного подхода (выявление аргументов с учётом, во-первых, лексико-грамматических маркеров и особых синтаксических конструкций для представления рассуждений, во-вторых, абстрактных логических моделей их организации, в-третьих, контекста реализации аргументов в целостной структуре рассуждений полного текста);
- 2) корпусной лингвистики (создание аргументационных корпусов, методики разметки, разметка текстов, оценка согласия аннотаторов);
- 3) обработки естественного языка (морфологический анализ, построение поисковых шаблонов для маркеров аргументации);
- 4) машинного обучения (классификация и кластеризация текстов и их элементов, извлечение аргументации).

Теоретическая (источниковая) база исследования включает работы двух групп. Одну представляют труды по исследованиям и моделированию аргументации с позиций различных дисциплин: в первую очередь, на основе лингвистических подходов, но также и в рамках иных направлений (философии, риторики, формальной логики). Другую группу составляют работы по компьютерной обработке текстов на естественном языке (как в рамках специализированных исследований аргументации, так и для общих смежных

задач, таких как автоматическая классификация текстов) и фундаментальные труды по анализу данных разных типов (например, работы с описанием алгоритмов кластеризации, применимых как к текстовым, так и иным данным).

Соответственно, **новизна** диссертационного исследования заключается в следующем:

1) Впервые применены методы формального аргументационного анализа к текстам научной коммуникации на русском языке. Даны строгие определения приема и стратегии аргументации, позволившие построить формальное представление аргументационной структуры текста, коллекции текстов, набора коллекций.

2) Впервые аргументация в текстах на русском языке обрабатывается компьютерными методами на уровне глубинной аргументационной семантики: через многоаспектный анализ элементарных и составных приёмов (в том числе с учётом их функционального сходства), методов убеждения (средств дополнения фактического содержания рассуждений апелляциями к авторитету и эмоциональным воздействием), совокупных особенностей последовательного обоснования тезисов в отдельных целостных текстах. Предложена функциональная классификация аргументов по структурно-семантическим особенностям их реализации в научных статьях, указаны четыре обобщённых типа таких аргументов («через детализацию», «через анализ причинно-следственный связей», «через рассуждения практического плана», «от авторитета»). Указана специфика реализации методов убеждения в научных текстах, подготовлен словарь их маркеров (в виде многокомпонентных лексико-грамматических шаблонов на основе регулярных выражений).

3) В рамках исследования подготовлен и размечен аргументационный корпус научных статей на русском языке, содержащий 11495 аргументов и 10123 утверждений (тезисов). Разметка текстов заключается в трёхуровневом моделировании их аргументационной структуры согласно стандарту Argument Interchange Format и классификации аргументационных схем Д. Уолтона. В процессе создания корпуса разработана методология многоаспектной разметки

аргументации в текстах научных статей, которая регулирует действия аннотаторов при разметке. Надёжность аргументационной разметки подтверждается коэффициентами согласия аннотаторов, подсчитанными на всех уровнях аннотирования аргументации. Для анализа расхождений в аргументационном моделировании на всех его уровнях разработан специализированный коэффициент аргументационного соответствия.

4) Разработан программный комплекс для многоаспектной компьютерной обработки аргументации в текстах на русском языке. Комплекс образован программными модулями для различных отдельных задач аргументационного анализа: распознавание тезисов в текстах, связей между смежными тезисами и отдельных частотных типов аргументов (для автоматизации экспертной разметки); характеристика аргументов в тексте по методам убеждения (этос, пафос); выявление составных приёмов аргументации, используемых авторами текстов научной коммуникации, из графов, построенных аннотаторами; кластеризация текстов научной коммуникации по сходству аргументационных структур и ранжирование групп текстов по формальной сложности аргументации (которая подсчитывается через интеграцию разноаспектных показателей аргументационной специфики текстов: структурной организации рассуждений, реализуемых приёмов аргументации, интенсивности выражения отдельных методов убеждения); жанровая классификация текстов на основе сведений о реализованных в них приёмах аргументации.

Таким образом, **теоретическая значимость** исследования обусловлена выявлением формальных особенностей построения аргументации в текстах научной коммуникации на русском языке. Определены тематические различия в использовании приёмов аргументации в работах по лингвистике и по информационным технологиям. Предложена функциональная классификация аргументов по структурно-семантическим особенностям их реализации в научных статьях, указаны четыре обобщённых типа таких аргументов («через детализацию», «через анализ причинно-следственных связей», «через рассуждения практического плана», «от авторитета»). Указана специфика

реализации методов убеждения в научных текстах, подготовлен словарь их маркеров (в виде многокомпонентных лексико-грамматических шаблонов на основе регулярных выражений). Результаты диссертационного исследования дополняют теоретическую базу для иных исследований в области теории аргументации (в рамках лингвистического и формального подходов) и лингвистической прагматики, предоставляют новые сведения об организации текстов (в первую очередь научных статей), их целостности и связности на уровне аргументационной структуры.

Практическая значимость диссертационного исследования заключается в разработке и реализации методов компьютерной обработки аргументации в текстах на русском языке. На основе размеченного корпуса экспериментально оценена эффективность различных алгоритмов для отдельных задач аргументационного анализа, изучены и типизированы ошибки на этапах компьютерного распознавания элементов аргументации (тезисов, аргументов по разным методам убеждения). Корпус с многоаспектной аргументационной разметкой научных статей может использоваться в иных исследованиях по формальному анализу аргументации (в частности, для подготовки методов машинного обучения). Иным аспектом практической значимости является учебно-методический: материалы исследования могут применяться в процессе педагогической деятельности в областях компьютерной лингвистики, лингвистики текста, риторики, академического письма (для развития у обучающихся компетенций в области анализа научной аргументации, критического восприятия информации, для практики навыков написания научных текстов с ясным представлением и обоснованием выдвигаемых тезисов).

Список используемой литературы содержит 153 источника, 98 из которых представлены на английском языке.

Положения, выносимые на защиту, сформулированы ниже.

1. Формальное представление аргументации позволяет охарактеризовать аргументационную специфику целостного текста в разных аспектах: взаимосвязи языковых средств выражения аргументов и абстрактных схем в основе их связей,

использовании приёмов аргументации и методов убеждения. Особенности аргументационной организации отдельных текстов проявляются при сопоставительном анализе их формальных моделей.

2. Авторы разных текстов научной коммуникации используют повторяющиеся между текстами приёмы аргументации, которые образованы небольшим числом типовых моделей рассуждения.

3. Представление и логический анализ фактов в научных статьях при обосновании исследовательских результатов периодически дополняется непрямым эмоциональным воздействием через использование оценочной лексики и модальных конструкций для привлечения внимания читателей и акцентирования отдельных аспектов работы.

4. Элементы аргументационных структур могут распознаваться методами компьютерной обработки текстов на естественном языке с достижением приемлемых показателей качества, что позволяет автоматически анализировать и неразмеченные тексты.

5. Использование приёмов аргументации и методов убеждения (общих способов аргументационного воздействия на читателя) в текстах научной коммуникации зависит от их жанровой специфики и предметной области. Представление текстов на уровне реализуемых в них приёмов обеспечивает возможность их надёжного разграничения как по жанру и тематике, так и по более специфичным аспектам содержания.

6. Компьютерный анализ характеристик аргументационной структуры в разных аспектах позволяет выявлять тексты научной коммуникации с похожей организацией рассуждений и указывает на взаимосвязь показателей разных уровней аргументации.

Достоверность результатов исследования обуславливается подготовкой теоретических и методологических оснований для его проведения на базе фундаментальных работ в областях традиционных исследований аргументации, компьютерной обработки текстов на естественном языке и формального аргументационного анализа. Все приводимые выводы подтверждаются

экспериментами, которые проведены в соответствии с общепринятыми стандартами. Для корпуса с аргументационной разметкой, используемого в экспериментах, подсчитаны формальные оценки согласия аннотаторов на разных уровнях моделирования аргументации. Другим свидетельством надёжности результатов является их публикация в ряде рецензируемых научных изданий. Результаты исследования также использованы в работе по грантам, финансируемым Российским научным фондом в рамках проектов № 23-21-00325 и № 23-11-00261.

Апробация результатов исследования включала их представление в рамках докладов на следующих конференциях и научных семинарах: Международная научная конференция «Корпусная лингвистика 2021 (2023)» (Corpora-21, Corpora-23) (Санкт-Петербург, 2021 г. и 2023 г.); Международная конференция «Марчуковские научные чтения 2021 (2022, 2024)» (Новосибирск, 2021 г., 2022 г., 2024 г.); 2022 (2024) IEEE 23rd (24th, 25th, 26th) International Conference of Young Professionals in Electron Devices and Materials (EDM) (республика Алтай, 2022 г., 2023 г., 2024 г. и 2025 г.); 2022 International Symposium on Knowledge, Ontology, and Theory (KNOTH) (Новосибирск, 2022 г.); XXV International Conference Data Analytics and Management in Data Intensive Domains (DAMDID/RCDL 2023) (Москва, 2023 г.); Международная научно-практическая конференция «Обработка естественного языка (NLP) для лингвистики, IT и бизнеса» (Томск, 2025 г.); регулярный объединённый семинар Института систем информатики Сибирского отделения Российской академии наук и кафедры программирования Новосибирского государственного университета «Интеллектуальные системы и системное программирование» (Новосибирск, 2021 г.); международный научно-практический семинар «Экспериментальные исследования языка и речи: Ресурсные базы и технологии» (E-SoLaS-2022) (Томск, 2022 г.); регулярный объединённый семинар Федерального исследовательского центра информационных и вычислительных технологий и Новосибирского государственного университета «Информационные технологии в задачах филологии и компьютерной лингвистики» (Новосибирск, 2022 г.).

Основное содержание работы отражено в 20 публикациях, из которых 4 статьи изданы в журналах из перечня российских рецензируемых научных журналов, в которых должны быть опубликованы основные научные результаты диссертаций на соискание ученой степени кандидата наук (список ВАК), а ещё 9 работ опубликованы в изданиях, индексируемых в базе данных Scopus.

Личный вклад соискателя может быть охарактеризован в следующих аспектах:

- 1) определения темы исследования, формулировки его целей и задач (как во всём масштабе работы, так и на отдельных этапах), изучения и выборе методологического аппарата, постановки и проверки гипотез;
- 2) подготовки аргументационного корпуса научных статей (отбор и разметка текстов), определения теоретических оснований и принципов для формального представления структур аргументации, составления методики разметки;
- 3) разработки и программной реализации методов автоматического анализа аргументации, их применения для формальной обработки структур рассуждения в текстах корпуса, интерпретации полученных данных;
- 4) подготовки публикаций с представлением результатов исследований, их презентации на международных конференциях и научных семинарах;
- 5) написании текста диссертационного исследования, формулировании выводов и положений, выносимых на защиту.

Структура диссертации образована введением, четырьмя главами (двумя теоретическими и двумя экспериментальными), заключением, библиографией (которая включает 153 наименования) и 6 приложениями. Работа иллюстрирована 5 рисунками, содержит 29 таблиц.

Глава 1. Теоретические основы представления аргументации

Первая глава представленной работы обращена к проблеме определения теоретических оснований для моделирования и анализа аргументации, выражаемой средствами естественного языка. Рассматриваются две группы подходов к аргументационным исследованиям: традиционные лингвистические и формальные.

В первом, вводном разделе Главы 1 детализируется значимость анализа аргументации в фундаментальном лингвистическом контексте. Второй раздел Главы 1 направлен на сопоставление различных лингвистических подходов к изучению аргументации. В пункте 1.2.1 рассматривается общая классификация взглядов на соотношение системы языка и аргументационных структур. В пунктах 1.2.2, 1.2.3, 1.2.4 представлены отдельные подходы к анализу аргументации с лингвистической точки зрения: теория семантических блоков М. Карел (расширение радикального аргументативизма Ж.-К. Анскомбра и О. Дюкро), лингвокультурологическое моделирование аргументационной семантики Э. Риготти и А. Роччи, социо-дискурсивная концепция Р. Амосси (на основе естественной логики Ж.-Б. Грiza). Особое внимание в пункте 1.2.5 обращено к развитию аргументационной теории в отечественных исследованиях (от работ по «языку революции» до когнитивной модели А. Н. Баранова и современного подхода Е. Р. Иоанесян).

В третьем разделе Главы 1 представлен обзор формальных средств для моделирования аргументации. Пункт 1.3.1 содержит описание модели С. Тулмина, позволяющей обозначать детализированный состав отдельных аргументов. В пункте 1.3.2 рассматривается теория риторической структуры за авторством У. Манна и С. Томпсон (как ранее разработанная лингвистическая модель для анализа риторической связности текстов, потенциально применимая и к аргументационным отношениям). В пункте 1.3.3 описывается специальный формат для записи именно аргументационных структур (Argument Interchange Format). Этот формат поддерживает моделирование аргументации в соответствии

с семантической теорией аргументационных схем Д. Уолтона, которая в настоящей работе охарактеризована в пункте 1.3.4.

1.1 Анализ аргументации как фундаментальная лингвистическая проблема

Анализ аргументационных структур в целостных текстах представляет интерес для традиционной лингвистики ввиду многосторонней природы текста как продукта коммуникации, одновременно обращённого и к системе языка (средствами которой он образуется), и ко внеязыковой действительности (к которой относятся и связи между описываемыми в тексте явлениями), и к речемыслительной деятельности говорящего (отдельной языковой личности). Успешное решение обширного круга текстовых проблем требует развития междисциплинарного подхода, интегрирующего достижения различных наук: как собственно лингвистики, так и семиотики, психологии, герменевтики, стилистики, наконец, риторики [Бернацкая, 2009].

Важность риторики в изучении текста подчёркивается лингвориторической парадигмой, согласно которой специфика языковой личности (ЯЛ) в классическом представлении Ю. Н. Караулова (всех трёх её структурных уровней, вербально-семантического, лингво-когнитивного и мотивационного) проявляется наиболее полно в организации речи: в композиции логоса, пафоса и этоса (как мыслительного, эмоционального и нравственного оснований речи) на разных риторических этапах выражения мысли (инвенции, диспозиции, элокуции) [Ворожбитова, 2014]. Как отмечает сам Ю. Н. Караулов, главенствующую роль в иерархии уровней ЯЛ играет мотивационный уровень, на котором выражаются её коммуникативные потребности, а к числу трёх ключевых типов этих потребностей относится воздейственная: аргументация же по своей природе выражает специфику воздействия одной ЯЛ на другую в процессе общения, в связи с чем анализ применяемой аргументации значим для изучения ЯЛ в её культурном контексте (например, изучение семиотического обращения текстов в речевом воздействии, проверка, к каким именно текстам производятся отсылки

разнообразных видов при обосновании тех или иных точек зрения, позволяет выявить прецедентные тексты культуры) [Караулов, 2010, с. 211–237]. Отдельно Ю. Н. Караулов акцентирует ключевые аспекты аргументации, обеспечивающие отражение в её акте отличительных характеристик соответствующей ЯЛ: диалогичность, адресованность (определённой личности или группе людей), субъективность (мотивированность частным интересом), а также её двойственная природа (реализация логической траектории языковыми средствами) [Караулов, 2010, с. 245–258]. Соответственно, значимость риторической перспективы в изучении текста как сложного целого, выстраиваемого по итогам речемыслительной деятельности в процессе коммуникации, влечёт актуальность исследования формальной организации текста на уровне реализованных, уже в нём материально представленных аргументационных структур (в рамках которых и отражается выражение логоса, пафоса и этоса).

Анализ аргументации как воздействия на мнение слушателя важен и в контексте исследования опосредованного языком социального взаимодействия в рамках лингвистического дискурс-анализа (в связи с антропоцентрическим фокусом современного языкознания) [Красина, 2018]. Исследования дискурса, в свою очередь, тесно связаны с лингвистикой текста, к ключевым проблемам которой относится изучение категорий связности и цельности текста [Минина, Карзунина, 2010] как двух его сторон, одновременно противопоставляемых и предполагающих друг друга (невозможности в тексте изолированных компонентов, не связанных с иными, при несводимости текста как образа ситуации к отдельным его компонентам) [Мурзин, Штерн, 1991, с. 11–18]. Одной из таких категорий выступает его аргументационная организация, которая может рассматриваться как один из аспектов связности текста, сочетающий логику изложения и композицию языковых средств, на уровне его макроструктуры [Негрышев, 2011] в классическом представлении Т. ван Дейка (как иерархического источника глобальной связности текста, или его семантического плана, регулирующего стратегию его представления на локальном и линейном уровне, например, в распределении эксплицитной и имплицитной информации,

установлении локальной текстовой связности) [ван Дейк, 2000, с. 41–67]. В свою очередь, разработанные Т. ван Дейком и В. Кинчем концепции макроструктуры и макропропозиции удобно использовать для представления понятий аргументационной структуры текста и отдельного аргумента: аргумент, образованный системой отдельных утверждений (посылок и заключения), может быть соотнесён с макропропозицией как пропозицией, выводимой из ряда пропозиций в результате свёртывания их смысловой структуры, а роль отдельных аргументов в аргументационной структуре близка в данном ракурсе соответствующей роли макропропозиций в обеспечении глобальной связности текста на уровне его макроструктуры (с учётом того, что две данные пары понятий принадлежат разным теориям и, соответственно, разным терминологическим подсистемам).

В дополнение к основным категориям, общим для всех текстов, современными исследователями подчёркивается способность реальных текстов проявлять категории, обусловленные коммуникативной спецификой различных сфер деятельности, причём изучение таких ситуационных категорий необходимо для всестороннего представления текста [Иргашева, 2011]. Рассмотрение аргументации как текстовой категории обуславливается и распространённым мнением о её реализации исключительно на уровне целостных текстов в процессе коммуникации, без возможности ограничения аргументации традиционными языковыми единицами с предшествующих уровней языка [Филиппов, 2003, с. 270] (однако существует и точка зрения о принадлежности аргументационного компонента системе языка, выраженная, в частности, в теории «радикального аргументативизма» [Anscombre, Ducrot, 1989]). Анализ способов ведения дискуссии, изложения доказательств и убеждения слушателей в научной коммуникации представляет естественный интерес при исследовании научного дискурса, но также следует отметить его прикладное педагогическое применение для усовершенствования методик развития языковой личности [Архипова, 2005].

В свою очередь, выявление аргументационной специфики разных жанров научной коммуникации и исследование различительной способности приёмов

аргументации в передаче жанровой специфики обретает особую значимость в контексте стилистики и конкретно проблемы разграничения и описания стилей, выявления связей между устойчивыми группировками текстов и изменчивыми наборами реализуемых в них языковых конструкций [Костомаров, 2005, с. 49–52]. Сведения об аргументационной организации текстов разных жанров, отражающей их смысловое структурирование в соответствии с прагматическими факторами в коммуникативной ситуации (ввиду направленности процесса убеждения на читателя), дополняют набор разноаспектных характеристик для типологизации текстов [Чаплина, 2009]. Типология текстов выступает одной из основных проблем лингвистики текста и даже выделяется в её самостоятельный раздел (по примеру лингвистической типологии языков) [Чернявская, 2009, с. 51–68]. Эта задача осложняется тем, что на характер текста влияют все его параметры в совокупности, в связи с чем научная типология текстов должна принимать во внимание признаки со всех уровней языковой системы (по аналогии с отмеченной типологией языков) [Левицкий, 2006, с. 202–205]. Анализ аргументационных структур в текстах научной коммуникации, направленных на доказательство излагаемых выводов, позволит оценить информативность показателей этого уровня в разграничении типов текстов.

Двуплановая природа аргументации (с одной стороны, как воздействие на мнение слушателя посредством доказательства некоторой идеи, с другой, как процесс достижения вывода путём построения логических связей в ходе мышления) обеспечивает её теоретическую актуальность в контексте обеих ключевых парадигм современной лингвистики: не только коммуникативной, но и когнитивной. Аргументационные структуры отражают мыслительные процессы, связанные с возникновением, языковой репрезентацией и передачей знания, которые находятся в фокусе внимания когнитологов [Кубрякова, 2012, с. 17]. Отдельное внимание когнитивной науки привлекает репрезентация знания в научной среде, специализирующейся на получении нового знания. В этом контексте исследуется как конструирование структур знания в научном дискурсе (в академическом, научно-популярном и научно-учебном типах) с учётом его

прагматической специфики в организации сообщаемой информации для убеждения адресатов [Манерко, 2017], так и осмысление научной коммуникации субъектами науки из разных социальных групп [Белоусов, Гатаулин, Ерофеева, 2017]. Оба указанных аспекта проявляются в аргументационных структурах, анализируемых в представленной работе.

Соответственно, анализ формальных аргументационных структур в текстах научной коммуникации (в том числе в аспекте их компьютерной обработки) представляется актуальным в контексте широкого круга разделов фундаментальной лингвистики: прагматики, лингвистики текста, стилистики, типологии текстов и когнитивной лингвистики.

1.2 Лингвистические подходы к анализу аргументации

В этом разделе рассматривается проблема соотношения аргументационных структур и выражающих их языковых средств (в терминах более общего вопроса о связях между лингвистикой и теорией аргументации). Указанная проблема проистекает из междисциплинарного характера аргументационных исследований, обуславливающего различия в определении базовых понятий [van Eemeren et al., 2014, с. 14–16]: так, теоретические перспективы для работы с аргументацией могут выстраиваться на основе риторики, формальной либо неформальной логики, когнитивной психологии и, наконец, лингвистики (в частности, в рамках дискурсивного анализа), чья методология оказывается особенно актуальной ввиду преимущественно языковой природы у средств выражения аргументации [Там же, с. 479–516].

В частности, уже центральное понятие аргументации определяется по-разному в зависимости от теоретического подхода. Например, прагматико-диалектический подход Ф. Еемерена основывается на деятельностном представлении аргументации с выделением трёх ключевых аспектов [van Eemeren, Grootendorst, 2004]: вербального (аргументация реализуется через применение языковых средств), социального (она предполагает обращение к иным участникам

коммуникации, явное или неявное) и рационального (ориентирована на построение в первую очередь интеллектуальных доводов). Согласно третьему аспекту Ф. Еемеерена, аргументационному анализу подлежит логическая структура доказательства (с прескриптивным разграничением корректных и некорректных доводов), тогда как эмоциональное воздействие в процессе убеждения считается нарушением интеллектуальной точности доводов. Указанное преобладание рационального компонента подразумевает сужение традиционной концепции Аристотеля о методах убеждения: согласно классификации, изложенной в «Риторике» [Аристотель, 1978], доказательство может выстраиваться как посредством цепочек рассуждения (что соответствует логическому аспекту, или логосу), так и через внушение доверия к говорящему (этос) либо воздействие на эмоции слушателей (пафос).

При этом Аристотель отдельно подчёркивает равную значимость всех трёх методов: уместность каждого метода (корректность или некорректность реализации в конкретной ситуации) определяется убеждаемой аудиторией и личностью говорящего. Так, эмоциональные аргументы сильнее воздействуют на юношей, чем логические, однако менее эффективны при воззвании к возрастным слушателям. Доводы от личного авторитета убедительнее звучат у почтенного говорящего, чем у юного. Соответствие аргументации контексту конкретной коммуникативной ситуации отражается в аспекте её своевременности (кайросе), который обретает особую значимость в рамках аристотелевской модели ввиду его непосредственного влияния на проявление трёх других аспектов, в большей мере эмоционального и этического, но также и логического [Kinneavy, Eskin, 2000].

Так, при воздействии на чувства аудитории следует учитывать не только её эмоциональное состояние, но и ситуативные факторы, обуславливающие это состояние (для двух слушателей, раздражённых в одинаковой степени по разным причинам, могут потребоваться два разных способа обращения к их раздражению: например, для чувствующего себя обделённым человека эффективно акцентирование несправедливости, тогда как для обычно спокойного слушателя полезно указать на вынужденный, неординарный характер его раздражения).

Сходным образом, при апелляции к статусности важно учитывать, кто и по каким причинам является авторитетом для аудитории, и даже при построении сугубо логических доказательств полезно оценивать степень восприимчивости к ним слушателей (например, в зависимости от уровня образованности аудитории могут быть наиболее иллюстративными логические построения меньшей или большей сложности). Таким образом, кайрос характеризуется определяющей ролью по отношению к итоговой эффективности аргументов (как собственно убедительности, так и, в зависимости от целей оратора, поучительности и полезности) в конкретной коммуникативной ситуации в определённый момент времени.

Ввиду рассмотрения эмоционального и этического компонента наравне с рациональным к классической концепции Аристотеля ближе определение аргументации, предлагаемое Х. Перельманом и Л. Олбрехт-Тытека в «Новой риторике»: система дискурсивных техник, позволяющих обусловить или усилить поддержку слушателями представленного тезиса для их согласия с ним [Perelman, Olbrechts-Tyteca, 1969, с. 4]. Однако при подобном понимании аргументации никак не учитывается вербальный компонент, который Ф. Емерен выделяет в первую очередь. Тем не менее, языковой аспект эксплицируется в ином определении в рамках дискурсивного подхода: Р. Амосси представляет аргументацию как «управление ракурсом для восприятия реальности, воздействие на точку зрения и поведение посредством широкого спектра вербальных средств» [Amossy, 2009]. При этом Р. Амосси подчёркивает, что при её формулировке аргументация как исследовательский объект оказывается полностью встроенной в расширенную лингвистическую парадигму (полную предметную область науки о языке без ограничения внутри дискурсивного анализа). Впрочем, в определении Р. Амосси подчёркивается применение языка в рамках речевой деятельности (активное использование вербальных средств для управления ракурсом представления реальности), тогда как отдельным лингвистическим подходам свойственно выделение аргументации внутри системы языка. Например, согласно лаконичной формулировке Л. Пуиг (в контексте теории семантических блоков,

разработанной под сильным влиянием структурализма), термин «аргументация» обращён к связям между высказываниями [Puig, 2012, с. 128].

Таким образом, любое определение аргументации как явления задаётся метаязыком теории, в терминах которой оно сформулировано.

1.2.1 Подходы к соотношению аргументации и языка как системы

Исследование аргументации с позиций лингвистической теории предполагает в первую очередь определение соотношения двух явлений: выражаются ли аргументационные структуры непосредственно на уровне языка (а следовательно, принадлежат его системе) либо же представляют собой отвлечённые логические или риторические конструкции, внешние по отношению к языковым средствам и их функционированию.

Сравнивая возможные подходы к обозначенному вопросу, французские лингвисты Ж.-К. Анскомбр и О. Дюкро разграничивают четыре класса взглядов на соотношение языка и аргументации, противопоставляемые по уровню признаваемой связи между обоими типами явлений [Anscombre, Ducrot, 1989]:

- 1) лингвистические структуры высказываний никак не связаны с их аргументационным функционированием (их использованием в построении и обосновании рассуждений);
- 2) аргументационные структуры не принадлежат системе языка, но их реализация зависит от базовых семантических свойств высказываний (к примеру, распределения пресуппозитивного и ассертивного содержания);
- 3) значение отдельных языковых выражений задаётся исключительно их использованием в аргументации (некоторые языковые конструкции наделены аргументационным содержанием, тогда как иные его лишены);
- 4) аргументационный аспект характеризует все языковые выражения: не существует полностью информативных высказываний, лишённых аргументационного значения (теория «радикального аргументативизма»).

1.2.1.1 Независимость аргументации от языкового выражения

Представление аргументации в отрыве от лингвистического аспекта отмечается, в частности, в подходах Я. Хинтика [Hintikka, 1989] и Ч. Хэмблина [Hamblin, 1970]. Так, Я. Хинтика возражает против междисциплинарного разграничения формальных приёмов аргументации и стратегий их применения (вынесение таких стратегий из логики, отнесение к риторике или лингвистике считает некорректным): указывает на предпочтительность описания обоих типов явлений на общем строго логическом основании. В свою очередь, Ч. Хэмблин анализирует логические ошибки (*fallacies*) в аргументации, в связи с чем рассматривает проблему демаркации аргументов (к которым применима проверка на отсутствие логических ошибок) и неаргументативного содержания (для которого такая проверка нерелевантна). Например, указание на ненадёжность человека, высказавшего некий тезис, не является логической ошибкой, если выдвигается как изолированный комментарий, а не как посылка к заключению о некорректности этого тезиса. Согласно Ч. Хэмблину, подобные отношения между аргументами проявляются на уровне их формальной организации, логическая точность которой лишь в ограниченной мере отражается языковым выражением (рассуждение, аккуратно организованное в лингвистическом плане, может являться логически несвязным, и наоборот). Соответственно, в рамках указанного подхода анализ аргументации предполагает переход на уровень логической структуры, абстрагирование от языкового оформления.

Таким образом, согласно первой группе подходов к представлению аргументационных структур в дискурсе предполагается, что эти структуры соотносятся непосредственно с фактами реальности, а не с высказываниями, отсылающими к этим фактам. Под высказыванием понимается реализация некоего предложения естественного языка в процессе коммуникации. Так, высказывание *U*, реализующее предложение *P*, приводит к заключению *C* в том случае, если *U* указывает на факт *F*, а участники дискурса имеют основания полагать, что *C* истинно в случае *F*. Соответственно, при таком представлении

между высказыванием U и заключением С нет прямой связи, а принадлежащее естественному языку предложение Р рассматривается только в аспекте его применения в процессе коммуникации.

Указанная точка зрения свойственна аналитической традиции в анализе аргументации: структуры убеждения рассматриваются как строго логические, а следовательно внеязыковые конструкции. При данном подходе различные проблемы языкового выражения аргументации (например, лингвистическая неоднозначность предложений, которые допускают несовместимые варианты интерпретации их соотношения с фактами реальности) отграничиваются от проблем собственно анализа аргументации и выводятся из рассмотрения.

1.2.1.2 Влияние базовых семантических свойств на аргументацию

Однако, как показывают Х. Перельман и Л. Олбрехт-Тытека в теории новой риторики, заключения могут основываться не только на самих фактах, но и на интерпретируемых намерениях говорящих коммуницировать данные факты посредством отсылающих к ним высказываний [Perelman, Olbrechts-Tyteca, 1969, с. 123–126]. Акцент на намерении как на основе аргументации (соответственно, и ситуации дискурса) проявляется особенно наглядно при реализации ряда риторических фигур, таких как аллюзия: эмфатическое выделение значимого события или культурного факта через их не прямое указание способно подчеркнуть духовную общность говорящего и слушающих, поддержать убедительность его выводов [Там же, 1969, с. 171–179]. В подобных аргументационных конструкциях способность U выражать F зависит по меньшей мере частично от лингвистической структуры предложения Р в основе U (как высказывания в процессе коммуникации, заданного выбранными языковыми выражениями).

1.2.1.3 Аргументационная роль отдельных языковых выражений

При этом Ж.-К. Анскомбр и О. Дюкро демонстрируют, что некоторые языковые выражения могут приносить аргументационное содержание в высказывания независимо от ситуативных особенностей их употребления. Так, коннективы вида «поэтому», «следовательно», «так как» указывают в любой коммуникативной ситуации, что факты, вводимые через одно высказывание, подразумевают истинность фактов, к которым отсылает другое утверждение на естественном языке. Такие коннективы также подчёркивают двойственную связь между обеими парами «высказывание и факт» (на уровне как синтаксиса, так и логической структуры). Следствием отмеченной тенденции является возможность разграничить конструкции с аргументационным значением и строго информативные выражения, не обладающие таким значением.

1.2.1.4 Теория радикального аргументативизма

Однако по итогам анализа предполагаемых языковых единиц из второй группы (строго информативных) Ж.-К. Анскомбр и О. Дюкро приходят к выводу, что в применении даже наиболее нейтральных из них (с предельно общим, элементарным значением) отмечаются явные аргументационные ограничения. В частности, союз «и» не может связывать два противоречащих друг другу высказывания, тогда как добавление артикля «un» способно изменять аргументационный потенциал утверждений (их роль в дальнейшем развитии дискурса): заявления «Pierre a peu travaillé» и «Pierre a un peu travaillé» передают противоположные оценки («Пьер мало поработал» и «Пьер чуть поработал»).

При этом сочетаемость выражений в аргументационных конструкциях задаётся выбором не только служебных, но и знаменательных слов. Например, хотя утверждения «Pierre est aussi grand que Marie» («Пётр столь же высок, как и Мария») и «Pierre a la même taille que Marie» («Пётр того же роста, что и Мария») передают одно фактическое содержание (о соотношении роста Петра и Марии),

только второе из них доступно для стилистически нейтрального применения в обосновании тезиса «*Pierre n'est vraiment pas grand pour son âge*» («Пётр не так высок для своего возраста»). В частности, во втором утверждении Пётр и Мария сопоставляются по нейтральной шкале роста, тогда как в первом им обоим приписывается позитивное выражение признака высокого роста. Как следствие, в рассуждении вида «[Пётр не так высок...], потому что [Пётр так же высок...]» возникает противоречие между посылкой, отсылающей к выражению признака, и заключением, опровергающим выражение этого же признака. Такое обоснование передаёт неотъемлемый иронический компонент, который маркирован стилистически и состоит в описании человека посредством качественного прилагательного, указывающего на несвойственный этому человеку признак.

В связи с выявленной тенденцией Ж.-К. Анскомбр и О. Дюкро формулируют концепцию радикального аргументативизма, согласно которой семантика предложений полностью задаётся выбором и представлением фактов с аргументационной точки зрения. Такой выбор проявляется уже на уровне применяемых слов: к примеру, два варианта описания объекта, как «дорогого» либо «недешёвого», отсылают к различным наборам понятий – «дороговизне» либо «дешевизне». При развитии дискурса для первого набора могут актуализироваться максимы вида «чем вещь дороже, тем она качественнее» или «чем вещь дороже, тем хуже сделка» (в зависимости от отношения говорящего), тогда как для второго набора акцентируется отрицание воззрений о дешевизне (которые основываются на максимах со сходной структурой, а именно на попарном сопоставлении оценочных шкал).

Такие подразумеваемые воззрения, не выражаемые в дискурсе явно, но обеспечивающие смысловую связность текста, Ж.-К. Анскомбр и О. Дюкро называют топосами (по аналогии с концепцией Аристотеля). При этом, согласно их определению, любой топос можно представить как соотношение нескольких оценочных шкал: например, «чем усерднее работаешь, тем больше зарабатываешь» или «чем усерднее работаешь, тем сильнее устаёшь». Смысловое содержание каждой шкалы обуславливается тогда её связями с другими шкалами

(в соответствии с парадигмой структурализма и концепцией «значимости» Соссюра [Соссюр, 1999], в частности). В свою очередь, лексическое значение единиц языка сводится к системе потенциально актуализируемых ими топосов: например, предикат «работать» может быть представлен через уточнение отношений внутри пространного набора шкал (вкладываемых усилий, получаемой выгоды, усталости, образованности и д.р.).

Таким образом, согласно концепции радикального аргументативизма, любой текст на естественном языке выражает аргументационное значение. Подобный неотъемлемый компонент проявляется при уточнении ракурса представления излагаемых сведений на уровне как выбираемых лексических средств, так и их интеграции в единое смысловое целое через предпочтение тех или иных синтаксических структур.

1.2.2 Теория семантических блоков

Развитием и уточнением концепции радикального аргументативизма является теория семантических блоков, разработанная М. Карел [Carel, 1995]. Эта теория основывается на переосмыслении понятия топоса как элементарной семантической единицы при аргументационном представлении любых лексических значений. Если в модели Ж.-К. Анскомбра и О. Дюкро исходным считается непротиворечивое соотношение смысловых шкал (противопоставленные воззрения «чем человек богаче, тем больше его любят» и «чем человек богаче, тем меньше его любят» соответствуют разным топосам), то М. Карел указывает на смысловую общность в основе подобных сопоставлений, различающихся по заявленной оценке. Например, если в значении слова «богатство» в системе языка выражается воззрение «чем человек богаче, тем он счастливее», то утверждение вида «чем сильнее Пьер богатеет, тем несчастнее он становится» не будет, согласно теории М. Карел, отсылать к принципиально иному значению: наоборот, оно продолжает реализовывать исходный смысловой компонент, поскольку при описании ситуации подразумевается её нетипичный

характер (который потерял бы актуальность для языка, где понятия богатства и счастья изначально не соотносятся). Иначе говоря, М. Карел акцентирует внимание на общем основании сопоставления, описывая через него семантический блок как всю совокупность потенциальных сопоставлений на этом общем основании, а Ж.-К. Анскомбр и О. Дюкро, игнорируя общее основание сравнения, фокусируются на случаях согласованности или противопоставленности оценочных шкал.

Таким образом, элементарной смысловой единицей в системе языка при подходе М. Карел становится семантический блок – набор потенциальных аргументационных связей между соотносимыми понятиями, содержание которых задаётся именно их взаимной зависимостью. Как следствие, описание значения отдельного слова, синтагмы или предложения подразумевает выявление аргументационных связей, допускаемых этой единицей (согласно определению Л. Пуиг [Puig, 2012], использующей теорию семантических блоков при анализе романских языков). Такие связи, выявляемые через анализ аргументации, и отражают смысловую зависимость языковых единиц: употребление одной подразумевает указание другой при раскрытии и неявном обосновании точки зрения на описываемые внеязыковые реалии.

Следует отметить, что отношения между внеязыковыми реалиями и лингвистическими единицами полагаются в теории семантических блоков исключительно опосредованными: лексические значения рассматриваются в их ограниченности аргументационным потенциалом (применимостью в различных схематизациях действительности через использование языковых конструкций: изложение, обоснование, оценку тех или иных точек зрения). Указанный подход к представлению лексических значений сближает теорию семантических блоков (на уровне философских оснований) с концепцией Л. Витгенштейна, согласно которой значение слова соответствует его употреблению в языке [Витгенштейн, 1994, с. 99].

Как следствие, теория семантических блоков, являясь развитием концепции радикального аргументативизма с переосмыслением элементарной смысловой

единицы (ввиду перехода от отдельного оценочного воззрения к набору потенциальных аргументационных связей), сохраняет идею Ж.-К. Анскомбра и О. Дюкро о тесной связи языковых выражений и аргументационных значений.

1.2.3 Анализ аргументации в выявлении ключевых слов культуры

Альтернативный подход к анализу лингвистической семантики в терминах аргументационной теории предлагают Э. Риготти и А. Роччи, италоязычные исследователи из Университета Лугано. В работе [Rigotti, Rossi, 2005] они обращаются к проблеме выявления ключевых слов для заданной языковой культуры, основываясь на формулировке задачи А. Вержбицкой. Хотя А. Вержбицкая утверждает, что объективной процедуры для определения ключевых слов в какой-либо культуре не существует [Wierzbicka, 1997, с. 6], Э. Риготти и А. Роччи описывают способ формальной проверки, является ли некое слово ключевым для соответствующей культуры, посредством анализа возможных ролей этого слова в аргументативных текстах.

Так, Э. Риготти и А. Роччи определяют исследуемый объект как лексические выражения, указывающие на внутреннее устройство культуры в целом, на её основополагающие воззрения, ценности, традиции и институты. Согласно их концепции, культурная значимость таких ключевых слов проявляется особенно явно в их расширенном потенциале (по сравнению с иными словами) для реализации в ценностных суждениях, а именно в их активном применении в роли среднего термина внутри энтимем для ведения доказательств, опирающихся на ценностные ориентиры общества. Так, под энтимемами здесь понимаются неполные простые категорические силлогизмы (логические структуры с построением заключения на основе двух посылок), в которых при этом одна из посылок либо заключение пропускается, то есть выражается неявно. В свою очередь, средний термин связывает обе посылки в структуре силлогизма (содержится в них обеих, однако не входит в заключение). Как следствие, применение в этой роли ключевого слова культуры (семантическое поле которого

расширяется чёткими ассоциативными связями, задаваемыми ценностными ориентирами общества) значительно облегчает восстановление слушателем пропущенной посылки (за счёт выявления соответствующей ассоциации для среднего термина) либо заключения (через распознавание точного характера связи между двумя явно выраженными посылками).

Примером аргумента такого типа, в котором неполнота выражения обуславливается предположением об очевидности подразумеваемой посылки, может служить выражение «X – предатель и, соответственно, заслуживает смерти». Анализ аргументационной структуры этого выражения позволяет выявить два явно представленных утверждения: «X – предатель» и «X заслуживает смерти». При этом переход от первого ко второму обеспечивается неявной апелляцией к посылке «предатели заслуживают смерти». Хотя такая характеристика не заключена в непосредственном лексическом значении слова «предатель» (среднего термина внутри аргумента), подразумеваемая очевидность посылки (выражение в пресуппозиции её культурно значимого содержания) свидетельствует об однозначной отрицательной направленности слова «предатель» (среднего термина в аргументе), а также о его культурной значимости ввиду неотъемлемой передачи ценностного компонента. Так, аргументационный потенциал ключевого слова и заключается в его тесной связи с таким оценочным компонентом (расширяющим значение этого слова для данной культуры). Важно отметить, что Э. Риготти и А. Роччи рассматривают случаи употребления энтимем для выражения имплицитно коммуницируемой информации, выводимой слушателем (импликатур): в этих случаях говорящий намеревается привлечь внимание слушателя к имплицитной части содержания (акцентируемой через неявное представление), а не старается скрыть манипулятивное воздействие через умалчивание некоторых сведений в основе логического перехода.

Следует отметить, что аргумент из приведённого примера окажется бессмысленным для культуры, в которой понятие «предатель» лишено подразумеваемого отрицательного компонента и не отсылает к общественным

ценностям. В этой связи тезис о расширенном аргументационном потенциале ключевых слов культуры (а соответственно, ограниченной употребительности слов, не являющихся ключевыми) находит поддержку в исследовании М. Стаббса [Stubbs, 1996]. Так, посредством анализа политических речей британских консерваторов М. Стаббс демонстрирует, что выстраиваемые ими аргументы представляют при строгом рациональном рассмотрении цепочки *non sequitur*, или несвязанных утверждений. Осмысленность же этих речей обеспечивается не логической корректностью эксплицитно представленных рассуждений, а активным обращением к подразумеваемым ценностям через их представление в неявных посылках за счёт частого повторения одних и тех же слов.

При этом Э. Риготти и А. Роччи подчёркивают, что актуализация отдельных понятий как ключевых (на уровне как общеязыковой культуры, так и для отдельных социальных групп) может сопровождаться явной формулировкой тезисов, которые впоследствии и подразумеваются при аргументационном применении этих слов. В качестве примера речи, трансформирующей значимость отдельного слова внутри целой культуры, исследователи приводят выступление Дж. Буша от 4 апреля 2002 г., посвящённое политике США на Среднем Востоке. Речь президента организована вокруг тезиса «*Terror must be stopped*», для обоснования которого Дж. Буш приводит утверждения вида «*No nation can negotiate with terrorists*» и «*For there is no way to make peace with those whose only goal is death*». Между тем заявленные американским президентом характеристики не выводятся напрямую из лексического значения слов «*terror*» и «*terrorist*». Таким образом, раскрытие функционирования этих лексических единиц в американской культуре подразумевает выяснение их аргументационной роли в изначальной речи Дж. Буша, а также и в общественном дискурсе впоследствии.

Указанная ориентированность на аргументационный анализ семантики непосредственно в коммуникативном пространстве отличает подход Э. Риготти и А. Роччи от модели Ж.-К. Анскомбра и О. Дюкро, в которой значения лексических единиц определяются на уровне их системных связей. Противопоставление структурализма и коммуникативной парадигмы, таким

образом, отмечается и для аргументационных исследований внутри лингвистической теории. Кроме того, позиция Э. Риготти и А. Роччи о качественном различии аргументационного потенциала ключевых слов культуры и слов, не являющимися ключевыми, близка третьей группе взглядов о соотношении языковых систем и аргументационных структур (по классификации Ж.-К. Анскомбра и О. Дюкро).

1.2.4 Дискурсивный подход к аргументации в естественном языке

Иной альтернативой для концепции радикального аргументативизма, включающего аргументационный компонент в семантику языковых единиц (в первую очередь на лексическом уровне), выступает дискурсивный подход к лингвистическому анализу аргументации. Источником этого подхода является естественная логика Ж.-Б. Гриза [Plantin, 2003], название которой обусловлено её противопоставлением формальной логике: абстрактность формальной логики («отсутствие субстанции») уточняется через обращение к лингвистическим формам (предложениям и текстам), посредством которых логические схемы и выражаются при повседневном применении языка.

Так, согласно Ж.-Б. Гризу, все когнитивные и лингвистические операции, задающие организацию текста, выражают аргументационный компонент: как и в модели Ж.-К. Анскомбра и О. Дюкро, любое высказывание признаётся выражающим некую точку зрения (которую говорящий схематизирует для слушателя). Тем не менее, для естественной логики аргументационные структуры проявляются именно при употреблении языка в реальных коммуникативных ситуациях: как следствие, дискурсивный подход к аргументации (развитие концепции Ж.-Б. Гриза в лингвистическом аспекте) ориентирован на анализ целостных текстов, или продуктов коммуникации, с учётом ситуативной специфики их создания.

В частности, как утверждает Р. Амосси при описании дискурсивного подхода [Amossy, 2009], аргументационный компонент не задаётся

семантическими отношениями внутри системы языка, а выражается только в коммуникативной ситуации: способность отдельных языковых выражений отсылать к другим (на основе их смысловой сопоставимости) ограничивается условиями общения (жанром дискурса, обсуждаемой темой, статусом говорящего и слушающего, их заявляемыми и скрываемыми намерениями, интертекстуальным влиянием иных дискурсов и т.п.). Свободно варьирующиеся условия общения не могут быть зафиксированы в системе языка, а в свою очередь, они так же активно влияют на выбор говорящими языковых единиц, на их представление своих точек зрения, как и системные значения используемых языковых выражений.

Кроме того, как указывает Р. Амосси, анализ аргументационных структур на уровне их текстовой реализации позволяет реконструировать формальную организацию применяемых доводов. Такое восстановление абстрактных схем, задающих переходы от одних утверждений к другим при раскрытии позиции говорящего, оказывается особенно важным при анализе связности текста как лингвистического целого. Однако и формальное строение аргументационных структур не может рассматриваться изолированно от языковых средств, их выражающих: согласно Р. Амосси, именно сочетание обоих типов конструкций (логических и лингвистических) обеспечивает представление в тексте мнений и отношений, неотъемлемо характеризующих описание любой ситуации.

Соответственно, дискурсивный подход к аргументации на естественном языке представляет собой умеренное сочетание двух парадигм: формально-логической и лингвистической. Кроме того, этот подход отличает проведение анализа на уровне целостных текстов, что особенно актуально для представленного в этой работе исследования академической коммуникации. Так, согласно исходному положению о выражении аргументационных структур на уровне законченных текстов, в качестве материала для проведённого формального анализа аргументационных стратегий взяты полные тексты научных статей: каждая статья рассматривается как целостное, неразрывное доказательство ключевого тезиса в её основе.

1.2.5. Развитие аргументационной теории в отечественной лингвистике

При обращении к истории аргументационных исследований в отечественной науке необходимо отметить сравнительно позднее включение естественноречевой аргументации в предметную область лингвистики, указываемое А. Н. Барановым. В его представлении вывод аргументационной проблематики из внимания отечественных исследователей был сопряжён с эксплицитной идеологизированностью, присущей советскому языкознанию. Так, после появления первых работ о средствах убеждения в языке революции, опубликованных в 20-е гг. XX века, проблематика речевого воздействия практически не получала дальнейшего освещения вплоть до 80-х гг. XX века [Баранов А.Н., 1990].

В рамках исследования языка революции А. М. Селищев рассматривает динамику французского языка во время Французской революции, сравнивает её с новыми на тот момент времени тенденциями в русском языке [Селищев, 1928]. Он анализирует экспрессивность и убедительность в языке публикаций, обращённых к широким слоям общества, изучает разные формы передачи эмоциональности ораторами революции (такие как тыканье, повторение, стремление к краткости речи через её прямолинейность). А. М. Селищев указывает на интенсификацию языковой деятельности в периоды социальных потрясений: разного рода патриотические общества и клубы революционной эпохи вовлекают всё больше участников, которые из идеологических потребностей регулярно упражняются в ораторском искусстве для привлечения новых сторонников и в ходе митингов, массовок, съездов и заседаний. Автор также отмечает историческое преимущество русского революционного языка над французским: последний характеризовался резким разрывом с предшествующим салонным языком (полному устоявшимся, практически канонизированным изящным форм), тогда как литературный русский язык XIX в. уже хорошо подходил для выражения самых сложных и утончённых общественных реалий и требовал только дальнейшего обогащения новыми выражениями. При таком

различии в отношении революционных языков к прежним доминирующим нормам А. М. Селищев подчёркивает преемственность средств языкового воздействия русскими революционерами от французских предшественников, в особенности в аспекте эмоционального влияния.

Актуальность изучения языка революции для формирующегося нового общества оказывалась сопряженной с задачами повышения убедительности агитации и противодействия неблагонадёжной полемике. Так, Л. П. Якубинский анализирует рефлексию В. И. Ленина о воздействии на общественную точку зрения [Якубинский, 1926] и концентрируется на «революционной фразе» как ключевом понятии в его полемическом арсенале. «Революционная фраза» рассматривается как этап исторического развития революционного лозунга, его перерождение и вырождение в отрыве от объективных обстоятельств. Важность связи лозунга и социальных реалий в авторской интерпретации В. И. Ленина обуславливает необходимость анализа политических высказываний в их конкретном применении, иными словами, учёта их прагматического компонента. Акцентируется деятельностный характер убеждающего лозунга: его применение представляет собой целенаправленное действие и неразрывно связано с конкретной ситуацией, на уровне которой и сформирована цель. Так, провозглашение лозунгов без учёта общественных потребностей уменьшает их эффективность в воздействии на общественную точку зрения и может привести к появлению отторжения к отстаиваемым тезисам. Привязанность лозунга к реальной ситуации проявляется в различных аспектах, в том числе экспрессивном: уменьшение резкости выражений и применение более аккуратных формулировок, балансирующих между интересами разных сторон, может привести к разрыву этой связи, когда лозунг перестаёт отвечать требованиям действительности. Однако сама по себе эмоциональность высказывания не тождественна его прагматической надёжности: наоборот, экспрессивные призывы, оторванные от социальных реалий, способны успешно вводить людей в заблуждение, ввиду чего явно обозначается актуальность выявления и противодействия неблагонадёжной агитации. В соответствии с этим

Л. П. Якубинский формулирует типы возможных уловок для введения в заблуждение: притворная самоидентификация оратора с аудиторией, умалчивающая эвфемистичность стиля и дистанцирование высказываемых лозунгов от контекста иных высказываний в общественной дискуссии.

Общественный спрос на изучение агитационных механизмов мотивировал и более общие теоретические исследования языковой стороны убеждения. Так, в работе [Якубинский, 1986] автор указывает на диалогическую природу ораторской речи и активную вовлечённость убеждаемой аудитории в процесс воздействия на её точку зрения. Подчёркивается роль невербальных условий коммуникации на эффективность аргументации: слушатель, лучше видящий оратора, скорее и более полно воспримет его тезисы.

Тем не менее, ввиду общественной ситуации и положения дел в гуманитарных науках отечественные исследования аргументации не получали системного развития на протяжении нескольких десятилетий. Впрочем, значимость аргументации как языкового явления обусловила постепенное возвращение исследовательского внимания к её отдельным аспектам. Так, Ю. Н. Караулов рассматривает языковые средства убеждения в терминах характеристики ими языковой личности [Караулов, 2010, с. 245–258]. Автор указывает, что логическая или содержательная ограниченность аргументов могут успешно компенсироваться их лингвистическим воплощением, например, за счёт сочетания разнообразных языковых средств оформления (как лексико-грамматических, так и более сложных, таких как изменение модальности или стилистической окраски, метафоризация, движение интонационного контура). При этом особенно информативными в характеристике языковой личности Ю. Н. Караулов полагает понятия ценностного поля, как актуализируемые в процессе аргументации, так и, наоборот, последовательно избегаемые на протяжении различных коммуникативных ситуаций.

Наиболее значимым трудом в истории отечественных исследований по естественно-языковой аргументации выступает докторская диссертация А. Н. Баранова, также являющаяся первой систематизирующей работой в данной

области. Его методологический подход к изучению аргументации сформирован в рамках когнитивной парадигмы: «...лингвистическое изучение аргументации следует рассматривать как часть общей модели мышления и коммуникативной деятельности человека...» [Баранов, 1990, с. 4]. В соответствии с этим, автор отмечает двойственную природу исследуемого явления и разграничивает аргументацию как комплекс языковых средств для воздействия на поведение человека и как отдельный тип дискурса с особыми коммуникативными и иллокутивными целями. Речевое воздействие понимается как совокупность процедур онтологизации знания (ввода новых знаний и модификации уже имеющихся в модели мира адресата) путём применения определённых языковых выражений в процессе коммуникации.

В рамках когнитивной методологии А. Н. Баранов выделяет два типа аргументации, способных сочетаться в одном речевом акте. Номинативная аргументация заключается в навязывании адресату концептуальных рамок восприятия ситуации, причём фиксирование этих рамок осуществляется используемыми языковыми выражениями. Интегративная аргументация предполагает иерархическое структурирование передаваемых смыслов по важности, что подразумевает, как правило, манипулирование распределением сведений между пропозицией и пресуппозицией. Таким образом, воздействие на адресата может обеспечиваться как иллокутивной семантикой внутри конкретного речевого акта, так и оценочными смыслами в лексическом значении используемых слов. Степень интенсивности оценочных смыслов разных типов (количественных, прототипических, целевых, общих) служит основанием для ранжирования языковых выражений посредством оценочной шкалы.

Отдельно следует отметить предложенную А. Н. Барановым классификацию аргументационных маркеров: обоснование тезиса может сопровождаться ссылкой на источник (авторитет, общеизвестность, предшествующий контекст), апелляцией к адресату, ссылкой на очевидность или ирреальной отсылкой к гипотетической ситуации. Автор выделяет и дополнительные классы аргументационных маркеров: структурирующие процесс

аргументации и экземплификаторы (показатели примеров). Такая классификация близка иным зарубежным подходам к семантико-функциональной классификации аргументов (например, классификации Д. Уолтона). В целом, характеризуя аргументационную теорию А. Н. Баранова, следует подчеркнуть его акцент на взаимодействии языковых структур в актах убеждения и общих механизмов мышления.

Среди современных отечественных работ по аргументационной теории ключевое место занимают труды Е. Р. Иоанесян. Её подход к представлению аргументации основан на работах французских лингвистов, в частности, Ж.-К. Анскомбра и О. Дюкро. В рамках этого подхода аргументативная сила высказывания рассматривается как неотъемлемый компонент его смысла: значение высказывания определяется влиянием на собеседника, ориентированием его внимания в определённом направлении (что близко упомянутому определению аргументации у Р. Амосси). Следует подчеркнуть приоритет аргументационного значения над информационным: «во многих случаях информативность высказывания является в действительности вторичной по отношению к аргументативному значению: часто целью локутора является не сообщение о каком-либо факте, а оказание давления на адресата, на его представления, мнения, действия» [Иоанесян, 2023, с. 186].

Е. Р. Иоанесян также акцентирует исключительную роль лингвистики среди иных дисциплин в анализе аргументации: аргументативное функционирование высказывания зависит от его лингвистической структуры, а не только заложенного информационного содержания; как следствие, при анализе аргументативного дискурса необходимо фокусироваться на лингвистическом оформлении высказываний, а не экстралингвистической риторике или логической организации рассуждений. Иными словами, «аргументация целиком и полностью принадлежит уровню речи» [Иоанесян, 2023, с. 165]. В соответствии с акцентированно лингвистическим ракурсом данного подхода, работы Е. Р. Иоанесян обращены к аргументативному функционированию конкретных

языковых единиц, таких как предикаты или союзы разных типов, кванторные единицы, маркеры статуса предложения.

Таким образом, отечественные подходы к исследованию аргументации характеризуются гармоничным сочетанием самобытной оригинальности и одновременным вниманием к международному опыту. Их специализация в рамках фундаментальной лингвистики ограничивает возможность опоры на рассмотренные работы в прикладном исследовании по компьютерной обработке аргументации – области, практически не представленной в отечественной науке, но активно развивающейся за рубежом.

1.3 Формальное моделирование аргументационных структур

Представление аргументации как сложной структуры – неразделимого сочетания утверждений, выражаемых языковыми средствами, и абстрактных схем убеждения для перехода между такими лингвистически оформленными пропозициями на уровне целостного связного текста – подразумевает анализ как отдельных тезисов, так и отношений между ними. Тезисы представляют собой утверждения в составе аргументации, собственное значение которых дополняется значением на основе контекста, образуемого системой аргументационных отношений (связей) между утверждениями. Тезисы, объединяемые аргументационными связями в общую структуру рассуждения, используются говорящим для представления определённой точки зрения. Соответственно, реконструкция аргументационных структур заключается в их формальном описании посредством абстрактной модели, причём построение такой модели делает возможной компьютерную обработку аргументации [Peldzus, Stede, 2003].

В частности, формальное представление аргументационных структур для их обработки компьютером может осуществляться в рамках решения двух методологически различных задач:

- 1) уточнение метаязыка описания аргументации с целью эффективной обработки рассуждений и доказательств в системах искусственного интеллекта, в том числе многоагентных;
- 2) компьютерный анализ аргументации в тексте на естественном языке с целью выявления целостной системы тезисов и организующих их связей, обеспечивающих эффект убеждения.

В контексте постановки и решения первой задачи особенно влиятельной является модель С. Тулмина, основанная на представлении отдельных аргументов посредством шести связанных компонентов [Toulmin, 2003]. Ещё одно направление сформировалось в развитии области прикладных исследований по компьютерному анализу тональности (сентимент-анализу) вследствие расширения исходной задачи (определения эмоциональной окраски текстов) посредством постановки акцента на раскрытии авторских обоснований выявляемых оценок.

1.3.1 Аргументационная модель С. Тулмина

1.3.1.1 Анализ аргументации как процесса обоснования мнения

Античным работам по логике и аргументации свойственно неявное предположение о структурном параллелизме убеждающей речи и процесса рассуждения (а этот процесс представляется как последовательное выведение новых заключений на основе некоторого числа исходных посылок): так, Аристотель говорит о предметной общности риторики и диалектики, основывая единство этих дисциплин на предполагаемой непосредственной связи между принципами логики и человеческого мышления, которые как универсальное средство постижения действительности организуют любое обсуждение или спор [Аристотель, 1978, с. 15–19]. Из возражения такому предположению исходит С. Тулмин, формулируя в своих аргументационных исследованиях тезис о ретроспективном характере аргументации: доводы в поддержку доказываемого

утверждения не задают предварительно его содержания, но определяются впоследствии на основе высказываемого мнения для его верификации [Toulmin, 2003, с. 6]. Соответственно, согласно модели С. Тулмина, такие доводы не являются функционально самостоятельными: они осознаются говорящим только в составе целостного аргумента на основе его понимания.

Таким образом, центральным для аргументационной модели С. Тулмина выступает понятие обоснования (*justification*, в противовес процессу вывода, *inference*), что наглядно отражено в представленной структуре аргумента. Так, убеждающее целое образуется шестью компонентами, разделяемыми по двум триадам согласно обязательности или факультативности их реализации при построении аргумента. Компоненты первой группы – утверждение, довод и основание; вторая группа содержит поддержку, опровержение и определитель.

1.3.1.2 Шесть составляющих аргумента по С. Тулмину

С. Тулмин определяет обязательные и факультативные компоненты в составе любого аргумента следующим образом [Toulmin, 2003, с. 87–131]:

- 1) утверждение (*claim*) понимается как непосредственно доказываемое содержание аргумента, а соответственно, выступает его центральной частью;
- 2) довод (*datum*) выявляется через анализ окружения утверждения как приводимое для его обоснования доказательство;
- 3) основание (*warrant*) обеспечивает переход от довода к утверждению за счёт объединения в своём смысловом поле элементов как первого, так и второго, причём, в отличие от них, может выражаться неявно (подразумеваться без словесной формулировки);
- 4) поддержка (*backing*) проявляется во взаимосвязи с основанием как его уточняющее раскрытие, предназначенное для усиления убедительности всего аргумента (в особенности если переход от довода к утверждению не выражается явно, например, во избежание избыточности);

5) опровержение (*rebuttal*) выражает признаваемые ограничения аргумента (к примеру, возможность альтернативного объяснения рассматриваемой ситуации, несовместимого с центральным утверждением аргумента);

6) определитель (*qualifier*) определяется как модальный показатель силы основания (а именно прочности логической связи между утверждением и доводом).

Взаимоотношение шести компонентов может быть проиллюстрировано следующим примером [Toulmin, 2003, с. 117], где выстраивается аргументационная модель для выражения вида *Поскольку Анна – сестра Джека, а все его сёстры рыжие, то у неё, скорее всего, рыжие волосы, если она их не перекрасила*:

- 1) утверждение: *у Анны рыжие волосы*;
- 2) довод: *Анна – сестра Джека*;
- 3) основание: *любая сестра Джека наделена рыжими волосами*;
- 4) поддержка: *ранее наблюдалось, что все сёстры Джека рыжие*;
- 5) опровержение: *если Анна не перекрасила волосы*;
- 6) определитель: *скорее всего*.

Приведённый пример показывает возможность реализации всех шести компонентов в составе конкретного аргумента. Между тем обязательными для функционирования аргумента является выражение лишь утверждения, довода, а также основания – которое, как подчёркивает С. Тулмин, может задаваться неявно: подразумеваться без языкового оформления. Так, конструкция вида *У Анны рыжие волосы, потому что она сестра Джека* представляет собой полноценный аргумент: говорящий обосновывает утверждение о цвете волос Анны через указание её родственной связи с Джеком, причём полагает связь двух этих характеристик очевидной.

Как следствие, возможность неявного выражения в аргументе одного из трёх обязательных компонентов означает, что в модели С. Тулмина аргументом является любое сочетание утверждения и довода в его поддержку: приведение

доказательства для любой заявляемой пропозиции подразумевает построение аргумента. При этом выделение четырёх иных компонентов позволяет более точно формализовывать содержание отдельных аргументов. Однако единицей анализа при подходе С. Тулмина служит именно изолированный аргумент: в его теории не задаются специальные структуры для моделирования связей между несколькими самостоятельными аргументами.

1.3.1.3 Критика аргументационной модели С. Тулмина

Шестикомпонентная модель аргумента, представленная С. Тулминым в 1958 г., оказала значительное влияние на привлечение исследовательского внимания к формальному анализу аргументации. Между тем разработка этой концепции проводилась преимущественно с теоретических позиций, тогда как прикладное применение модели С. Тулмина для моделирования аргументации (в том числе её компьютерной обработки) продемонстрировало ряд сложностей в представлении аргументационных структур, встречающихся в текстах на естественных языках. Так, А. Пелджус и М. Сид характеризуют три функциональных неясности в компонентах С. Тулмина.

1. Во-первых, разграничение основания и поддержки оказывается на практике избыточным. Поддержка в модели С. Тулмина дополняет переход от довода к утверждению – указывает на корректность такого перехода внутри конкретного аргумента. Однако в естественных текстах релевантность приводимого доказательства может раскрываться рекурсивно: как через введение новой поддержки (либо параллельной исходной, либо, наоборот, её уточняющей), так и через ограничение области действия изначальной (посредством реализации внутреннего опровержения или определителя, который функционально соответствует аналогичному компоненту на уровне полного аргумента). Тогда при обработке подобных сложных структур удобнее представлять поддержку как элемент иного полнозначного аргумента (где она служит доводом, а основание исходного аргумента – утверждением). Впрочем, доступность этого

представления для расширенной поддержки означает и его применимость для стандартного компактного случая, выделенного С. Тулминым: поскольку поддержка обосновывает применимость основания внутри аргумента, их отношение функционально близко связи между утверждением и доводом. Некоторое различие в структурном контексте (основание и поддержка полагаются вспомогательными к утверждению и доводу, которые передают непосредственное содержание аргумента) уравнивается предпочтением единообразного представления поддержки при её как компактной, так и рекурсивной реализации.

2. Во-вторых, функциональные компоненты С. Тулмина не поддерживают разграничение двух видов прагматического воздействия на аргумент: со стороны говорящего либо его собеседника. В рассматриваемой модели не учитывается возможная активность иных участников коммуникации, что значительно затрудняет моделирование диалогического дискурса. В частности, опровержение аргумента может соотноситься как с собеседником (реальным либо вымышленным) и его внешней атакой на доказательство (автор которого вследствие признаёт ограниченность аргумента), так и с самим говорящим в рамках сложной убеждающей структуры: аргументация в естественных текстах может строиться через приведение возможных возражений с их последующим опровержением от самого автора. В исходной же модели С. Тулмина не предусмотрены средства для уточнения специфики опровержения.

3. В-третьих, при анализе аргументов как самостоятельных конструкций из некоего отдельного утверждения и его обоснования (через иные компоненты, ограниченные доказательством) из внимания ускользает более сложная макроструктура аргументации – аргументационное окружение, а именно связи между аргументами в процессе коммуникации. Две ранее обозначенные сложности являются следствиями ограничения фокуса внимания только изолированным аргументом как единицей анализа: рекурсия в процессе доказательства, как и диалоговая специфика отдельных компонентов, проявляются на уровне макроструктуры аргументации, которая в свою очередь выражается лишь на уровне связного целостного текста [Peldzus, Stede, 2003].

Как следствие, анализ А. Пелджуса и М. Стида указывает на ограниченную применимость модели С. Тулмина в формальном представлении аргументации на уровне полных текстов. Детальное описание отдельных аргументов через разграничение шести компонентов выступает особенно информативным при подготовке ёмких доказательств, локализации логических ошибок в процессе убеждения. Однако анализ аргументации на основе положений дискурсивного подхода обуславливает предпочтительность других формальных моделей, которые поддерживают представление аргументов с учётом их взаимных связей внутри более сложной общей структуры.

1.3.2 Теория риторической структуры в анализе аргументации

1.3.2.1 Ключевые положения теории риторической структуры

Формальное представление структур аргументации на уровне полных текстов может строиться как в соответствии со специализированной моделью, разработанной с учётом их языковой специфики, так и на основе имеющейся лингвистической теории для описания текстовой связности (при условии её применимости для аргументационного анализа). Такую потенциально сходную модель, предназначенную для формального представления функциональных отношений между сегментами текста в прагматическом аспекте, представляет теория риторической структуры (TPC) У. Манна и С. Томпсон [Mann, Thompson, 1988]. Эта теория применяется для анализа структурной организации текстов на уровне отношений их дискурсивных единиц. Несмотря на то, что хотя ТРС активно применяется на практике для анализа текстов на различных языках [Taboada, Mann, 2006a, 2006b], опыт использования этого метода на материале русского языка сравнительно невелик.

В частности, М. К. Тимофеева, исследуя стили рассуждения в русскоязычных поэтических текстах в терминах ТРС, указывает на возможность определения данной теории как «метода формального представления структуры

рассуждения, реализованного в тексте» [Тимофеева, 2020, с. 116–120]. При этом структуры рассуждения реконструируются через анализ функционального назначения отдельных сегментов текста: для каждого такого целостного сегмента может быть определена цель его построения автором, а также его намеренный вклад как в содержание, так и в общую структурную связность текста (на уровне иерархических отношений между дискурсивными единицами, элементарными и составными) [Там же, с. 109–137]. Таким образом, авторская интенция проявляется именно в совместной организации сегментов текста, их объединении внутри целостной структуры посредством риторических отношений из заданного однородного набора: одни и те же отношения могут применяться для связи и элементарных мыслей (в их языковом оформлении), и их сложными комплексами.

Точный набор риторических отношений (как ключевой компонент ТРС) определяется конкретной реализацией исходной концепции У. Манна и С. Томпсон, а именно её адаптацией под решаемую задачу. При исследованиях на материале русскоязычных текстов ТРС применяется в различных направлениях: анализе когнитивных структур, например, для выявления речевых аномалий [Кибрик, Подлесская, 2009], изучении стилей рассуждения в художественных текстах [Timofeeva, 2021], а также компьютерной обработке языковых данных с целью уточнения моделируемого содержания посредством распознавания маркеров дискурса [Kononeko, Sidorova, Akhmadeeva, 2020], в том числе в текстах научного жанра [Бакиева, Батура, 2017; Chistova et al., 2019]. В соответствии с задачей для когнитивного и стилистического анализа выступает значимым детализированное представление риторических отношений (с расширением их списка в зависимости от жанра моделируемых текстов), тогда как для компьютерной обработки свойственно ограничение разнообразия в выделяемых структурах (в целях надёжного разграничения сходных связей).

Отмеченная вариативность ТРС как инструмента была запланирована У. Манном и С. Томпсон изначально. Естественным следствием служит анализ её применимости как для формального представления аргументационных связей в терминах риторических отношений, так и для упрощения аргументационной

аннотации текстов на основе сведений об их риторической структуре (которая доступна для построения компьютером).

Оценивая возможную применимость ТРС к анализу аргументации, А. Пелджус и М. Сид подчёркивают удобное разграничение в исходной концепции У. Манна и С. Томпсон двух типов риторических отношений: предметных (отражающих связи между фактами реального мира, к примеру, отношение причины и следствия) и презентационных (которые автор текста применяет для активного воздействия на мнение адресата либо его отношение к некоему вопросу). Примерами отношений из второй группы можно указать приведение доказательства (*evidence*) либо раскрытие мотивации (*motivation*), причём именно презентационные отношения представляют основной интерес для моделирования аргументационных связей [Peldzus, Stede, 2003]. А. Пелджус и М. Сид отмечают также предпочтительность ТРС над другими моделями текстовой связности (такими как Segmented Discourse Representation Theory (SDRT), Discourse Lexicalized Tree-Adjoining Grammar и Dynamic Discourse Model) ввиду её направленности на прагматическое описание (обращения к интенциям автора без строгой зависимости от синтаксических структур и семантики, на уровне которых и осуществляется анализ связности в иных теориях – что менее актуально для моделирования аргументации как прагматического явления).

В частности, Э. Поттер дополнительно сопоставляет ТРС и SDRT в рамках аннотации текстов для компьютерного анализа аргументации [Potter, 2022]. Согласно его наблюдениям, характерное для SDRT представление дискурса через теоретико-истинностную семантику не позволяет достаточно полно описать интенции автора, лежащие в основе организации текста (которые являются ключевыми как для ТРС, так и для дискурсивного подхода к аргументационному анализу). По той же причине менее предпочтительным, чем ТРС, для анализа интенционального аспекта в структуре связного текста является и подход на основе анализа эксплицитно выраженных дискурсивных маркеров преимущественно в лексическом аспекте (такой подход реализован при моделировании текстов в Penn Discourse Treebank) [Demberg, Schoman, Asr, 2019].

1.3.2.2 Соотношение аргументационной и риторической структуры

В свою очередь, аннотирование текстов на основе ТРС с целью перехода к аргументационному анализу предполагает ответ на вопрос о соотношении риторических и аргументационных структур (как в целостном представлении, так и на уровне отдельных сегментов текста и их отношений). Сформулировав отмеченную проблему в работе [Kononenko, Sidorova, Akhmadeeva, 2020], группа исследователей разграничивает эти структуры в системе трёх уровней дискурса:

- 1) жанрового (наивысший уровень, или суперструктура: соответствует композиционной и семантической организации полного текста);
- 2) риторического (отражает объединение в целостный текст линейных формальных сегментов с уточнением их функциональных связей);
- 3) аргументационного (представляет текст как реализацию аргументации, или системы обосновываемых тезисов в их поддержку или опровержение, содержание которых раскрывается взаимными отношениями).

Таким образом, риторическая структура текста анализируется в аспекте его линейной связности (на уровне цепочек смежных сегментов) и покрывает текст полностью, тогда как аргументационная структура выделяется поверх некоторого подмножества сегментов (чьи отношения характеризуются более узкой спецификой, а именно ролью воздействия на мнение адресата). Кроме того, как демонстрирует указанный коллектив авторов, процессы построения риторической аннотации (согласно принципам ТРС) и аргументационной (под конкретную задачу компьютерного анализа научно-популярного дискурса) различаются во многих значимых аспектах.

Во-первых, хотя оба типа аннотации основываются на сегментировании текста на элементарные дискурсивные единицы (ЭДЕ: предложения, клаузы, минимальные сегменты с пропозициональным содержанием), только ТРС предполагает объединение элементарных единиц в их более сложные комплексы (которые связываются теми же отношениями, что и элементарные единицы). Этот процесс итеративного вложения меньших единиц в большие обеспечивает

построение риторической структуры в виде дерева (в математическом понимании) – ациклического связного графа, где между любыми двумя вершинами (текстовыми сегментами) доступен только один путь. Наоборот, при аргументационной аннотации все единицы выделяются на одном уровне, без объединения меньших сегментов текста в более крупные (поскольку их связи характеризуют аргументацию в содержательном аспекте, а не линейную организацию текста, как в ТРС). Кроме того, аргументационный граф не обязательно является деревом: одно утверждение может служить посылкой для двух разных заключений, тогда как заключение способно поддерживаться несколькими посылками (что обуславливает возможное существование множественных путей между двумя утверждениями в аргументации).

Во-вторых, между риторическими и аргументационными сегментами в общем случае нет однозначного соответствия. Дискурсивные единицы без аргументационного содержания могут игнорироваться при разметке аргументации. Цельный компонент аргумента (к примеру, отдельный тезис) может на риторическом уровне разделяться на несколько единиц (если риторическое отношение между ними носит предметный характер, не отражает убеждающего воздействия). Элементарная же единица риторического уровня, наоборот, может разделяться на разные компоненты аргумента (например, в аргументе от экспертного мнения одной посылкой выступает указание авторитетного источника, тогда как другой – выражение им некоего мнения, причём эти посылки могут принадлежать одной клаузе).

В-третьих, в двух видах аннотации по-разному представляются маркеры дискурсивных отношений. При моделировании риторических структур они включаются в дискурсивные единицы, однако могут выводиться из текстовых сегментов при разметке аргументации (если не относятся к пропозициональному содержанию посылки или заключения). Поэтому даже формально совпадающие дискурсивные единицы в риторической и аргументационной структурах могут различаться границами в тексте.

В-четвёртых, отдельные дискурсивные маркеры могут различаться по релевантности для риторических и аргументационных структур. Так, частицы вида «именно», «только», «по крайней мере» не указывают на риторические отношения (в исходной версии ТРС), однако передают потенциальное аргументационное значение.

Данный список различий между двумя видами прагматических структур может быть дополнен замечаниями от А. Пелджуса и М. Стида:

1. Хотя риторическая структура покрывает текст целиком (без допустимых пропусков), её отношения могут связывать несмежные сегменты текста только опосредованно (в составе более сложных сегментов, смежных между собой). В аргументационной же структуре, поддерживающей пропуски частей текста, допустимы непосредственные связи между линейно отдалёнными сегментами (к примеру, если тезис предваряется двумя разными доводами).

2. Риторическая структура покрывает оформленное содержание текста, а следовательно, составляется явно выраженными сегментами. Структура же аргументации может включать подразумеваемые пропозиции, не сформулированные напрямую (например, в случае энтимем: раскрытие перехода в их основании, от неполной посылки к заключению, требует восстановления невысказанного компонента) [Peldzus, Stede, 2003].

Как показывает сопоставление двух типов структур, при их сходстве в прагматическом рассмотрении отмечаются и значимые различия в аспекте их функционирования. Такие различия проявляются при практическом описании риторической и аргументационной структуры для одного текста.

1.3.2.3 Применение ТРС для моделирования аргументационных связей

Исследуя возможность компьютерного построения аргументационных структур на основе риторической разметки, Э. Поттер выявил две сложности теоретического плана при практическом решении данной задачи [Potter, 2022].

Во-первых, актуальные методы компьютерной обработки риторических структур основываются на некотором упрощении исходных принципов ТРС. Так, автоматизация разметки может достигаться через постановку акцента на лексико-синтаксических характеристиках сегментируемого текста (в частности, через обработку дискурсивных маркеров). Сочетание расширенного словаря маркеров и методов синтаксического анализа позволяет выстраивать риторические структуры, достаточно близкие к экспертной аннотации: риторически значимые сведения теряются лишь в ограниченной мере. Тем не менее оптимизация разметки в целях автоматизации приводит к выраженной потере содержания, полезного именно для аргументационного анализа (в первую очередь, интенциональных характеристик). В этой связи для автоматической обработки аргументации более эффективна экспертная разметка риторических структур, которая, однако, отличается высокой трудозатратностью (ввиду чего аналогичное усилие экспертов может быть обращено к аргументационной аннотации напрямую).

Во-вторых, различия в функциональном представлении риторических и аргументационных отношений влекут расхождения в их направленности внутри общей структуры текста. Так, большинство риторических связей являются ассиметричными: среди связываемых ими сегментов текста один является главным (ядром), другой – подчинённым (сателлитом). Однако в аргументационной структуре взаимные роли сегментов могут как сохраняться, так и инвертироваться (в зависимости от направления развития рассуждения при убеждающем воздействии). Так, отношение *Solutionhood* в ТРС связывает проблему и предлагаемое решение, где первая является сателлитом, а второе – ядром (по определению У. Манна и С. Томпсон). При этом в аргументации возможен как вывод решения из проблемы (если ключевое внимание направлено именно к решению как предлагаемому действию), так и проблемы из решения (если автор желает подчеркнуть сложность этой проблемы, для чего указывает критические недостатки в доступных решениях, а затем обращается к ирреальной модальности с предписанием изначально не допускать возникновения такой

сложной ситуации). Иным примером инвертирования ролей может выступить риторическое отношение *Elaboration*: его ядром является уточняемый объект, тогда как сателлитом – уточняющее свойство. Хотя свойство и второстепенно по отношению к объекту, в процессе аргументации оно может акцентироваться в большей степени, чем объект. Подобные случаи неясного направления связи, как демонстрирует Э. Поттер, требуют для корректного представления анализа структурного контекста (и интерпретации интенций автора), что затруднительно для компьютерных методов.

С иных позиций обращаются к задаче аргументационного анализа через риторические структуры А. Пелджус и М. Сид [Peldzus, Stede, 2003]: они рассматривают не автоматизированное построение аннотаций одного типа на основе другого, а непосредственное формальное представление аргументационных структур в терминах ТРС и приходят к выводу, что ввиду исходной концептуальной гибкости и методологической близости к задаче обработки аргументации ТРС подходит для той лучше аналогичных лингвистических инструментов. Однако А. Пелджус и М. Сид отмечают две значимые проблемы в таком применении ТРС, требующие адаптации её базовых положений.

1. ТРС ориентирована на линейное представление сегментов текста, тогда как аргументационным структурам линейность вовсе не свойственна. Она нарушается уже при обосновании одного заключения несколькими параллельными доводами (которые не объединяются общей связью, а вводят различное содержание для поддержки тезиса в разных аспектах). Кроме того, структура аргументации заметно усложняется с глубиной: доводы в поддержку некоего тезиса могут дополнительно раскрываться через их обоснование, возражение же способно использоваться для защиты доказываемой идеи через приведение контр-опровержения. При этом одна посылка может применяться для доказательства нескольких промежуточных заключений. Особую сложность представляет линейная отдалённость связанных сегментов (в связи с возможностью связи таких единиц в ТРС лишь опосредованно, в составе более

сложных сегментов), которая может достигать размеров всего текста (например, если он начинается доказываемым тезисом, а завершается прямым расширением центральной идеи).

2. Отдельные виды значимых для аргументации структур не могут быть надёжно представлены в терминах ТРС: к примеру, контр-возражения получится формально задать только как опровержение возражения, без обозначения их возможной связи с исходным тезисом (два разных вида контр-возражений – поддерживающие исходный тезис, как и частично возражающие ему для вывода компромиссного варианта – будут выглядеть одинаково при описании в терминах ТРС).

Таким образом, ТРС поддерживает моделирование аргументационной связности текстов, причём подходит для этой задачи лучше иных аналогичных инструментов ввиду ориентированности на прагматическое функциональное описание. Однако для полной передачи специфических черт аргументации как языкового, в частности дискурсивного явления предпочтительно применение специализированной модели, учитывающей методологические особенности автоматической обработки текстов.

1.3.3 Моделирование текстов через Argument Interchange Format

В работе [Rahwan, Reed, 2009], обращённой к формальному описанию аргументационных структур, И. Рахван и К. Рид показывают недостаток специализированных моделей, разработанных под конкретную программную реализацию: такие средства представления (Argument Markup Language [Reed, Rowe, 2004], ClaiMaker [Shum et al., 2007]) задают строгие теоретические ограничения на спецификацию аргументов и их отношений (в соответствии с концепцией, выбранной создателями), отчего выстраиваемые с их помощью структуры поддерживаются только исходным приложением (и недоступны для обработки иными). Кроме того, приведённые модели ориентированы на удобство визуализации аргументов (наглядность их представления), что подразумевает как

появление искусственных ограничений на построение запутанных сложных структур, так и ограниченную поддержку методов формальной логической обработки для аргументационных моделей.

В связи с отмеченными недостатками иных языков описания И. Рахван и К. Рид представляют собственный стандарт для формального представления аргументации: Argument Interchange Format (AIF). AIF как модель отличается высокой степенью абстракции и ориентирован на два следующих ключевых принципа, обозначенных при разработке:

- 1) применимость в многоагентных системах (открытых или закрытых) с поддержкой процессов рассуждения и взаимодействия, основанных на аргументации посредством общего формализма;

- 2) совместимость между отдельными инструментами для обработки и визуализации аргументов (иллюстрацией к этому принципу служит база данных AIFdb [Lawrence, Reed, 2014], которая поддерживает загрузку пользователями любых аргументационных корпусов, размеченных в соответствии с форматом AIF вне зависимости от использованного инструмента).

AIF основывается на формальном представлении аргументации в виде графов, или математических структур с двумя типами элементов: вершинами, которые соответствуют отдельным объектам, и рёбрами – парными связями между вершинами (отражение объектных отношений). Отличительной чертой AIF выступает дополнительное разграничение двух типов вершин:

- 1) информационные вершины отражают пропозициональное содержание для отдельного утверждения (текстового сегмента либо неявного высказывания, соответствующего цельному компоненту некоторого аргумента);

- 2) вершины-схемы указывают на реализацию аргументационных схем, или типовых моделей рассуждения, посредством которых и осуществляется переход между компонентами аргумента (например, вывод заключения на основе посылок через апелляцию к экспертному мнению).

Выделение вершин-схем как самостоятельных структурных элементов обуславливается их значимостью в аргументации (убеждающее воздействие

достигается не только собственно языковым содержанием утверждений, но и организующими их связями в процессе коммуникации). Такое введение отдельного структурного типа для вершин-схем (отражающих связи между посылками и заключениями) соответствует выделению в ТРС реляционных пропозиций, соответствующих связям между сегментами (У. Манн и С. Томпсон отмечают важность реляционных пропозиций в структуре рассуждения, а развитием их концепции служит модель Э. Поттера: связи между сегментами описываются посредством предикатов высших порядков, аргументами которых являются сегменты текста, связываемые тем или иным отношением).

В свою очередь, для информационных вершин задаётся ограничение на связность: одно ребро в аргументационном графе не может соединять две вершины этого типа (любой переход от одного аргументативного компонента к другому рассматривается как реализация некой схемы). При этом полностью допускается соединение вершин разного типа либо же двух вершин-схем (что позволяет моделировать особый тип атак на аргументы: не через опровержение посылок, а посредством возражения об адекватности вывода заключения из потенциально корректных посылок).

Соответственно, в AIF различается два типа вершин-схем: выражающие отношение поддержки (одного тезиса неким доводом) либо возражения (атаки на аргумент). При этом высокая степень абстрактности для этой модели (её ориентированность на широкую применимость для различных задач анализа аргументации с совместимостью структур) проявляется в поддержке любой классификации схем рассуждения. Одной такой классификацией, отмечаемой авторами AIF, служит компендиум аргументационных схем Д. Уолтона.

1.3.4 Классификация аргументационных схем Д. Уолтоном

Компендиум Д. Уолтона [Walton, Reed, Macagno, 2008] представляет собой систему типовых моделей рассуждения, предназначенную для семантического анализа связей между утверждениями в аргументации. Данная классификация не

является специализированной, поскольку содержит в целях прикладной применимости схемы убеждения разных типов: свойственные как повседневному общению, и научному языку, так и, например, юридическому дискурсу.

При этом число моделей обоснования выбрано, согласно Д. Уолтону, в некотором роде произвольно: учтены только основные схемы аргументации с возможностью значительного расширения классификации более редкими при её адаптации для конкретной задачи. Тем не менее, каждая приведённая схема раскрывается критическими вопросами к ней (которые позволяют проверить уместность её употребления в конкретном случае), а ряд многоаспектных моделей уточняется через выделение их подтипов, или противопоставляемых типовых вариантов их реализации (которые могут дробиться далее).

Так, представление схем аргументации в классификации Уолтона можно проиллюстрировать на примере доказательства от практического рассуждения (*Practical Reasoning*) [Walton, Reed, Macagno, 2008, с. 323-325]. Для данной схемы выделяются шесть подтипов: практический вывод (*Practical Inference*), схема обязательного условия (*Necessary Condition Schema*), схема достаточного условия (*Sufficient Condition Schema*), практическое рассуждение через ценность (*Value-Based Practical Reasoning*), аргумент от цели (*Argument from Goal*) и аргумент от итогового результата и средств (*Argument from Ends and Means*). Каждый из этих подтипов определяется посредством семантической характеристики его посылок и заключений, что можно продемонстрировать через сопоставление первого и третьего из выделенных шести:

1) практический вывод поддерживает две посылки: обозначение цели (вида «у меня есть цель G ») и обозначение действия для достижения этой цели («выполнение действия A позволяет достичь цели G ») с выводом из них заключения о предпочтительности указанного действия в свете обозначенной цели («с практической точки зрения мне следует выполнить действие A »);

2) более специализированная схема достаточного условия допускает пять посылок (посылку цели «моей целью является достижение G », посылку альтернатив «выполнение хотя бы одного условия из $[A_1, \dots, A_n]$ необходимо для

достижения G», выбора «достаточное условие A_i наиболее предпочтительно для достижения G», практичности «перед выполнением A_i нет непреодолимых препятствий» и посылку побочных эффектов «достижение G более предпочтительно, чем невыполнение A_i » с построением на них заключения «мне следует выполнить A_i ».

Реализация какой-либо схемы из классификации Д. Уолтона не требует эксплицитного выражения в тексте всех поддерживаемых ею посылок, поскольку некоторые из них могут только подразумеваться (как для достижения большей убедительности, так и для обеспечения краткости). С целью же иллюстрации критических вопросов в представлении схем можно указать таковые к первому подтипу практического рассуждения:

- 1) Имеются ли у меня иные цели, несовместимые с целью G?
- 2) Какие действия, отличные от A, также позволяют достичь цели G?
- 3) Какое действие из A и его альтернатив наиболее эффективно?
- 4) Какие имеются основания полагать практическую выполнимость A?
- 5) Каковы иные последствия от выполнения A?

Таким образом, классификация аргументационных схем Д. Уолтона позволяет подробно представлять содержательный аспект отношений между утверждениями в аргументационной структуре текста. Она отличается значительной гибкостью и может быть применена для формального описания текстов различных жанров (как монологических, так и диалогических). Как следствие, моделирование аргументационной семантики по классификации Уолтона обеспечивает возможность сопоставительного аргументационного анализа текстов из разных отраслей (в рамках самостоятельного исследования либо через соотнесение результатов). Наконец, подробное описание Д. Уолтоном принципов выявления аргументационных схем способствует единообразию в моделировании аргументационных структур разными исследователями.

Выводы

В рамках обзора теоретических оснований для исследований аргументации представлены и сопоставлены актуальные лингвистические подходы к её анализу, в первую очередь свойственные континентальной европейской традиции. Учтены также аргументационные исследования в отечественной лингвистике, отличающиеся гармоничным сочетанием самобытной оригинальности и внимания к международному опыту. Исследования с позиций англо-американской традиции, согласно которой аргументация является строго внеязыковым явлением, не рассматриваются ввиду их преимущественно формально-логической направленности. Указаны расхождения в понимании аргументации как языкового явления (компонента лексической семантики или источника текстовой связности), вызванные междисциплинарным влиянием теоретического контекста (учитывает ли исследователь модели формальной семантики или обращается к аргументации под частичным влиянием риторики либо психологии). Показано, что аргументационные исследования проводятся с позиций преимущественно двух парадигм современной лингвистики: структуралистской и коммуникативной. Отдельный акцент посвящён детализированию значимости аргументационного анализа в его связях с фундаментальной лингвистикой.

Представленное в настоящей работе исследование аргументационных стратегий в академической коммуникации соответствует трём теоретическим положениям из дискурсивного подхода (в формулировке Р. Амосси): во-первых, о проявлении аргументационных значений на уровне целостных текстов, во-вторых, о возможности реконструкции формального строения структур убеждения на основе выражающих их языковых средств, а в-третьих, о равной значимости лингвистических средств и абстрактных схем убеждения (отражающих переходы от одних утверждений к другим) в аргументационном воздействии.

После определения ключевых теоретических принципов для проводимого исследования решается задача выбора модели для формального представления аргументационных структур на уровне полных текстов. Рассматривается

концепция С. Тулмина, значимая ввиду исторического влияния на аргументационные исследования в рамках формального подхода. Оценивается возможность адаптации существующей лингвистической модели (теории риторических структур как инструмента для реконструкции текстовой связности в функциональном прагматическом аспекте) к моделированию аргументации.

По результатам сопоставительного анализа моделей принято решение использовать для формального представления аргументации формат AIF, ориентированный на компьютерную обработку выявляемых структур, с его семантическим расширением через компендиум аргументационных схем Д. Уолтона. Эта классификация позволяет моделировать отношения между компонентами аргументов в содержательном аспекте и поддерживает анализ текстов разных жанров, а подробное формулирование Д. Уолтоном принципов выявления отдельных схем способствует единообразию аргументационной разметки у разных аннотаторов.

Глава 2. Методы анализа аргументационных приёмов и стратегий

Во второй главе представленной работы описываются методы компьютерной обработки аргументационных структур в их формальном представлении. Методы сгруппированы по функциональному критерию (согласно решаемым ими задачам аргументационного анализа). В свою очередь, представление методов включает как обзор существующих подходов к решению определённой задачи, так и уточнение алгоритмов, реализованных непосредственно в рамках данного исследования.

В частности, в разделе 2.1 обозначаются принципы разметки аргументации, которые дополняют базовые правила стандарта AIF и сформулированы с учётом жанровой специфики доказательств в научных статьях. Эти принципы применены в разработке методологии разметки и подготовке аргументационного корпуса, обрабатываемого в рамках данного исследования. В разделе 2.2 рассматриваются методы частичной автоматизации аргументационного аннотирования, используемые во вспомогательной роли для повышения эффективности ручной разметки. Особенно подробно представлены методы для компьютерного распознавания утверждений с аргументацией, реализованные и экспериментально проверенные в данной работе. В разделе 2.3 описываются коэффициенты согласия аннотаторов, посредством которых оценена надёжность разметки в подготовленном аргументационном корпусе. Представлен как коэффициент, специально разработанный в ходе исследования с учётом многоуровневой специфики аргументационной разметки, так и два стандартно используемых с учётом случайного согласия между аннотаторами. Раздел 2.4 посвящён методам формального выявления приёмов аргументации – структур из одного и более аргументов определённых типов. В разделе 2.5 анализируются методы для оценки убедительности аргументации в разных аспектах ведения доказательств (через построение логических связей на основе доступных фактов, через воздействие на эмоции либо апелляции к авторитету). Наконец, в разделе 2.6 представлены методы для выявления текстов со сходной организацией доказательств

посредством кластеризации их аргументационных структур с учётом реализуемых приёмов аргументации и выраженности отдельных аспектов убеждения (из обозначенных трёх).

2.1 Экспертный подход к разметке аргументации

Аннотирование аргументации в корпусе научных статей, применяемом в представленном исследовании, проводится в соответствии со стандартом AIF на основе классификации моделей рассуждения Д. Уолтона. Поскольку размеченные тексты предназначены для компьютерной обработки, процесс аннотирования регулируется набором дополнительных принципов. Четыре ключевых принципа из этого набора перечислены ниже.

1) Для каждого текста определяется строго один главный тезис, выражающий основную идею научной статьи. Выявление этого тезиса основывается на авторском структурировании разделов текста. Как правило, главный тезис формулируется в заключении статьи (как основной вывод исследования, от которого могут зависеть второстепенные выводы). В некоторых случаях он приводится при постановке цели работы или описании проверяемой гипотезы (тогда возможно повторное указание главного тезиса в завершении статьи, но в сокращённой версии, через отсылку к уже представленной полной формулировке). Например, если целью исследования, описываемого в научной статье, является создание специализированного корпуса, то его отличительные черты могут обозначаться уже при постановке цели (а при подведении выводов кратко отмечается факт успешной разметки необходимого количества текстов). Если в статье демонстрируется корректность заявленной гипотезы, то итоговый вывод может заключаться в указании на факт её подтверждения. В редких случаях главный тезис может приводиться уже в названии публикации, если он не представляется более полно в тексте работы (к примеру, статья может быть посвящена анализу определённого метода или группы методов в различных

аспектах, однако не содержать итогового обобщающего комментария обо всём методе либо их группе в целом).

Если в тексте статьи обнаруживаются несколько предполагаемых главных тезисов, то, как правило, один из них является наиболее общим, объединяющим остальные (которые, соответственно, ведут цепочки рассуждения к этому тезису, выступают вспомогательными по отношению к нему и не являются главными). Этот тезис следует обозначить как главный. Научных статей, содержащих несколько равноправных главных тезисов, в корпусе не выявлено.

2) Каждый текст представляется в виде аргументационного графа, который является ориентированным, связным (слабо-связным), ациклическим и корневым.

Ориентированность графа означает, что любому ребру между его вершинами присваивается направление (что соответствует переходу от одной вершины к другой). Указание направлений рёбер важно для обозначения ролей утверждений внутри аргумента (какие являются посылками, а какое заключением).

Связность графа означает, что для любой пары вершин a_1 , a_2 из множества вершин существует путь между a_1 и a_2 . Требование связности графа, обозначенное для разметки текстов в представленном исследовании, сопряжено с принципом выявления единственного главного тезиса для каждой статьи (граф не может не являться связным, так как тогда как минимум одно утверждение не будет иметь отношения к главному тезису, что противоречит его определению как ключевой идее текста, доказываемой всеми тезисами в составе аргументации). Свойство связности также подразумевает, что сегменты, не связанные с главным тезисом текста, не считаются аргументативными (и не вводятся в граф аргументации). Поскольку граф также является ориентированным, то понятие связности ограничивается до *слабой связности*: при проверке существования пути между двумя вершинами не учитываются направления рёбер. Так, если две посылки ведут к общему заключению, то они считаются связанными (опосредованно через это заключение). Наоборот, понятие *сильной связанности* (с учётом направления рёбер при проверке существования пути) не применимо к

аргументационным графам: существование направленного пути от заключения ксылке возможно только в случае «порочного круга» (заключение поддерживаетсясылкой, которая обосновывается этим же заключением).

Запрет на построение «порочных кругов» в процессе разметки отдельно фиксируется через требование ацикличности графа, или отсутствия в нём циклов. С учётом ориентированности графа под циклом понимается существование пути (согласно направлениям рёбер) из некоей вершины обратно в неё саму. Циклами не являются случаи, когда некоторое заключение обосновывается двумя разнымисылками, которые в свою очередь поддерживаются одной общейсылкой (так как при проверке на ацикличность учитываются направления рёбер).

Аргументационный граф является корневым. Корневая вершина соответствует главному тезису текста. В реализуемом подходе каждое ребро в графе обращено направлением к корневой вершине (напрямую либо опосредованно через иные рёбра и вершины).

3) Основной вариант ввода утверждений и связей в граф – итеративно от главного тезиса по расширяющейся окрестности с учётом авторского структурирования текста (на уровне разделов научной статьи, абзацного членения, сегментации абзацев на предложения). В этом случае сперва выявляютсясылки к главному тезису, затем определяются обоснования для этихсылков, после чего процесс повторяется до покрытия текста связной аргументационной структурой. Разметка текста экспертом завершается, когда на очередном шаге не выявляется ни одного нового тезиса, доступного для присоединения к построенному аргументационному графу. Преимущество такого подхода заключается в ориентированности на глобальную структуру аргументации в размечаемом тексте, а соответственно, её более надёжной передаче. Другой возможный подход основан на выявлении в первую очередь локальных связей между отдельными утверждениями по убыванию их прозрачности (от явных связей по очевидным схемам к менее явным связям с несколькими возможными схемами), с итоговым объединением разрозненных структур из разных частей текста в одну.

4) К аргументационному графу присоединяются только явно выраженные тезисы. Имплицитные утверждения, которые подразумеваются автором, но не представлены в тексте в точной формулировке, не добавляются в граф. Причинами пропуска таких неявных тезисов являются ориентированность на моделирование текстового выражения аргументации (в особенности для компьютерной обработки) и отсутствие точного критерия для разграничения имплицитных тезисов по степени их неявности. Например, при кореференции часто не оговаривается явно, что два разных наименования указывают на одно и то же явление, тогда как авторы научных статей могут обращаться к базовым сведениям своей области, которые полагают очевидными для иных исследователей и не разъясняют в тексте отдельно. Соответственно, восстановление неявных утверждений сопряжено с рисками их некорректной интерпретации (ввиду отсутствия точной формулировки в тексте), избыточного переполнения ими аргументационного графа (вплоть до превышения числа явных тезисов) и искажения авторской организации текста (при написании научной статьи её автор совершает осознанный выбор, какие сведения представлять явно, а какие не выражать в тексте).

Четыре обозначенных ключевых принципа для аргументационной разметки, проведённой в представленном исследовании, дополняются рядом более частных вспомогательных. Такие принципы покрывают, к примеру, определение границ тезисов в зависимости от конкретных аргументационных маркеров, выявление моделей рассуждения в отдельных аргументах. В частности, для распознавания этих моделей (аргументационных схем) экспертам-разметчикам предлагается использовать дерево вопросов, фрагмент которого представлен в виде таблицы 3. В столбце «№» обозначается номер вопроса, в столбцах «+» и «-» – номер следующего вопроса либо название конкретной модели рассуждения для положительного и отрицательного ответа на текущий вопрос (в столбце «Вопрос») соответственно.

Таблица 3 – Фрагмент дерева вопросов для определения схем аргументации

№	Вопрос	+	-
(1)	Выражает ли аргумент причинно-следственное отношение?	(2)	(7)
(2)	Содержит ли аргумент компонент практического соображения?	(3)	(6)
(3)	Приведена ли положительная или отрицательная оценка результатов?	(4)	(5)
(4)	Является ли эксплицированная в аргументе оценка положительной?	<i>Positive Consequences</i>	<i>Negative Consequences</i>
(5)	Сосредоточен ли практический компонент в представлении метода, которым достигнуты результаты (да), или проявляется в сочетании цели и способа её достижения (нет)?	<i>Applied Method</i>	<i>Practical Reasoning</i>
(6)	Строится ли рассуждение от видимых явлений к их причинам?	<i>Cause to Effect</i>	<i>Correlation to Cause</i>
(7)	Основан ли аргумент на обращении к некоей точки зрения?	(8)	(10)
(8)	Явно ли обозначен источник, предлагающий эту точку зрения?	(9)	<i>Popular Opinion</i>
(9)	Позиционируется ли источник этой точки зрения как авторитет?	<i>Expert Opinion</i>	<i>Position to Know</i>
(10)	Содержит ли заключение указание действия, корректность которого обосновывается распространённой практикой в посылке?	<i>Popular Practice</i>	(11)

Окончание таблицы 3

(11)	Обращается ли рассуждение к изучаемым объектам через их классификацию?	<i>Verbal Classification</i>	(12)
(12)	Основана ли связь на иллюстрации заключения через рассмотрение частного случая?	(13)	(15)
(13)	Описывает ли заключение более общую ситуацию по отношению к случаю в посылке (да), или же в рассуждении акцентируется сходство частных ситуаций (нет)?	(14)	<i>Analogy</i>
(14)	Приводится ли указанный частный случай как конкретный пример?	<i>Example</i>	<i>Sign</i>
(15)	Возможна ли интерпретация сущности в посылке как составной части сущности в заключении?	<i>Part to Whole</i>	Схема не из списка

Более подробно методы экспертной разметки аргументации представлены в работе «Методология многоаспектной разметки аргументации в научных статьях» ([https://uniserv.iis.nsk.su/arg_files/МетодологияРазметки_v1\(2024\)mod.pdf](https://uniserv.iis.nsk.su/arg_files/МетодологияРазметки_v1(2024)mod.pdf)), которая разработана в рамках представленного исследования и содержит также статистические характеристики итоговой версии созданного корпуса научных статей. Отдельные вспомогательные материалы для использования в процессе разметки представлены в приложениях к представленной работе. Приложение А содержит маркеры главного тезиса, выявленные в текстах размеченного корпуса, с указанием их позиционного расположения в тексте. В приложении Б приведены маркеры аспектов содержания научных статей, отражающие связь аспектов содержания с типовыми аргументационными схемами, которые реализуются при

представлении этих аспектов. Приложение В содержит набор из 16 наиболее частотных в корпусе аргументационных схем, где для каждой представлено её формализованное и неформализованное описание, приведён пример реализации схемы в аргументе.

2.2 Методы автоматизации аргументационной разметки

2.2.1 Обзор подходов к формальному выявлению аргументации

Работа по глубокой разметке текста, каковой является аннотирование аргументации, очень трудоемка, поэтому вручную размеченных текстов часто недостаточно для качественного обучения систем распознавания. В этой связи для увеличения объёма размеченных данных может производиться автоматизация аргументационной разметки (что подразумевает компьютерное выявление элементов аргументации в качестве вспомогательного инструмента для эксперта-разметчика, ускорения аннотирования вручную).

Компьютерное построение аргументационных структур обычно проводится по следующей схеме, отдельные этапы которой соответствуют шагам при ручном (экспертном) аннотировании [Lawrence, Reed, 2019]:

- 1) сегментация исследуемого текста;
- 2) разделение выделенных сегментов на аргументативные (утверждения, также тезисы) и неаргументативные;
- 3) выявление роли утверждений в структуре аргументации – главного тезиса, посылок и заключений;
- 4) установление связности утверждений с использованием знаний о схемах рассуждений и / или тематической структуры текста.

Для текстов с преобладанием аргументативного содержания (например, определенного жанра – дискуссии, или темы – юридические тексты) пункт (2) может быть пропущен.

В рамках формального выявления аргументации много внимания среди узкотематических текстов уделяется юридическим [Moens et al., 2007; Mochales Palau, Moens, 2009], иные часто исследуемые жанры – эссе [Stab, Gurevych, 2014], дебаты [Madnani et al., 2012; Hua, Wang, 2017] (иногда аргументы для дебатов готовят заранее и извлекают из интернет-ресурсов [Levy et al., 2014]), научные статьи [Lawrence et al., 2014], новостные публикации [Sardianos et al., 2015], онлайн комментарии [Park, Cardie, 2014].

Иногда исследователи, решающие задачу разделения выделенных сегментов на аргументативные и неаргументативные, не выделяют в анализируемом тексте на первом этапе потенциальные утверждения в виде клауз, а проводят анализ на уровне отдельных предложений. Аргументативные предложения как поисковые объекты фигурируют в работах [Moens et al., 2007; Hua, Wang, 2017]. Часто интерес представляет извлечение только тех аргументов, которые обеспечивают поддержку (или опровержение) тезиса, например, [Hua, Wang, 2017; Park, Cardie, 2014]. Часть работ посвящается обнаружению аргументативных предложений / клауз, играющих роль тезиса и / или посылки, например, [Daxenberger et al., 2017; Stab, Gurevych, 2014].

Прогностические модели обсуждаются в работах [Accuosto, Saggion, 2019; Passon et al., 2018]. Первая ориентирована на оценку перспективности научных докладов при их отборе на конференции, во второй анализируется качество отзывов о товарах.

Для поиска утверждений с аргументацией традиционно применяют такие методы машинного обучения, как наивный байесовский классификатор (MNB) [Lawrence et al., 2014; Sardianos et al., 2015], максимум энтропии [Moens et al., 2007], метод опорных векторов (SVM) [Park, Cardie, 2014], логистическая регрессия [Sardianos et al., 2015], деревья решений [Fischeva, Kotelnikov, 2019], нейронные сети [Ajjour et al., 2017].

Используемые системы признаков, как правило, включают несколько типов характеристик: структурные, лексические, синтаксические, контекстные, индикаторы. В качестве признаков рассматриваются n -граммы ($n \geq 1$), пары слов

(включая разрывные), теги отдельных частей речи, грамматические характеристики (время, лицо глагола и др.), число знаков препинания (Park, Cardie, 2014), позиции в предложении (основная или вспомогательная, например, в придаточном предложении), дискурсивные маркеры [Madnani et al., 2012]. В дополнение к перечисленным признакам используют информацию о типах поисковых объектов (аргументов или доводов в поддержку) [Levy et al., 2014], о риторической структуре [Madnani et al., 2012], тематической структуре текста [Lawrence et al., 2014] и пр. Наилучшие комбинации признаков и методов машинного обучения, как правило, выбираются экспериментально. В основном используются бинарные значения признаков, но иногда реализуется взвешивание, например, по $TF*IDF$, как в [Passon et al., 2018] и др.

Применение методов машинного обучения (MNB, максимум энтропии, SVM) к юридическим текстам для обнаружения аргументативных предложений [Mochales Palau, Moens, 2009] показывает, что наилучшей комбинацией «система признаков + метод» является алгоритм MNB и максимум энтропии, обученные на n -граммах ($n = 1, 2, 3$), парах слов, отдельных частях речи, ключевых словах, пунктуации и текстовой статистике. Такой показатель качества, как точность, зависит от обучающей коллекции и составляет 73% и 80% для юридических текстов. Следует заметить, что юридические тексты отличаются лапидарностью языка, что менее характерно для научных текстов.

Активное распространение больших языковых моделей (large language model, LLM), обучающихся на объёмных коллекциях разножанровых текстов и способных решать широкий спектр задач обработки естественного языка, побуждает исследователей изучать их применимость и к таким сложным проблемам компьютерной лингвистики, как формальный анализ аргументации. Основания для использования LLM при решении узкоспециализированных задач заключаются в их эмерджентных способностях: языковые модели, обучаемые технически прямолинейной процедуре предсказывать n -граммный контекст, начинают решать различные задачи обработки текста без какой-либо настройки под эти задачи, а только за счёт масштабного увеличения обучающих данных

(крупных коллекций неразмеченных текстов) и расширения обрабатываемого контекстного окна ценой привлечения значительных вычислительных мощностей [Wei J. et al., 2022]. При этом возможность решения новых, ранее не поддававшихся задач никак не выводится из экстраполяции работы модели на меньшем объёме данных, что затрудняет предсказание изменений в работе LLM при дальнейшем масштабировании данных.

Так, в работе [Chen et al., 2024] оценивается эффективность LLM трёх разных семейств (Chat GPT-3.5-Turbo, Flan-UL-2, Llama-2-13B) в обработке аргументации на естественном языке. Авторы применяют два подхода к формулировке задачи: стандартный для традиционных методов формального анализа аргументации с разграничением четырёх подзадач для поочерёдного решения каждой моделью и специализированный, состоящий в генерации с помощью LLM полемического текста с оспариванием доказательства в обрабатываемом (что позволяет косвенно охарактеризовать адекватность «восприятия» моделью этого доказательства). Для проверки моделей на четырёх подзадачах стандартного подхода, выделенных авторами, использованы разные корпуса кратких текстов, специально созданных для испытания методов формального анализа аргументации и напоминающих студенческие эссе с высказыванием и обоснованием позиции по определённой проблеме. В эксперименте LLM достигают точности в 64–84% для установления главного тезиса, 35–65% для определения авторской позиции по рассматриваемой проблеме, 65–71% для выявления утверждений в поддержку главного тезиса и 58–73% для определения общего типа поддержки (приведение статистики или фактов, отсылка к авторитету или примеру). Эти значения уступают показателям качества традиционных методов формального анализа аргументации, что авторы объясняют неестественностью пошагового разграничения подзадач для больших языковых моделей, воспринимающих аргументацию в обрабатываемых текстах целиком в широком контексте ввиду своего устройства.

Для демонстрации синкретичного характера обработки аргументации большими языковыми моделями в исследовании дополнительно сопоставляются

два подхода к генерации оспаривающих текстов: LLM предлагается либо породить опровержение исходной аргументации сразу, либо сперва сгенерировать краткий пересказ доказательства в обрабатываемом тексте, а затем оспорить аргументацию по созданному пересказу. Тексты, сгенерированные моделями, оцениваются экспертами-людьми по показателям последовательности (как логической, так и грамматической) и убедительности (fluency и persuasiveness, по шкалам от 1 до 5), а также по проценту покрытия исходных аргументов в ходе опровержения. Подход с генерацией опровержения напрямую превосходит подход с промежуточным пересказом исходной аргументации по всем трём показателям: 4.32 против 3.56 по средней последовательности, 3.8 против 2.8 по убедительности, 95% против 78% по полноте покрытия исходной аргументации. В соответствии с полученными результатами авторы исследования указывают на способность LLM создавать логичные и убедительные тексты, достаточно полно передающие суть рассуждений в обрабатываемых, и рекомендуют не запутывать модели введением понятных для человека, но заметно более сложных для LLM промежуточных шагов.

Эффективность LLM в порождении связных убедительных текстов, в том числе на основе других через переформулирование исходной аргументации, подчёркивается и в работе [Ziegenbein et al., 2024]. Авторы решают задачу автоматического перефразирования неприемлемых комментариев в онлайн-дискуссиях (излишне эмоциональных или агрессивных) с сохранением логики аргументации из исходного сообщения. В ходе эксперимента отмечено, что хотя по оценке экспертов модель корректно передаёт структуру доказательства из обрабатываемых текстов в большинстве случаев, в отдельных ситуациях (без указания частот) LLM искажает исходные аргументы либо через изменение авторской позиции (например, из-за замены отдельных слов, как «not» на «is»), либо через дописывание утверждений, противоречащих доказываемому тезису. Авторы работы объясняют такие ошибки малым объёмом обрабатываемых комментариев (не больше 220 слов, не больше 1100 символов), что затрудняет учёт моделью контекста.

В статье [Sun et al., 2024] также рассматривается применение LLM для решения частных задач в извлечении аргументов (разграничение тезисов и посылок, проверка аргументативной связи в конкретных парах утверждений, определение общего характера связи в терминах «поддержка / атака»). Представлен метод на основе автоматической тонкой настройки запросов к LLM по мере обработки отдельных пар утверждений, которая динамично продолжается до получения непротиворечивого представления полного текста (с возвращением к ранее построенным/пропущенным связям в случае появления структурных аномалий в аргументации). Экспериментально выявлен выигрыш в 4% по F-мере над другими подходами, решающими аналогичные задачи на том же корпусе. При этом авторы отмечают нетипично частые ошибки модели с построением связей между утверждениями на значительном позиционном отдалении внутри текста при игнорировании контактных и близких связей. Авторы объясняют эти ошибки разницей в ширине контекстного окна при обработке аргументации между LLM и человеком.

Таким образом, следует отметить явный потенциал быстро развивающихся больших языковых моделей в компьютерной обработке аргументации, особенно при отсутствии подходящего под задачу размеченного корпуса для подготовки традиционных методов машинного обучения. Тем не менее, хотя LLM показывают хорошие результаты на содержательно более простых задачах анализа аргументации (с лингвистической точки зрения), отмеченные ошибки в их работе могут свидетельствовать о резком ухудшении результатов на более сложных. Так, в работах [Sheng et al., 2023] и [Schaeffer, Miranda, Koyejo, 2023] подвергается сомнению наличие эмерджентных способностей у больших языковых моделей: по мнению авторов, видимость их проявления обуславливается проверкой на задачах определённого круга и не повторяется на глубинных проблемах вне него, а также обеспечивается избирательным подбором метрик для проверки эмерджентности. Соответственно, даже при дальнейшем масштабировании LLM определённые типы задач могут оказаться недоступными для их эффективного решения. Отмеченная проблема выходит за рамки

представленного исследования, которое направлено на реализацию традиционных методов машинного обучения для пошагового анализа аргументации в текстах на русском языке. Это исследование на базе размеченного корпуса предоставит ориентировочные оценки качества (*baseline*) для сравнения с иными методами, в частности, на основе более сложных нейронных сетей.

2.2.2 Компьютерное распознавание утверждений с аргументацией

В представленном исследовании реализована частичная автоматизация аргументационной разметки посредством компьютерного распознавания в тексте предложений, потенциально используемых в построении аргументов.

Предложение считается аргументативным (тезисом), если в нем содержатся утверждения из структуры рассуждений (посылка, заключение), представленных в тексте с целью доказательства или опровержения главного тезиса. Распознавание аргументативных предложений в текстах сводится к решению задачи классификации множества предложений $S = \{s_i\}$ текстов коллекции на два класса: $C = \{C_a, C_n\}$, где C_a – метка класса для предложений, содержащих посылку или заключение, а C_n – метка класса, к которому относятся предложения без аргументации ($C_a \cap C_n = \emptyset$). Требуется построить на обучающей коллекции классификатор для точной классификации $F': S \times C \in \{0, 1\}$, максимально близкий к целевой функции $F(s_i, C_j) = \{0, \text{если } s_i \notin C_j, 1 \text{ если } s_i \in C_j\}$ ($j = 1, 2$).

Обучение и распознавание проводится на множестве размеченных экспертом примеров (значения целевой функции F известны): $S = S^L \cup S^T$, где S^L – обучающее множество, S^T – тестовое. Это позволяет вычислять показатели качества (точность P , полноту R , и F -меру) для S^T автоматически.

Предобработка текстов заключается в нормализации слов с помощью программы PyMorphu2 [Korobov, 2015] и устранении слов, не содержащих кириллицу. Используемая модель текста – мешок слов, представление – вектор нормализованных униграмм.

Поскольку формирование набора лингвистически обоснованных признаков требует значительного времени, в данной работе реализуется статистическая фильтрация признаков. Для этого применяются процедуры взвешивания компонентов вектора по мерам TF*IDF и ожидаемой информационной выгоды (EMI), а также отбор униграмм согласно критериям Variance и χ^2 [Manning, Schutze, 2000].

Мера TF*IDF имеет высокие показатели, если термин часто встречается в небольшом количестве документов и низкие, если термин встречается во многих документах, что позволяет отсеять служебную и общеупотребительную лексику (в нее могут входить и маркеры и / или их составляющие, как правило, слабые, не слишком точно идентифицирующие аргументативную лексику). Величина информационной выгоды EMI позволяет отфильтровать леммы, часто встречающиеся в разных классах, так как измеряет количество информации, которое обеспечивается присутствием / отсутствием признака в классе; максимум достигается, когда признак встречается в предложениях из одного класса.

С помощью критерия Variance (Var) удаляют признаки с низкой дисперсией (меньше заданного порога), то есть те, что имеют близкие значения для всех предложений в обучающей коллекции. Величина критерия χ^2 позволяет судить о том, насколько ожидаемая и наблюдаемая вероятности появления признака и появления класса отклоняются друг от друга в предположении о независимости этих двух событий. Использование критерия отсеивает признаки, которые с наибольшей вероятностью будут независимыми от класса, то есть те не имеют отношения к классификации.

Статистическая фильтрация признаков реализуется с помощью программ из библиотеки Scikit-Learn). Пороги для четырёх указанных выше мер и критериев выбираются экспериментально, в соответствии с результатами распознавания.

Для классификации предложений выбраны часто упоминаемые в литературе при решении аналогичных задач методы MNB (наивный байесовский классификатор), SVM (метод опорных векторов) и MLP (многослойный перцептрон). Глубокая настройка алгоритмов не проводится для избежания

переобучения классификаторов на обрабатываемой коллекции (и потенциального понижения их эффективности при масштабировании на иные тексты).

2.2.3 Методы машинного обучения для распознавания тезисов и схем

Алгоритм MNB. Для оценки вероятности принадлежности предложения s_i классу C_a или C_n алгоритм MNB использует формулу Байеса на обучающей коллекции. Каждое предложение представляется мультимножеством лемм (используются бинарные значения отобранных признаков), а наилучшим классом для него является тот, что имеет наибольшую апостериорную вероятность. Предполагается, что признаки (леммы) встречаются в тексте независимо друг от друга. В формуле подсчета вероятностей учтено сглаживание по Лапласу [Lidstone, 1920] для корректной реакции на неполноту обучающей коллекции.

Алгоритм SVM. Алгоритм ищет в векторном пространстве документов разделяющую поверхность между двумя классами, максимально удалённую от всех точек обучающего множества. В качестве ядра выбрана радиальная базисная функция RBF: $\exp(-\gamma ||x - x'||^2)$ [Buhmann, 2003]. Экспериментально проверено, что полиномиальная квадратичная функция показывает худшие результаты. При обучении SVM с ядром RBF учитываются два параметра: C (устанавливает баланс между ошибкой классификации и простотой поверхности принятия решений) и γ (определяет влияние одного обучающего примера).

Алгоритм MLP. Классификатор реализует алгоритм MLP, который обучается с использованием метода обратного распространения ошибки. Компоненты векторов s_i ($x_1, x_2, \dots, x_m$) составляют входной слой. В скрытых слоях каждый нейрон преобразует значения из предыдущего слоя со взвешенным линейным суммированием $w_1x_1 + w_2x_2 + \dots + w_mx_m$. В применяемой нами реализации используется функция активации $f(x) = \max(0, x)$. Персептрон по умолчанию имеет 100 слоев, применён решатель из семейства квази-Ньютоновых методов (минимизируется функция логарифмических потерь при реализации стохастического градиентного спуска). Программа выполняет итерации до

сходимости, задаваемой параметром (в эксперименте равен 0.0001), или до тех пор, пока количество итераций не достигнет порогового значения p .

2.3 Методы оценки качества разметки через согласие аннотаторов

Субъективный характер экспертной разметки обуславливает необходимость дополнительной проверки её надёжности и применимости для корректного обучения классификаторов. Надёжность разметки может оцениваться в терминах согласия аннотаторов (сходства их решений при разметке одних и тех же текстов) посредством автоматического подсчёта количественных коэффициентов.

Известно большое количество коэффициентов согласия разного типа, причём, как отмечено в работе [Олейник и др., 2014], при выборе одного из них следует учитывать структурную специфику анализируемых текстов, их доступность для различных интерпретаций (или, наоборот, направленность на передачу конкретного недвусмысленного сообщения), а также цель их содержательного моделирования (разметки). Необходимо принимать во внимание ограничения коэффициента, предположения в его основе (например, учитывается ли возможность случайного согласия, назначаются ли разные или одинаковые веса размечаемым классам с разной частотой), а также такие параметры, как сравниваемое число аннотаторов и полнота разметки (каждый ли из текстов коллекции размечен каждым аннотатором).

Согласие разметчиков в разработанном корпусе оценено посредством трёх различных коэффициентов, одного специализированного и двух универсальных (каппа Коэна и альфа Криппендорфа). Первый коэффициент, коэффициент аргументационного соответствия (КАС), создан в рамках представленного исследования непосредственно для задачи аргументационной разметки с учётом многоуровневой специфики такой разметки и подробно представлен в работе [Пименов, 2023].

2.3.1 Подсчёт коэффициента аргументационного соответствия

Значение КАС подсчитывается отдельно по каждому тексту, для двух вариантов его разметки (с последующим усреднением значений по всему корпусу), посредством представленного ниже алгоритма.

Пусть $R = \{r_i^j\}$ – множество графов аргументации, построенных j экспертами ($i = 1, \dots, I$, I – число текстов в корпусе). В общем случае допустимо принимать значение $j = 2$ (если некоторый текст размечен $n > 2$ аннотаторами, его можно рассматривать как m разных текстов, где m – число всех пар разметчиков из n). Рассмотрим граф r_i^j , который представляется множеством вершин $\{s_i^j\}$ и множеством соединяющих эти вершины помеченных ребер $\{e_i^j, ch(e_i^j)\}$, где e_i^j – ребра (связи), а соответствующие им метки $ch(e_i^j)$ – схемы аргументации.

Подсчет согласия по КАС (сходства графов аргументации) проводится последовательно в три этапа: отдельно для $\{s_i^j\}$, подмножества $\{e_i^j\}$ и $\{ch(e_i^j)\}$. На первом этапе определяется доля утверждений $K(S_i^{12})$, совпавших в графах r_i^1 и r_i^2 , от суммарного числа утверждений в объединении вершин этих графов: $K(S_i^{12}) = |S_i^{12}|/|\{s_i^1\} \cup \{s_i^2\}|$, где $S_i^{12} = \{s_i^1\} \cap \{s_i^2\}$. K здесь и далее обозначает коэффициент на соответствующем уровне разметки, $|S_i^{12}|$ означает мощность S_i^{12} .

На втором этапе вычисляется доля совпавших связей от числа всех связей между совпавшими вершинами из множества S_i^{12} : $K(\hat{E}_i^{12}) = |\hat{E}_i^{12}|/|E_i^{12}|$, где $\hat{E}_i^{12}$ – множество совпавших связей, а E_i^{12} – множество всех связей между вершинами из S_i^{12} .

На третьем этапе вычисляется доля совпавших меток (схем аргументации) у выявленных на втором этапе совпавших связей из $\hat{E}_i^{12}$: $K(\hat{Ch}_i^{12}) = |\hat{Ch}_i^{12}|/|\hat{E}_i^{12}|$, где $\hat{Ch}_i^{12}$ – множество совпавших схем аргументации.

Для уточнения характера расхождений в выборе схем дополнительно подсчитывается значение $K'(\hat{Ch}_i^{12})$, в котором не учитываются разногласия насчёт содержательно и функционально сходных схем (например, «от мнения эксперта»,

«с позиции знающего», «от популярного мнения»). Так, в случае трёх отмеченных схем аннотаторы фактически согласны, что доказательство строится от авторитетного мнения, а расходятся лишь в оценке природы этого авторитета. Всего выделено четыре группы функционально близких схем из всех схем, использованных аннотаторами при разметке корпуса, а их подробное представление доступно в разделе 4.4.

Диапазон значений КАС в корпусе вычисляется как $\min_i \{K(S_i^{12})\}$, $\min_i \{K(\hat{E}_i^{12})\}$, $\min_i \{K(\hat{C}h_i^{12})\}$, $\min_i \{K'(\hat{C}h_i^{12})\}$ и $\max_i \{K(S_i^{12})\}$, $\max_i \{K(\hat{E}_i^{12})\}$, $\max_i \{K(\hat{C}h_i^{12})\}$, $\max_i \{K'(\hat{C}h_i^{12})\}$. Среднее значение КАС определяется следующим образом: $\sum_i K(S_i^{12})/I$, $\sum_i K(\hat{E}_i^{12})/I$, $\sum_i K(\hat{C}h_i^{12})/I$, $\sum_i K'(\hat{C}h_i^{12})/I$.

2.3.2 Оценка согласия через каппу Коэна и альфу Криппендорфа

Каппа Коэна [Cohen, 1960] – статистическая метрика, которая количественно определяет уровень согласия между двумя аннотаторами, классифицирующим объекты по взаимоисключающим категориям. Она традиционно применяется для оценки надёжности размеченных данных самого разного целевого назначения, в том числе при разметке аргументации, например, по модели С. Тулмина [Teruel et al., 2018]. Этот коэффициент (κ) высчитывается по парам разметчиков по следующей формуле, обеспечивающей коррекцию получаемого значения с учётом вероятности случайного совпадения меток [Artstein, Poesio, 2008]: $\kappa = (A_0 - A_e) / (1 - A_e)$, где A_0 – доля объектов с совпавшими метками категорий, назначенных аннотаторами, а A_e – ожидаемое случайное совпадение меток. Так, $A_0 = \sum_{i=1}^I \text{agr}_i / I$, где I – число дискретных объектов в размечаемом наборе данных, а agr_i равняется 1 (аннотаторы указали одинаковые категориальные метки для объекта) или 0 (объект охарактеризован разными метками). В свою очередь, $A_e = \sum_{k=1}^K n_{c_1 k} n_{c_2 k} / I^2$, где k – категория из списка возможных категорий K , а $n_{c_i k}$ – частота объектов с категорией k среди всех объектов, размеченных аннотатором c_i .

Иными словами, каппа Коэна придаёт больший вес согласию аннотаторов на редких, им не свойственных метках, и меньший – на частотных метках коллекции. Верно и обратное: расхождение аннотаторов на редких метках не так выражено понижает значение коэффициента, как разногласие по частотным категориям.

Значение k в стандартной формулировке принадлежит диапазону $[-1, 1]$: 1 указывает на полное согласие аннотаторов по категориям всех объектов, а -1 – на обратную корреляцию их решений в разметке (когда один аннотатор приписывает объекту одну категорию, другой обязательно отмечает другую). 0 в этом диапазоне характеризует ситуацию полной неопределённости, когда по решению одного аннотатора нельзя сказать ничего о решении другого (он может как согласиться, так и указать иную категорию). В статье [Landis, Koch, 1977] предлагается следующая интерпретация положительных значений: от 0.01 до 0.20 соответствует небольшому согласию, между 0.21 и 0.40 означает удовлетворительное согласие, 0.41 и 0.60 – умеренное согласие, 0.61 и 0.80 – существенное согласие, а 0.81 и 1 – практически идеальное согласие.

Таким образом, подсчёт каппы Коэна требует полной разметки всего списка объектов каждым разметчиком, а также явного и дискретного характера объектов. В случае же с многоуровневой разметкой аргументации по стандарту AIF условие дискретности не выполняется уже для выявления утверждений: аннотаторы работают с целостным текстом и самостоятельно определяют границы аргументативных сегментов в его составе. Следовательно, применение каппы Коэна для оценки согласия аннотаторов в размеченном корпусе возможно лишь путём вспомогательного представления данных в доступном для подсчёта виде. Предлагается следующий подход.

Пусть T – размечаемый текст, образованный предложениями t ($T = \{t_{ij}\}$, $j = 1, \dots, J$, J – число предложений в тексте). Рассмотрим каждое предложение t как отдельный дискретный объект и определим его метку u аннотатора c_i через обращение к списку утверждений S , выявленных этим разметчиком: если существует такое s из S , что t является подстрокой s либо s является подстрокой t ,

то это предложение считается аргументативным сегментом; если же такого s нет, то оно признаётся неаргументативным сегментом. Сегменты t' , отмеченные аргументативными у обоих разметчиков, используются затем для оценки согласия на связях по списку аргументов: для каждой пары предложений t'_1 и t'_2 из числа всех возможных пар t' проверяется наличие аргумента из списка аргументов, выявленных разметчиком, посылка которого соответствует одному предложению из этой пары, а заключение – другому (здесь соответствие понимается по аналогии с прошлым шагом как вхождение подстроки). Если такой аргумент есть в разметке аннотатором, то пара сегментов t'_1 и t'_2 считается аргументативно связанной, иначе делается вывод об отсутствии связи между этими сегментами. Наконец, оценка согласия на уровне аргументационных схем для пар сегментов, одинаково связанных у обоих аннотаторов, тривиальна: меткой для пары считается схема, указанная в разметке для соответствующего аргумента. В отличие от подсчёта k на уровнях утверждений и связей, где выделяется по две возможных категории (аргументативные и неаргументативные сегменты, связанные и несвязанные пары аргументативных сегментов), на уровне схем аргументации число категорий равно числу уникальных схем в разметке всех текстов коллекции обоими аннотаторами.

Иными словами, дискретными объектами для подсчёта коэффициента выступают целостные предложения исходного текста на уровне утверждений и их пары на уровнях связей и схем.

Итоговые значения k в представленном исследовании подсчитываются по всей совокупности объектов каждого уровня в коллекции (например, по всем предложениям всех текстов, без деления по текстам). Для подсчёта капши Коэна применена программная реализация алгоритма из библиотеки Scikit-Learn для языка Python. Поскольку двойная разметка 100 статей корпуса проведена тремя экспертами, каждый из которых разметил неполную часть этих 100 текстов, то значение k подсчитывается отдельно для каждой пары аннотаторов по текстам, размеченным ими обоими.

Альфа Криппендорфа [Krippendorff, 2004] работает по аналогичным принципам, что и каппа Коэна, но применима к любому числу аннотаторов (не только к двум) и допускающим неполное покрытие данных (отсутствие метки кого-то из аннотаторов на некоторых объектах). Поскольку её подсчёт также предполагает дискретность размечаемых объектов, в исследовании задействована аналогичная предобработка данных, что и в случае с каппой Коэна. Использование альфы Криппендорфа позволяет при этом получить общее значение согласия аннотаторов на каждом из трёх уровней разметки без усреднения попарного согласия. В исследовании задействована программная реализация альфы Криппендорфа из библиотеки Fast Krippendorff для языка Python [Castro, 2017].

Существуют и иные показатели согласия аннотаторов, например, каппа Флейсса [Fleiss, 1971] или коэффициент корреляции Мэтью [Matthews, 1975], однако принципы их подсчёта близки к двум рассмотренным, ввиду чего включение дополнительных коэффициентов представляется избыточным (в ходе эксперимента также выявлено, что значение альфы входит в диапазон значений Каппы).

2.4 Методы выявления составных приёмов аргументации

2.4.1 Обзор подходов к обработке приёмов аргументации

Некоторые исследователи в области Argument Mining предлагают включать в разметку приемы аргументации (типовые структуры из нескольких аргументов определённых моделей), чтобы упростить исследование ее эффективности. Так, в работе [Anand et al, 2011] показано, что знания о приемах аргументации оказывают существенную помощь в классификации блогов по убедительности.

Моделирование приемов (либо более сложных стратегий аргументации, традиционно понимаемых как сочетание нескольких приёмов) актуально для

вычислительного синтеза аргументации. Примерами такого рода исследований могут служить работы [Wachsmuth et al., 2018] и [Al-Khatib et al., 2017].

В работе [Wachsmuth et al., 2018] представлены результаты анализа аргументационных стратегий в редакционных новостных статьях. На размеченных текстах изучено применение в текстах разной тематики приемов аргументации, образуемых аргументами трёх типов (статистика, анекдот, свидетельство от наблюдателя, эксперта, организации). Приведены примеры 15 встретившихся последовательно приемов (цепочек длиной от 1 до 7 аргументов трёх указанных типов), использованных авторами. Выявлена корреляция между темами и используемыми приемами. Область применения результатов – синтез, идентификация аргументов, классификация текстов.

Авторы работы [Al-Khatib et al., 2017] строят модель стратегии, которая опирается на каноны риторики, сформулированные Аристотелем. Оценка модели проводится экспертами. Модель реализуется в три этапа, которые выполняются вручную. Для каждого заданного тезиса выбираются аргументативные дискурсивные единицы (ADU) из базы. Из них составляются структуры аргументации (цепочки из посылок, заключений). Выбранные утверждения оформляются в определенном стиле согласно принимаемой стратегии. Под стратегией в этой работе понимается пропорция используемых средств, например, логос – 70%, этос – 10%, пафос – 20%. Результаты синтеза коротких аргументативных текстов совпадают у разных экспертов примерно на 50%. Изучение структурных особенностей в организации последовательности рассуждений, реализуемой в текстах, не входило в задачу авторов. Знания о такой организации могут быть полезны на этапе формирования структуры. Они могли бы отчасти нивелировать различия в аргументации, синтезированной разными экспертами.

Авторы работы [El Baff et al., 2019] предлагают алгоритм, согласно которому аргументационные цепочки строятся автоматически вне зависимости от темы из базовых ADU. База ADU содержит тезисы, аргументы «за» и «против» с указанием реализованной в них стратегии (логос, пафос). Выбранные ADU

упорядочиваются согласно связности, определяемой путем вычисления семантического расстояния между ними. Совпадение с построенными экспертами структурами аргументации составило порядка 50% на уровне двух подряд следующих аргументов. Детализация связей между аргументами и учет организации схем в структуры могут быть полезным дополнением данному методу.

2.4.2 Выявление приёмов аргументации через анализ подграфов

В представленном исследовании под приемами аргументации понимается применение отдельных повторяющихся моделей рассуждений и образуемых ими повторяющихся структур в аргументации, реализуемой авторами в текстах. Задача поиска приемов аргументации в таком случае – это извлечение из заданной коллекции всех повторяющихся подграфов (графов, образованных некоторым подмножеством вершин-схем графа аргументации и некоторым подмножеством смежных им рёбер).

Такого сорта задачи могут решаться методами частотного анализа подграфов FSM (Frequent subgraph mining) [Jiang, Coenen, Zito, 2004], который включает генерацию подграфов, представляющих интерес, а также подсчет частоты встречаемости этих подграфов в заданном наборе данных. В данной работе генерацию кандидатов заменяет разбиение графа аргументации текста на подграфы с фиксированным числом вершин. Для подсчета частоты требуется провести сравнение полученных подграфов. Эта процедура известна как проверка подграфов на изоморфизм. Самые известные алгоритмы, позволяющие получить точное решение, предложены в работах [Ullmann, 1976; Cordella et al., 2001]. В исследовании применяется последний, использующий стратегию поиска в глубину и набор правил для сокращения времени поиска и потребляемой памяти компьютера.

Пусть $A = \{a_i\}$ – множество аргументационных аннотаций коллекции текстов. Каждая аннотация a_i представляет собой связный граф G_i ,

$G_i = \langle V_i^{\text{inf}} \cup V_i^{\text{ch}}, E_i \rangle$, где V_i^{inf} – множество информационных вершин, V_i^{ch} – множество вершин-схем, E_i – множество ребер, соединяющие вершины. На данном этапе нами проводится анализ структур (графов) $\hat{G}_i$, образуемых из графов G_i с сохранением только вершин-схем: $\hat{G}_i = \langle V_i^{\text{ch}}, \hat{E}_i \rangle$. Ребра $\hat{E}_i$ строятся следующим образом: 1) в случае удаления информационной вершины при цепочечной их организации – путем замены входящего и исходящего ребра на ребро, соединяющее две схемы и сохраняющее направление; 2) в случае конвергенции информационных вершин аналогичным образом строятся ребра при удалении каждой из них; 3) в случае дивергентного связывания информационных вершин строятся ребра, объединяющие каждую схему с входящим в информационную вершину ребром с каждой схемой, соединенной исходящими из информационной вершины ребрами. Аннотацию a_i удобно представлять совокупностью частотных характеристик

$\Phi(a_i) = \{\Phi_1(a_i), \dots, \Phi_n(a_i), \dots, \Phi_{N_{\max}}(a_i)\}$, где $\Phi_1(a_i)$ – статистика отдельных схем (разных вершин из V_i^{ch} и частот их встречаемости $F_{\text{abs}}(V_i^{\text{ch}}) > 1$). Характеристика $\Phi_n(a_i)$ ($n > 1$) аккумулирует множество подграфов $\{sg_{ik}^n\}$ графа $\hat{G}_i$ с фиксированным числом вершин равным n ($n = 2, \dots, N_{\max}$, $N_{\max}$ – задаваемая константа, ограниченная сверху значением $|V_i^{\text{ch}}|$): $\Phi_n(a_i) = \{sg_{ik}^n\}$, $k = 1, \dots, K_n$ – число разных подграфов с n вершинами. Для каждого подграфа sg_{ik}^n вычисляется число его вхождений $F_{\text{abs}}(sg_{ik}^n)$ в граф $\hat{G}_i$. В $\Phi(a_i)$ включаются подграфы с $F_{\text{abs}}(sg_{ik}^n) > 1$. Будем называть $\Phi(a_i)$ частотным спектром приемов, реализованных в тексте.

Представление множества A в виде частотного спектра включает аналогичные характеристики и является совместным частотным спектром коллекции: $\Phi(A) = \{\Phi_1(A), \Phi_2(A), \dots, \Phi_n(A), \dots, \Phi_{N_{\max}}(A)\}$, где $N_{\max}$ – максимальное число вершин в подграфах, общих хотя бы для двух графов аргументации разных текстов: $\Phi_n(A) = \{sg_{ik}^n\}_{i=1}^{|A|}$ ($|A|$ обозначает мощность множества A). Для каждого подграфа sg_{ik}^n вычисляется $F_{\text{abs}}(sg_{ik}^n)$ и $F_{\text{txt}}(sg_{ik}^n)$ –

абсолютная и текстовая частота в коллекции A . В частотный спектр входят подграфы с $F_{\text{txt}}(sg_{ik}^n) > 1$. Под текстовой частотой понимается число текстов из коллекции A , в графах которых встречается данный подграф. Все приемы из $\Phi(A)$ могут повторяться внутри одного текста и/или в разных текстах. Совокупность внутритекстовых приемов обозначим $\Phi^{\text{In}}(A)$, а межтекстовых – $\Phi^{\text{Ex}}(A)$.

Если множество текстов A по некоторому критерию может быть разбито на отдельные подмножества (например, подмножества текстов одинакового жанра), тогда $A = \bigcup_{r=1, \dots, R} A_r$, где R – число подмножеств (жанров). Жанровым приемам фиксированного по r подмножества соответствуют подграфы sg_{mk}^n из $\{sg_{mk}^n\}_{m=1}^{|A_r|}$ с текстовой частотой в этом подмножестве – $F_{\text{txt}}(sg_{mk}^n) > 1$. Среди жанровых приемов также будем рассматривать межтекстовые $\Phi_r^{\text{Ex}}(A_r)$ и внутритекстовые приемы $\Phi_r^{\text{In}}(A_r)$ с соответствующими частотами: $F_r(sg_{mk}^n) > 1$. Межжанровая частота подграфа $F(sg_{ik}^n)$ равна числу жанров r ($0 < r < R + 1$), среди подграфов которых встречается данный подграф.

Стратегию аргументации текста i определим как совокупность приемов $\Phi(a_i)$, которые могут входить как в $\Phi_r^{\text{Ex}}(A_r)$, $\Phi_r^{\text{In}}(A_r)$, так и в $\Phi^{\text{Ex}}(A)$ и $\Phi^{\text{In}}(A)$. Таким образом, формирование спектра $\Phi(A) = \bigcup_{n,i} \Phi_n(a_i)$ сводится к выявлению структурно идентичных общих подграфов как в отдельной аннотации a_i , так и во всей коллекции A , а также в ее подмножествах. Как уже было сказано во Введении, для определения идентичности структур двух графов используется понятие изоморфизма графов. По определению два графа изоморфны, если у них одинаковое число вершин (n) и вершины каждого из них можно переименовать так, что в первом графе две вершины соединены ребром тогда и только тогда, когда вершины с такими же именами соединены во втором графе. Для установления изоморфизма используется реализация алгоритма Корделлы VF2 (NetworkX) [Hagberg, Schult, Swart, 2008], которая для каждой комбинации значений индексов i, j, k, n ; $i \neq j$ определяет, не содержит ли граф $\hat{G}_i$ подграф, изоморфный подграфу $sg_{ik}^n \subset \hat{G}_i$. Кроме того, из всей совокупности приемов удаляются все вложенные подграфы.

2.4.3 Построение частотного спектра для повторяющихся подграфов

Построение совместного частотного спектра $\Phi(A)$ производится в три этапа:

1. *Этап предобработки.* Основные функции этого этапа – извлечение из каждого аргументационного графа, считываемого из описания в формате .json, множества применяемых схем и множества его подграфов согласно заданному числу n вершин-схем $\{sg_{ik}^n\}$ ($n > 1$). Для этого выполняются следующие шаги:

1) Преобразование текстового представления графа из .json-формата в матрицу смежности.

2) Построение по матрице смежности компактного представления графа $\hat{G}_i$, состоящего только из вершин-схем (без информационных вершин с текстовым содержанием отдельных утверждений, но с сохранением всех связей, проходящих через вершину-схему). Для обеспечения связности графа $\hat{G}_i$ (в условиях отсутствия информационных вершин) добавляется техническая вершина MainThesis, обозначающая главный тезис исходного графа (с указанием всех переходов от ведущих к нему вершин-схем).

Удаление из рассмотрения информационных вершин обуславливается абстрагированием на данном этапе исследования от содержания тезисов при анализе приёмов аргументации (приёмы определяются на уровне типовых моделей рассуждения). В этом случае число вершин графа уменьшается почти два раза, что важно для сокращения трудоёмкости поиска изоморфных подграфов.

3) Извлечение пар связанных вершин.

4) Извлечение подграфов $sg_{ik}^n \subset \hat{G}_i$ для всех $i = 1, \dots, |A|$; $k = 1, \dots, K_n$. Нижняя граница диапазона n соответствует наименьшему числу вершин, при котором проявляется структурная вариативность между подграфами (все подграфы с $n = 2$ образуют цепочку, три же вершины могут соединяться последовательно, либо через конвергентный переход от двух к третьей, либо через дивергенцию от одной к двум другим). Верхняя граница изменения n

определяется в ходе работы алгоритма на этапе исследования подграфов с $n + 1$ вершиной.

Извлечение всех подграфов sg_{ik}^n осуществляется полным перебором: из множества вершин графа поочерёдно выбирается n вершин ($n > 2$), для которых по матрице смежности графа $\hat{G}_i$ проводится проверка связности вершин.

2. Построение совместного частотного спектра $\Phi(A)$.

- 1) Вычисление характеристик $\Phi_1(A)$ на базе V_i^{ch} для всех $i = 1, \dots, |A|$.
- 2) Вычисление характеристик $\Phi_2(A)$ по парам связанных вершин из $(V_i^{ch}, \hat{E}_i)$ для всех $i = 1, \dots, |A|$.
- 3) Вычисление характеристик $\Phi_n(A)$, $n > 2$.

Поиск изоморфных подграфов. На этом этапе производится попарный анализ преобразованных графов $\hat{G}_i$ всех текстов коллекции с помощью алгоритма VF2. Выявляемые изоморфизмы sg_{ik}^n фиксируются в $\Phi(A)$ с указанием вычисленных $F_{abs}(sg_{ik}^n)$ и $F_{txt}(sg_{ik}^n)$ больше единицы.

Итеративный поиск изоморфных подграфов с числом вершин $n + 1$ ($n > 2$) реализуется только при обнаружении для пары текстов хотя бы одного изоморфизма на n вершинах. Из анализа исключаются вершины, не вошедшие в изоморфные подграфы с n вершинами, они не могут входить ни в какой изоморфный подграф размера $n + 1$: в этом случае они бы содержались в его вложенном подграфе размера n , также изоморфном для двух текстов.

3. Постобработка результатов.

1) Как было отмечено, алгоритм VF2 выявляет изоморфизмы, то есть графы одинаковой структуры (без учёта имен вершин). На данном этапе исследования при анализе аргументационных графов мы ограничиваемся рассмотрением случая точного совпадения имен вершин-схем. Исследование одинаковых конфигураций, образуемых изоморфными подграфами с разными именами вершин, несомненно, представляет интерес и будет проведено в следующих работах. Постобработка изоморфных подграфов осуществляется с целью фильтрации подграфов с несовпадающими именами вершин.

2) Фильтрация вложенных подграфов проводится по наборам приёмов, встречающихся в одних и тех же текстах (наборы текстов, реализующих эти приёмы, полностью совпадают), через попарные проверки вложенности меньших приёмов в большие (по мере увеличения числа вершин).

Если проводится жанровый анализ текстов, то процедура построения совместного частотного спектра выполняется для каждого подмножества A_r отдельно, а затем рассматривается пересечение отдельных частотных спектров: $\Phi(A) = \bigcap_{r=1, \dots, R} \Phi_r(A_r)$. Вычисленные абсолютные и текстовые частоты позволят различать приемы из $\Phi^{\text{In}}(A)$ и $\Phi^{\text{Ex}}(A)$.

2.5 Методы оценки убедительности аргументации

2.5.1 Анализ рассуждений по выраженности логоса, пафоса, этоса

Традиционно (согласно канонам Аристотеля) считается, что аргументация проявляется в трех аспектах: интеллектуальном (логос), эмоциональном (пафос) и этическом (этос). Три перечисленных аспекта характеризуют самостоятельное содержание аргументации, рассматриваемое отдельно от контекста её реализации. Связь выстраиваемой аргументации с прагматическим контекстом, с конкретной коммуникативной ситуацией отражается в кайросе, аспекте её своевременности (уместности, степени соответствия содержания доводов моменту их приведения). Соответственно, анализ кайроса в коллекции текстов требует разностороннего учёта экстралингвистической специфики каждого текста: к примеру, для научной статьи полезно охарактеризовать как степень её соответствия тематике конкретного журнала, в котором эта статья опубликована, так и своевременность описанного исследования в более широком контексте состояния соответствующей научной области на момент публикации работы (акцентирует ли автор применение современных передовых методов или довольствуется традиционными; ссылается ли на недавние или классические труды других исследователей; как характеризует актуальное состояние предметной области,

вклад в неё проделанной работы и перспективы прикладного использования полученных результатов). Желательно также учитывать временной разрыв между проведением исследования и публикацией статьи, причём этот разрыв может значительно различаться для разных работ. Таким образом, многосторонний анализ кайроса представляется значительно более сложным по сравнению с анализом трёх аспектов собственного содержания аргументации, в особенности для формальных методов, в связи с чем он не проводится в представленном исследовании.

Поскольку логос, пафос и этос определяют разные механизмы воздействия, аргументы разделяют на три типа: логические, эмоциональные и этические. Главная составляющая убеждения – логическая, но важна и подача аргументов, сочетающая логические рассуждения с апелляциями к этике и эмоции. Учет механизмов воздействия актуален для построения текстов, рассчитанных на разные аудитории.

В частности, современный анализ профессиональных текстов показывает, что границы между текстами научного и научно-популярного жанров стираются. В работе [Ilynska, Platonova, Smirnova, 2016] демонстрируется проявление этой тенденции на текстах по архитектуре. Авторы объясняют это явление многообразием рассматриваемой проблематики. В силу чего «помимо передачи информации и убеждения читателей многие специализированные тексты также реализуют экспрессивные функции для привлечения внимания к представленным сведениям» [Там же, с. 33].

Количественная оценка убедительности текстов представляет интерес как для анализа уже созданных текстов, так и для построения планируемых. Убедительность текстов может приравниваться к их «логичности», степени истинности вывода. Например, в работе [Zagorulko et al., 2020] количественная оценка убедительности текста подсчитывается через итерационную процедуру путем вычисления весов аргументов по графу аргументации с помощью операций нечеткой логики. Начальные веса задаются исходя из знаний об аудитории.

Для синтеза текстов важен выбор стратегии, учитывающей не только логику, но и другие механизмы влияния. Авторы работы [El Baff et al., 2019] строят модель стратегии, которая опирается на каноны риторики, сформулированные Аристотелем. Оценка модели проводится экспертами. Модель реализуется в три этапа, выполняемые вручную. Для каждого заданного тезиса выбираются аргументативные дискурсивные единицы (ADU) из заранее сформированной базы. Из них составляются структуры аргументации (цепочки из посылок, заключений). Выбранные утверждения оформляются в определенном стиле согласно принимаемой стратегии. Под стратегией в указанной работе понимается пропорция используемых утверждений, содержащих логос, этос, пафос (например, в соотношении 70 : 20 : 10). Варианты текстов, составленных разными экспертами по определенной заранее стратегией, сближаются.

Авторы работы [Lukin et al., 2017] сравнивают влияние логоса и пафоса на аудиторию американских социальных сетей. Они анализируют восприимчивость разных типов личности к разным способам убеждения (посредством фактов либо эмоционального воздействия). Для выделения личностных типов авторы применяют пятифакторную модель черт характера (открытость опыту, нейротизм, экстраверсия, доброжелательность, сознательность). Для нейтральных аргументов (без эмоционального воздействия либо интенсивной апелляции к фактам) демонстрируется равная эффективность со всеми типами личности. Сознательные люди оказываются наиболее восприимчивы к эмоциональным аргументам, тогда как доброжелательные сильнее других склонны соглашаться с интенсивным перечислением фактов.

Наконец, распознавание этоса и пафоса также представляет интерес для компьютерного анализа текстов разных жанров с целью установления успешности коммуникации. В работе [Duthie, Budzynska, Reed, 2016] показано, что основные события и тенденции в политическом ландшафте отражаются или предвосхищаются путем анализа парламентского отчета (UK parliamentary record, Hansard, 70 тыс. слов, 253 спикера) с помощью этос-аналитики. Распознавание наличия и полярности выражений этоса, проведенное с помощью методов

машинного обучения при использовании в качестве признаков n -грамм, содержащих лексикон этоса ($n = 1, 2, 3$), улучшает F -меру на 10–20% в зависимости от используемого алгоритма машинного обучения.

2.5.2 Анализ методов убеждения на основе маркеров и схем

Оценка методов убеждения (пафос, этос) в представленной работе строится на выделении в утверждениях с аргументацией соответствующей лексики (маркеров пафоса и этоса) и учете доли каждого из методов в схемах. Оценки для схем задаются исходя из функционально-сопоставительного анализа специфики их реализации авторами научных статей.

2.5.2.1 Формирование словарей пафоса и этоса

Для создания словарей лексики, маркирующей пафос и этос в тезисах, заданы критерии, по которым эти словари сформированы.

Будем считать языковое выражение маркером *пафоса*, если оно соответствует хотя бы одному из следующих критериев:

1) *избыточности*: языковое выражение не дополняет излагаемые сведения ни в смысловом содержании, ни в организации, а при его удалении не изменится ни содержание тезиса, ни его выражаемая связь с другими утверждениями (например, частицы *лишь, же, даже, только, вовсе* и др.);

2) *деонтической модальности*: выражение содержит сему долженствования (*нужно, требуется, должен, вынужден* и т.д.);

3) *эксплицитной семантики оценки с отражением высокой интенсивности*: выражение может заменяться на нейтральный синоним (без явной оценки либо с меньшей интенсивностью) без изменения смысла утверждения (*изумительный, чудовищный, потрясающий, прекрасный* и т.п.).

Следует отметить, что отбор маркеров в словари осуществлён с учётом специфики их реализации преимущественно в статьях двух анализируемых

областей (лингвистики и компьютерных технологий). Некоторые из применяемых маркеров могут обладать специализированным значением, не связанным с методами убеждения, в иных предметных областях. Например, конструкция «тогда и только тогда» является традиционной и нейтральной структурирующей связкой в математических определениях. Тем не менее, её реализация в работах по лингвистике или информационным технологиям не является типичной для этих двух областей, а потому может выражать в тексте компонент пафоса.

Этос в данном исследовании предлагается понимать шире, чем подразумевает определение Аристотеля: как обоснование тезиса через его поддержку некоторым авторитетом. Этим авторитетом способен являться конкретный индивид (например, исследователь или специалист), обобщённая группа (коллектив учёных), общественное мнение либо сам автор текста (например, в случае иллюстрации тезиса произвольным примером, который выбирается автором для повышения убедительности рассуждений: релевантность примера как частного случая к общей ситуации поддерживается в том числе авторитетом исследователя и фактом намеренного предпочтения конкретного примера из многих доступных). Такое расширение понятия этоса обуславливается тем, что доказательства тезисов от авторитета не принадлежат в общем случае логосу или пафосу, но при этом распространены в научных публикациях и близких им научно-популярных работах. Соответственно, маркеры этоса определяются через критерий авторизации: языковое выражение считается маркером, если указывает на источник сведений (*автор считает; по мнению эксперта; в работе утверждается, что; примером может служить* и т.п.). Маркерами считаются и ссылки на литературу, представленные обычно в квадратных или круглых скобках.

Словарное представление маркеров построено на базе регулярных выражений в виде шаблонов, применяемых непосредственно для поиска. Для каждого маркера указаны его грамматические характеристики (согласно меткам в морфологическом анализаторе Rymorphy2) и необходимый контекст (лексический, пунктуационный). При необходимости объемный перечень

вариантов (в примере ниже они записаны через разделитель «|») задается списком. Допустимы ограниченные пропуски в последовательности выявляемых элементов маркера. Проиллюстрируем структуру шаблонов на примере двух следующих маркеров:

< V^P >: [(единица|десятка|сотня|тысяча|миллион) // plur] [– // gent].

< V^E >: [по // PREP] [...] [данные // datv]

Первый маркер указывает на выражение пафоса и состоит из двух элементов, первый из которых задаётся и лексическими, и грамматическими константами (допускает словоформы множественного числа для пяти лемм), тогда как второй обозначает любое слово в родительном падеже (иные грамматические признаки, такие как часть речи или род, не являются значимыми для этого элемента). Данный маркер соответствует эмфатическому обращению к численности неких сущностей при контекстуальном подчёркивании их изобилия либо, наоборот, практического отсутствия. Второй маркер отмечает выражение этоса и допускает вставку произвольного количества словоформ (вплоть до длины тезиса) между первым и третьим элементами (где оба соответствуют строго заданной лексической константе с заданным грамматическим признаком). Он свидетельствует о ссылке на особый вид данных, уточняемый либо одиночным прилагательным (по новым/актуальным/последним данным), либо более сложной многословной конструкцией (недавно подтверждённым, экспериментально полученным и т.п.).

Словари маркеров этоса и пафоса в представленном исследовании собраны специально для анализа аннотированного корпуса научных статей (с целью их как можно более полного покрытия). Построение шаблонов осуществлено экспертно через итеративное компьютерное распознавание конструкций этоса и пафоса на множестве всех утверждений корпуса: после каждого запуска процедуры компьютерного распознавания вручную проверялись как новые случаи срабатывания маркеров (для выявления ошибочных конструкций, не выражающих пафос либо этос, и уточнения шаблонов), так и тезисы без

выявленных элементов проверяемого метода убеждения (для обнаружения ложных пропусков, добавления новых шаблонов).

2.5.2.2 Оценка моделей рассуждения по влиянию пафоса и этоса

Сравнительный функциональный анализ схем аргументации позволяет дать количественную оценку представленности отдельного метода рассуждения (пафос, этос, логос) в каждой из них в диапазоне $[0,1]$. Количественные характеристики наиболее частотных моделей приведены в таблице 4.

Таблица 4 – Количественная оценка доли методов рассуждения в схемах

Схема	Логос	Пафос	Этос
CauseToEffect	0.8	0.2	0
<i>Example</i>	0	0.33	0.66
<i>ExpertOpinion</i>	0	0	1
NegativeConsequences	0.66	0.33	0
<i>PopularOpinion</i>	0	0.33	0.66
<i>PositionToKnow</i>	0	0.2	0.8
PositiveConsequences	0.66	0.33	0
PracticalReasoning	0.75	0.25	0
VerbalClassification	1	0	0

Определение количественных оценок произведено за счёт сопоставления моделей рассуждения с их группировкой по смысловой близости и последующим сравнением похожих схем в пределах одной группы. Так, по способу доказательства все схемы можно разделить на две группы: выстраивающие рассуждение в пределах содержания утверждений (указаны в таблице обычным шрифтом) либо же через привлечение внешнего авторитета (выделены курсивом). Соответственно, в схемах первой группы сильнее выражен компонент логоса, в

схемах второй – этоса. При этом для отдельных схем основной компонент может дополняться иными методами рассуждения с разной степенью влияния: так, пропорции трёх методов частично отражают смысловые различия между похожими схемами. Иными словами, хотя точные количественные показатели схем являются условными, они применяются не для подсчёта абсолютных значений в отдельных текстах, а для сравнения текстов по соотношению методов. Веса схем одинаковы для всех текстов, они отражают их ранжирование по выраженности методов убеждения.

Так, среди схем с преобладанием логоса его наибольшая доля отмечается для модели *VerbalClassification* (с воспроизведением структурных отношений между элементами классификации). Близкие *VerbalClassification* причинно-следственные связи в *CauseToEffect* и *CorrelationToCause* (обратные по направлению) допускают большую вариативность в организации тезисов, что позволяет дополнительное проявление пафоса. Такое эмоциональное воздействие выражается более явно при рассуждении от поставленной цели (*PracticalReasoning*) и оказывается особенно сильным при эксплицитной оценке последствий (*Positive / Negative Consequences*).

При доказательстве от авторитета этос выражается наиболее полно при апелляции к конкретному эксперту (*ExpertOpinion*). Если у авторитета отсутствует профессиональная квалификация, то обоснование доверия к нему может привести компонент пафоса (как для *PositionToKnow*). Убедительность авторитета становится менее явной при отсылке к обезличенному источнику (в *PopularOpinion*), что пропорционально усиливает эмоциональность аргумента и приближает его к доказательству через пример (когда авторитет проявляется не через суждение, а через релевантность некоего явления для утверждения).

Назначение числовых значений для схемы отражает предполагаемое соотношение между компонентами трёх методов убеждения с учётом корреляции этих методов между разными схемами (на основе сопоставительного анализа схем по функциональной специфике их применения). Так, и *PositionToKnow*, и *PopularOpinion* выражают преобладающий компонент этоса со вторичным

влиянием пафоса. Но такое влияние становится более значимым при обращении к неопределённому источнику. Соответственно, для этой схемы отношение пафоса к этосу определено как 1 : 2, тогда как для *PositionToKnow* – 1 : 4. Аналогичное распределение применяется для разграничения *CauseToEffect* и, в первую очередь, *PracticalReasoning*, а при усилении разницы, *Positive/Negative Consequences*.

2.5.3 Оценка выраженности методов убеждения в целостных текстах

Оценка механизмов воздействия, использованных в тексте, предполагает рассмотрение аргументационной разметки текста как множества отдельных аргументов: множества информационных вершин), связанных через вершиную схему, без учета их взаимовлияния в общей структуре рассуждений. Для механизмов этоса и пафоса это влияние не так значимо ввиду их выражения в пределах одного аргумента, а соответственно малой зависимости от внешнего контекста связываемых утверждений (в отличие от логоса, при котором убедительность доказательства зависит от убедительности каждой из посылок, убедительность которых, в свою очередь, определяется тезисами в их поддержку). Такое допущение позволяет оценивать неразмеченные тексты с распознанными в них тем или иным способом (например, с применением шаблонов маркеров аргументов или машинного обучения) утверждениями и моделями рассуждения без построения всей аргументационной структуры текста. Последнее является отдельной сложной задачей, как на уровне экспертной разметки, так и для компьютерного распознавания.

Пусть A – аннотация текста из коллекции, трактуемая как множество, содержащее k аргументов: $A = \{\arg^k\}$. Каждый аргумент может быть представлен тройкой $\arg^k = \langle p^k, sch^k, c^k \rangle$. Обозначим $M = \{E, P\}$ – исследуемые методы (этос, пафос). Будем считать, что присутствие элемента $m \in M$ в любой из составляющих аргумента усиливают его. Количественная составляющая для каждого метода $m \in M$ в заключении аргумента определяется числом утверждений, содержащих соответствующую лексику:

$$Q^m(c^k) = 1, \text{ если утверждение содержит лексику метода } m \\ \text{(задается словарем маркеров } V^m) \\ = 0, \text{ если такая лексика отсутствует.}$$

Подобным образом вычисляем представленность метода m в посылках:

$Q^m(p^k) = \sum_{i=1}^I Q^m(p_i^k)$, где I – число посылок в аргументе, $Q^m(p_i^k)$ вычисляется по правилу (1), где вместо $Q^m(c^k)$ подставляется значение $Q^m(p_i^k)$.

Количественная оценка метода m в схеме назначается по таблице 4 ($T = \{t_{ij}\}$) согласно индексу столбца m ($m = 2, 3$) и строки l , идентифицирующей схему sch^k : $Q^m(sch^k) = t_{lm}$. Тогда оценка аргумента имеет вид:

$$Q^m(arg^k) = 0, \text{ если } Q^m(p^k) = Q^m(c^k) = Q^m(sch^k) = 0, \\ = (k_p Q^m(p^k) + k_c Q^m(c^k) + 1) * (Q^m(sch^k) + 1) \text{ иначе.}$$

Коэффициенты k_p и k_c применяются для сглаживания весов утверждений, реализуемых в одинаковой роли в разных аргументах (когда одна посылка поддерживает несколько заключений либо одно заключение доказывается через параллельную реализацию нескольких моделей рассуждения). $k_p = 1 / F_p$, $k_c = 1 / F_c$, где F_p соответствует числу аргументов, задействующих утверждение как посылку, F_c – включающих его в роли заключения.

Оценка текста задается формулой:

$$Q^m(A) = \sum_{k=1}^K Q^m(arg^k), \text{ где } K \text{ – число аргументов в тексте.}$$

Величины $Q^m(A)$ для разных текстов сравнимы только для одного и того же метода m и при условии нормализации на число утверждений в тексте A . Нас интересует пропорция выраженности методов $Q^P: Q^E$, определяемая отношением:

$Q^P(A) / (Q^E(A) + \alpha)$, где α – малая техническая константа (например, 0,001), которая не позволяет знаменателю дроби обратиться в ноль.

Проиллюстрируем процесс интеграции оценок отдельных аргументов на примере подсчёта убедительности для фрагмента разметки, представленного на Рис. 2. Фрагмент содержит четыре аргумента (A_{19} , A_{18} , A_9 , A_{16}), которые сперва оцениваются поочерёдно, независимо друг от друга. Синие и красные рамки внутри тезисов обозначают маркеры этоса и пафоса, выявляемые по словарю.

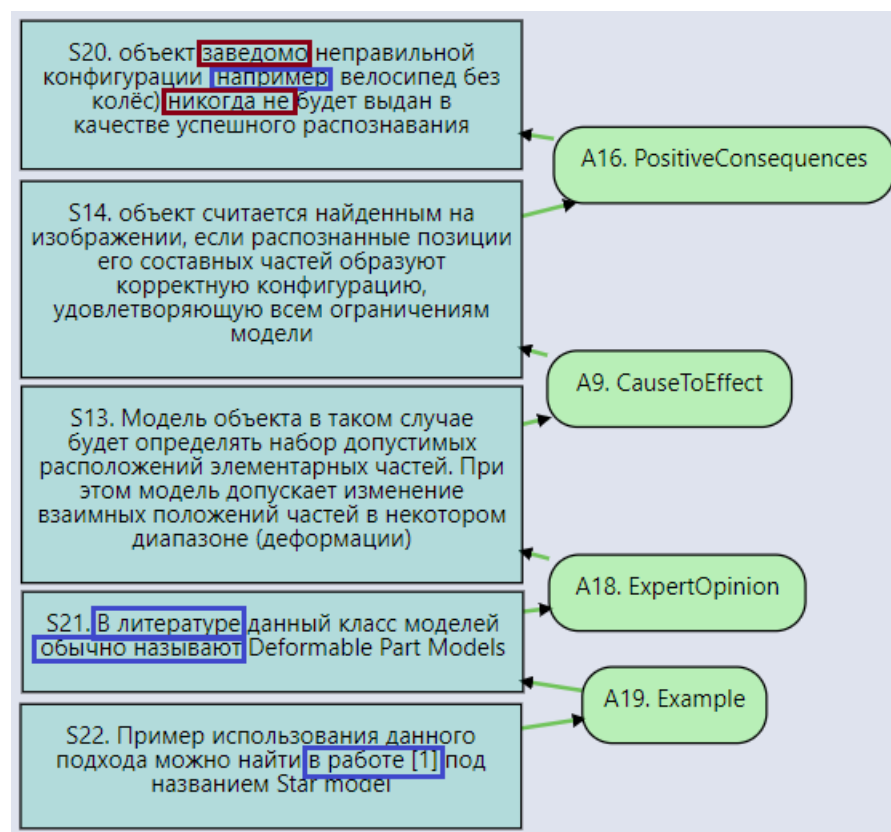


Рисунок 2 – Пример фрагмента графа для оценки этоса и пафоса

Для аргумента A19 маркеры этоса выявлены и в его посылке, и в заключении, тогда как ни один из тезисов не содержит маркеров пафоса. Реализованная модель рассуждения, *Example*, считается выражающей этос в два раза сильнее, чем пафос (соответствующие значения равны 0.66 и 0.33). Согласно формуле (2), значение этоса для A19 достигает $(1 + 1 + 1) * (0.66 + 1)$, или 4.98, а значение пафоса равняется $(1 + 0 + 0) * (0.33 + 1)$, или 1.33. Аналогично значения этоса и пафоса для A18 достигают, соответственно, 4 $[(1 + 0 + 1) * (1 + 1)]$ и 0 (ни в посылке, ни в заключении нет маркера пафоса, тогда как схема *ExpertOpinion* рассматривается как образцовое выражение этоса). Эти же значения для A9 равны 0 и 1.2 $[(0 + 0 + 1) * (0.2 + 1)]$, а для A16 – 2 $[(0 + 1 + 1) * (0 + 1)]$ и 2.66 $[(0 + 1 + 1) * (0.33 + 1)]$. Суммирование значений по всем четырём аргументам выдаёт 10.98 весом этоса, 5.19 весом пафоса. Итоговое соотношение $Q^P : Q^E$ равняется 0.47, ввиду чего проанализированный фрагмент признаётся выражающим этос в два раза более интенсивно, чем пафос.

2.6 Кластеризация текстов по сложности аргументации

Моделирование аргументационных структур в виде графов обеспечивает возможность формального выявления текстов с похожей организацией рассуждений. Такая проблема сводится к задаче *кластеризации* графов, где под кластеризацией понимается выявление групп похожих объектов среди заданного фиксированного множества [Батура, 2016, с. 113]. Объектами в нашем случае служат аргументационные графы, представляемые векторами их отличительных признаков. Такие признаки отражают специфику графов как в структурном, так и содержательном аспектах (например, через указание средней длины путей рассуждения, среднего объёма информационных вершин в словах). Построение векторов для обрабатываемых объектов на основе отобранных признаков позволяет сгруппировать эти объекты по степени сходства через применение алгоритмов кластеризации. Эффективность работы отдельных алгоритмов зависит от структуры коллекции объектов (к примеру, от возможности выделения центров кластеров, то есть эталонных объектов для выделяемых групп, либо от наличия уникальных объектов, значительно отличающихся от любых других). В этой связи практическое решение задачи кластеризации может предполагать применение нескольких алгоритмов разных типов (пример классификации доступен в [Батура, 2016, с. 131–139]) и сравнение результатов методами оценки качества кластеризации.

2.6.1 Традиционные показатели сложности текстов

В представленном исследовании решается задача кластеризации текстов по сложности аргументации. Эта задача является частной по отношению к проблеме оценки сложности текстов, где сложность определяется как доступность текстов для понимания и лёгкость их восприятия [Солнышкина, Кисельников, 2015] или как сумма всех элементов текста, влияющих на понимание темы текста, скорость его прочтения и уровень интереса к нему [Лапошина, 2017].

Как правило, наборы таких элементов включают единицы разных уровней языка, от фонетического до семантического, причём они могут дополняться экстралингвистическими показателями. Примерами таких расширений выступают отслеживание кожно-гальванической реакции испытуемых при чтении текста (эта реакция свидетельствует о степени активации головного мозга, интенсивность чего зависит от сложности текста) [Самсонов, Чмыхова, Давидов, 2015] либо анализ Колмогоровской сложности текста (сравнение объёма исходного текста с объёмом его сжатой версии после обработки архиватором, то есть определение длины описания этого текста на мета-языке). В частности, компьютерный анализ текста по характеристике Колмогорова применяется при оценке сложности текстов для перевода [Кутузов, 2009] и классификации текстов [Ryabko, Gus'kov, Selivanova, 2017].

Тем не менее исследования текстовой сложности почти не затрагивают более сложные лингвопрагматические явления, такие как риторическая или аргументационная организация текста, несмотря на их значимость для восприятия и понимания сообщаемых сведений. Причиной является трудоёмкость обработки структур прагматического уровня языка. Вместе с тем исследования собственно аргументационной составляющей текстов традиционно обращаются к оценке их убедительности, а не сложности. Примером такого оценивания убедительности русскоязычных научно-популярных текстов по графу аргументации может служить работа [Zagorulko et al., 2020]. Результирующая оценка вычисляется путем суммирования весов утверждений и схем рассуждений на основе правил нечеткой логики согласно связям в графе (от листовых вершин к главному тезису), а начальные веса задаются исходя из знаний об аудитории.

Задача анализа аргументационной сложности текстов может решаться через выявление текстов со сходной организацией рассуждений, их группировку и ранжирование групп по сложностным признакам. Формально она представляется как задача анализа коллекции аргументационных аннотаций $A = \{a_i\}$ ($i = 1, \dots, N_a$) с целью разбиения их на группы $\{c_j\}$ ($j = 1, \dots, K_c$) таким образом, чтобы внутри одной группы оказались аннотации, близкие по сложности аргументации, а

различающиеся попадали в различные группы. Обучающее подмножество множества A с заранее известными метками $C = \{c_j\}$ кластеров отсутствует. Задача сводится, во-первых, к выбору множества признаков (приемов аргументации и характеристик всей аргументационной структуры текста), по которым проводится кластеризация, во-вторых, к поиску множества C , каждый элемент которого содержит близкие по сложности аргументации тексты и оценке качества кластеров, а в-третьих, к интерпретации полученных кластеров.

2.6.2 Кластеризация аргументационных графов

Как обозначено ранее, эффективность работы различных алгоритмов кластеризации зависит от структуры обрабатываемой коллекции (причём её специфика заранее неизвестна: например, наличие в коллекции уникальных объектов, не похожих ни на один другой, выясняется только на этапе кластеризации). В этой связи в исследовании применены три алгоритма разных типов, функционирование которых основывается на различных принципах: центроидный *K-means*, иерархический *Ward* и графовый *Spectral*. Программные реализации этих алгоритмов взяты из библиотеки Scikit-learn для языка Python.

K-means делит набор N_a аннотаций из A на K_c непересекающихся кластеров, каждый из которых описывается центроидом (μ_j) этого кластера. Алгоритмом выбираются центроиды, которые минимизируют сумму квадратов расстояний от аннотаций a_i до центроидов μ_j , к которым они относятся. В используемой реализации алгоритма *k-means++* [Arthur, Vassilvitskii, 2007] инициализация начальных центроидов производится из максимально удаленных a_i . Для их нахождения используется 10 попыток. Шаги алгоритма (не более 300) повторяются, пока перемещение центроидов больше заданного значения (0.0001).

Алгоритм *Ward*, как и *K-means*, минимизирует сумму квадратов разностей во всех кластерах, но решение ищется с помощью агломеративного иерархического подхода [Ward, 1963]. Объединение кластеров на каждом шаге

производится на основе минимизации дисперсии объединяемых кластеров. Останов алгоритма происходит по достижению заданного числа кластеров (в нашем случае их четыре).

Для работы алгоритма спектральной кластеризации *Spectral* [Luxburg, 2007] определяется матрица схожести аннотаций: $Adj_{i,j} = - ||x_i - x_j||^2$. Эта матрица описывает полный граф с вершинами-аннотациями и рёбрами между каждой парой аннотаций с весом, соответствующим степени сходства этих вершин. Процесс кластеризации заключается в решении «Normalized cuts problem», а именно разделении получившегося графа на две части так, чтобы аннотации в двух графах были в среднем более похожи на другие внутри получившейся той же части графа, чем на аннотации в другой части.

2.6.3 Оценка качества кластеризации

Выбор оптимальной кластеризации текстов из вариантов их группирования разными алгоритмами на разных наборах признаков предполагает оценку качества получаемых кластеров в каждом отдельном случае. Для подсчёта такой оценки по разным аспектам кластеризации в исследовании совмещаются три меры: основанный на коэффициенте Жаккарда (*Jaccard-Based*) подсчет суммарного числа совпадений в элементном составе кластеров и две внешние меры, *V-measure* и *FM-score* (*Fowlkes-Mallows*), которые определяют надёжность гипотезы об эталонности одного из вариантов кластеризации через сопоставление разных вариантов.

При подходе *Jaccard-Based* проводится попарное сравнение наборов кластеров, построенных разными алгоритмами. Сравнение двух наборов кластеров производится в цикле по каждому кластеру (сперва из одного набора, затем из второго) следующим образом: 1) кластер из одного набора сравнивается с каждым кластером из другого набора: число аннотаций в пересечении кластеров делится на число аннотаций в их объединении; 2) для данного кластера определяется и фиксируется наибольшая мера близости с кластерами другого

набора. По завершении цикла строится массив из наибольших мер близости по отдельным кластерам. Для построенного массива подсчитывается средний показатель сходства кластеров, который и определяется как итоговое значение близости кластеров, построенных двумя алгоритмами.

V-measure представляет собой гармоническое среднее *однородности* (кластер содержит только a_i одного класса) и *полноты* (все a_i класса относятся к одному кластеру). Данные меры формально определяются с использованием функций энтропии и условной энтропии: $h = 1 - H(C|K)/H(C)$, $c = 1 - H(K|C)/H(K)$, где H – функция подсчёта энтропии, K – результат кластеризации, C – эталонное разбиение на классы согласно гипотезе (при двух сравниваемых вариантах кластеризации поочерёдно каждый из них принимается за эталонный). Итоговое значение оценки подсчитывается как $V = 2hc/(h+c)$.

FM-score определяется как среднее геометрическое для попарной точности и полноты: $FM = TP/(TP+FP)(TP+FN)$, где TP – True Positive, то есть количество a_i , принадлежащих одним и тем же кластерам как в разбиении, которое мы принимаем за эталонное, так и в том, которое мы считаем предсказанными; FP – False Positive, то есть количество a_i , принадлежащих к кластерам с эталонными метками, но не предсказанными, а FN является количеством ложно-отрицательных (False Negative) случаев, когда a_i не принадлежат к кластерам с эталонными метками, но в предсказании считаются к ним принадлежащими.

Выводы

В Главе 2 диссертационного исследования представлены методы формального анализа аргументации в рамках последующих экспериментов. Описание методов включает как обзор существующих подходов к решению отдельных задач в области компьютерной обработки аргументации (в том числе анализ применимости современных больших языковых моделей), так и уточнение алгоритмов, реализованных в диссертационном исследовании.

В частности, в Главе 2 обозначаются принципы разметки аргументации, дополняющие базовые правила AIF с учётом жанровой специфики доказательств в научных статьях. На этих принципах разработана методология многоаспектной разметки аргументации, которая применена в подготовке аргументационного корпуса, обрабатываемого в экспериментальной части диссертационного исследования. Рассматриваются методы частичной автоматизации аргументационной разметки, где основное внимание уделяется компьютерному выявлению тезисов (по аналогии, и отдельных приёмов) через обучение классификаторов тремя алгоритмами (наивный байесовский классификатор, MNB; метод опорных векторов, SVM; многослойный перцептрон (базовая нейронная сеть) MLP) и векторное представление тезисов через статистическую фильтрацию незначимых для аргументации лемм (посредством компьютерного морфологического анализа, подсчёта мер $TF*IDF$, *Mutual Information* и критериев χ^2 и Variance на основе встречаемости различных лемм в аргументативных и неаргументативных сегментах). Описаны методы подсчётов коэффициентов согласия аннотаторов для оценки надёжности разметки: представлены как стандартно применяемые каппа Козна и альфа Криппендорфа (с уточнением специфики их подсчёта на трёх уровнях разметки), так и коэффициент аргументационного соответствия, специально разработанный в ходе исследования с учётом многоуровневой специфики моделирования аргументации.

Во второй главе представлены и методы выявления **составных приёмов аргументации**, определяемых как применяемые авторами разных текстов структуры из нескольких элементарных приёмов (отдельных аргументационных схем). Обнаружение составных приёмов аргументации достигается методами автоматического извлечения изоморфных подграфов из аргументационных графов (изоморфные подграфы соответствуют идентичным подструктурам в полных графах).

Две другие группы методов ориентированы на переход от анализа приёмов к **стратегиям аргументации** (конструкциям высшего уровня, образуемым сочетанием отдельных приёмов аргументации и общими структурно-

содержательными особенностями организации рассуждений на протяжении целостных текстов). Описываются методы компьютерной оценки выраженности **методов убеждения** (трёх общих аспектов аргументационного воздействия: через рациональные доказательства, влияние на эмоции, апелляции к авторитетам) посредством анализа аргументационных схем и лексико-грамматических маркеров. Представлены методы выявления текстов со сходными аргументационными структурами и их ранжирования по формальной сложности аргументации (через интеграцию показателей сложности разных типов). Эта задача решается применением алгоритмов кластеризации трёх функциональных типов (центроидного *K-means*, иерархического *Ward* и графового *Spectral*) и формальных оценок качества кластеризации (*V-measure*, *Fowlkes-Mallow score*, оценки на основе коэффициента Жаккара).

Представленные в Главе 2 методы используются в экспериментах, описываемых в третьей и четвёртой главах диссертации.

Глава 3. Выявление элементов аргументационных стратегий в графах аргументации

В Главе 3 рассматриваются эксперименты по формальному выявлению элементов аргументационных стратегий (приёмов аргументации и отдельных тезисов) в аннотированных аргументационных структурах научных статей. Под **приёмами аргументации** в данном исследовании понимается применение повторяющихся отдельных моделей рассуждения и образуемых ими повторяющихся структур (подграфов) аргументации, реализуемых в различных научных статьях разными авторами. Соответственно, приёмы аргументации выражаются конструкциями ограниченного размера в полной структуре аргументационных схем, связывающих тезисы отдельного текста.

В соответствии с обозначенной задачей в разделе 3.1 описывается анализ элементарных приёмов аргументации, состоящих из одной схемы. Такой анализ предполагает проведение аргументационной разметки на корпусе текстов. Приведены значения коэффициентов согласия разметчиков для построенного корпуса, подсчитанные на всех уровнях (выявления тезисов, связей, схем).

В разделе 3.2 экспериментально оценивается компьютерное распознавание тезисов и его применимость в частичной автоматизации разметки. В продолжение выявления утверждений представлены дальнейшие эксперименты по извлечению аргументационных связей и аргументов отдельных типов из неразмеченных текстов (по аналогии, методами бинарной классификации).

3.1 Выявление элементарных приёмов аргументации

Моделирование аргументации в целостных текстах согласно стандарту AIF подразумевает выделение аргументов (наборов тезисов, оформленных средствами естественного языка и связанных через реализацию типовой модели рассуждения) как составных элементов общей структуры. Такой подход позволяет эксплицитно представить в аргументационном графе элементарные приёмы аргументации,

образованные одной аргументационной схемой. Соответственно, задача выявления подобных приёмов сводится к проведению аргументационной разметки на корпусе текстов, причём выстраиваемые в результате аргументационные графы применимы и для решения многих иных задач формального анализа аргументации, напрямую не связанных с обнаружением приёмов.

3.1.1 Характеризация аргументационного корпуса научных статей

Разметка аргументации в текстах проведена тремя экспертами в соответствии со стандартом AIF, в частности, с выделением детализированных моделей рассуждения на основе компендиума Д. Уолтона, посредством использования специализированного инструмента для аннотирования аргументации ArgNetBank Studio [Сидорова и др., 2020].

Размеченные тексты характеризуются объёмом от 800 до 2000 слов, ориентированы на доказательство одной центральной идеи (общего вывода, обосновываемого на протяжении научной статьи). Представленные статьи относятся к двум предметным областям: лингвистике (преимущественно корпусной лингвистике и лексической семантике) и компьютерным технологиям (различным прикладным разделам, таким как информационная безопасность, виртуальное моделирование, обработка изображений и др.). Выбранные для разметки статьи опубликованы в научных изданиях, что указывает на достижение ими академического уровня, достаточного для прохождения рецензирования.

Согласованность аргументационной разметки обеспечена инструкцией для аннотаторов, которая содержит набор вспомогательных правил для определения главного тезиса текста, выявления связанных утверждений и установления аргументационной схемы, реализуемой в их связи (с помощью дерева вопросов). Ключевые принципы инструкции и дерево вопросов представлены в главе 2, разделе 2.1.

Пример фрагмента аргументационной разметки текста приведен на Рис. 3. Вершины-утверждения изображены прямоугольниками, вершины-схемы – эллипсами. Утверждения и схемы обозначены символами S и A соответственно и снабжены номерами (согласно линейному порядку утверждений в тексте). Связи между аргументами на уровне общих тезисов позволяют объединить их в общую структуру текста, причём на рисунке представлены разные типы таких связей: параллельные (Convergent Arguments, через разные аргументы) (S33, S35 → S32), сцепленные (Linked Arguments, внутри одного аргумента) (S24 & S30 → S34), расходящиеся (Divergent Arguments) (S24 → S34, S35), последовательные (Sequential Arguments) (S24 → S35 → S32).



Рисунок 3 – Фрагмент аргументационной разметки текста

Пример полной аргументационной разметки научной статьи (с представлением аргументов, иллюстрацией структуры аргументационного графа, указанием неаргументативного содержания) представлен в приложении Г.

3.1.2 Оценки согласия аннотаторов на размеченном корпусе

Значения коэффициентов согласия аннотаторов на трёх уровнях разметки для построенного корпуса представлены в таблицах 5 и 6. Подсчёт значений

осуществлѐн в соответствии с принципами, указанными в разделе 2.3. Уровень схем представлен двумя строками: собственно детализированные схемы и их обобщѐнные функциональные группы, которые представлены в разделе 4.4.

Таблица 5 – Детализированное согласие аннотаторов на разных уровнях разметки

Уровень	КАС	Каппа Коэна	Альфа Криппендорфа
Утверждения	0.87	0.30 – 0.60	0.51
Связи	0.54	0.49 – 0.67	0.66
Схемы	0.52	0.29 – 0.46	0.43
Функциональные группы схем	0.71	0.45 – 0.58	0.55

Таблица 6 – Согласие аннотаторов после автоматической модификации разметки

Уровень	КАС	Каппа Коэна	Альфа Криппендорфа
Утверждения	0.93	0.58 – 0.76	0.73
Связи	0.54	0.49 – 0.68	0.67
Схемы	0.78	0.51 – 0.66	0.65
Функциональные группы схем	0.87	0.66 – 0.77	0.71

Таблица 5 содержит показатели для исходной версии итогового корпуса в 100 коротких научных статей (версии, сохраняющей стилистические особенности разметки аннотаторами), таблица 6 – для версии, полученной на основе автоматической модификации типовых случаев расхождения. Процедура автоматической модификации (полностью проводимой посредством компьютерной обработки) включала такие шаги, как удаление уникальных листовых утверждений (которые выявлены лишь одним аннотатором и не поддерживаются другими утверждениями в его варианте разметки, а при этом отнесены другим аннотатором к неаргументативной части) и фиксирование одной

из спорных схем, выделенных аннотаторами для одного аргумента, согласно иерархии правил, отражающей смысловую специфику этих схем. Подробное описание процедуры доступно в статье [Pimenov, Salomatina, 2024b].

Раздельный подсчёт КАС по текстам корпуса на всех уровнях разметки позволяет выявлять тексты с низким согласием для их корректировки аннотаторами. Как показал анализ расхождений в работе [Пименов, 2023], их критичным источником, способным привести к ухудшению согласия сразу на всех трёх уровнях разметки, выступает различное определение главного тезиса научной статьи. Например, при разметке текста, в котором описывалось применение метода частотного анализа к нескольким европейским языкам по разнородным лексическим группам и выведение частных наблюдений при сопоставлении этих языков, один аннотатор воспринял статью как иллюстрирующую применимость метода с указанием выявленных ограничений при вспомогательной роли языковых примеров, тогда как другой посчитал её направленной на определение взаимной близости языков при вторичности метода. Различное определение главного тезиса вызвало значительные расхождения и в выделении тезисов, и в определении их связей, и в интерпретации реализуемых моделей рассуждения.

Соотношение значений между показателями отражает специфику учёта случайного согласия: так, специализированный коэффициент аргументационного соответствия (КАС) указывает на достаточно высокое согласие на уровне выявления утверждений, однако значения каппы Коэна и альфы Криппендорфа оказываются заметно ниже. Это обусловлено непропорциональным соотношением двух классов для коротких научных статей ввиду их жанровой специфики: ограничение на объём текстов побуждает их авторов к тезисному ёмкому стилю, изложению доказываемых идей без сторонних отвлечений, ввиду чего совокупность неаргументативных частей оказывается, как правило, значительно короче аргументативных. Однако при подсчёте каппы Коэна и альфы Криппендорфа такая диспропорция категорий приводит к тому, что указание

аргументативности некоего сегмента обоими аннотаторами оценивается как случайное совпадение и практически не повышает значения коэффициентов.

Подобные ограничения на применимость двух этих мер отмечаются в работе [Delgado, Tibau, 2019]: применимость коэффициентов со штрафами за случайное совпадение зависит от ряда факторов, таких как сбалансированная представленность меток разных категорий в размечаемом корпусе, сопоставимая значимость разных меток для разметки, возможность чёткого разграничения категорий. Если указанные факторы не выполняются ввиду постановки задачи и/или специфики размечаемой коллекции, то расчёт согласия разметчиков через, к примеру, каппу Коэна может приводить к некорректным результатам (коллекция с более качественной разметкой, может получать меньшее значение каппы Коэна, чем эта же коллекция с менее качественной разметкой, где качество разметки устанавливается экспертным анализом с учётом специфики задачи).

Следует также отметить, что формально вычисляемые коэффициенты согласия лишь косвенно характеризуют надёжность размеченных данных, так как ориентированы на соответствие мнений аннотаторов относительно категорий изолированных объектов без учёта, например, насколько последовательны метки одного аннотатора на сходных объектах. Как следствие, повышение значений коэффициентов не обязательно гарантирует повышения качества алгоритмов распознавания на размеченных данных, как экспериментально показано в работе [Pimenov, Salomatina, 2024b].

3.1.3 Анализ встречаемости отдельных аргументационных схем

Для анализа употребления в научных статьях корпуса элементарных приёмов аргументации, содержащих одну типовую схему, подсчитаны абсолютные (F_{abs}) и текстовые (F_{text}) частоты таких схем. Под абсолютной частотой схемы понимается общее число её реализаций в аргументационных графах, под текстовой – число графов, включающих хотя бы одно вхождение этой схемы. Относительные частоты схем для каждой из двух тем уточнены через

подсчёт характеристик F_{comp} и F_{ling} (соответственно доли реализации схемы от общей характеристики F_{abs} в графах для лингвистических статей и графах для работ по компьютерным технологиям).

Эксперимент по анализу реализации элементарных приёмов аргументации проведён на ранней версии корпуса, а именно подкорпусе из 30 текстов, который содержит 1301 тезис и 1174 аргумента.

Частоты аргументационных схем, встречающихся в более чем 20% текстов подкорпуса ($F_{text} > 5$), а соответственно регулярно применяемых разными авторами, приведены в таблице 7 с упорядочением по абсолютной частоте F_{abs} . Для удобства восприятия F_{abs} каждой схемы сопровождается указанием её доли от общего числа аргументов в корпусе (1174, F_{abs} (%)).

Таблица 7 – Частоты встречаемости отдельных аргументационных схем

Схема	Обозначение	F_{abs}	$F_{abs}\%$	F_{text}	F_{comp}	F_{ling}
От примера	<i>Example (E)</i>	210	18%	29	37%	63%
Через классификацию	<i>VerbalClassification (VC)</i>	184	16%	26	56%	44%
От взаимосвязи к причине	<i>CorrelationToCause (CtC)</i>	149	13%	26	36%	64%
От причины к следствию	<i>CauseToEffect (CtE)</i>	146	12%	27	51%	49%
От части к целому	<i>PartToWhole (PtW)</i>	112	10%	16	57%	43%
От практической цели	<i>PracticalReasoning (PR)</i>	72	6%	22	60%	40%
От знака к означаемому	<i>Sign (S)</i>	58	5%	14	31%	69%
От экспертного мнения	<i>ExpertOpinion (EO)</i>	51	4%	18	39%	61%
Через опровержение	<i>LogicalConflict (C)</i>	34	3%	15	44%	56%
От применяемого метода	<i>AppliedMethod (AM)</i>	32	3%	16	44%	56%
От положительных результатов	<i>Positive Consequences (PC)</i>	23	2%	8	100%	0%

Окончание таблицы 7

От негативных результатов	<i>Negative Consequences (NC)</i>	18	2%	8	100 %	0%
От распространённого мнения	<i>Popular Opinion (PO)</i>	15	1%	7	60%	40%
По аналогии	<i>Analogy (A)</i>	13	1%	10	46%	54%
С позиции знающего	<i>Position To Know (PtK)</i>	12	1%	6	50%	50%
<i>Суммарно</i>	-	1129	96%	-	-	-

Самой частотной в текстах коллекции является схема с поддержкой тезиса примером (*Example*). Хотя эта модель применяется почти во всех работах (29 из 30), в лингвистических статьях она встречается значительно чаще, чем в статьях по информационным технологиям (63% против 37%). Так, в работах по лингвистике удобно иллюстрировать тезисы языковым материалом. Кроме того, статьям этой тематики свойственен анализ глубинных языковых структур через их наглядные внешние проявления (по схемам *Correlation To Cause* для общих языковых тенденций и *Sign* на уровне частных семиотических переходов; две указанные модели встречаются в лингвистических текстах в 64% и 69% случаев соответственно). Рассуждения от экспертного мнения также более характерны для статей по лингвистике (61%). Менее активная реализация трёх этих схем в текстах по компьютерным технологиям объясняется акцентом данной части корпуса на авторских разработках, что уменьшает потребность в цитировании иных исследователей и удобство представления созданных программ от их устройства к особенностям функционирования (в обратном направлении относительно описания языковых явлений).

В свою очередь, прикладная направленность работ по информационным технологиям влечёт предпочтение практических доказательств как через анализ решаемой задачи (*Practical Reasoning*, 60%), так и за счёт эксплицитной оценки

возможных результатов (*PositiveConsequences* и *NegativeConsequences*, причём рассуждения по этим моделям не появляются ни в одной лингвистической статье корпуса). Наконец, обеим темам присущи прямые причинно-следственные связи (*CauseToEffect*) и основательная систематизация описываемых явлений по схеме *VerbalClassification*.

Соотношение текстовой и абсолютной частот характеризует, насколько избирательно отдельные схемы применяются авторами. Например, рассуждения от экспертного мнения применяются в большем количестве текстов (18), чем модели *PartToWhole* и *Sign* (16 и 14), которым близко по текстовой частоте развитие доказательства через опровержение промежуточного тезиса (15). Однако схемы *ExpertOpinion* и *LogicalConflict* уступают двум другим указанным моделям по абсолютному числу реализаций (51 и 34 против 112 и 58). Это объясняется спецификой их употребления в научных текстах: при доказательстве главного тезиса достаточно привести ограниченное число чужих авторитетных суждений, на основе которых затем развить подробный самостоятельный анализ (в котором логические переходы между частью и целым, знаком и означающим способны выстраиваться неоднократно). Конфликты же могут применяться в обосновании главного тезиса через указание возможных возражений и их встречное опровержение, при этом рассмотрение большого числа таких возражений нетипично ввиду отвлечения внимания от основного доказательства. Сходным образом при рассуждении от практической цели (на уровне всего исследования или отдельного этапа) могут подробно анализироваться пути её достижения, что отражается в соотношении частот схемы (*PracticalReasoning* применяется 72 раза в 22 текстах, тогда как для *PartToWhole* отмечается в полтора раза больше реализаций при меньшем количестве включающих её работ).

Следует отметить, что 15 моделей рассуждения, применяемых в более чем 20% текстов подкорпуса, покрывают 96% всех его аргументов (1129). Это свидетельствует о сравнительной однородности доказательств в научных статьях корпуса: авторы исследованных текстов предпочитают строить рассуждения на основе узкого круга типовых моделей и демонстрируют сдержанность в

творческом применении экзотических схем, нетипичных для жанра (например, аргументов от нравственных или эстетических установок, исключительных случаев, социальных комментариев с апелляцией к страху либо личностным убеждениям, а тем более *ad hominem* критики альтернативных воззрений либо иных исследователей).

Для уточнения аргументационной специфики научных работ было проведено дополнительное исследование на материале корпуса коротких научно-популярных статей как текстов смежного жанра [Пименов, 2021]. Данный корпус содержит 95 статей, размеченных с целью изучения отдельных аргументационных схем и характеризующихся меньшей интенсивностью доказательств (в среднем в отдельной научно-популярной статье корпуса выделяется по 20 тезисов и 14 аргументов при аналогичном объёме текста). Выявленные различия в реализации аргументационных схем между научными и научно-популярными текстами (задающие их жанровую специфику) объясняются расхождением по целевой направленности. Так, поскольку научно-популярные статьи в большей степени ориентированы на лёгкость восприятия представляемых сведений, для них не столь значима точная систематизация описываемых явлений (обеспечиваемая разноплановыми классификациями, которым соответствует схема *VerbalClassification* – одна из наиболее частотных в группе научных текстов, но покрывающая только 1% аргументов в размеченных научно-популярных статьях). Работам этого жанра также более свойственно построение рассуждений через привлечение авторитетных оценок (по схеме *ExpertOpinion*, 11% против 4% для корпуса научных текстов), поскольку ввиду их жанровой специфики выступает допустимой организация аргументации в тексте посредством исключительного апеллиций к авторитетным источникам (тогда как в научном жанре предполагается расширение цитат от иных экспертов собственным комментарием автора).

Другим направлением развития представленного анализа являются работы, обращённые к компьютерному распознаванию аргументов отдельных схем. Эта проблема решается как задача классификации по аналогии с компьютерным

извлечением утверждений, представленным в разделе 3.3 (посредством алгоритмов классификации, указанных в 2.2.2).

Так, в статье [Pimenov, Salomatina, 2024a] описано компьютерное извлечение аргументов, образованных по схемам *Example*, *Practical Reasoning* и *Expert Opinion* (680, 386 и 172 аргументов соответственно). Это исследование направлено на оценку эффективности признаков разных типов (лемм, маркеров дискурса, индикаторов методов убеждения с учётом или без учёта позиционной специфики, их встречаемости в посылке или заключении) в сравнении друг с другом. В связи с этим полученные оценки качества рассматриваются как ориентировочный *baseline* для дальнейших работ (достижение как можно более высоких оценок качества распознавания не входило в цель исследования на данном этапе). Тем не менее, достигнуты показатели F-меры, равные 68% для извлечения аргументов по модели *Expert Opinion* (точность P равна 65%, полнота R – 72%), 58% для *Example* ($P = 65\%$, $R = 51\%$), 54% для *Practical Reasoning* ($P = 67\%$, $R = 45\%$). Показано также, что, как правило, учёт маркеров дискурса в аргументах повышает точность распознавания схем, тогда как учёт индикаторов методов убеждения – полноту. В частности, применение маркеров дискурса повышает точность в среднем на 10% (тогда как сопутствующую F-меру – на 7%) по сравнению с распознаванием исключительно по тематическим леммам. Отмеченные ошибки извлечения схем сводятся к следующим основным причинам: 1) пересечение аргументов в структуре рассуждения на уровне утверждений (одно утверждение может являться посылкой в одном аргументе и заключением в другом, в связи с чем содержать показатели схемы как первого, так и второго); 2) реализация в аргументах многозначных маркеров (соответствующих нескольким схемам либо преимущественно используемым для некоей одной, но в конкретном контексте выражающим иную); 3) реализация аргументов без явных показателей исследуемых схем (более характерно для *Practical Reasoning*). Из трёх указанных моделей наиболее явные маркеры характеризуют схему *Expert Opinion*, как отдельно показано в [Salomatina et al., 2020], где для компьютерного распознавания аргументов от экспертного мнения

успешно применяются вручную созданные лингвистические шаблоны. Применимость маркеров дискурса и специально создаваемых шаблонов для анализа аргументации рассматривается также в работе [Пименов, Саломатина, Засыпкин, 2024] в контексте компьютерного распознавания связей в аргументах, образованных по схемам *Example*, *Analogy* и *Verbal Classification*: точность шаблонов на основе маркеров дискурса достигает 80%, 88% и 95% соответственно для этих схем. Аналогично, маркеры в дополнении особенностями явленной структуры текста используются в работе [Саломатина, Пименов, 2024] для построения лингвистических правил, посредством которых определяются главные тезисы отдельных абзацев и всего текста (с достижением F-меры в 80%). В свою очередь, компьютерное выявление аргументативных связей между смежными клаузами и предложениями без опоры на словарь маркеров (исключительно методами машинного обучения) представлено в статье [Саломатина, Сидорова, Пименов, 2024], где анализу подвергаются тексты разных жанров научной коммуникации с достижением F-меры в 58%.

В свою очередь, в статье [Pimenov, Salomatina, 2024b] представлен эксперимент по компьютерному распознаванию аргументов по схемам *Part to Whole*, *Verbal Classification*, *Correlation to Cause*. Данный эксперимент проведён для оценки эффективности процедуры компьютерной модификации разметки (реализованной для устранения типовых расхождений между аннотаторами) и ориентирован на выявление динамики оценок эффективности распознавания (до и после модификации разметки). Поскольку интерес в рамках работы представляет именно динамика показателей качества, а не их абсолютные значения, извлечение аргументов основано на базовой модели «мешка слова» (важно использование однотипных признаков в представлении аргументов для корпуса до и после модификации). Однако даже при такой простой модели отмечены достаточно хорошие оценки распознавания уже на изначальном корпусе до модификации: достигнута F-мера в 65% для *Verbal Classification* (1176 аргументов, $P = 58\%$, $R = 73\%$), 64% для *Part to Whole* (1040 аргументов, $P = 53\%$, $R = 80\%$), 61% для *Correlation to Cause* (786 аргументов, $P = 49\%$, $R = 81\%$). Низкая точность для

этой тройки схем по сравнению с прошлой обусловлена, во-первых, использованием более простого признакового представления аргументов, а во-вторых, их менее явной маркированностью (даже относительно *Practical Reasoning*, уступающей по явности маркеров и *Example*, и *Expert Opinion*). Перечисленные выше работы посвящены распознаванию разных типов элементов аргументационных структур на отдельных этапах, а подход с объединением разных этапов в компьютерной обработке аргументации представлен в статье [Zasyrkin, Pimenov, Salomatina, 2023].

3.2 Компьютерное распознавание предложений, содержащих аргументацию

Анализ организации рассуждений в тексте предполагает моделирование его аргументационной структуры, что является особенно трудоёмкой задачей ввиду сложности аргументации как явления прагматического уровня языка. В этой связи актуальна по крайней мере частичная поддержка экспертной разметки тезисов и их связей посредством компьютерной обработки неструктурированных текстов. Такая обработка может осуществляться на разных уровнях, в частности, для компьютерного распознавания предложений, содержащих аргументацию.

Проблема компьютерного распознавания в тексте предложений, содержащих тезис из некоторого аргумента, сводится к задаче автоматической бинарной классификации коллекции предложений (их разбиения на две группы: относящиеся и не относящиеся к аргументационной структуре). Методы обучения подобных классификаторов представлены в главе 2, разделе 2.2. Результаты эксперимента по распознаванию утверждений в научных и научно-популярных текстах также опубликованы в работе [Salomatina, Sidorova, Pimenov, 2021].

3.2.1 Обучение классификатора с учётом специфики корпуса текстов

В нашем случае решение задачи классификации предполагает выполнение следующих этапов:

- 1) сбора коллекции текстов;
- 2) проведения аргументационной разметки текстов (с целью обучения классификатора и оценки его качества для последующего использования на текстах без известной структуры аргументации);
- 3) формирования обучающей и тестовой коллекций текстов;
- 4) фильтрации признаков (в нашем случае лемм текста);
- 5) обучения классификатора с помощью выбранного метода;
- 6) проверки классификатора на тестовой коллекции;
- 7) вычисления оценок качества классификации.

Коллекция текстов, используемая в рамках компьютерного распознавания тезисов, содержит 10 научных статей (5 лингвистических, 5 по компьютерным технологиям) и 12 научно-популярных. Предложения из этих текстов извлечены посредством компьютерной сегментации по пунктуационным знакам и образуют обучающую и тестовую коллекцию для настройки классификатора. Распределение предложений по двум данным группам регулируется принципами, во-первых, сохранения целостности текстов для избежания тематического влияния между частями текста и завышения оценок качества (все предложения из одного текста включаются либо в обучающую, либо в тестовую коллекцию, но не могут распределяться по обеим); во-вторых, примерно равного распределения предложений обеих групп по коллекциям с целью корректного обучения и оценки качества классификатора.

Построенная обучающая коллекция (S^L) включает предложения из 16 текстов (10 научно-популярных и 6 научных: 3 по компьютерным технологиям, 3 лингвистических), тестовая (S^T) – из 6 (2 научно-популярных, 2 компьютерных, 2 лингвистических). Количественные характеристики обеих коллекций на уровне предложений (с учётом соотношения между классами C_a и C_n , аргументативных и неаргументативных соответственно) представлены в таблице 8.

Таблица 8 – Распределение предложений для подготовки классификатора

	C_a	C_n	Всего
S^L	392	390	782
S^T	91	89	180

Предобработка предложений заключается в их пословном разбиении, лемматизации слов с помощью морфологического анализатора PyMorphy2 и устранении слов, не содержащих кириллицу. Используемая модель текста – мешок слов, представление – вектор нормализованных униграмм (лемм).

Для выбора признаков в исследовании экспериментально проверяется эффективность их статистической фильтрации. Для этого применяются процедуры взвешивания компонентов вектора по мерам $TF*IDF$ и ожидаемой информационной выгоды (EMI), а также отбора униграмм согласно критериям Variance (Var) и χ^2 . Методы статистической фильтрации признаков представлены в разделе 2.2.2.

Для классификации предложений выбраны часто упоминаемые в литературе при решении аналогичных задач методы MNB и SVM, а также алгоритм MLP. Функционирование этих методов описано в разделе 2.2.3.

3.2.2 Статистическая фильтрация лемм на основе их частоты

На этапе обучения для сокращения признакового пространства проводится взвешивание лемм двумя способами (по мерам $TF*IDF$ и EMI), используются критерии Var и χ^2 , а также различные их комбинации. Процедура обучения и тестирования реализуется на признаках (леммах), отобранных на основании весов и критериев. После фильтрации признаков в обучении и распознавании используются их бинарные значения.

Объём обучающей коллекции без фильтрации составляет 3302 разных леммы. Рассмотрено 15 вариантов V_i ($i = 1, \dots, 15$) фильтрации признаков. Объёмы

отдельных наборов признаков, условия, согласно которым проводится фильтрация, а также пороги (p_1 , p_2 , p_3) представлены в таблице 9.

Таблица 9 – Количество лемм при разных способах фильтрации.

Фильтры	Критерий	V_i
–	–	3302
TF*IDF	$p_1 > 0.5$	3164
EMI	$p_2 > 0.001$	800
TF*IDF, EMI	$p_1 > 0.5 \& p_2 > 0.001$	718
Var	$p_3 > 0.01$	333
Var, TF*IDF	$p_3 > 0.01 \& p_1 > 0.5$	202
Var, EMI	$p_3 > 0.01 \& p_2 > 0.001$	168
Var, TF*IDF, EMI	$p_1 > 0.5 \& p_2 > 0.001 \& p_3 > 0.01$	88
χ^2	25% лучших	825
χ^2 , TF*IDF	25% & $p_1 > 0.5$	791
χ^2 , EMI	25% & $p_2 > 0.001$	200
χ^2 , TF*IDF, EMI	25% & $p_1 > 0.5 \& p_2 > 0.001$	179
χ^2	10% лучших	330
χ^2 , TF*IDF	10% & $p_1 > 0.5$	317
χ^2 , EMI	10% & $p_2 > 0.001$	80
χ^2 , TF*IDF, EMI	10% & $p_1 > 0.5 \& p_2 > 0.001$	71

Выбор порогов (p_1 , p_2 , p_3) произведён экспериментально. Из-за заполненности каждого вектора нулями в среднем на 99% фильтрация с помощью меры TF*IDF осуществляется путем проверки ее значения для каждого предложения. Признак (лемма) считается информативным, если порог преодолевается хотя бы для одного предложения из всех, составляющих вектор. Таким же образом разреженность векторов обуславливает выбор сравнительно низкого значения для порога EMI: значение EMI каждой леммы подсчитывается

на всём векторе в соответствии с частотами его компонентов в аргументативных и неаргументативных предложениях. Порог для отбора лемм по критерию Variance настроен на фильтрацию лемм, встречающихся в менее чем 99% предложений (что в 2 раза превышает среднюю долю реализуемых в отдельных предложениях лемм). Порог χ^2 настроен на отбор заданной доли лемм, встречаемость которых в предложениях наиболее сильно зависит от их класса. При комбинировании фильтров значения выбранных порогов сохраняются. В ближайшей окрестности выбранных пороговых значений алгоритмы показывают примерно одинаковые показатели качества, которые снижаются с удалением от выбранного порогового значения.

3.2.3 Результаты автоматической классификации предложений

Показатели качества получены для каждого из алгоритмов (MNB, SVM, MLP) на всех наборах признаков (V_i). Основные результаты экспериментов отражены в таблице 10 (показатели качества указаны в процентах). Полужирным шрифтом выделены два лучших показателя точности (P), полноты (R) и F-меры для каждого алгоритма.

Таблица 10 – Качество распознавания тезисов на разных наборах лемм

V_i	P			R			F-мера		
	MNB	SVM	MLP	MNB	SVM	MLP	MNB	SVM	MLP
3302	56.86	53.75	55.95	63.74	47.25	51.65	60.10	50.29	53.71
3164	58.33	47.17	55.29	61.54	54.95	51.65	59.89	50.76	53.41
825	54.55	48.53	47,13	59.34	36,26	45.05	56.84	41.51	46.07
800	53.06	55.07	60.53	57.14	41.76	50.55	55.03	47.50	55.09
333	56.67	51.43	56.10	56.04	39.56	50.55	56.35	44.72	53.18
330	57.14	49.21	59.74	65.93	34.07	50.55	61.22	40.26	54.76

Окончание таблицы 10

791	54.21	48.11	55.56	63.74	56.04	60.44	58.59	51.78	57.89
718	52.59	57.14	59.77	67.03	35.16	57.14	58.94	43.54	58.43
317	56.82	55.93	59.42	82.42	36.26	45.05	67.26	44.00	51.25
168	52.58	62.90	72.50	56.04	42.86	63.74	54.26	50.98	67.84
80	51.68	58.11	60.27	84.62	47.25	48.35	64.17	52.12	53.66
179	56.83	56.67	64.15	86.81	37.36	37.36	68.70	45.03	47.22
71	53.29	60.66	64.58	89.01	40.66	34.07	66.67	48.68	44.60

В третьей строке таблицы размещены результаты распознавания аргументативных предложений без фильтрации признаков. В следующих пяти строках – результаты работы алгоритмов после применения каждого отдельного фильтра к набору признаков, далее – результаты применения комбинированных фильтров. Число признаков после фильтрации указано в первом столбце (критерии фильтрации см. в таблице 9). Фильтрация признаков по одному из фильтров улучшает точность суммарно для всех алгоритмов в половине случаев, полноту и F-меру в трети случаев. Лучшие показатели точности для MNB достигаются при фильтрации по одному из фильтров: TF*IDF или χ^2 (10% лучших), полноты и F-меры – по совокупности фильтров. При сокращении признакового пространства по одному из фильтров, полнота, как правило, падает, исключение представляют результаты работы SVM на признаках, полученных фильтрацией по TF*IDF, и MNB – по χ^2 (10% лучших). Выигрыш по F-мере при применении MNB возможен при использовании χ^2 (10% лучших); результаты SVM улучшает фильтрация по TF*IDF; результаты MLP улучшаются при фильтрации по χ^2 (10% лучших) или EMI.

Одновременное применение нескольких фильтров улучшает показатели качества избирательно. Влияние комбинированных фильтров на точность, полноту и F-меру разных алгоритмов представлено в списке ниже.

1. Комбинирование фильтров не улучшает точность работы MNB, но практически все комбинации улучшают точность SVM и MLP (максимум – Var & EMI, а также χ^2 (10% лучших) & TF*IDF & EMI).
2. Полнота повышается при применении MNB с фильтрацией по χ^2 (10% лучших) & TF*IDF & EMI и χ^2 (10% лучших) & EMI; SVM с фильтрацией по TF*IDF или χ^2 (25% лучших) & TF*IDF; MLP – χ^2 (25% лучших) & TF*IDF или Var & EMI.
3. F-мера улучшается, если использовать MNB в сочетании с фильтрацией по χ^2 (25% лучших) & TF*IDF & EMI или χ^2 (10% лучших) & EMI; SVM – Var & EMI или χ^2 (10% лучших) & EMI; MLP – Var & EMI или TF*IDF & EMI.

Подводя итог, можно сказать, что фильтрация признаков в целом полезна, но требует подбора критериев для каждого алгоритма и для каждого показателя качества.

Как правило, классификатор хорошо работает либо с аргументативными, либо неаргументативными предложениями. Близкое качество в распознавании одним классификатором предложений, относящихся к обоим классам, достигается очень редко. Алгоритм MNB лучше работает с аргументативными предложениями, SVM и MLP – с неаргументативными. Лучшие показатели по *точности* распознавания (72.50%) демонстрирует алгоритм MLP на 168 признаках (при этом показывая и лучшие для него полноту и F-меру); по *полноте* (89.01%) – алгоритм MNB на наборе из 71 признака, при этом показывая лучшую F-меру (68.70%) на 179 признаках.

Анализ ошибок показывает, что основную часть неправильно распознанных предложений составляют предложения, не содержащие маркеров или содержащие неоднозначные маркеры (*значит, говорить о том, что* и др.), встречающиеся как в предложениях с аргументацией, так и без нее. Аргументативные предложения с отсутствующими в обучающей коллекции моделями, тем не менее, были распознаны в тестовой благодаря, в частности, содержащимся в них

риторическим и аргументативным маркерам (*может быть; тем самым; отсюда можно найти; оказывается, что; для того, чтобы; написано, что* и т.п.).

Полученные в проведенных экспериментах результаты сложно сравнивать с представленными в литературе в силу имеющихся языковых и жанровых различий в данных, а также использования других показателей качества. Однако по точности результаты исследования сопоставимы, например, с представленными в [Moens, 2007]: 0.73 и 0.72.5 (MLP, 168 признаков).

Выводы

В третьей главе диссертационного исследования описаны эксперименты по решению отдельных задач формального выявления элементов аргументационных стратегий. К этим задачам относятся анализ использования элементарных приёмов аргументации (из одной модели рассуждения) и компьютерное распознавание тезисов в исходных текстах научных статей. Отдельное внимание уделено проблеме оценки качества аргументационной разметки: представлен анализ согласия аннотаторов для созданного корпуса, используемого в экспериментах, на основе количественных мер.

Количественная оценка качества разметки в подготовленном корпусе подсчитывается посредством трёх разных коэффициентов: специализированного с учётом многоуровневого характера аргументационной разметки (коэффициент аргументационного соответствия, разработан в рамках исследования) и двух универсальных с учётом случайного согласия между аннотаторами (каппа Коэна и альфа Криппендорфа, подсчёт которых для многоуровневой аргументационной разметки требует незначительного преобразования формата данных). Отмечено, что показатели согласия для подготовленного корпуса достигают диапазона «умеренного» согласия на всех уровнях разметки, а при её автоматической модификации на основе типовых расхождений повышаются до «существенного». При этом указываются ограничения количественных метрик согласия, которые

могут искажать информативность высчитываемых значений при указании на надёжность разметки.

Эксперимент по частотному анализу элементарных приёмов аргументации показывает, что 16 типовых моделей рассуждения, применяемых более чем в 20% текстов проанализированной коллекции, покрывают 96% её аргументов. Это свидетельствует о сравнительной однородности доказательств в научных статьях корпуса: авторы исследованных текстов предпочитают строить рассуждения на основе узкого круга типовых моделей и избегают активного применения схем, нетипичных для жанра. Выявлена и тематическая зависимость в использовании элементарных приёмов для двух сравниваемых предметных областей: лингвистики и информационных технологий. Так, лингвистическим статьям свойственно активное приведение примеров для иллюстрации языковых фактов (67% от всего числа реализаций схемы *Example*), изучение глубинных языковых структур через их наглядные внешние проявления (по схемам *Correlation to Cause* и *Sign*, соответственно 64% и 69% аргументов с которыми применяются именно в публикациях по языкознанию) и обращение ко взглядам иных исследователей (по модели *Expert Opinion*, 61%). Прикладная же направленность работ по информационным технологиям влечёт предпочтение практических доказательств через анализ решаемой задачи (*Practical Reasoning*, 60%) либо за счёт эксплицитной оценки возможных результатов (*Positive Consequences* и *Negative Consequences*, причём рассуждения по этим моделям не появляются ни в одной лингвистической статье корпуса). Общая жанровая специфика статей обеих тематических групп проявляется в одинаково интенсивном анализе прямых причинно-следственных связей и основательной систематизации описываемых явлений (по моделям *Cause to Effect* и *Verbal Classification*, которые используются в публикациях по лингвистике и информационным технологиям с одинаковой частотой, а при этом покрывают соответственно 16% и 12% аргументов всей коллекции при учёте любых аргументационных схем).

Развитием анализа элементарных приёмов аргументации выступают эксперименты по компьютерному извлечению аргументов, построенных по

отдельным шести схемам, покрывающим более 55% аргументов корпуса (*Example, Expert Opinion, Practical Reasoning, Verbal Classification, Part to Whole, Correlation to Cause*). Проведённые эксперименты по компьютерному распознаванию этих схем являются вспомогательными при решении иных задач (определение эффективности признаков разных типов, проверка процедуры компьютерной модификации разметки), тем не менее, в них достигнуты показатели качества распознавания в 65–68% F-меры (хотя на этом этапе исследования цель получить как можно более высокие оценки компьютерного извлечения схем не ставилась).

Эксперимент по компьютерному распознаванию тезисов методами машинного обучения показывает оправданность использования таких алгоритмов для частичной автоматизации ручной разметки и применимость статистической фильтрации значимых для аргументации лемм (на основе их частот в аргументативных и неаргументативных сегментах текста, без экспертного составления маркеров). Качество компьютерного распознавания тезисов оценивается в терминах точности, полноты и усредняющей их F-меры. Лучших показателей по точности распознавания достигает многослойный перцептрон (72.50%, а его полнота и F-мера равны 63.74% и 67.84%), наибольшая же полнота и F-мера отмечается для полиномиального байесовского классификатора (86.81% и 68.70% при точности в 56.83%). Анализ ошибок показывает, что основную часть неправильно распознанных предложений составляют предложения, не содержащие маркеров аргументации или содержащие неоднозначные маркеры. Тезисы в аргументах с отсутствующими в обучающей коллекции (для подготовки алгоритмов) моделями рассуждения, тем не менее, распознаются в тестовой при наличии в них риторических и аргументативных маркеров (которые выявляются компьютером за счёт методов формальной фильтрации).

Глава 4. Стратегии аргументации как совокупность общих приёмов

В четвёртой главе данной работы решается задача определения **стратегий аргументации**, реализуемых авторами отдельных научных статей для организации рассуждений в целостный текст. Так, раздел 4.1 обращён к проблеме выявления **составных приёмов аргументации**, соответствующих частотным подструктурам в аргументационных графах из нескольких элементарных приёмов. Представлен анализ таких приёмов как в собственно научных статьях, так и в текстах разных жанров научной коммуникации (научных и аналитических научно-популярных статьях *Nabr*, новостях науки). Демонстрируется высокая различительная способность приёмов аргументации в представлении текстов, их эффективность в разграничении текстов как по жанру, так и по тематике (на основе информации только о реализованных в тексте приёмах аргументации).

В разделе 4.2 описывается выявление **методов убеждения** на основе классической модели Аристотеля (с разграничением логоса, пафоса, этоса): методы убеждения характеризуют полный текст в аспекте предпочитаемых его автором способов доказательства (посредством логических рассуждений, апелляций к авторитету, воздействия на эмоции). При этом ввиду специфики исследуемого жанра (научных статей) анализ методов убеждения в данной работе сфокусирован на авторитетности и эмоциональности отдельных текстов как вспомогательных категориях к логическому доказательству, которые отражают не убедительность аргументов согласно истинностным значениям тезисов, а способы оформления фактического содержания для его более продуктивного восприятия предполагаемой аудиторией. На основе анализа методов убеждения в текстах с аргументационной разметкой решается задача компьютерного распознавания аргументов по интенсивности выражения пафоса и этоса, что представлено в разделе 4.3.

В разделе 4.4 анализируется сочетаемость элементарных приёмов аргументации в полных путях рассуждения (от посылок, приводимых без обоснования, к главному тезису научной статьи) с учётом функциональной

близости отдельных схем. Наконец, раздел 4.5 представляет исследование **стратегий аргументации** в терминах аргументационной сложности целостных научных текстов (их формального сравнения по специфике организации доказательств) посредством интеграции результатов аргументационного анализа в различных аспектах (использования элементарных и составных приёмов, реализации методов убеждения, сочетания аргументов из разных функциональных групп, структурных особенностей графов аргументации). Показанная в разделе 4.5 эффективность сведений об аргументационных структурах в кластеризации текстов дополнительно подтверждается представленным в разделе 4.6 экспериментом по жанровой атрибуции текстов многожанровой коллекции научной коммуникации с использованием исключительно сведений о реализуемых в этих текстах приёмах аргументации разной сложности.

4.1 Выявление составных приёмов аргументации

Составным приёмам аргументации, включающим не менее двух схем, соответствуют подграфы двух типов. Совокупность связанных вершин в них может образовывать дерево (именуемое далее «цепочкой»), в котором каждая вершина имеет только одно входящее и исходящее ребро, кроме корневой (нет исходящего ребра) и листовой (нет входящего ребра). Такое определение дерева отличается от стандартного (согласно которому в корневой вершине нет входящих рёбер, а в листовых нет исходящих), однако стандартное представление может быть получено через изменение направления для всех связей (тогда связь вида $A \rightarrow B$ будет означать не «А поддерживает В», а «А поддерживается В»), что не вызовет структурных или содержательных искажений для приёма аргументации.

Любой же приём аргументации, содержащий ровно две схемы, может быть представлен в виде цепочки. Подграфы другого типа, в которых хотя бы одна вершина имеет более одного входящего либо исходящего ребра), именуются

ветвящимися структурами. Как следует из определения, такие подграфы включают не менее трёх аргументационных схем.

4.1.1 Составные приёмы аргументации в научных текстах

Исследование составных приёмов аргументации проведено на материале ограниченной части корпуса (25 научных статей: 14 по лингвистике, 11 по компьютерным технологиям). Результаты исследования представлены также в работе [Пименов, Саломатина, Тимофеева, 2022].

Для анализа встречаемости таких приёмов построен частотный спектр подграфов обоих типов в соответствии с процедурой, описанной в разделе 2.3. Количество всех встретившихся в двух и более текстах коллекции разных цепочек (из характеристики $\Phi_n(A)$, $n = 2, \dots, 5$) и разных ветвящихся структур (из характеристики $\Phi_n(A)$, $n = 3, \dots, 8$) – 231 и 221 соответственно.

Статистика применения наиболее частых приёмов-цепочек из n вершин ($n = 2, \dots, 5$) приведена в таблице 11 с указанием для каждого приёма числа текстов, где он реализуется хотя бы один раз (F_{txt}), дополнительно уточняемого по темам (F_{Comp} и F_{Ling}). Схемы в составе приёмов обозначаются аббревиатурами, приведёнными в таблице 7. Обозначение *MainThesis* соответствует достижению цепочкой главного тезиса текста.

Аналогичная статистика для ветвящихся приёмов из n аргументационных схем ($n = 3, \dots, 8$) приведена в таблице 12.

Таблица 11 – Частоты встречаемости цепочек аргументов

N	F_{txt}	F_{Comp}	F_{Ling}	Цепочки
$n = 2$				
1	14	7	7	$CtoE \rightarrow PR$
2	13	4	9	$CtoC \rightarrow CtoE$
3	12	3	9	$E \rightarrow CtoC$
4	12	2	10	$CtoE \rightarrow CtoC$
5	11	5	6	$E \rightarrow CtE$

Окончание таблицы 11

$n = 3$				
1	8	1	7	$E \rightarrow CtoC \rightarrow CtoE$
2	7	4	3	$CtoE \rightarrow PR \rightarrow CtoE$
3	6	0	6	$CtoC \rightarrow CtoE \rightarrow CtoC$
4	5	1	4	$CtoC \rightarrow CtoE \rightarrow MainThesis$
5	5	4	1	$PR \rightarrow CtoE \rightarrow PR$
$n = 4$				
1	5	4	1	$CtoE \rightarrow PR \rightarrow CtoE \rightarrow PR$
2	4	0	4	$E \rightarrow CtoC \rightarrow CtoE \rightarrow CtoC$
3	3	0	3	$CtoE \rightarrow CtoC \rightarrow CtoE \rightarrow MainThesis$
4	3	0	3	$EO \rightarrow CtoE \rightarrow CtoC \rightarrow CtoE$
$n = 5$				
1	2	2	0	$E \rightarrow CtoE \rightarrow PR \rightarrow CtoE \rightarrow PR$
2	2	0	2	$CtoC \rightarrow CtoE \rightarrow CtoC \rightarrow CtoE \rightarrow MainThesis$
3	2	0	2	$CtoE \rightarrow E \rightarrow CtoC \rightarrow CtoE \rightarrow MainThesis$
4	2	1	1	$CtoE \rightarrow PR \rightarrow CtoE \rightarrow PR \rightarrow MainThesis$

Таблица 12 – Частоты ветвящихся приёмов аргументации.

N	F_{txt}	F_{Comp}	F_{Ling}	Деревья
$n = 3$				
1	6	0	6	$CtoC \searrow$ $CtoC \rightarrow CtoE$
2	5	3	2	$CtoC \searrow$ $CtoC \rightarrow MainThesis$
3	5	0	6	$E \searrow$ $E \rightarrow CtoC$
4	4	0	4	$CtoE \searrow$ $CtoE \rightarrow PR$
$n = 4$				
1	3	0	3	$PtoW \searrow$ $E \rightarrow PtoW \rightarrow PtoW$
2	3	3	0	$E \searrow$ $E \rightarrow CtoE \rightarrow PR$
3	3	1	2	$E \searrow$ $CtoE \rightarrow PR \rightarrow CtoE$
$n = 5$				
1	3	0	3	$CtoC \rightarrow CtoC \searrow$ $CtoC \rightarrow CtoC \rightarrow MainThesis$
2	2	2	0	$VC \searrow$ $VC \rightarrow CtoC$ $PtoW \rightarrow VC \nearrow$
3	2	1	1	$VC \searrow$ $VC \rightarrow PtoW \rightarrow CtoE$ $VC \nearrow$

Окончание таблицы 12

<i>n</i> = 6				
1	2	2	0	$ \begin{array}{l} E \searrow \\ E \rightarrow \\ E \rightarrow CtoE \rightarrow PR \\ CtoC \nearrow \end{array} $
2	2	0	2	$ \begin{array}{l} E \rightarrow CtoC \searrow \\ CtoC \rightarrow CtoE \rightarrow CtoC \rightarrow MainThesis \end{array} $
3	2	1	1	$ \begin{array}{l} CtoC \searrow \\ CtoC \rightarrow \\ CtoC \rightarrow MainThesis \\ CtoE \rightarrow CtoC \nearrow \end{array} $
<i>n</i> = 7				
1	2	2	0	$ \begin{array}{l} CtoC \searrow \\ CtoC \rightarrow \\ CtoC \rightarrow MainThesis \\ CtoE \rightarrow PR \rightarrow CtoE \nearrow \end{array} $
2	2	0	2	$ \begin{array}{l} CtoC \rightarrow CtoE \searrow \\ CtoE \rightarrow E \rightarrow CtoC \rightarrow CtoE \rightarrow MainThesis \end{array} $
3	2	1	1	$ \begin{array}{l} VC \searrow \\ VC \rightarrow \\ VC \rightarrow PR \\ CtoE \rightarrow PR \rightarrow CtoE \nearrow \end{array} $
<i>n</i> = 8				
1	2	0	2	$ \begin{array}{l} E \rightarrow CtoC \searrow \quad CtoC \searrow \\ CtoE \rightarrow CtoE \rightarrow CtoC \nearrow \quad CtoC \rightarrow MainThesis \end{array} $

Анализ приёмов аргументации (цепочек и ветвящихся структур) позволил выявить зависящие и не зависящие от темы особенности организации рассуждений в научных текстах коллекции. Эти особенности перечислены ниже.

Во-первых, если в узлах дерева возникает ветвление, то вне зависимости от уровня узла в большинстве случаев схемы в узлах следующего уровня на разных ветвях будут одинаковы, то есть построение параллельных доводов чаще происходит по одной и той же модели вне зависимости от темы текста (см. примеры в табл. 3: схема *MainThesis* ($n = 3, 5, 6, 7, 8$) поддерживается несколькими *CorrelationToCause*; *CorrelationToCause* ($n = 5$), *PracticalReasoning* ($n = 7$) и *PartToWhole* ($n = 5$) – несколькими *VerbalClassification*; *CauseToEffect* поддержан двумя *CorrelationToCause* ($n = 3$) и несколькими *Example* ($n = 4, 6$) и др.

Такое построение свидетельствует о частом использовании «интенсивной аргументации» – через организацию рассуждения от нескольких аргументов одного типа. Доказательство тезиса через несколько аргументов разных типов является менее регулярным.

Иными словами, если авторы научных текстов выстраивают рассуждение от нескольких доводов, то они регулярно организуют эти аргументы по одинаковой модели рассуждения (при возможном применении нескольких: *CorrelationToCause*, *Example*, *VerbalClassification*). Сочетание двух доводов через разные модели является менее типичным. Отмеченная тенденция может объясняться удобством восприятия: читателю проще осмыслить доказательство от нескольких доводов, когда они все соотносятся с тезисом одинаково. И наоборот, использование нескольких доводов разного типа может оказаться избыточным: более убедительным будет применение одного довода из нескольких доступных.

Во-вторых, корневая (техническая) вершина в графе *MainThesis* (см. п. 2.1) чаще является корнем в дереве, чем конечным звеном в цепочке: в 63-х подграфах из 221 (28.5%) аргументы направлены непосредственно на доказательство главного тезиса, тогда как среди цепочек только в 25-ти из 122 (с тремя или более вершинами), то есть 20%, примыкают к главному тезису. Важно подчеркнуть, что цепочки встречаются и внутри подграфов (при рассмотрении их отдельных путей), в то время как обратное структурно невозможно. Таким образом, для научных текстов в окрестности главного тезиса чаще отмечаются сочетания нескольких аргументов, чем одного завершающего в цепочке.

Причиной этому служит предпочтение в научных текстах многоуровневых доказательств главного тезиса: ключевые выводы анализируются в различных аспектах (возможно, разрозненных, если пути рассуждения не пересекаются вне главного тезиса), а их обоснование обеспечивает аргументационную связность (и, следовательно, логико-семантическую целостность) всего текста. Так как главный тезис служит связующим элементом для всех путей рассуждения, типичным является выделение смысловых подблоков текста в непосредственной с ним связи.

Наглядный пример для двух указанных наблюдений приведён на Рис. 4. На нем представлен главный тезис текста (без исходящих связей к каким-либо иным утверждениям), в котором утверждается преимущество одного алгоритма над двумя другими. Три довода в поддержку главного тезиса приведены по одной и той же модели (от положительных результатов), а их обоснование состоит в подробном представлении каждого из алгоритмов по отдельности.

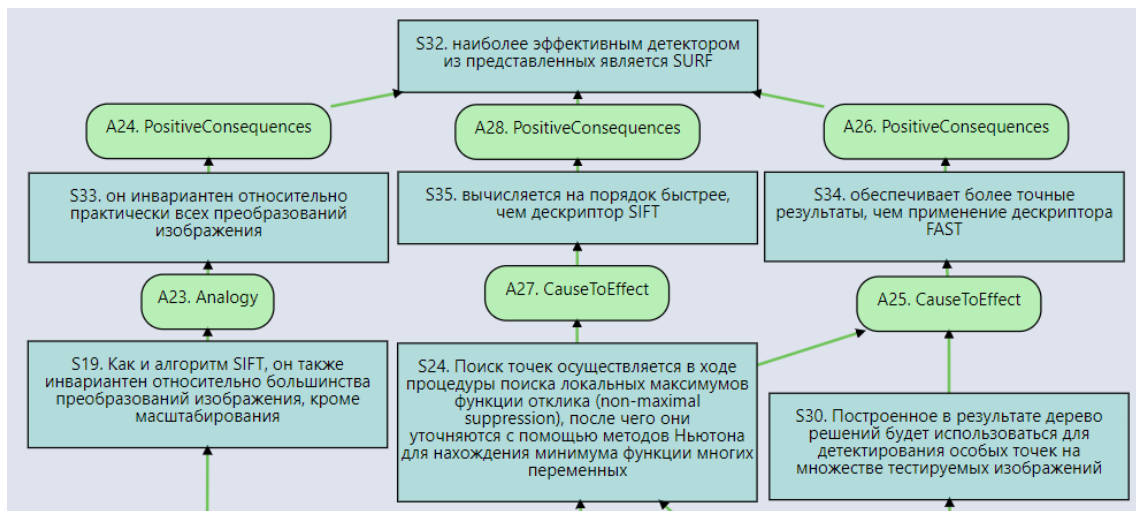


Рисунок 4 – Фрагмент графа с повторением схемы при ветвлении

Наконец, в употреблении приёмов аргументации отмечается заметное влияние тематической области. Так, среди 15 наиболее частотных приёмов (реализуемых в 5 или более текстах) треть используется исключительно в лингвистических работах. К этим приёмам относятся, например, $[E, E \rightarrow CtC]$ (такая линейная запись соответствует ветвящейся структуре с поддержкой аргумента по модели *CorrelationToCause* двумя моделями *Example*, приведённой в таблице 12: $n = 3, N = 3$), $[CtE \rightarrow CtC \rightarrow CtE]$, $[CtC, CtC \rightarrow CtE]$, $[CtC \rightarrow CtE \rightarrow CtC]$ и $[E \rightarrow CtC \rightarrow CtE]$. Несмотря на высокую частоту каждой из этих схем по отдельности (в том числе внутри текстов по информационным технологиям: по абсолютным значениям частоты указанные модели уступают в них только классификационной схеме *VC*), их комбинации применяются только в текстах одной группы.

Такая особенность обуславливается тематическими ограничениями на сочетаемость схем: модели *Example*, *CauseToEffect* и *CorrelationToCause* применяются в работах по информационным технологиям не совместно друг с другом, а с добавлением иных схем. К примеру, цепочка $[CtE \rightarrow PR \rightarrow CtE]$, отличающаяся на одну замену от $[CtE \rightarrow CtC \rightarrow CtE]$, встречается в трёх статьях по лингвистике и в четырёх по информационным технологиям.

Тематические особенности аргументации проявляются и в конфигурации деревьев, что особенно заметно при $n > 4$. Деревья, характеризующие приемы аргументации, применяемые в лингвистических текстах, значительно реже демонстрируют интенсивное ветвление в вершинах (кроме корневой вершины), в отличие от текстов компьютерной направленности. Исключение составляют случаи применения моделей *Example* и *VerbalClassification*, которые встречаются при интенсивном ветвлении в вершинах деревьев, характеризующих аргументационные приемы в текстах обеих тематик.

4.1.2 Приёмы аргументации в текстах разных жанров

Развитием первоначального исследования приёмов аргументации в научных статьях выступает их анализ на материале большего числа жанров, а именно текстов из области научной коммуникации [Pimenov, Salomatina, 2024c; Пименов, Саломатина, Тимофеева, 2025]. Новое исследование обращено как к научным, так и научно-популярным статьям (в рамках этой работы под ними понимаются аналитические статьи значительно большего объёма, чем короткие научно-популярные тексты в разделе 3.1.2), а также новостям науки. Научные статьи (S) задействованного корпуса включают 25 статей, использованных в исследовании из пункта 4.1.1, и 25 новых текстов, размеченных теми же аннотаторами. Две тематические области, лингвистика и информационные технологии, представлены поровну среди 50 статей (по 25 текстов обеих областей). Научно-популярные статьи (H) взяты с онлайн-ресурса Habr (Habr.com), тогда как новостные тексты (N) найдены с помощью научно-информационного портала «Поиск», на котором

представлены актуальные новости мировой и отечественной науки (poisknews.ru). И использованные научные новости покрывают различные области научного знания: археологию, биологию, математику, психологию, физику, химию, экологию, исследования искусственного интеллекта.

Двойная разметка текстов трёх жанров произведена совокупно пятью аннотаторами по аналогичным принципам, что и разметка корпуса научных статей. Два аннотатора размечали научные статьи, два работали с научными новостями и статьями *Nabr*, один аннотатор размечал тексты всех трёх жанров. Количественные характеристики полученного корпуса представлены в таблице 13 (Text – число текстов, Stat – число утверждений, Arg – число аргументов, Agr_Stat – согласие аннотаторов на уровне утверждений, Agr_Sch – согласие аннотаторов на уровне аргументационных схем). Согласие разметчиков в научных статьях подсчитано посредством альфа Криппендорфа, в статьях *Nabr* и новостных текстах – с помощью специализированного коэффициента, приводимого в работе (Сидорова и др., 2023).

Таблица 13 – Характеристики корпуса научной коммуникации

Жанр	Text	Stat	Arg	Agr_Stat	Agr_Sch
S	50	2502	2155	0.67	0.44
H	20	1681	1415	0.31	0.67
N	28	859	758	0.48	0.70
<i>Всего</i>	98	5042	4328		

Выявление приёмов в текстах трёх жанров проведено теми же методами, что и на ограниченной части научных статей. Единственное исключение: анализу подвергнуты исключительно вершины-схемы, без выделения главного тезиса текста в техническую вершину *Main Thesis* (для удобного сопоставления приёмов между разными жанрами: в научно-популярных статьях, написанных в более свободном стиле по сравнению с научными, и в некоторых новостных текстах

выделяется сразу несколько потенциально главных тезисов без возможности установить наиболее важный).

Анализ составных приёмов аргументации в текстах корпуса корректно предварить рассмотрением элементарных схем, соединения которых и формируют более сложные приёмы. Частотный спектр схем с указанием их относительных частот в текстах каждого жанра приведён в таблице 14. Представлены 15 наиболее частых схем в совокупности по всем жанрам, полужирным шрифтом выделены наиболее частые схемы для каждого жанра.

Таблица 14 – Частоты отдельных схем в текстах научной коммуникации

Схема	S	H	N
<i>Verbal Classification</i>	0.13	0.03	0.01
<i>Example</i>	0.12	0.09	0.05
<i>Practical Reasoning</i>	0.12	0.05	0.06
<i>Cause to Effect</i>	0.11	0.20	0.11
<i>Positive Consequences</i>	0.10	0.03	0.07
<i>Correlation to Cause</i>	0.09	0.02	0.03
<i>Part to Whole</i>	0.08	0.13	0.04
<i>Applied Method</i>	0.05	0.02	0.07
<i>Sign</i>	0.05	0.06	0.05
<i>Negative Consequences</i>	0.05	0.03	0.01
<i>Expert Opinion</i>	0.04	0.07	0.24
<i>Logical Conflict</i>	0.01	0.06	0.03
<i>Evidence to Hypothesis</i>	0.01	0.04	0.07
<i>Popular Opinion</i>	0.01	0.03	0.03
<i>Effect to Cause</i>	-	0.02	0.06
<i>Все перечисленные схемы</i>	0.97	0.88	0.93

Статистика реализации отдельных схем отражает различия в аргументации между текстами трёх жанров на самом базовом уровне. Авторы научных текстов пользуются для обоснования тезисов наименьшим числом уникальных схем (15 схем, указанных в таблице, покрывают 97% аргументов в этой части корпуса, содержащей больше всего текстов). Такая последовательность аргументации может определяться существующими стандартами в написании научных статей. Наоборот, авторы статей *Nabr* следуют неформальному стилю с наибольшим разнообразием аргументационных схем для поддержания интереса читателей и стимулирования обсуждений в комментариях. Среднее разнообразие в применении схем в научных новостях объясняется нацеленностью журналистов на представление результатов исследований широкой аудитории в нейтральной и простой для восприятия манере.

Ограниченный набор схем, характерный для научных статей, отличается сходной интенсивностью реализации наиболее частых из них: три преобладающие схемы (*VC*, *PR*, *E*) практически не отличаются по частоте и близки следующим трём (*CtE*, *PC*, *CtC*), а вместе эти шесть схем покрывают 67% аргументов в научных текстах. Научно-популярные статьи *Nabr* сочетают наибольшее разнообразие аргументации по числу схем с доминированием двух основных, *CtE* и *PtW*: выведение причинно-следственных связей и анализ целого по компонентам образуют нейтральный базис доказательств, дополняемый многообразием более редких экзотических схем. Стоит отдельно отметить выраженность полемической аргументации в статьях *Nabr* по сравнению с двумя другими жанрами: схема логического противоречия отмечена в них частотой в 6% и уступает, помимо двух основных моделей, только схемам «от экспертного мнения» и «от примера». Атаки на аргументы встречаются в два раза реже в новостях (3%) и, как правило, соответствуют случаям расхождения между цитируемыми экспертами, а коротким научным статьям вовсе не свойственны (1%), так как ограничения на их объём способствуют развитию авторских доказательств без эксплицитного оспаривания мнений других исследователей. Аргументация же в научных новостях, обращённых к широкой публике, строится

во многом на апелляции к авторитету учёных: схема «от экспертного мнения» встречается почти в четверти всех аргументов (24%) с большим отрывом по частоте от следующей модели (CtE , свойственной также и научно-популярным, и научным статьям).

Различия на уровне отдельных аргументационных схем влекут и расхождения в составных приёмах. Так, в корпусе выявлено 767 различных составных приёмов, встретившихся лишь в текстах одного жанра, и только 88, реализующихся в текстах двух или трёх жанров. Число приёмов в обоих случаях указано с учётом фильтрации вложенных: для каждой пары приёмов с разным числом вершин, совпадающих по текстовой частоте, проверяется, не является ли меньший подграфом в большем, а если является, то исключается из списка приёмов как функционирующий только в составе расширенной структуры. В таблице 15 отражено разбиение всех 855 приёмов корпуса по их встречаемости: либо в текстах только одного жанра (S, H, N), либо в текстах некоторой комбинации жанров.

Таблица 15 – Число составных приёмов в текстах каждого жанра

	S	H	N	S + H	S + N	H + N	S + H + N
Приёмы	579	132	56	32	22	13	21

Приёмы, реализованные в текстах всех жанров, представлены цепочками из двух схем, большинство из которых содержат модель причинно-следственной связи (исключениями являются цепочки $PtW \leftarrow EO$, $E \leftarrow EO$, $PtW \leftarrow PC$, $VC \leftarrow E$). Наибольшей суммарной текстовой частотой по всему корпусу характеризуются цепочки $CtE \leftarrow CtE$ (17, 9, 8 в S, H, N соответственно), $CtE \leftarrow E$ (17, 10, 6) и $CtE \leftarrow EO$ (12, 4, 12). Самая распространённая из них встречается в трети текстов корпуса. Первая и вторая из указанных цепочек наиболее частотны в научных статьях (среди всех цепочек, реализованных в текстах трёх жанров), вторая – также в статьях $Набр$, а третья – в новостных текстах. На уровне относительных текстовых частот (по отношению к общему числу текстов жанра в

корпусе) отмечается тяготение цепочек к отдельным жанрам. В целом частота цепочки зависит как от частот схем в её составе, так и от их позиции: приём $CtE \leftarrow PO$ с поддержкой причинно-следственной связи её общепризнанностью реализован в 12 текстах (4, 5, 3), тогда как обратная цепочка $PO \leftarrow CtE$, где объясняются причины распространённости некоего мнения, встретила лишь в 7 (2, 3, 2). Выраженность схемы CtE в приёмах, общих для всех жанров, объясняется её достаточно высокой частотой в каждом из них (10%, 20%, 11%), тогда как, к примеру, модель PtW , активно реализуемая в научных и научно-популярных статьях, практически не свойственна новостям.

Приёмы, встречающиеся в текстах двух жанров, отражают наибольшую близость аргументации в научных и научно-популярных статьях. Наиболее частым из них выступает приём $PtW \leftarrow (PtW, PtW, PtW)$, где скобки отражают параллельное присоединение трёх вершин PtW к корню PtW . Такая конструкция отвлечённого и скрупулёзного анализа целого по частям тяготеет к текстам *Nabr* (встречается в 7 статьях *H* из 20 и лишь в 5 *S* из 50), тогда как коротким научным статьям, ввиду большей строгости стиля и ограничений на объём, более свойственны приёмы с конкретными схемами (например, через обращение к классификациям).

Среди приёмов, общих для новостных и научно-популярных текстов, более частотные также характеризуются тяготением к одному жанру по сравнению с менее частотными, реализуемыми в обоих со сходной интенсивностью: цепочка $EtH \leftarrow EO$ с приведением экспертных свидетельств в поддержку некой гипотезы выявлена в 11 текстах (9 новостных, 32%, и 2 научно-популярных, 10%), тогда как пара $LC \leftarrow E$, с атакой тезиса через пример противоположного, встретила в 5 статьях *Nabr* и лишь в 2 новостях, а, например, приём $CtE \leftarrow (CtE, EO)$ и цепочка $PR \leftarrow ConflictingGoals$ отмечены лишь в 2 и 3 текстах из *H* и *N* соответственно.

Наконец, общие для научных статей и новостей приёмы преимущественно образованы схемами практического плана, применимыми для представления актуальных исследовательских результатов в обоих жанрах, а именно PR , AM , PC : например, $PR \leftarrow CtE$ (56% текстов из *S*, 14% текстов из *N*), $AM \leftarrow PR$ (21% и

34%), $PR \leftarrow EO$ (10%, 18%), $CtE \leftarrow (PC, PC)$ (12%, 7%). Лишь 2 приёма из 22 не содержат ни одной из этих схем: $CtC \leftarrow CtE$ (28%, 14%) и его менее частотное расширение $CtC \leftarrow (CtE, EO)$ (8%, 7%). Три отмеченных схемы сближают аргументацию в двух жанрах, тогда как иные модели, как правило, либо не характерны для какого-то одного, либо проявляются и в научно-популярных статьях. Общие приёмы тяготеют к научным статьям, отмечаются заметно меньшей относительной текстовой частотой в новостях науки: цепочка $PR \leftarrow EO$ является единственным исключением. Это может свидетельствовать об одностороннем стилистическом влиянии научных статей на новости, что не наблюдается в обратную сторону.

Жанровая специфика аргументации проявляется и на уровне приёмов, ограниченных текстами лишь одного жанра, в том числе в аспекте их разнообразия. Наибольшее число приёмов в научных статьях (579 против 132 и 56 для Н и N) сопряжено с устойчивой вариативностью доказательств: реализуется меньше уникальных схем, которые при этом близки по частоте без преобладания некоей одной. Количество текстов и аргументов в текстах жанра влияет слабее: текстов в N больше, чем в Н, но приёмов в них меньше в два с лишним раза, тогда как в сравнении S с N, аргументов в научных статьях больше в 1.5 раза, а составных приёмов – в четыре. В свою очередь, наименьшее разнообразие приёмов в новостях обуславливается спецификой схем в этих приёмах: из 56 приёмов, уникальных для N, 21 содержат модель EO , 25 – EtH , причём семантика этих схем (ссылка на экспертное мнение либо обобщённое приведение свидетельств некоторой гипотезы без уточнения, как именно эти свидетельства её подтверждают) сопутствует их реализации в кратких путях доказательства без расширения иными схемами. Как следствие, приёмы в новостях характеризуются меньшим объёмом по сравнению с приёмами других жанров: 43% приёмов в N состоят из двух схем, 45% из трёх (и лишь 12% включают четыре и более схемы), тогда как для S и Н эти показатели равны 13% и 40%, 31% и 34% соответственно.

Для наглядной иллюстрации различий приёмов разных жанров по сложности представлен пример общих подграфов с наибольшим числом вершин,

выявленных в аргументационных графах корпуса для каждого жанра. Самый сложный повторяющийся подграф в научных новостях состоит из пяти вершин-схем: *Evidence to Hypothesis* $\leftarrow$ [*Expert Opinion*, *Applied Method*, *Positive Consequences* $\leftarrow$ {*Evidence to Hypothesis*}]. Данная запись означает, что посылка в аргументе по схеме *EtH* (*Evidence to Hypothesis*) поддерживается параллельно тремя посылками по схемам *EO*, *AM* и *PC*, причём посылка в третьем аргументе (по схеме *PC*) обосновывается доводом по схеме *EtH*. Эта структура реализуется в двух новостных текстах при выражении содержания, представленного в двух следующих абзацах. Нумерация утверждений соответствует их порядку в исходном тексте.

Текст 1

(1) «Древние жители Аляски были пресноводными рыбаками 13 тыс. лет назад» $\leftarrow$ ***Evidence to Hypothesis*** $\leftarrow$ (3) «коренные американцы, живущие на территории современной центральной Аляски, могли начать ловить рыбу в пресноводных водоемах около 13 000 лет назад во время последнего ледникового периода» $\leftarrow$ [***Expert Opinion*** $\leftarrow$ «Группа исследователей обнаружила самые ранние из известных свидетельств того, что»; ***Applied Method*** $\leftarrow$ (5) «Для исследования Поттер и соавторы использовали комбинацию анализа ДНК и изотопов для идентификации 1110 образцов рыб, извлеченных из шести мест поселений людей, в том числе в бассейнах рек Танана, Кускоквим, Суситна и Коппер, на территории, которая когда-то была восточной Берингией (центральная Аляска)»; ***Positive Consequences*** $\leftarrow$ (4) «Исследование дает представление о том, как ранние люди использовали меняющийся ландшафт, и может дать представление современным людям, столкнувшимся с аналогичными изменениями» { $\leftarrow$ ***Evidence to Hypothesis*** $\leftarrow$ (6) «“Хотя мы не знаем, почему использование водоплавающих птиц уменьшилось, мы знаем, что климат менялся, – отметил Поттер. – Один из способов, которым люди смогли адаптироваться, – это интегрировать новые виды и новые технологии“» }]

Текст 2

(1) «В синаптической передаче также участвуют клетки, поддерживающие нервные волокна, – глия» ← **Evidence to Hypothesis** ← (3) «разработали нейронную сеть, которая смоделировала активное участие астроцитов в передаче сигналов между нейронами головного мозга» » ← [**Expert Opinion** ← (2) «Исследователи из Нижегородского государственного университета (Нижний Новгород)»; **Applied Method** ← (4) «За основу разработки ученые взяли спайковую нейронную сеть»; **Positive Consequences** ← (5) «Построив новую модель, группа ученых исследовала механизм астроцитарной регуляции работы нейронов в ответ на сенсорный стимул» { ← **Evidence to Hypothesis** ← (6) «Подаваемый сенсорный стимул вызывал переключение нейронной сети из режима передачи сигнала между одиночными нейронами в состояние, когда множество нейронов синхронизируются и передают сигнал в формате «волны» спайков, называемой пачкой».

Реализация одинаковой структуры в новостных текстах разных тематик (археология и математическая биология) связана с их организацией по общему шаблону: в новости представляется новая гипотеза, приводятся свидетельства в её поддержку с указанием как исследователей, выдвинувших эту гипотезу на основе обработанных ими данных, так и исследовательской методологии в обработке соответствующих данных и практической пользы исследования, которая подтверждается дополнительными свидетельствами. В ряде других новостных текстов корпуса отмечены вложенные варианты этого приёма из меньшего числа вершин (четырёх или трёх) ввиду отсутствия какого-либо вида сведений в локальном контексте графа (например, деталей методологии или вспомогательного обоснования практической пользы). Ни данный приём, ни его вложенные варианты (с аналогичной корневой вершиной) не встречаются в текстах иных жанров: ни в научных, ни в научно-популярных статьях корпуса не выявлены повторяющиеся структуры с корневой вершиной *Evidence to Hypothesis* (обоснование результатов в научных статьях основано на глубинном анализе

материала, а не на поверхностном указании свидетельств, тогда как статьям Хабра не свойственно акцентированное изложение гипотез в отрыве от контекста их практической значимости, ввиду чего схема *Evidence to Hypothesis* реализуется в них на уровне частных посылок на периферии аргументационной структуры).

Наиболее объёмный приём в научно-популярных текстах образован уже восемью схемами и имеет вид *CauseToEffect* $\leftarrow$ [*PartToWhole*, *PartToWhole*, *PartToWhole* $\leftarrow$ {*CauseToEffect*, *CauseToEffect*, *CauseToEffect*, *CauseToEffect*}]. Эта структура, включающая всего две уникальные схемы, содержит три уровня с активным ветвлением (посылка в корневой причинно-следственной связи обосновывается тремя меронимическими аргументами, один из которых раскрывается подробно через четыре причинно-следственных связи). Текстовая реализация данной структуры приведена в двух следующих абзацах.

Текст 1

(1) «Почему форумы продолжают жить» $\leftarrow$ ***Cause to Effect*** $\leftarrow$ (2) «Преимущества форумов перед другими форматами комьюнити» $\leftarrow$ [***Part to Whole*** $\leftarrow$ (3) «1. Асинхронный (более спокойный) формат обсуждения»; ***Part to Whole*** $\leftarrow$ (4) «2. Механизм репутации здесь эффективнее фильтрует контент, чем в других комьюнити»; ***Part to Whole*** $\leftarrow$ (5) «4. Форумы удобнее в технических дискуссиях» $\leftarrow$ { ***Cause to Effect*** $\leftarrow$ (6) «1. Длительный срок жизни обсуждений. Некоторые вопросы получают ответы спустя несколько месяцев или лет»; ***Cause to Effect*** $\leftarrow$ (7) «2. Удобная публикация фрагментов кода, цитат из документации, скриншотов, прикрепленных файлов»; ***Cause to Effect*** $\leftarrow$ (8) «3. Меньшая аудитория»; ***Cause to Effect*** $\leftarrow$ (9) «5. Уникальный профессиональный контекст» }].

Текст 2

(9) «Это помощник в создании креативных иллюстраций, который решает проблему с недостатком качественных изображений для статей и творческих проектов, но не заменит работу медийщиков, дизайнеров и иллюстраторов» $\leftarrow$

Cause to Effect ← (1) «Мы протестировали работу 9 самых популярных сервисов, рисующих картинки по текстовому запросу и сделали выводы» ← [**Part to Whole** ← (8) «Общее впечатление: достойный конкурент Midjourney по качеству изображений, но уступает по функциональности и стоимости тарифов»; **Part to Whole** ← (2) «Общее впечатление: качество изображений среднее, хотя встречаются интересные фотографии. Прорисовка людей в определенных стилях оставляет желать лучшего. «Кандинский» плохо прорисовывает пальцы, лица и профиль человека. Кроме того, в фотографиях по одному и тому же запросу меняется только ракурс, а идея остается исходной»; **Part to Whole** ← (3) «Общее впечатление: хорошее качество изображений, можно получить стоящие результаты, если подобрать удачные указания в запросе. Система может выдать как странную абстракцию, так и работу с объёмными детализированными объектами. Интерфейс сайта и галерея фотографий требуют доработки, так как сейчас они неудобные» ← { **Cause to Effect** ← (4) «Приложение работает на основе трёх алгоритмов: первый создаёт более фантазийные и абстрактные изображения (он называется Altair); второй – более реалистичные (Orion); третий – специализируется на рендеринге (Argo)»; **Cause to Effect** ← (5) «Дополнительно к тексту запроса можно добавить желаемый стиль изображения, либо загрузить готовую картинку, которую ИИ использует в качестве отправной точки, а также указать количество вариаций и уровень проработки»; **Cause to Effect** ← (6) «Возможность улучшать фото: можно улучшать фотографии и генерировать аналоги по изображениям»; **Cause to Effect** ← (7) «Вариации стилей/разрешений: нет ограничений по стилям, можно увеличить разрешение за дополнительные кредиты. Бесплатно доступны пять вариантов разрешений, четырех из которых вертикальные» }].

Два этих примера отражают свойственное авторам некоторых текстов Nabr явное тезисное структурирование излагаемых сведений для удобства восприятия технической аудиторией. Автор первого фрагмента использует неполные предложения и нумерованные списки, второй же фрагмент организован через

прямое повторение маркеров выстраиваемой структуры («общее впечатление» в заключении каждого блока с описанием опыта работы с каждой из рассматриваемых нейронных сетей). Подобная прямолинейность в оформлении текста (практически на манер программного кода) отличает научно-популярные статьи корпуса от научных новостей и научных статей (в том числе и научных статей по ИТ). В характеристике рассматриваемого приёма следует отметить, что его ключевой чертой является сочетание активного ветвления посылок (при не такой значительной наибольшей длине пути, равной таковой для наиболее сложного подграфа в научных новостях) и использования всего двух уникальных схем с достаточно общей семантикой, тогда как точное количество вершин на каждом уровне ветвления (три или четыре) менее значимо. Что касается специфики рассматриваемого приёма относительно других жанров, то новостным текстам не свойственен экстенсивный анализ целого по большому числу компонентов (схема *Part to Whole* встречается лишь в 4% аргументах), тогда как научным статьям более характерно варьирование типов аргументов, ввиду чего подобные ветвящиеся структуры со схемами *Part to Whole* и *Cause to Effect* включают в них и иные модели (впрочем, конструкция из одной схемы, *Part to Whole* $\leftarrow$ [*Part to Whole*, *Part to Whole*, *Part to Whole*], встречается в четырёх разных научных статьях).

Наконец, наиболее сложная повторяющаяся структура в научных статьях образована девятью вершинами и имеет вид *AppliedMethod* $\leftarrow$ [*VerbalClassification*, *VerbalClassification*, *VerbalClassification* $\leftarrow$ {*Example*}, *PracticalReasoning* $\leftarrow$ {*CauseToEffect* $\leftarrow$ (*ExpertOpinion*, *ExpertOpinion*)}]. Этот приём соответствует детализированной характеристике метода через классификацию (с иллюстрацией одного из классов примером) с практическим обоснованием предпочтения этого метода, которое, в свою очередь, раскрывается через причинно-следственную связь, обозначаемую с апелляцией к двум экспертам. В приёме отмечается активное ветвление по схеме детализированного анализа (по схеме *Verbal Classification*) с одновременным разнообразием

реализованных схем (шесть уникальных). Текстовые реализации этого приёма вынесены в Приложение Д в связи с их значительным объёмом.

Уникальные приёмы в научно-популярных статьях Habr характеризуются активной реализацией полемической аргументации, в том числе по специализированным редким схемам (таким как *General Acceptance Doubt* или *Sign From Other Events*). Более четверти из выявленных приёмов (37 из 132) содержат как минимум одну конфликтную схему. В шести приёмах отмечено сразу две схемы, которые могут объединяться как параллельно (атака одного аргумента сразу двумя возражениями), так и последовательно (упреждающее опровержение гипотетического возражения). В двух текстах из Н выявлена цепочка из трёх конфликтных схем $LC \leftarrow LC \leftarrow LC$, реализуемая как атака на аргумент с упреждающим опровержением возможного сопротивления этой атаке. Такая цепочка конфликтов усиливает эмоциональную вовлечённость читателя, что отражает актуальность экспрессивной полемической аргументации для статей Habr.

Для проверки того, насколько точно сведения о реализуемых в тексте приёмах аргументации характеризуют его жанровую отнесённость, проведён эксперимент по автоматической кластеризации 98 текстов корпуса [Пименов, 2025]. Для каждого текста строится его векторное представление в пространстве признаков, в качестве которых выступают приёмы из заданного набора (если приём реализован в тексте, то на соответствующей позиции в векторе указывается единица, иначе – ноль). Затем осуществляется группировка текстов на основании сходства векторов алгоритмами кластеризации Ward и K-Means из библиотеки Scikit-Learn. Экспериментально проверены четыре варианта формирования набора приёмов: $\Phi_1(A)$, $\Phi_2(A)$, $\Phi_n(A)$ ($n > 2$), $\Phi_n(A)$ ($n > 0$), или соответственно одиночные схемы, пары схем, структуры из трёх и более схем либо полный набор всех выявленных приёмов корпуса. В итоге установлено, что наилучшее качество кластеризации достигается при учёте всех выявленных приёмов одновременно ($\Phi_n(A)$, $n > 0$) и числе кластеров $k = 5$, при этом кластеры, построенные обоими алгоритмами, практически полностью совпадают по текстовому составу (в том

числе одинаково выделяются аномальные кластеры). Состав кластеров по текстам каждого жанра и комментарии к каждому кластеру представлены в таблице 16.

Таблица 16 – Кластеризация текстов корпуса по полному набору приёмов

Кластер	S	N	H	Комментарий
1	27	-	1	Н: обзор открытого ПО в задачах макроэкономики, обстоятельная аргументация, близкая научным текстам. S: статьи по ИТ, а также лингвистические статьи по компьютерной и корпусной лингвистике.
2	11	-	-	Статьи по иным разделам лингвистики.
3	9	-	-	Статьи по передовым и популярным информационным технологиям (3 текста по облачным технологиям, 2 по блокчейну, 2 по искусственному интеллекту).
4	-	-	13	Научно-популярные статьи по собственно ИТ.
5	3	28	6	Научные статьи с нетипично высокой долей неаргументативного содержания. Научно-популярные статьи по смежным социальным вопросам, касающимся работников ИТ.

Как следует из таблицы, в каждый кластер попали тексты преимущественно одного жанра, тогда суммарное число текстов по преобладающим в кластерах жанрах равняется 88 (доля ошибок при кластеризации равна 10% от объёма корпуса, причём отнесение текстов к кластерам иных жанров обуславливается спецификой этих текстов). Кроме того, реализация приёмов аргументации передаёт сведения не только о жанровых, но и о тематических особенностях текстов. Совокупность приёмов аргументации, характеризующих тексты одного кластера, может рассматриваться как определённая стратегия аргументации, особенно с учётом дополнительных структурных характеристик аргументационных графов (что рассматривается более подробно в главе 4.5).

4.2 Выявление методов убеждения в аргументационных структурах

Задача выявления аргументационных стратегий не сводится только к анализу структурной организации рассуждений (как отдельных элементарных и составных приёмов, так и особенностей их сочетания автором отдельной статьи), но также предполагает изучение общих подходов к подаче доводов, проявляющихся на уровне полного текста. Такое оформление излагаемых доказательств позволяет дополнить логическую составляющую аргументов иными способами воздействия на аудиторию (в частности, манипуляцией эмоциями, апелляциями к авторитету), помогающими вызвать доверие у людей разных психотипов и тем самым повысить убедительность всего текста.

Два данных направления вспомогательного воздействия при доказательстве могут быть соотнесены с двумя методами убеждения из классической концепции Аристотеля: пафосом (обращением к эмоциям) и этосом (отсылкам к авторитету). Выделяемый наравне с ними логос (обозначающий рациональную основательность доводов) выступает преобладающим в текстах научного жанра, так как регулирует фактическую связность и обоснованность тезисов. По этой причине представленное исследование направлено на анализ выражения пафоса и этоса, их соотношения друг с другом в научных статьях.

4.2.1 Оценка этоса и пафоса по аргументационным графам

Поскольку моделирование аргументационной структуры текста по стандарту AIF означает построение организованной системы лингвистических и логических единиц (тезисов и абстрактных моделей рассуждения), анализ выраженности двух методов убеждения в тексте производится интеграцией оценок по элементам обоих типов. Базовой единицей выступает отдельный аргумент как организованное целое, внутри которого реализуется связь между тезисами в процессе доказательства по некоторой модели рассуждения. В свою очередь, оценка этоса и пафоса в аргументе подразумевает анализ как лексико-

грамматического оформления тезисов (посылок и заключения, уникальных составляющих данного текста) по словарю маркеров, так и выраженности этих методов убеждения в типовых моделях рассуждения (общих между разными текстами) посредством их функционального сопоставления.

Принципы составления словаря маркеров пафоса и этоса представлены в главе 2, разделе 2.5.2.1. В разделе 2.5.2.2 описана процедура назначения весов методов убеждения для отдельных моделей рассуждения. Формулы для оценки аргументов (по выраженности пафоса либо этоса) и интеграции весов аргументов для целостных текстов приведены в разделе 2.5.3.

Построенные словари маркеров содержат 64 разных шаблона (25 для этоса, 39 для пафоса). Точность работы маркеров равна 87% на размеченных текстах, а причиной ошибок является полисемия отдельных констант в шаблонах. Например, конструкция «по новым данным» распознаётся как выражение этоса и в значении «согласно новым данным» (корректное срабатывание шаблона), и при употреблении в роли «используя новые данные» (некорректное, встречающееся во фразах вида *как только модель обучена, она может принимать решения по новым данным*). Дальнейшее уточнение шаблонов в словарях, позволяющее избежать оставшихся случаев ошибочного распознавания, приводит к частому пропуску искомых конструкций, выражающих этос либо пафос.

Проверка точности построенных словарей позволяет также определить наиболее часто употребляемые маркеры. Так, самым реализуемым шаблоном (в 88 тезисах) является отсылка к цитируемой работе через квадратные скобки с библиографическим номером внутри (маркер этоса). Близок к нему по частоте маркер пафоса, соответствующий группе одноэлементных эмфатических частиц (таких как *лишь, даже, именно* и др.; 75 утверждений). Следующие шаблоны в частотном упорядочении характеризуются менее активным употреблением: один из них образован одиночными предикатами-предписаниями (29 утверждений: *необходимо, надо, нужно* и др.), другой соответствует любым прилагательным (вне зависимости от леммы) в превосходной степени (26 тезисов). Эти два маркера указывают на выражения пафоса, тогда как сразу за ними следует *например* в

роли маркера этоса (24 утверждения). Двумя наиболее частыми многоэлементными шаблонами пафоса являются конструкция *не только / просто ..., а / но* и составная форма превосходной степени *самый + любое прилагательное* (18 и 16 тезисов соответственно). В свою очередь, самый частый многоэлементный шаблон этоса образован неопределённо-личным указанием на авторитетную группу с глаголом когнитивной семантики (например, *авторы считают*, в 5 тезисах).

Отличительной чертой аргументационных маркеров, заданных на лексико-грамматическом уровне, выступает их сравнительно редкое употребление в научных текстах: как показывает проверка словарей, 67% тезисов не содержат явных маркеров для методов убеждения (что отмечается и при экспертном анализе утверждений без срабатывания шаблонов). Эта тенденция объясняется лингвистической сложностью аргументации как явления прагматического уровня языка: обоснование тезисов в основном производится не применением особых лексических и грамматических констант, а через организацию утверждений с разнообразным содержанием в целостную аргументационную структуру. Высокая роль связей между тезисами для обеспечения убеждающего воздействия означает и важность оценивания отдельных типовых моделей рассуждения по выраженности в них пафоса или этоса.

Набор шаблонов для этоса и пафоса приведён в приложении Е в их формальном представлении для компьютерной обработки.

4.2.2 Оценка авторитетности и эмоциональности текстов

Анализ подсчитанных оценок убедительности для текстов корпуса позволяет выделить четыре группы текстов (представленные в таблице 17) с разными соотношениями этоса и пафоса [Pimenov, Salomatina, Timofeeva, 2022]. Выявление этих групп методами кластеризации демонстрирует значительное влияние предметной области на пропорцию пафоса и этоса в статье. Так, группа с наибольшим показателем $Q^P : Q^E$ содержит только работы по компьютерным

технологиям, тогда как в следующей группе по этому соотношению их в два раза больше, чем лингвистических статей. Группа же с равной интенсивностью обоих методов убеждения содержит лишь 20% от всех компьютерных статей корпуса, но более половины лингвистических.

Таблица 17 – Группировка текстов по методам убеждения с учётом тематики

Группа	Число текстов		$Q^P: Q^E$
	Лингвистика	ИТ	
1	0	2	5:1
2	2	4	4:1
3	5	6	2:1
4	8	3	1:1

Указанная тенденция обуславливается сравнительной редкостью обращений к этосу в работах по компьютерным технологиям. Редкость апелляций к авторитету вызвана прикладной ориентированностью этих статей: анализ одиночных моделей рассуждения показывает малую частоту схем с доказательством от авторитета (*ExpertOpinion*, *PopularOpinion*, *PositionToKnow* вместе покрывают лишь 6% от всех аргументов в работах по компьютерным технологиям), тогда как схемы практических рассуждений (*PracticalReasoning*, *Positive* и *NegativeConsequences*, каждая из которых выражает значительный компонент пафоса) реализуются в 16% аргументативных связей (и, наоборот, менее 5% в размеченных лингвистических статьях).

Соответственно, поскольку лингвистические статьи корпуса направлены на теоретическое описание языковых явлений, качественная оценка получаемых результатов (и использованных методов) оказывается для них менее важной, чем для прикладных работ по компьютерным технологиям, в которых отдельные программы и алгоритмы (в областях информационной безопасности, распознавания изображений, применения компьютерных технологий в медицине) анализируются преимущественно в аспекте их преимуществ и недостатков. Как

следствие, прикладная направленность этих статей влечёт активное употребление как положительной и негативной оценочной лексики (эмфатической и стилистически маркированной), так и конструкций с эксплицитной деонтической модальностью (подчёркивающих необходимость преодоления определённых сложностей либо использования неких конкретных методов).

4.3 Компьютерное распознавание этоса и пафоса в текстах

Компьютерная оценка текстов по выраженности двух методов убеждения (эмоционального воздействия, апелляций к авторитету), представленная в разделе 3.3, предполагает наличие точных сведений об аргументационной структуре таких научных статей: представленных в них тезисах, связях между этими утверждениями с учётом реализуемых моделей рассуждения. Однако выражение той или иной абстрактной модели рассуждения проявляется на уровне избираемых автором лексических средств, как потенциальных маркеров, так и менее явных конструкций, в том числе служебных единиц. Поскольку типовые модели рассуждения передают компоненты эмоционального воздействия и/или апелляций к авторитету (либо их значимое отсутствие, как в случае с *VerbalClassification*), то связь между моделью, лексикой в соединяемых ею тезисах и весами этих утверждений по этосу и пафосу может обрабатываться и в обратном направлении: не только через оценку методов убеждения в аргументе с известной схемой, но и через определение потенциальной модели аргумента на основе его предполагаемой оценки по этосу и пафосу. Такой подсчёт оценки без знания реализуемой схемы может осуществляться средствами автоматической классификации тезисов по их лексическому составу по аналогии с распознаванием утверждений двух классов (относящихся и не относящихся к аргументации).

4.3.1 Построение бинарных классификаторов для пафоса и этоса

Задача компьютерного распознавания аргументов, выражающих пафос либо этос, может быть сведена к двум аналогичным частным подзадачам бинарной классификации: выявлению отдельно аргументов с этосом, отдельно аргументов с пафосом [Саломатина, Пименов, Сидорова, 2022]. Это подразумевает подготовку двух классификаторов, обучающихся и проверяемых на разных коллекциях (исходные тексты по-разному распределяются по классам в обеих коллекциях).

Этапы обучения классификаторов для распознавания аргументов с этосом и пафосом полностью аналогичны таковым для компьютерного выявления тезисов, представленным в 2.2. Следует отметить, что в данном эксперименте испытан дополнительный алгоритм классификации LGR (логистическая регрессия, модель вероятностного прогнозирования события по логистической кривой).

Распределение аргументов по S^L и S^T (обучающей и тестовой коллекции) основано на результатах, полученных при анализе аннотированных текстов по методам убеждения m (размеченные тексты нужны для обучения алгоритмов, их применение также позволяет подсчитывать ориентировочное качество классификации при масштабировании на иные тексты, отсутствующие в коллекции). С целью сбалансированного распределения аргументов по классам пороговые значения для их разграничения задаются по среднему весу метода m в аргументах коллекциях.

Количественные характеристики S^L и S^T для этоса и пафоса приведены в таблице 18 (символы «+» и «-» обозначают число аргументов, вес которых по методу m соответственно достигает или не достигает среднего веса по этому методу среди всех аргументов коллекции). Стоит отметить, что при распределении аргументов по классам строго соблюдается принцип сохранения целостности исходных текстов, а для выравнивания объёма классов корпус был расширен двумя новыми научными статьями (одной лингвистической, одной по компьютерным технологиям).

Таблица 18 – Распределение аргументов с выражением этоса и пафоса

	Этос		Пафос	
	+	-	+	-
S^L	472	480	364	588
S^T	163	138	138	163

4.3.2 Результаты распознавания аргументов с этосом и пафосом

Результаты работы классификаторов на тестовых коллекциях аргументов (с известными весами этоса и пафоса, которые недоступны алгоритму распознавания) позволяют автоматически оценить качество классификации в терминах точности P , полноты R и F-меры. Значения этих характеристик для обоих методов убеждения по четырём алгоритмам приведены в таблице 19. Полужирным шрифтом отмечены лучшие значения каждой характеристики для обоих методов.

Таблица 19 – Оценки распознавания аргументов по методам убеждения

	Этос			Пафос		
	P	R	F-мера	P	R	F-мера
MNB	0.59	0.72	0.65	0.62	0.65	0.64
SVM	0.56	0.67	0.61	0.82	0.33	0.47
MLP	0.54	0.57	0.56	0.74	0.49	0.59
LGR	0.53	0.58	0.56	0.89	0.46	0.61

Проведённый эксперимент показывает, что аргументы с этосом и пафосом в научных текстах могут распознаваться с F-мерой более 60% (65% для этоса, 64% для пафоса), причём наилучших результатов по F-мере достигает алгоритм MNB. Однако при необходимости как можно более точного извлечения

аргументов с пафосом (без учёта полноты их покрытия) предпочтительными выступают три других алгоритма, в особенности LGR.

Ошибки в распознавании аргументов проанализированы отдельно по каждому классификатору с сопоставлением четырёх возможных сочетаний классов (один аргумент может выражать оба метода убеждения (E+, P+), либо только пафос (E-, P+), либо только этос (E+ , P-), либо не выражать ни одного метода (E-, P-). Как следствие, можно выделить четыре типа ошибок при выявлении каждого метода. Проиллюстрируем их на примере распознавания этоса: классификатор может найти этос в аргументе, где он отсутствует, причём этот аргумент способен как выражать пафос (E-, P+), так и не выражать (E-, P-); с другой стороны, пропуски этоса в аргументах, действительно его выражающих, могут проявляться с разной частотой в зависимости от наличия в нём пафоса (E+, P+) либо отсутствия (E+, P-).

Анализ распределения ошибок разных типов позволяет выделить следующие наблюдения:

1) Эмоциональные отсылки к авторитету выявлять легче, чем нейтральные, даже если пафос не учитывается явно при обучении. Так, все четыре алгоритма выявляют этос надёжнее, если он сопровождается пафосом: ошибок вида False Negative (пропуск положительного случая, игнорирование этоса в этом контексте) для аргументов «E+, P+» значительно меньше, чем для «E+, P-».

2) Если этоса нет в аргументе, наличие или отсутствие пафоса не влияет на корректность алгоритма (число ошибок вида False Positive для «E-, P+» и «E-, P-» близко во всех алгоритмах при распознавании этоса).

3) Наоборот, пафос сложнее выявлять в аргументах, реализующих апелляцию к авторитету. Так, для всех алгоритмов False Negative значительно меньше в «E-, P+», чем «E+, P+».

4) Однако если пафоса нет в аргументе, то наличие/отсутствие этоса также не влияет на корректность алгоритма.

Таким образом, выражение пафоса облегчает выявление этоса (что вполне естественно для эмоциональных апелляций к авторитету), но при этом выражение

этого размыкает для алгоритмов выявление пафоса (при эмоциональных отсылках к авторитету оказывается, что классификатор выносит решение преимущественно на основе ссылок на авторитет, тогда как эмоциональные выражения оказывают лишь второстепенное влияние на работу алгоритма).

Анализ распределения ошибок по аргументам разных типов может быть уточнён указанием типичных случаев неправильного распознавания в зависимости от лексического содержания тезисов. Ошибки в распознавании этоса вызываются следующими факторами:

1) смешение отсылок к иным исследованиям с авторскими результатами (*установлено, что, на данный момент авторы чаще не учитывают, анализ показал и т. п.*);

2) нетипичные конструкции, отсутствующие в S^L (*часто можно прочитать, что, находится в центре внимания и т. п.*);

3) указание на конкретного автора по его имени (*раскрывает Гусейнов, О. П. Ермакова конкретизирует и т. п.*);

4) неявное выражение этоса, раскрываемое через контекст других аргументов (*сегодня во многих вузах существуют как скрытая посылка от распространённой практики для обоснования актуальности исследования, не рассматривается вовсе в силу некоторых особенностей*).

В свою очередь, причины ошибок в распознавании пафоса можно свести к двум ключевым:

1) применение маркеров, отсутствующих в S^L (*вынужден, неизбежно и т. п.*);

2) приведение литературных примеров для языковых явлений.

Методы предотвращения таких ошибок заключаются в следующем:

1) пополнении словаря маркеров (также при расширении обрабатываемых данных, увеличении объёма размеченного корпуса);

2) реализации в процедуре распознавания дополнительного блока для разграничения авторского и неавторского текста.

Таким образом, ошибки при выявлении этоса являются более вариативными, чем в случае с пафосом, что согласуется с соотношением точности и полноты между методами убеждения. Все четыре алгоритма часто находят этос, но при этом больше ошибаются; однако находят не так много аргументов с пафосом, но определяют его надёжно.

4.4 Классификация аргументов по функциональной близости

Анализ аргументационных структур на уровне элементарных и составных приёмов аргументации позволяет установить специфичные тенденции в употреблении авторами научных статей типовых моделей рассуждения. Однако при моделировании таких структур на основе одноуровневой классификации аргументационных схем не так явно представляются различия между моделями рассуждения по их функциональному сходству друг с другом.

Подобные различия косвенно проявляются при обработке аргументационных аннотаций, в частности, при изучении частот отдельных приёмов и их зависимости от тематики текста: так, значимое различие между статьями по компьютерным технологиям и лингвистическими работами на уровне элементарных приёмов заключается в частом применении в первых схем *PracticalReasoning*, *Positive* и *Negative Consequences*, которые практически не употребляются в текстах второй группы. Сопоставление моделей по выраженности методов убеждения также демонстрирует возможность ранжировать их по интенсивности пафоса либо этоса, а следовательно, демонстрирует близость одних схем другим и удалённость от иных.

Подобные опосредованные зависимости в реализации приёмов аргументации могут быть выражены явно через классификацию моделей рассуждения на основе функциональной специфики их применения. Последующее изучение сочетаемости аргументов из разных функциональных групп в полных цепочках рассуждений (от начальных посылок, не обосновываемых далее, до главного тезиса текста) позволяет выявить общие

тенденции организации доказательств в научных статьях корпуса, а также предоставит возможность сравнения текстов в аспекте их аргументационного сходства.

4.4.1 Функциональное сопоставление моделей рассуждения

Представление аргументационной структуры текста посредством модели AIF предполагает семантическую интерпретацию связи между утверждениями в составе каждого аргумента. Такое определение реализуемой схемы рассуждения может осуществляться посредством как одноэтапного указания её детальной специфики (например, через выбор подходящей модели из компендиума Д. Уолтона), так и через вспомогательное указание обобщённых функциональных особенностей перехода в качестве промежуточного шага. Первый вариант сопряжён с большими практическими сложностями при аргументационной разметке коллекции текстов, тогда как определение общих функциональных черт для впоследствии уточняемых моделей рассуждения позволяет соотносить их по степени сходства или различия.

Для обозначения набора функционально близких моделей аргументации, представляемых через одну обобщённую модель на промежуточном этапе, мы применяем понятие функциональной аргументационной группы (ФАГ). Состав каждой группы определяется через сопоставительный анализ детализированных аргументационных схем (из классификации Д. Уолтона) по содержательным особенностям их использования: явность или неявность убеждающего воздействия, ракурс представления сведений, сочетаемость с иными схемами в процессе рассуждения, позиционная вариативность в структуре полного текста (зависимость/независимость расположения в общей структуре от позиции, повторяемость схемы между последовательными или параллельными связями).

Статистика использования детализированных моделей рассуждения в 1174 аргументах корпуса (из 30 научных статей) приведена в разделе 3.1.3 в таблице 7. Показано, что 15 схем, встречающихся более чем в 20% текстов, характеризуют

практически полный набор аргументов коллекции (96%). Эти схемы, отобранные для функциональной классификации (за исключением схемы Logical Conflict), распределяются согласно вышеперечисленным особенностям их использования по четырем различным классам, которые представлены в таблице 20.

Таблица 20 – Состав функциональных аргументационных групп

ФАГ	Аббр.	Схемы аргументации
От авторитета (From Authority)	FA	Expert Opinion, Popular Opinion, Position to Know
От практической значимости (From Practical Value)	PV	Applied Method, Practical Reasoning, Positive / Negative Consequences
Анализ причинно-следственных связей (Causal Analysis)	CA	Cause to Effect, Correlation to Cause
Через детализацию (Through Elaboration)	TE	Analogy, Example, Part to Whole, Sign, Verbal Classification

Первая группа (аргументы от авторитета, FA) включает три схемы, доказательство посредством которых обращается вне рассматриваемых реалий ко мнениям на их счёт. Прямое указание на авторитет эксплицирует убеждающее воздействие таких моделей. В рассмотренных научных текстах для рассуждений от авторитета характерно употребление посылок, не обосновываемых иными доказательствами: указание на весомость источника делает менее актуальным воспроизведение детализирующих шагов при достижении этим источником цитируемого вывода (а раскрытие приводимых рассуждений снижает потребность в акцентировании авторитетности источника).

Тем не менее на Рис. 5 приведены два примера нетипичного употребления схемы *PopularOpinion* (с развитием посылки о распространённом мнении через её дальнейшее уточнение). В одном случае (в аргументе, отмеченном вершиной A26) такое уточнение ограничивается приведением дополнительного примера, тогда

как в другом аргументе (A19) ссылка на распространённое мнение сопровождается указанием на частные расхождения среди сторонников этого мнения. Автор статьи объясняет причины такого расхождения, что приводит к заметному удлинению цепочки рассуждений и её разветвлению в аргументах A24 и A30.

Вторая группа (аргументы от практической значимости, FPV) образована схемами, при которых переход между утверждениями строится через сопоставление собственных особенностей представляемой сущности со внешними ориентирами предпочтительности или нежелательности актуализации таких свойств. Оценочный компонент, характеризующий применимость описываемых средств или методов, также усиливает убеждающее воздействие соответствующих схем. При этом для посылок, приводимых при реализации данных схем, характерно дальнейшее обоснование с раскрытием факторов, влияющих на применимость метода или средства. На Рис. 5 приведён пример реализации схемы *PracticalReasoning* (A32), наиболее частотной в этой ФАГ, с построением рассуждения от эксплицитно обозначенной задачи, актуальность которой описывается дополнительно.

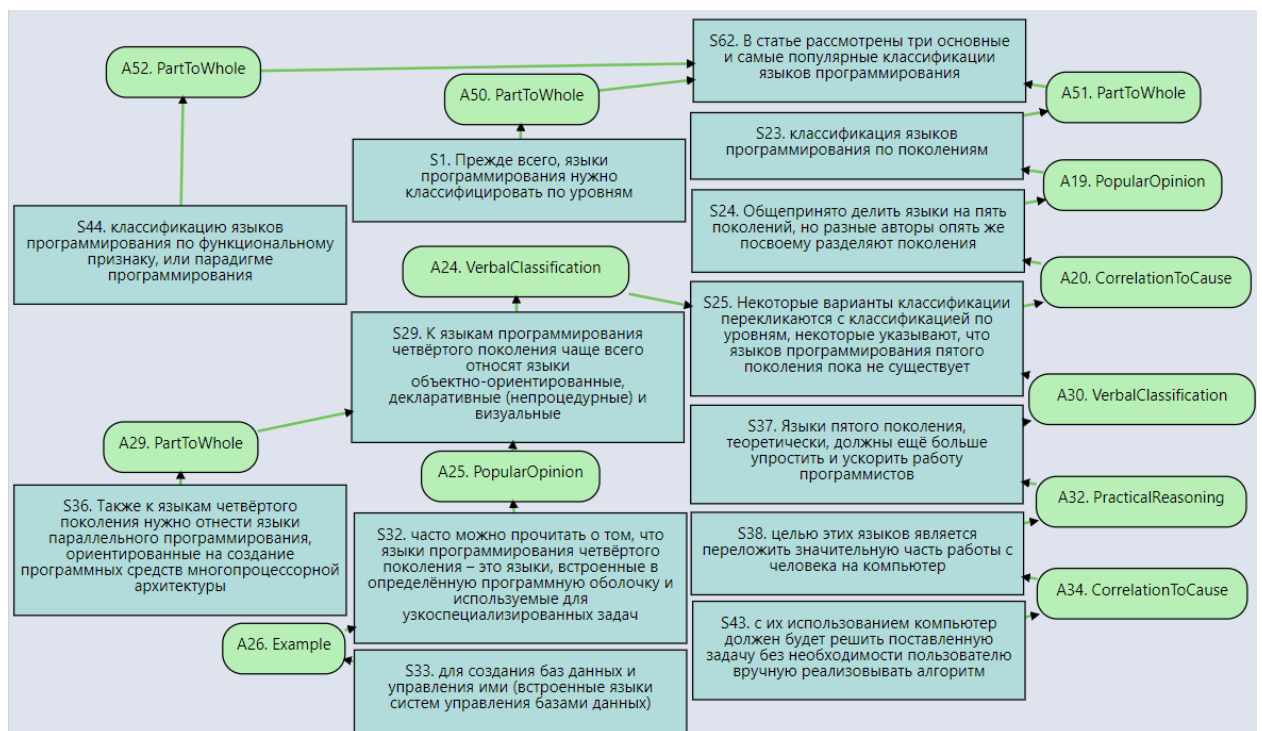


Рисунок 5 – Пример фрагмента графа со схемами разных функциональных групп

Третья группа (анализ причинно-следственных связей, СА) содержит две частотные схемы, различающиеся направлением перехода. Использование каждой из них предполагает разграничение двух или более явлений, самостоятельных на уровне выражения (причина и следствие соответствуют отдельным событиям, корреляция проявляется между сопоставляемыми явлениями). Высокая общность значений этих схем обеспечивает их наибольшую позиционную гибкость при расположении в аргументационной структуре полного текста. Например, на Рис. 5 указаны два случая реализации схемы *CorrelationToCause*: как при присоединении исходной посылки, не обосновываемой далее (A34), так и в центре цепочки рассуждения (A20).

Наконец, четвёртая группа (через детализацию, ТЕ) объединяет схемы, основанные на самостоятельном представлении некоторого отдельного явления через освещение его сторон. Это представление может производиться посредством уточнения связей между частью и целым, частным и общим, означающим и означаемым, а также иллюстрации явления с помощью примера или аналогии. Как и в случае с предыдущей группой, для рассуждений по этим схемам нехарактерно эксплицитное убеждающее воздействие. При этом в текстах коллекции схемам именно этой группы наиболее свойственно активное параллельное повторение: так, раскрытие одной части или класса дополняется контрастом с иной частью или классом. На Рис. 5 представлены примеры употребления схем *PartToWhole*, *VerbalClassification* и *Example*, где первые две реализуются именно при таком параллельном повторении (соответственно A50, A51, A52 и A24, A30).

4.4.2 Представление цепочек аргументов через функциональные блоки

Построенная классификация схем по их функциональной специфике может применяться и для анализа сочетаемости разных типов моделей рассуждения в текстах коллекции. В сравнении с характеристикой по детализированным схемам, при анализе текстов по укрупнённым группам (ФАГ) менее точно представляется

специфика отдельных статей, что компенсируется учётом функционального сходства разных схем, а соответственно позволяет выявлять тенденции более общего плана в организации аргументов. В этой связи функциональная классификация особенно значима для компьютерного определения сходства текстов по аргументационной структуре.

В представленной работе сочетаемость ФАГ анализируется посредством обработки полных цепочек рассуждений, элементы которых преобразованы в функциональные аргументационные блоки (ФАБ). Такое преобразование производится в соответствии со следующим набором шагов.

1) Исходные цепочки рассуждений извлекаются из размеченных текстов и отображают последовательность аргументационных схем в каждом полном пути (от начальной посылки, приводимой без дальнейшего обоснования, к главному тезису текста). Поэтому в цепочках учитываются только модели рассуждения без содержания соединяемых ими утверждений.

2) Последовательные детализированные схемы в цепочке заменяются блоками (ФАБ) соответствующих им функциональных групп (ФАГ).

3) Редкие схемы, не включенные в классификацию, не учитываются в преобразованной цепочке. Также пропускается схема конфликта (опровержения одного тезиса другим), встречающаяся в 2% аргументов коллекции, ввиду контрастивной специфики её употребления (что отличает конфликт от иных схем, отражающих поддержку доказательства посылками в процессе рассуждения).

4) Одинаковые функциональные блоки (соответствующие одной ФАГ) в смежных позициях преобразованной цепочки объединяются в один ФАБ. Это производится для учёта непосредственно качественных, а не количественных особенностей организации рассуждений на уровне функциональных групп.

Процесс преобразования цепочек рассуждения в последовательность ФАБ удобно проиллюстрировать на примере фрагмента аргументационной структуры, приведённого на Рис. 5. Так, в этом фрагменте представлены три полных цепочки рассуждения, принимающие следующий вид на уровне аргументационных схем:

1) *CorrelationToCause* → *PracticalReasoning* → *VerbalClassification* → *CorrelationToCause* → *PopularOpinion* → *PartToWhole* (от утверждения S43 к главному тезису S62);

2) *Example* → *PopularOpinion* → *VerbalClassification* → *CorrelationToCause* → *PopularOpinion* → *PartToWhole* (от S33);

3) *PartToWhole* → *VerbalClassification* → *CorrelationToCause* → *PopularOpinion* → *PartToWhole* (от S36).

Одноэлементные цепочки от S1 и S44 не учитываются, так как в исходном тексте они развиваются дополнительными доводами (по аналогии с S23, однако на рисунке приведены компактно). Цепочка (1) преобразуется в последовательность *CA* → *PV* → *TE* → *CA* → *FA* → *TE* (включает 6 ФАБ, реализующих все 4 ФАГ из классификации), цепочка (2) – в *TE* → *FA* → *TE* → *CA* → *FA* → *TE* (6 ФАБ, 3 ФАГ), а цепочка (3) – в *TE* → *CA* → *FA* → *TE* (4 ФАБ, 3 ФАГ). Таким образом, и в (1), и в (2) по две ФАГ представлены несколькими ФАБ (соответственно два *TE* и два *CA*, три *TE* и два *FA*), которые чередуются с функционально различающимися доводами, причём в первой цепочке реализуются все четыре выделенных ФАГ. В свою очередь, в цепочке (3) две смежные схемы из одной ФАГ (*PartToWhole*, *VerbalClassification*) объединяются в один ФАБ.

4.4.3 Сочетаемость функциональных групп в цепочках рассуждений

Употребление моделей рассуждения разных функциональных групп проанализировано на наборе из 1030 цепочек аргументов, извлечённых из коллекции аргументационных аннотаций для 30 научных текстов. Результаты анализа отдельно представлены в публикации [Пименов, 2022]. Цепочки моделей рассуждения, анализируемые в эксперименте, соответствуют полным путям от исходных посылок каждого текста (утверждений, приводимых без дальнейшего доказательства) к его главному тезису. Такие цепочки представляются в виде последовательностей аргументационных схем (между связанными утверждениями

без учёта их содержания) и варьируются длиной от одной схемы (если одно из утверждений, напрямую поддерживающих главный тезис, приведено без дальнейшего обоснования) до двенадцати, а их средняя длина равна 5.43 схемам. Для анализа сочетаемости семантически близких и отдалённых моделей рассуждения эти цепочки преобразовываются в последовательности ФАБ.

Так, при объединении семантически близких схем из 1030 исходных цепочек получено 120 различных последовательностей ФАБ (длиной от 1 до 10 элементов со средним показателем в 4.6 блока). Уменьшение средней длины цепочек связано с объединением смежных функционально близких схем (из одной ФАГ) в общий блок. Сокращение средней длины обуславливается объединением сходных схем, тогда как выбрасывание редких моделей влияет на результаты лишь в ограниченной мере (редкие схемы покрывают 4% аргументов от 1174 в экспериментальном корпусе, то есть, всего 47). Только одна исходная цепочка выведена из рассмотрения ввиду преобразования в пустую последовательность (эта цепочка состоит из единственной редкой схемы «От исключительного случая», *ExceptionalCase*, с частотой в 0.3%, не учитывающейся при функциональной группировке). В таблице 21 приведены количественные характеристики для 120 цепочек ФАБ (Sequencies) с их группировкой по числу представленных ФАГ (Groups) и указанием количества исходных путей (Paths), реализующих соответствующее число уникальных групп.

Таблица 21 – Распределение цепочек по числу функциональных групп

Groups	Sequencies	Paths
1	4	67
2	26	398
3	69	394
4	21	170
Суммарно	120	1029

Как показывает таблица 21, проанализированным научным статьям свойственно построение многоаспектных доказательств (с реализацией 2 и более групп), сочетающих аргументационные схемы разных ФАГ даже внутри последовательно развиваемых цепочек рассуждения (не только между параллельными цепочками). Так, в 93% полных цепочек (от исходной посылки до главного тезиса текста) употребляются аргументационные схемы нескольких ФАГ, тогда как все четыре группы реализуются в 17% путей рассуждения (каждом шестом). Приведённые характеристики уточняются в таблицах 22 и 23, в первой из которых отражены частоты однородных цепочек, то есть состоящих из одной ФАГ (F_a – число путей, соответствующих этой цепочке, F_t – число текстов, содержащих такие пути).

Таблица 22 – Частоты однородных цепочек ФАБ

Цепочка	F_a	F_t
CA (Causal Analysis)	30	11
TE (Through Elaboration)	22	4
PV (From Practical Value)	14	7
FA (From Authority)	1	1

Примечательной особенностью размеченной коллекции является выделение единственного пути рассуждения из 1030, содержащего лишь схемы аргументации через авторитетное мнение (без дополнения апелляций к авторитету доводами иных типов). Этот путь отмечен в обзорной статье о машинном переводе идиоматических выражений, в главном тезисе которой утверждается сохраняющаяся сложность перевода таких конструкций (доказываемая на протяжении текста через разбор возникающих частных затруднений). При этом главный тезис дополнительно поддерживается сторонним аргументом о предпочтительности человеческой проверки и редактирования любых автоматически переведённых текстов: данный довод основан на цитировании авторитетного источника и не требует дополнительного раскрытия из-за выхода

вовне поля зрения работы (посвящённой частной проблеме машинного перевода, а не всей задаче в целом). Вся цепочка из единственной авторитетной схемы имеет следующий вид (с учётом содержания утверждений): *«Несмотря на интенсивные исследования в сфере машинного перевода и перспективность данного направления, до сих пор существует ряд проблем и задач, стоящих перед разработчиками программного обеспечения, и одной из основных проблем является перевод идиоматических выражений»* ← {ExpertOpinion} ← *«Осуществление перевода в существующих системах требует вмешательства человека в качестве предредактора, интерредактора или постредактора [Зубов, 2007, с. 254]»*.

Следует отметить, что схемы ФАГ «FromAuthority» встречаются в 29% исходных цепочек рассуждения (297 из 1030). Таким образом, почти полное отсутствие однородных цепочек рассуждений исключительно от авторитета обусловлено жанровыми особенностями: для научных статей нетипично приведение чужого мнения без его расширения авторским комментарием.

Кроме того, для однородных цепочек отмечается сравнительно низкая частота наиболее активно реализуемой из них: так, рассуждения на основе единственной ФАГ «CausalAnalysis» встречаются только в трети текстов. Эта тенденция указывает на предпочтение авторами многоаспектных доказательств (цепочек рассуждения с сочетанием нескольких ФАГ).

В таблице 23 указано распределение длин цепочек ФАБ для многоаспектных последовательностей, то есть содержащих схемы из двух и более ФАГ (UT – число разных ФАГ, Length – количество ФАБ в цепочке, N(Paths) – число таких цепочек и соответствующих исходных путей).

Таблица 23 – Распределение длин многоаспектных цепочек рассуждения.

UT	Length	N(Paths)	UT	Length	N(Path)s	UT	Length	N(Paths)
2	2	9 (249)	3	3	11 (83)	4	4	1 (1)
2	3	7 (69)	3	4	23 (156)	4	5	5 (60)

Окончание таблицы 23

2	4	6 (61)	3	5	15 (90)	4	6	8 (77)
2	5	3 (16)	3	6	13 (37)	4	7	5 (30)
2	6	1 (3)	3	7	3 (6)	4	9	1 (1)
-	-	-	3	8	1 (8)	4	10	1 (1)
-	-	-	3	9	2 (13)	-	-	-
-	-	-	3	10	1 (1)	-	-	-

Другой особенностью научных текстов является активное чередование аргументов разных ФАГ при развитии цепочек рассуждений. Как показано в таблице 23, даже при ограничении двумя группами возможно их чередование на протяжении пяти-шести блоков. Такое варьирование поддерживается для трёх разных групп: четыре отмеченные цепочки имеют вид:

« $CA \rightarrow PV \rightarrow CA \rightarrow PV \rightarrow CA \rightarrow PV$ »,

« $PV \rightarrow CA \rightarrow PV \rightarrow CA \rightarrow PV$ »,

« $TE \rightarrow PV \rightarrow TE \rightarrow PV \rightarrow TE$ »,

« $CA \rightarrow TE \rightarrow CA \rightarrow TE \rightarrow CA$ ».

В частности, первая цепочка (тройное повторение $CA \rightarrow PV$) в одном из текстовых выражений соответствует представлению общего решения для крупной задачи через описание последовательных этапов с их решениями и влиянием прошлого частного решения на новый этап (схемы из группы *CausalAnalysis* передают следственные связи, из группы *FromPracticalValue* – практические ориентиры при выборе решения). С учётом содержания утверждений (приведённого сокращённо) цепочка имеет вид «На разных уровнях управления АСУ имеют неодинаковый характер и преследуют различные цели...» → {*CorrelationToCause*} → «Проблема возникает при разработке...» → {*PracticalReasoning*} → «Необходимо хранить пересекающуюся информацию...» → {*CauseToEffect*} → «Дублирование данных приводит к...» → {*PracticalReasoning*} → «Дублирование можно исключить при...» →

$\{CauseToEffect\} \rightarrow$ «Возникает задача исключения...» $\rightarrow \{PracticalReasoning\} \rightarrow$ «При проектировании корпоративных АСУ необходимо...».

В свою очередь, цепочка из пяти блоков с чередованием «СА→ТЕ» в одном из текстов характеризует вывод следствия из классификации системных причин для языкового явления, анализируемого через примеры с детальным представлением внутренних причинных связей («Обоснование причины происходит с помощью...» $\rightarrow \{CauseToEffect\} \rightarrow$ «Реагирующая реплика с reason, I can see several possible reasons why, акцентирует внимание...» $\rightarrow \{Example\} \rightarrow$ «При анализе микротекстов обнаруживается, что функция reason не всегда одинакова...» $\rightarrow \{CorrelationToCause\} \rightarrow$ «В реагирующих репликах она реализует несколько функций...» $\rightarrow \{VerbalClassification\} \rightarrow$ «Риторическую, обращенную к общим фоновым знаниям...» $\rightarrow \{CauseToEffect\} \rightarrow$ «Общим и для иницирующих и для реагирующих реплик является...»). Цепочка же из повторяющихся ТЕ $\rightarrow$ PV появляется в тексте с описанием методов в бизнес-процессе на разных уровнях.

Характеризуя общее распределение числа реализуемых ФАГ и длины цепочки в ФАБ, следует подчеркнуть тенденцию к повторению сходных схем в многоаспектных рассуждениях (цепочках с тремя или четырьмя функциональными группами). Так, для цепочек с двумя ФАГ свойственно единоразовое применение каждой без их дублирования (249 цепочек из 398 с двумя группами, 63% случаев), а повторение одного блока столь же частотно, сколько и повторение обоих. Наоборот, среди цепочек ровно с тремя ФАГ преобладают последовательности с повторением любого одного блока (156 из 394, 40%), а разовое употребление каждой из сочетаемых групп отмечается в каждом пятом случае. Такая компактная реализация всех применяемых ФАГ ещё менее характерна цепочкам, реализующим все четыре выделенные группы: лишь в одной из 170 каждая ФАГ представлена ровно одним блоком, тогда как наиболее частые цепочки такого типа (6 блоков) содержат либо повторение одного ФАБ дважды, либо повторение двух по одному разу.

Таким образом, отмеченные тенденции характеризуют следующие особенности цепочек рассуждения в проанализированных научных текстах. В первую очередь, авторам научных текстов свойственно построение многоаспектных доказательств с активным чередованием аргументов разных ФАГ в последовательном развитии доводов. Во-вторых, однородные цепочки рассуждения (с реализацией одной ФАГ) являются достаточно редкими, причём среди них почти не встречаются последовательности исключительно из апелляций к авторитету: хотя отсылки к экспертным мнениям (в частности, указания на цитируемую литературу) содержатся в 30% проанализированных цепочек, они неизменно дополняются авторскими комментариями. В-третьих, высокая вариативность аргументов разных ФАГ сопряжена с неравенством частот отдельных схем внутри цепочек рассуждения: для последовательностей с тремя и более ФАГ скорее характерно повторение довода одного типа, чем употребление всех с равной частотой. В-четвёртых, выявленная вариативность цепочек доказательства на уровне ФАБ свидетельствует о применимости тех как для аргументационного аннотирования текстов (на начальном этапе, перед детализацией схем рассуждения в соответствии с углублёнными семантическими классификациями по типу компендиума Д. Уолтона), так и для отдельных прикладных задач обработки текстов (например, для их кластеризации на основе аргументационного сходства). Кроме того, достаточная информативность аргументационного анализа на уровне функциональных блоков может снимать потребность в углублении к уровню детализированно размеченных схем для отдельных задач.

4.5 Кластеризация аннотаций по структурной сложности

Анализ аргументационных стратегий, реализуемых авторами в отдельных научных статьях, может осуществляться по формальным признакам их разноуровневой организации посредством сопоставления текстов и выявления их отличительных черт. Такое сопоставление предполагает выявление текстов как со

сходным построением рассуждений, так и со значительно различающимся, а соответственно влечёт упорядочение и группировку текстов по некоторой обобщающей характеристике их аргументационного устройства, объединяющей разноаспектные формальные показатели аргументационной структуры. Этой характеристикой в представленном исследовании выступает сложность аргументации, выстраиваемой в научной статье: аргументация как явление прагматического уровня языка обеспечивает связность и целостность текста, а её сложность отражает ориентированность и адаптацию публикации для восприятия определённой аудиторией.

Решается задача анализа коллекции аргументационных аннотаций $A = \{a_i\}$ ($i = 1, \dots, N_a$) с целью разбиения их на группы $\{c_j\}$ ($j = 1, \dots, K_c$) таким образом, чтобы внутри одной группы оказались аннотации, близкие по сложности аргументации, а различающиеся попадали в различные группы. Обучающее подмножество множества A с заранее известными метками $C = \{c_j\}$ кластеров отсутствует. Задача сводится к, во-первых, выбору множества признаков (приемов аргументации и характеристик всей аргументационной структуры текста), по которым проводится кластеризация, во-вторых, поиску множества C , каждый элемент которого содержит близкие по сложности аргументации тексты и оценке качества кластеров, в-третьих, к интерпретации полученных кластеров. Решение указанной задачи представлено также в публикации [Pimenov, Salomatina, 2022].

4.5.1 Оценка аргументационной сложности отдельного текста

Наборы признаков, отражающие сложность аргументации, сформированы из представленных ниже пяти характеристик графов аргументации.

1) Un – элементарные приёмы аргументации (23 модели рассуждения из классификации Уолтона, использованные аннотаторами в ходе разметки текстов). Анализ элементарных приёмов аргументации представлен ранее в разделе 3.1.

2) *In* – составные приемы аргументации (511 подграфов аргументации: 247 цепочек и 264 деревьев, общих для разных аннотаций текстов и выявленных с помощью алгоритма Корделлы VF2). Эти 511 подграфов извлечены из 30 научных статей корпуса, а предварительное исследование специфики их употребления на 25 текстах приведено в разделе 3.2.

3) *Full* – 11 признаков, характеризующих граф аргументации в целом (в том числе и информационные вершины). Данные признаки перечислены в таблице 24 вместе с обозначающими их аббревиатурами (аббр.).

Таблица 24 – Структурные признаки графов аргументации

№	Признак полного аргументационного графа	Аббр.
1	среднее число слов в утверждениях	WA
2	отношение числа схем к числу утверждений	SS
3	число уникальных моделей рассуждения	UM
4	число аргументов с несколькими посылками	SPS
5	число листовых вершин	L
6	максимальная длина полного пути в графе	MPL
7	средняя длина полного пути в графе	APL
8	максимальная полустепень входа в информационную вершину	MID
9	средняя полустепень входа в информационную вершину	AID
10	максимальная полустепень выхода из информационной вершины	MOD
11	средняя полустепень выхода из информационной вершины	AOD

Как и в иных разделах исследования, под **утверждением** (тезисом) понимается сегмент текста, соответствующий одной информационной вершине в графе. Границы утверждения определяются по его способности исполнять роль посылки или заключения в аргументах. Утверждения выделяются в составе аргументов через их разграничение на посылки и заключения, при этом они различаются по длине: утверждения могут соответствовать отдельным

предложениям, их частям (например, простым предложениям в составе сложного) либо последовательностям из нескольких предложений (если они лишь совместно выражают некую посылку или заключение и не образуют отдельный аргумент с реализацией какой-либо типовой модели рассуждения). Так, поскольку авторы разных научных статей по-разному организуют представление сведений в тексте, значение средней длины утверждений в словах (WA) варьируется между статьями.

Под **полным путём** в графе понимается цепочка тезисов от утверждения, приводимого без обоснования (соответствует листовой вершине), к главному тезису текста (обозначенному корневой вершиной).

Значения **полустепеней входа и выхода** для информационных вершин в аргументационных графах определяются соответственно как число посылок, обосновывающих тезис в этой вершине, и как число заключений, поддерживаемых им. Подсчитывается общее число посылок и заключений по всем аргументам, в которых данный тезис реализуется в заданной роли (заключения для полустепени входа, посылки для полустепени выхода).

4) P_E – 4 признака, характеризующие эмоциональность и авторитетность текста: доли утверждений (от числа информационных вершин в графе) с маркерами пафоса (P) и с маркерами этоса (E). Некоторые утверждения могут принадлежать к обоим классам (содержащие маркеры обоих типов), некоторые – ни к одному. Распознавание явных маркеров в тезисах проведено по лексико-грамматическим шаблонам из словарей пафоса и этоса согласно процедуре, описанной в разделе 3.3.

5) $Comp$ – 30 составных приёмов аргументации, представленных на уровне функциональных блоков (из аргументов четырёх функциональных групп: *From Authority*, *CausalAnalysis*, *ThroughElaboration*, *FromPracticalValue*). Эти 30 приёмов получены из 247 приёмов-цепочек (из набора признаков *In*) путём их обобщения по функциональным группам согласно процедуре, представленной в разделе 4.2. Как будет показано дальше, такой переход сохраняет достаточную

информацию о сложности аргументации, существенно сокращая признаковое пространство.

4.5.2 Сопоставление текстов корпуса по аргументационной сложности

После определения значений выбранных признаков для каждой отдельной научной статьи корпуса произведено автоматическое группирование этих текстов посредством трёх алгоритмов кластеризации (K-means, Ward, Spectral), описанных в разделе 2.6. Выбор оптимального сочетания признаков разных типов основан на подсчёте оценок сходства кластеров между алгоритмами для каждого из семи наборов признаков: $\{Un\}$, $\{In\}$, $\{Full, P_E\}$, $\{Comp\}$, $\{Comp, Full, P_E\}$, $\{Un, In, Full, P_E\}$, $\{Un, Comp, Full, P_E\}$. Наибольшее согласие алгоритмов достигнуто на наборе $\{Comp, Full, P_E\}$ (при котором тексты представляются через составные приёмы аргументации, общие структурные признаки аргументационного графа, выражение пафоса и этоса) и принадлежит диапазону 0.64–0.71 (в зависимости от выбранной пары алгоритмов и оценки сходства из подсчитываемых трёх). В свою очередь, наибольшее *среднее* сходство кластеров на разных наборах признаков для отдельного алгоритма характеризует Spectral, что отражено в таблице 25.

Таблица 25 – Среднее сходство кластеров между наборами признаков.

Мера сходства	K-means	Ward	Spectral
Jaccard-based	0.52	0.45	0.56
V-measure	0.51	0.51	0.62
FM-score	0.40	0.40	0.54
Average	0.48	0.45	0.57

Соответственно, для анализа групп текстов по аргументационной сложности выбраны результаты кластеризации от алгоритма *Spectral* на признаках $\{Comp, Full, P_E\}$. Текстовый состав четырех построенных кластеров не демонстрирует

какой-либо зависимости от аннотатора и тематики научной статьи. Данная особенность свидетельствует как о достаточном сходстве разметки между аннотаторами (чем подтверждает формально подсчитанные оценки согласия в 60–70%, указанные в разделе 3.2), так и об ограниченном влиянии приёмов аргументации на аргументационную сложность текстов (учёт элементарных приёмов уменьшает согласие алгоритмов кластеризации на 5.5% в среднем по всем парам сравниваемых алгоритмов). Иными словами, использование отдельных моделей рассуждения не столь явно соотносится с иными параметрами аргументационной сложности текста, в том числе составных приёмов аргументации при их функциональном представлении (которые косвенно отражают применение элементарных приёмов, но при этом, как их композитные сочетания, сильнее влияют на общую аргументационную сложность текста согласно нашему определению).

4.5.3 Выявление отличительных признаков между группами текстов

В таблице 26 приведены значения признаков $\{Comp, Full, P_E\}$, характерные для каждого из четырех построенных кластеров. Для удобства восприятия данные значения представлены в диапазонах от высокого значения («+») к среднему (« $\bar{+}$ ») и низкому («-»). Диапазоны определяются для каждого признака отдельно. Стоит отметить, что не все значения признаков, помеченные символом «+» («-»), означают высокий (соответственно низкий) уровень сложности аргументации: например, большое число листьев (L) говорит о присутствии в аргументации многих тезисов, неподкрепленных доказательствами. Кроме того, применение 30 функциональных составных приёмов отражено кратко двумя признаками: их разнообразием в тексте M и средней длиной блоков ФАГ, обозначаемой LB . Общая аргументационная сложность научных статей в кластере определяется совокупностью значений всех признаков. Кластеры в таблице упорядочены по сложности, от наиболее сложного (1) к наименее (4).

Таблица 26 – Количественные оценки сложности аргументации

№	{Full}								{Comp}		{P E}	
	MLP	APL	L	WA	IDM	ODM	SPS	SS	M	LB	P	E
(1)	+	+	$\overline{+}$	$\overline{+}$	–	+	+	+	+	+	$\overline{+}$	–
(2)	$\overline{+}$	$\overline{+}$	$\overline{+}$	$\overline{+}$	+	$\overline{+}$	$\overline{+}$	+	$\overline{+}$	$\overline{+}$	+	$\overline{+}$
(3)	$\overline{+}$	$\overline{+}$	$\overline{+}$	+	$\overline{+}$	–	–	$\overline{+}$	$\overline{+}$	–	$\overline{+}$	–
(4)	–	–	+	+	+	$\overline{+}$	$\overline{+}$	–	–	–	$\overline{+}$	+

Три кластера из четырех построенных выделяются значительным объемом: (2) состоит из 10-ти текстов, (1) содержит 9 текстов и (3) – 7. В самый маленький по объему кластер (4) попали 4 текста, графы которых отличаются в среднем небольшой длиной полного пути, большим числом листовых вершин (то есть большим числом тезисов без поддержки какими-то ни было аргументами), большим числом слов в информационных вершинах (пространной формулировкой тезисов), частой поддержки заключения несколькими посылками, большим числом маркеров этоса. Разнообразие применяемых приемов (разных цепочек) для текстов этого кластера невелико (от 4 до 8), превалирующими среди них являются двухблочные: *CausalAnalysis* $\leftarrow$ *FromAuthority*; *CausalAnalysis* $\leftarrow$ *ComponentAnalysis*; *FromPracticalValue* $\leftarrow$ *CausalAnalysis*.

Самый сложный по аргументации кластер – (1). Разнообразие приемов аргументации в статьях этого кластера изменяется от 6 до 16, большая часть из них содержит трехблочные приемы, а половина – четырехблочные, например, *CausalAnalysis* $\leftarrow$ *FromPracticalValue* $\leftarrow$ *CausalAnalysis* $\leftarrow$ *ThroughElaboration*, два текста – пятиблочные: *FromPracticalValue* $\leftarrow$ *CausalAnalysis* $\leftarrow$ *FromPracticalValue* $\leftarrow$ *CausalAnalysis* $\leftarrow$ *ThroughElaboration*. Графы аргументации отличаются в среднем большей длиной пути, высокой степенью дивергенции (один тезис поддерживает несколько других), а также содержат небольшое число листьев и малое число аргументов с несколькими посылками. Число маркеров этоса невелико, пафос выражается со средней интенсивностью.

В самый большой кластер (2) вошли 10 текстов. Аргументация в научных статьях этого кластера менее сложна, чем в (1), разнообразие применяемых

приемов в одной аннотации ниже – (3 – 12). В основном, это трех- и двухблочные приемы, среди которых одним из частотных и представленным практически в каждой аннотации является прием в виде цепочки: *ThroughElaboration* ← *CausalAnalysis*. Кроме того, состав блоков (уникальные схемы) отличается значительным разнообразием. И пафос, и этос выражаются в научных статьях из этого кластера заметно более интенсивно, чем в текстах из (1).

В оставшийся кластер (3) попали 7 текстов, сложность аргументации которых выше, чем в (4) и ниже, чем в остальных кластерах. Доля двухблочных цепочек выше, чем трехблочных, четырехблочная цепочка встретила только в одной научной статье. Самая частая цепочка, встречающаяся в каждой аннотации кластера неоднократно – *FromPracticalValue* ← *ThroughElaboration*. Разнообразие применяемых приемов для аннотаций этого кластера изменяется в пределах от 4 до 10. Маркеры пафоса и этоса не являются высокочастотными.

Исследование признаков, полезных для оценки сложности аргументационных структур, дает возможность на следующем этапе построить алгоритм, результатом вычисления которого будет число, характеризующее сложность аргументации. Этап кластеризации позволяет выявить допустимые границы варьирования числового значения сложности для текстов одного кластера. Один из критериев адекватности работы алгоритма – отсутствие расхождения значений сложности для научных статей из одного кластера.

В свою очередь, содержательная интерпретация кластеров позволяет выявить взаимосвязь значений формальных признаков для разных аспектов аргументации. Так, интенсивность апелляций к авторитету, средняя длина путей доказательства и доля утверждений, приводимых без обоснования, характеризуют разные стороны организации аргументов в тексте. Однако для всех текстов кластера (4) признаки из этого набора демонстрируют практически идентичную пропорцию: активные обращения к авторитету, короткие пути доказательства, большое количество тезисов без поддержки (точные количественные значения по разным текстам также очень близки). Сходным образом во всех текстах кластера (1) сочетаются значительная длина цепочек доказательства,

обстоятельное обоснование тезисов через эксплицитное указание всех посылок, предполагаемых моделью рассуждения, и практическое отсутствие эмоционального воздействия в процессе аргументации. Подробное и явное приведение посылок разных типов соответствуют значительной «ширине» графа: можно было бы предположить, что такая ширина и длина цепочек рассуждения связаны в общем случае обратной пропорцией (в рамках противопоставления экстенсивного и интенсивного подхода к аргументации: через обозначение большого числа аргументов без их подробного разбора по отдельности либо через обстоятельное развитие малого числа доводов), однако во всех текстах кластера (1) оба признака достигают высоких значений (и превосходят по ним тексты из иных кластеров).

В рамках взаимосвязи формальных аргументационных показателей следует отметить, что кластеризация текстов осуществляется по набору признаков разных типов. Вклад 11 структурных признаков и оценок методов убеждения ограничен одновременным влиянием приёмов аргументации, которые характеризуют глубинную семантику отдельных аргументов (в опосредованном абстрагировании от содержания тезисов). Однако взаимосвязь формальных признаков разных типов проявляется и при многоаспектной кластеризации. Это свидетельствует о достаточно точной передаче аргументационной специфики отдельных текстов (на уровне реализуемых в них стратегий аргументации как систем доказательств, интегрирующих различные аспекты аргументации в основе связности текста) при моделировании и формальном анализе их аргументационных структур.

4.6 Использование приёмов аргументации в жанровой атрибуции

Эффективность формальных аргументационных структур в передаче содержательной специфики текстов может быть проверена и посредством эксперимента по жанровой атрибуции. Задача установления жанра текста является классической для компьютерной лингвистики [Pungeršek, Ljubešić, 2023] и близка проблеме определения авторства в рамках более общего вопроса о

выявлении стилистических характеристик текста, указывающих на его категориальную принадлежность (отличающих его от текстов иных групп и при этом объединяющих его со сходными текстами в некой заданной). В зависимости от постановки задачи такой группой могут выступать, к примеру, тексты одного жанра, одного автора или авторов из некоторой социальной общности, таких как представители одного поколения или носители одного языка [Argomon, Koppel, Avneri, 1998]. В свою очередь, вопрос определения различительных формальных характеристик конкретных текстов на уровне прикладных задач (как правило, в рамках компьютерной лингвистики) отражает фундаментальную теоретическую проблему разграничения и описания стилей, выявления связей между устойчивыми группировками текстов и изменчивыми наборами реализуемых в них языковых конструкций [Костомаров, 2005, с. 49–52].

4.6.1 Высокоуровневые признаки в распознавании жанров

Согласно авторам статьи [Лагутина, Лагутина, Бойчук, 2021], качество автоматической жанровой идентификации повышается с использованием более сложных признаков, принадлежащих к высшим уровням языковой системы, для представления обрабатываемых текстов. Так, в указанной работе описывается жанровая классификация на основе ритмических характеристик текста, образуемых повторением слов и словосочетаний (учитываются такие фигуры речи, как анафора, эпифора, диакота и апозиопеза). Вывода о предпочтительности высокоуровневых признаков достигают и авторы работы [Kuzman, Ljubešić, 2022], в которой описан сравнительный анализ шести наборов признаков для представления текстов при жанровой классификации, трёх лексических и трёх грамматических: отдельные словоформы в исходном морфологическом оформлении (с и без фильтрации стоп-слов), их леммы, частеречные теги, морфосинтаксические дескрипторы (частеречные теги в сочетании с граммемами таких категорий, как число и падеж) и, наконец, структуры синтаксических

зависимостей (без какой-либо лексической информации), на которых и достигается наибольшая F-мера (до 0.9 для отдельных жанров).

Иным свидетельством большей эффективности высокоуровневых признаков выступают результаты эксперимента по распознаванию четырёх малых литературных жанров в статье [Vujlova, 2018]: как установлено автором, наиболее простые статистические характеристики (такие как длины слов и предложений) уступают лексическим маркерам и морфологическим характеристикам, однако и те не обеспечивают надёжного разграничения художественных текстов близких жанров. Достаточное качество ($F = 0.88$) достигается переходом на синтаксический уровень и классификацией с учётом распределения различных конструкций глагольного управления и специфики их реализации (в прототипической, усечённой или расширенной форме). Различительная способность глагольных конструкций разных типов объясняется (в расширенном комментарии в работе [Буйлова, Ляшевская, 2018]) их зависимостью от структур нарратива, свойственных разным литературным жанрам: согласно авторскому примеру, любовным романам характерно чередование эмоционально интенсивных сцен со взаимодействием персонажей (с предельным заполнением зависящих от глаголов слотов) и сцен-филлеров с более простым контрастирующим синтаксисом глагольных групп.

4.6.2 Представление текстов через совокупность приёмов

В свете указанных работ с сопоставительным анализом эффективности признаков разных типов проведённый эксперимент по распознаванию жанров на основе аргументационных структур включает сравнение приёмов разной сложности [Pimenov, Salomatina, 2025]. Представление текстов в рамках эксперимента содержит сведения исключительно о выявленных в них аргументационных приёмах, никакие иные признаки (лексические, морфологические и т. д.) не используются.

Задействованные в эксперименте тексты относятся к трём близким жанрам научной коммуникации: научным статьям (S), научно-популярным текстам с онлайн-платформы Nabr (H) преимущественно по тематике информационных технологий и научные новости (N) с научно-информационного портала «Поиск» (ориентированного на повышение интереса к достижениям отечественной науки и повышение её престижа среди широкой аудитории). Стоит подчеркнуть, что тематическая близость текстов этих жанров значительно затрудняет их распознавание на базе низкоуровневых признаков (в том числе лексических); различие жанров связано с разной организацией излагаемого материала с учётом специфики аудитории (академической, узкопрофессиональной из области информационных технологий либо предельно разноплановой). В этой связи эксперимент ориентирован на проверку информативности приёмов аргументации в характеристике жанровой специфики изложения.

Для представления обрабатываемых текстов использованы четыре набора признаков, соответствующие приёмам аргументации разной сложности: отдельные схемы, пары схем (биграммы) из связанных аргументов, более сложные приёмы из трёх и более схем (с разнообразными структурами, возможностью ветвления на разных уровнях) и, наконец, полный спектр всех выявленных приёмов аргументации (сочетание простых и сложных). В рамках эксперимента конструкции из трёх и более схем будем называть *комплексными приёмами* (приёмы из двух схем исключены из этого определения, так как выведены в отдельную категорию для наглядности сравнения). Веса для схем и биграмм соответствуют их относительным частотам в тексте (от числа всех схем либо биграмм в данном тексте), тогда как для комплексных приёмов они являются бинарными: вес комплексного приёма в тексте признаётся за единицу при его наличии в тексте, за ноль при отсутствии. Применение относительных частот для комплексных приёмов непрактично ввиду их большей структурной вариативности.

Задействованный в эксперименте корпус состоит из 98 текстов: 50 научных статей, 20 научно-популярных, 28 научных новостей. Подробные сведения о

корпусе приведены выше в разделе 4.1.2 (в частности, в таблицах 13, 14 и 15). Всего в текстах корпуса выявлены 895 приёмов, из которых 40 соответствуют отдельным схемам (и как элементарные приёмы не учтены в таблице 15), 213 – биграммам, 642 – комплексным приёмам.

Проблема жанровой атрибуции текста сводится к задаче классификации. В эксперименте использованы ранее описанные традиционные алгоритмы MNB, SVM и MLP. Ввиду ограниченного объёма данных на уровне текстов для создания тестовых коллекция применён метод K-Fold [Kohavi, 1995]: 98 текстов корпуса разбиты на 5 непересекающихся тестовых коллекций (ТК) со сбалансированным распределением жанров (в каждой ТК по 10 научных статей, 6 научных новостей и 4 научно-популярных текста, за единственным исключением ТК с 4 новостями). Для тематической сбалансированности научных статей дополнительно соблюдена одинаковая для всех ТК пропорция лингвистических работ и текстов по информационным технологиям (по 5 статей каждой тематики среди 10 научных работ). Распределение конкретных текстов по пяти ТК фиксировано для всех четырёх наборов признаков.

Процедура экспериментального жанрового распознавания идентична для каждого из четырёх наборов признаков. Для каждой из пяти ТК образуется обучающая коллекция (ОК) из текстов, не входящих в данную ТК. Для каждого текста строится векторное представление на уровне приёмов аргументации, соответствующих используемому набору признаков. На текстах из ОК производится обучение алгоритмов классификации (MNB, SVM, MLP), которые затем применяются для жанровой атрибуции текстов из текущей ТК. Показатели качества (точность, полнота и F-мера) подсчитываются отдельно для каждого из трёх жанров и преобразовываются в общие показатели для текущей ТК посредством микро- и макро-усреднения. Общие показатели для пяти ТК затем усредняются в итоговые оценки качества для каждого набора признаков.

4.6.3 Идентификация жанров посредством приёмов разных видов

В таблице 27 указаны оценки качества, полученные в эксперименте по жанровой атрибуции посредством разных видов приёмов аргументации. Для наглядности и компактности представления указаны оценки, полученные наиболее эффективным алгоритмом для каждого набора признаков, в микро-усреднении по пяти ТК.

Таблица 27 – качество распознавания жанров по приёмам разной сложности

Вид приёмов	Лучший метод	Точность	Полнота	F-мера
Схемы	SVM	0.84	0.81	0.82
Биграммы	SVM	0.90	0.89	0.89
Комплексные приёмы	MLP	0.92	0.90	0.91
Полный спектр приёмов	MLP	0.97	0.93	0.95

Как показано в таблице 27, качество жанровой атрибуции повышается с усложнением приёмов аргументации, посредством которых представляется аргументационная структура обрабатываемых текстов. Самые низкие показатели достигнуты на отдельных схемах, однако в абсолютных значениях эти оценки (82% F-меры, 84% точности, 81% полноты) свидетельствуют о вполне приемлемом качестве жанровой идентификации в условиях близости жанров. И точность, и полнота монотонно возрастают при переходе к парам связанных схем, к конструкциям из более чем двух схем и, наконец, к полному спектру всех выявленных приёмов (вплоть до достижения F-меры в 95%). Отмеченная динамика согласуется с наблюдениями других исследователей о повышении качества жанровой идентификации с применением более сложных признаков.

В таблице 28 приведено число ошибок разных типов на каждом наборе признаков (для наиболее на нём эффективного алгоритма классификации). Типы ошибок соответствуют 6 возможным вариантам неправильной жанровой атрибуции и обозначаются двумя буквами: первая соответствует настоящему

жанру текста (S для научных статей, N для научных новостей, H для научно-популярных статей Habr), а вторая – ошибочной метке, назначенной алгоритмом. Полужирным шрифтом выделены наиболее частые типы ошибок для каждого вида приёмов аргументации.

Таблица 28 – ошибки в жанровой атрибуции по приёмам разной сложности

Вид приёмов	SN	SH	NS	NH	HS	HN
Схемы	1	4	1	1	5	3
Биграммы	4	1	4	1	0	1
Комплексные приёмы	4	0	2	0	3	0
Полный спектр приёмов	1	0	1	0	3	0

Как отражено в таблице 28, при использовании различающихся по сложности аргументационных приёмов выявлены разные типы ошибок. На уровне отдельных схем часто смешиваются друг с другом научные и научно-популярные статьи (9 раз), тогда как при учёте локального контекста на уровне биграмм из связанных аргументов такая ошибка отмечается лишь единожды. Однако при классификации по биграммам 8 раз смешиваются научные статьи и научные новости, хотя на уровне схем эта ошибки произошла лишь дважды. Стоит подчеркнуть, что и на схемах, и на биграммах наиболее эффективен SVM, так что различие ошибок между наборами признаков не может быть связано с различием алгоритмов. Оно не связано и с распределением текстов между разметчиками: за исключением аннотатора, работавшего со всеми тремя жанрами, одна пара размечала научные статьи, другая – новости науки и научно-популярные тексты, ввиду чего можно было бы предположить явное разграничение научных статей от текстов двух других жанров, часто смешиваемых между собой. Тем не менее, этого не наблюдается: научные новости и научно-популярные тексты смешиваются друг с другом только на уровне схем и, лишь дважды, на биграммах, тогда как при оперировании комплексными приёмами и полным спектром приёмов такие ошибки не встречаются вовсе. Соответственно, различие

ошибок на разных наборах признаков обусловлено не внешними факторами (применяемыми алгоритмами или почерком разметчиков), а непосредственно содержательной спецификой приёмов аргументации разной сложности в передаче особенностей аргументации в отдельных текстах.

Различие ошибок при классификации по разным приёмам отражено и в таблице 29, в которой уточняется распределение неверных жанровых меток по пяти зафиксированным ТК, одинаковым для всех наборов признаков. Как видно из статистики ошибок, при представлении аргументации посредством приёмов разной сложности некорректная идентификация жанра характеризует разные тексты из разных ТК. Содержательный анализ конкретных случаев неправильной жанровой атрибуции отдельных текстов представлен после таблицы в порядке усложнения аргументационных приёмов.

Таблица 29 – распределение ошибок классификации по тестовым коллекциям

Вид приёмов	ТК1	ТК2	ТК3	ТК4	ТК5
Схемы	2	1	4	4	4
Биграммы	1	4	2	2	2
Комплексные приёмы	1	0	3	2	3
Полный спектр приёмов	0	0	2	2	1

Ошибки в жанровой атрибуции на уровне отдельных схем

Ввиду большей доступности отдельных схем для компьютерного распознавания (относительно более сложных приёмов аргументации) оценка их применимости для жанровой атрибуции представляет особенный интерес. На этом уровне алгоритм SVM практически не совершает ошибок в выявлении научных новостей, однако приписывает некорректные жанровые метки 8 научно-популярным текстам из 20. По итогам детального анализа определены две основных причины ошибок, каждая из которых соответствуют одной из двух неверных меток: в некоторых статьях Habr с высокой частотой реализованы схемы, более характерные научным новостям (причём такие схемы различаются

между некорректно помеченными текстами), тогда как иные тексты Nabr принимаются за научные статьи ввиду стилистической нейтральности аргументации, присущей академическим текстам (отсутствия экспрессивных моделей для вовлечения читателей, малого числа реализуемых уникальных схем). Распределение схем в текстах научной коммуникации отражено ранее в подразделе 4.1.2 (в частности, в таблице 14).

Обе ошибки SVM на TK1 заключаются в причислении двух текстов Nabr к научным новостям. Одна из этих статей содержит нетипично частые отсылки к авторитетным экспертам (в 17% аргументов), при этом схема *Expert Opinion* встречается лишь в 7% аргументов в научно-популярных текстах (и вовсе не появляется в 9 текстах из 20), но реализуется в каждой научной новости без исключений (и покрывает в них 24% аргументов). Второй текст посвящён экономическому прогнозированию с частыми отсылками к статистическим данным (в 19% аргументах против 5% в среднем по статьям Nabr), что сближает его с одним из подтипов новостей, направленным на количественное представление отдельных явлений (например, социальных групп в новостях о психологических исследованиях; новостям такого типа свойственно частое использование схемы *Sign*, тогда как в других новостях она реализуется реже).

Единственная ошибка SVM на TK2 связана с отнесением к научным статьям научно-популярного текста с наименьшим числом уникальных схем (всего десяти, каждая из которых типична для академического стиля). Этот текст представляет сравнительный анализ нейронных сетей для генерации изображений, ввиду чего близок по стилю к научным обзорам. Сходным образом на TK3 один текст Nabr причислен к академическому жанру, однако другие три ошибки состоят в неправильной атрибуции научных статей. Одна из них отнесена к новостям из-за частого цитирования иных исследователей (цитаты составляют значительную часть текста и покрывают 31% его аргументов). В свою очередь, одна из научных статей, сгруппированных с научно-популярными текстами, характеризуется выраженным прикладным акцентом при анализе проблем в промышленном использовании технологий виртуального моделирования, тогда как другая

организована через активное противопоставление рассматриваемых тезисов по модели *Logical Conflict*, более свойственной популярному стилю с яркой полемической аргументацией и эксплицитными атаками на чуждые аргументы (что нехарактерно для академического жанра).

На ТК4 классификатор путает метки для одного текста *Nabr* и одной научной статьи, но более примечательна неверная атрибуция двух новостей, не встречающаяся на других ТК. Одна из новостей, в которой реализовано лишь 8 уникальных схем (стилистически нейтральных), отнесена к научным статьям, тогда как другая с детализованным описанием новейших открытий в области процессов репарации ДНК причислена к научно-популярному классу (обстоятельная конкретизация описываемых явлений несвойственна кратким новостям, ориентированным на широкую публику; интенсивный анализ причинно-следственных связей в её рамках больше присущ текстам *Nabr*). Наконец, на ТК5 классификатор причисляет одну научную статью к научно-популярному классу, три настоящих текста из которого отмечены неверно: два отнесены к академическим статьям из-за построения их авторами стилистически нейтральной аргументации без использования маркированных схем, свойственных текстам *Nabr*, третий отнесён к новостям ввиду частых апелляций к экспертному мнению (в 21% аргументов).

Ошибки в жанровой атрибуции на уровне отдельных биграмм

На уровне пар схем из связанных аргументов SVM вновь работает лучше других классификаторов, однако допускает иные ошибки по сравнению с представлением текстов на уровне изолированных схем. Как установлено при анализе результатов классификации, учёт локального контекста в приёмах аргументации (уже на уровне только ближайших, соседних моделей) позволяет классификаторам надёжнее выявлять жанровые нюансы в использовании схем ценой уменьшения сходства между текстами одного жанра (ввиду их большей вариативности на уровне биграмм, численно превосходящих схемы).

Так, хотя на уровне отдельных схем многие научно-популярные тексты отнесены к другим классам из-за отсутствия специализированных схем онлайн-

дискуссий (реализуемых в текстах *Nabr* и не встречающихся в иных жанрах), комбинирование в них стилистически нейтральных схем в достаточной мере отличается от сочетаний аналогичных схем в текстах иных групп, что позволяет классификаторам выявить разницу. Сходным образом, частые апелляции к экспертам или статистике используются авторами научно-популярных статей корпуса в ином функциональном аспекте, чем в новостях. Так, в 9 различных новостных текстах (из 28, почти трети) отмечается реализация схемы «*Expert Opinion*» при обосновании свидетельств, приводимых в поддержку некоторой гипотезы (представлению которой и посвящена новость) по схеме «*Evidence to Hypothesis*». Вторая модель обобщённого указания свидетельств в поддержку некоего теоретического предположения менее свойственна практически ориентированным статьям *Nabr* (4% аргументов против 7% в новостях, как отмечено в таблице 14 в разделе 4.1.2) и реализуется в них преимущественно в контексте анализа причинно-следственных связей (по схеме «*Cause to Effect*»). Цепочка же «*Expert Opinion* → *Evidence to Hypothesis*» выявлена лишь в 2 научно-популярных статьях из 20: каждая из этих моделей не является характерной для текстов *Nabr* ввиду отвлечённой теоретической семантики (отсылка к авторитету вместо доказательства по существу в одном случае, обобщённое указание эмпирических свидетельств без глубинного анализа явлений или учёта прикладных обстоятельств в другом), что ограничивает их сочетаемость друг с другом. В контексте научных новостей, наоборот, такая поверхностная абстрактность двух схем вполне целесообразна для представления новых научных открытий широкой публике, что и влечёт их высокую сочетаемость: цепочка «*Expert Opinion* → *Evidence to Hypothesis*» является второй по частоте среди всех биграмм, выявленных в научных новостях (уступает только биграмме «*Expert Opinion* → *Cause to Effect*» из двух самых частотных в новостях схем). Последующие примеры иллюстрируют типовую организацию новостных текстов через акцентирование освещаемого научного наблюдения, ёмкое указание свидетельств, на основе которых произведён представленный вывод, и отсылка к

исследователям-авторам этого вывода (часто через прямую или косвенную цитату).

[(0) «По словам Бена Поттера и его коллег-сотрудников Университета Аляски,» → *Expert Opinion* →] (1) «предки коренных жителей этого региона, средства к существованию многих из которых до сих пор зависят от пресноводной рыбы (лосось, например) могли начать натуральное рыболовство в ответ на сокращение пищевых ресурсов во время долгосрочного изменения климата» ← ***Evidence to Hypothesis*** ← (3) «люди, жившие между 13 000 и 11 500 лет назад на территории нынешней Внутренней Аляски, питались пресноводной рыбой – налимом, сигом и щукой» ← ***Expert Opinion*** ← (2) «Исследование, опубликованное в журнале *Science Advances*, показывает, что».

Нумерация утверждений соответствует их порядку в исходном тексте. Утверждение под номером 0 приведено для раскрытия контекста с двойным акцентированием отсылки к опубликованной статье, содержание которой кратко излагается в новости, как напрямую («исследование показывает, что»), так и через представление статьи как речевого акта коллектива её авторов («по словам Бена Поттера и коллег»). утверждение (1) параллельно поддерживается утверждениями (0) и (3), биграмма «*Evidence to Hypothesis* ← *Expert Opinion*» реализуется в цепочке утверждений (1) ← (3) ← (2). Примечательно, что рассматриваемая биграмма реализуется в тексте повторно, но уже с вводом слов первого автора статьи в виде прямой цитаты (причём в новости не указано явно, извлечены ли эти слова из описываемой статьи или были высказаны первым автором в виде комментария к статье в ходе интервью), как показывает следующий пример.

(1) «Коренные американцы, живущие на территории современной центральной Аляски, могли начать ловить рыбу в пресноводных водоемах около 13 000 лет назад во время последнего ледникового периода» ← ***Evidence to Hypothesis*** ← (2) «Кости были найдены внутри домов и очагов и, как правило, ассоциировались с базовыми лагерями, а не с краткосрочными охотничьими лагерями. Отсутствие рыболовных крючков или копий на этих участках

предполагает, что ранние жители Аляски, вероятно, использовали сети и плотины для ловли рыбы» ← Expert Opinion ← (3) ««Это убедительный, основанный на фактических данных случай пресноводной рыбалки в конце последнего ледникового периода», – отметил Поттер».

Для наглядной демонстрации сходной организации научных новостей корпуса (характерного построения в них аргументации, не свойственного статьям Nabr) приведены ещё два примера реализации рассматриваемой цепочки из других новостных текстов.

(1) «Одна инъекция экспериментальной генной терапии может быть всем, что требуется, чтобы навсегда остановить размножение самки кошки» ← Evidence to Hypothesis ← (2) «Мы показываем, что этот векторный контрацептив предотвращает овуляцию, вызванную размножением, приводит к полному бесплодию и может представлять собой безопасную и долговременную стратегию контроля воспроизводства у домашней кошки» ← Evidence to Hypothesis ← (3) «пишут исследователи».

В приведённом примере отмечается предельно обезличенная ссылка на цитируемых учёных (стоит отметить, что в полном тексте исходной новости никто из авторов ни разу не упоминается по имени, возможно, ввиду тематики статьи).

[(0) «Проведя новое исследование, ученые пришли к выводу, что» → Expert Opinion →] (1) «никакой косолапости не было» ← Evidence to Hypothesis ← (3) «“Когда я изучала Тутанхамона, я не заметила каких-либо доказательств того, что он был инвалидом. Я видела мумии с косолапостью, но здесь этого нет. Мы называем эти псевдопатологическими изменениями”» ← Expert Opinion ← (2) «Одна из авторов исследования, египтолог-биомедик Софья Азиз отметила:».

Структура данного примера полностью совпадает со структурой первого примера, в том числе на уровне порядка утверждений в исходном тексте (параллельная поддержка утверждения (1) утверждениями (0) и (3), нелинейная (относительно порядка сегментов в исходном тексте) реализация биграммы «Evidence to Hypothesis ← Expert Opinion» в цепочке (1) ← (3) ← (2), аналогичное

двойное акцентирование ссылки на цитируемую статью, от более отвлечённой ссылки к конкретной, оформленной в виде слов персонализированного автора). Важно подчеркнуть, что два этих примера взяты из аргументационных графов разных разметчиков.

В свою очередь, в научно-популярных текстах Nabr ни одна из этих двух схем не образует частотных сочетаний с другими. Примечательно, что наиболее частотная биграмма со схемой «*Expert Opinion*» образована её оспариванием по редкой полемической модели «*Expertise Inconsistency*» с указанием на противоречия в экспертных мнениях (в 4 текстах из 20), а не какой-либо поддерживающей схемой. Примеры реализации цепочки «*Expert Opinion* ← *Expertise Inconsistency*» представлены ниже.

(2) «отличия в метаболизме покоя между людьми составляет 5-8%» ← ***Expert Opinion*** ← (1) «Чтобы было о чем поговорить, предлагаю остановиться на этом обзоре исследований. Из него мы узнаем, что» ← ***Expertise Inconsistency*** ← (3) «Однако в обзоре анализировались все люди, независимо от пола, роста и возраста».

(1) «В юности я фанател от Булгакова, и мне импонировали разлетевшиеся на цитаты слова Воланда:» (2) ← ***Expert Opinion*** ← «“Никогда и ничего не просите! Никогда и ничего, и в особенности у тех, кто сильнее вас. Сами предложат и сами всё дадут!”» ← ***Expertise Inconsistency*** ← (3) «Часто забывают или не знают, что Воланд пародирует Евангелие от Матфея: “Просите и дано будет вам; ищите, и найдете; стучите, и отворят вам. Ибо всякий просящий получает, и ищущий находит, и стучащему отворят”». Этот пример взят из текста, посвящённого вопросам карьерного роста в ИТ-сфере (в частности, получению большей заработной платы).

С другой стороны, увеличивающаяся вариативность между текстами одного жанра (на уровне приёмов аргументации) подразумевает возможность реализации в тексте нетипичной комбинации схем, вполне типичных для его жанра (причём такая комбинация может быть более свойственна другому жанру). Так, три из четырёх научных статей, отнесённых к новостям при атрибуции по биграммам,

структурированы как обзоры (проблем в информационной безопасности, применений технологий виртуальной реальности в медицине, категоризации коронавируса в разных вариантах английского языка), ввиду чего близки в сочетаемости схем новостным текстам с обобщением взглядов разных исследователей. В свою очередь, некоторые из научных новостей объединены скрупулёзным разбором описываемого материала, ввиду чего близки и текстам академического жанра. Любопытно, что три из четырёх таких новостей посвящены биологическим открытиям и освещают соответствующие процессы в подробных деталях (одной из них является текст о репарации ДНК, ошибочно классифицированный и на уровне схем).

Ошибки в жанровой атрибуции на уровне комплексных приёмов

При жанровой классификации по конструкциям из трёх и более схем отмечаются те же тенденции, которые проявились при учёте локальных аргументационных контекстов в случае с парами схем в связанных аргументах. Одновременно этот уровень характеризуется ухудшением качества алгоритма SVM и улучшением MLP в связи с увеличением числа проверяемых приёмов аргументации (с 213 до 642): метод опорных векторов склонен к переобучению в векторных пространствах с большой размерностью, тогда как базовая нейронная сеть, наоборот, значительно ограничена малым числом признаков при классификации по схемам или биграммам. Тем не менее, MLP повторяет некоторые из ошибок SVM: оба новостных текста, неправильно отмеченных на уровне комплексных приёмов, так же отнесены другим алгоритмом к научным статьям на биграммах, а одна из статей Nabr, отнесённая к научному жанру на схемах, вновь ошибочно причислена к этому же классу. Другие шесть ошибок, впрочем, ранее не встречались, но и они сводятся к отмеченной причине нехарактерного для жанра сочетания схем (уже в более широком контексте, чем пары связанных схем).

Ошибки в жанровой атрибуции на полном спектре приёмов

При жанровой классификации по полному спектру приёмов достигнуты наиболее высокие показатели качества. Все пять ошибок, допущенных на этом

уровне алгоритмом MLP, повторяются с уровня комплексных приёмов. Четыре другие ошибки успешно предотвращены за счёт дополнительного учёта самостоятельных частот отдельных схем и биграмм. В частности, отнесение научных статей к новостям, выявленное лишь единожды на уровне схем, тоже встречается только один раз на полном спектре приёмов. С другой стороны, сохраняется ошибка с отнесением некоторых научно-популярных текстов к академическим статьям (хотя таких ошибок практически не было на уровне схем, они достаточно частотны на двух других уровнях). Примечательно, что хотя научный текст, ошибочно причисленный к новостям, отличается от аналогично потерянного текста с уровня схем, он также характеризуется частыми отсылками к экспертным и даже просто распространённым мнениям (соответственно 14% и 4% аргументов).

Общие комментарии к ошибкам в жанровой атрибуции

Таким образом, по совокупному анализу ошибок классификации по разным видам приёмов аргументации следует акцентировать две ключевых причины некорректной идентификации жанра текста: либо 1) частая реализация в тексте приёмов, свойственных другому жанру, либо 2) редкая реализация приёмов, свойственных иным текстам этого же жанра. По мере усложнения учитываемых приёмов уменьшается число ошибок первого типа, однако ошибки второго типа становятся более частотными. Впрочем, уменьшение первых ошибок является более выраженным, и качество жанровой атрибуции повышается с использованием более сложных приёмов. Наилучшее качество достигается при использовании полного спектра приёмов, так как приёмы разной сложности отражают различные аспекты аргументационной структуры текстов (как самостоятельные частоты отдельных моделей рассуждения, так и их сочетаемость в контекстах разной ширины), за счёт чего нивелируют ошибки друг друга (ввиду одновременного учёта как нюансов различий между жанрами, так и общих для каждого жанра характеристик). Наибольшее число ошибок в эксперименте характеризует научно-популярные статьи, наименьшее же – научные новости. Это объясняется их жанровой спецификой в рамках исследуемого корпуса научной

коммуникации: научные новости представляют собой наиболее своеобразный жанр из трёх ввиду ориентированности на широкую публику, направленности научной коммуникации вовне, а не к академической либо узкопрофессиональной аудитории. В свою очередь, научно-популярные тексты с онлайн-платформы близки обоим другим жанрам: они лишены строгих стилистических ограничений, свойственных научным статьям, но при этом посвящены не менее скрупулёзному анализу излагаемого материала.

Выводы

В четвёртой главе диссертации описаны эксперименты по компьютерному анализу аргументационных структур в более сложных аспектах: извлечению составных приёмов аргументации из графов, выявлению методов убеждения на уровне отдельных аргументов, оценке полных текстов по выраженности эмоционального воздействия и апелляций к авторитету. Отдельный эксперимент заключается в кластеризации текстов по формальной сложности аргументации через интеграцию элементов аргументационных структур разных типов (на основе иных проведённых экспериментов), результатом чего становится разграничение обобщённых подходов к аргументации (стратегий) в научных статьях из проанализированной коллекции. Эффективность приёмов аргументации в передаче содержательной специфики текстов дополнительно демонстрируется экспериментом по жанровой атрибуции текстов из разных областей научной коммуникации (научных и научно-популярных статей, а также научных новостей).

Анализ составных приёмов аргументации (объединённых структур из нескольких моделей рассуждения) в научных статьях показывает, что построение параллельных доводов чаще происходит по одной и той же модели вне зависимости от темы текста. Если авторы научных текстов строят рассуждение от нескольких доводов, то они преимущественно организуют эти аргументы по одинаковой модели рассуждения (сочетание нескольких доводов через разные

модели является менее типичным). Другой особенностью в употреблении приёмов аргументации является заметное влияние тематической области на сочетаемость между схемами (среди 15 наиболее частотных приёмов треть реализуется только в лингвистических статьях, хотя образующие их отдельные схемы активно используются и в текстах по компьютерным технологиям).

В свою очередь, анализ приёмов в текстах разных жанров позволяет выявить жанровую специфику аргументации для научных статей, научно-популярных работ и научных новостей. Собственно научные статьи в корпусе научной коммуникации характеризуются наибольшим разнообразием и устойчивостью составных приёмов аргументации, научно-популярные отличаются активностью полемических моделей, тогда как новостям свойственно низкое разнообразие аргументационных структур ввиду специфики преобладающих схем. Информативность представления текстов через совокупность реализованных приёмов аргументации продемонстрирована экспериментально: кластеризация по полному спектру приёмов позволяет сгруппировать тексты как по жанровым, так и тематическим особенностям.

В рамках эксперимента по выявлению методов убеждения (этоса и пафоса) составлены словари их маркеров в виде лексико-грамматических шаблонов. Точность поиска равна 78% (причиной ошибок является полисемия отдельных констант в шаблонах). Выраженность методов убеждения в аргументационных схемах оценивается через их сопоставительный функциональный анализ. На этапе распознавания методов убеждения в отдельных аргументах достигнуты значения F-меры в 65% для этоса и 64% для пафоса. При обработке полных текстов отмечено, что прикладная направленность статей по информационным технологиям влечёт активное использование в них оценочной лексики (положительной и негативной) и конструкций с модальностью долженствования. Теоретический же характер лингвистических работ проявляется в нейтральности их стиля, но также частом обращении к авторитетам иных исследователей (как современных, так и исторических основоположников отдельных подходов в языкознании).

В эксперименте по выявлению текстов со сходной аргументационной структурой показано, что лучшие показатели кластеризации достигаются при многоаспектном сочетании формальных показателей аргументационной специфики. На исследованной коллекции выявлена взаимосвязь между значениями разных показателей (например, при разнообразии приёмов аргументации в тексте отмечается приведение малого числа тезисов без обоснования, тогда как обилие таких тезисов характеризует статьи с частыми апелляциями к авторитету).

Посредством заключительного эксперимента по жанровой атрибуции подтверждена эффективность формальных аргументационных структур в передаче содержательной специфики текстов. Установлено, что качество жанровой атрибуции возрастает с увеличением сложности используемых в классификации приёмов: от 82% F-меры по микро-усреднению для отдельных схем до 89% для биграмм, 91% для комплексных приёмов и 95% для полного спектра приёмов аргументации. В результате анализа ошибок выявлено, что причины некорректной жанровой атрибуции могут быть сведены ко двум основным: 1) реализация в тексте приёмов, характерных текстам иных жанров; 2) отсутствие в тексте приёмов, свойственных другим текстам этого же жанра. По мере усложнения приёмов аргументации увеличивается число ошибок второго вида, тогда как ошибки первого вида становятся заметно менее частотными. Среди трёх жанров наибольшее количество ошибок распознавания отмечено для научно-популярных статей, которые близки научным статьям в степени детализации представляемого материала, а наименьшее – для научных новостей ввиду их организации с ориентиром на широкую публику, а не академическую или узкопрофессиональную аудиторию.

Заключение

Ключевые итоги диссертационного исследования, полученные в результате экспериментов по формальному анализу приёмов и стратегий аргументации в коллекции текстов научных статей, представлены в нижеизложенном списке.

1. Подготовлен корпус с многоуровневой аргументационной разметкой научных статей на русском языке, включающий 100 текстов из двух предметных областей, где для каждого текста представлена двойная разметка разными аннотаторами. Вклад соискателя заключается в а) разметке 92 статей, в которых выявлено 5297 утверждений и 4563 аргументов; б) создании методологии многоаспектной разметки аргументации в научных статьях, регулирующей согласованную разметку текстов разными аннотаторами; в) первоначальном отборе текстов для разметки; г) разработке программных скриптов для автоматической модификации разметки и устранению типовых расхождений между аннотаторами. Моделирование аргументационных структур данных текстов осуществлено в соответствии с дискурсивным подходом и стандартом Argument Interchange Format. Разметка аргументации включает как выявление тезисов в оформлении языковыми средствами, так и указание структурных связей между ними с обозначением абстрактных моделей рассуждения в основе этих связей (согласно классификации аргументационных схем Д. Уолтона). Корпус может применяться для решения различных задач в области компьютерной обработки аргументации, как показано в проведённых в ходе исследования экспериментах. Сочетание языковых средств оформления аргументов и их абстрактно-логических моделей при формальном представлении аргументации обеспечивает возможность многоаспектного аргументационного анализа.

2. Показано, что аргументационные структуры в обработанных текстах образованы небольшим числом базовых моделей рассуждения (15 типовых аргументационных схем покрывают 96% всех аргументов корпуса). Отмечается регулярное использование авторами разных текстов одних и тех же составных приёмов аргументации, которые соответствуют конфигурациям из нескольких

отдельных моделей (элементарных приёмов). Выявлена тематическая зависимость в реализации как элементарных, так и составных приёмов аргументации от жанра и предметной области, к которым относится текст.

3. Отмечено, что прикладной характер статей корпуса по информационным технологиям влечёт активное употребление в них оценочной лексики и модальных конструкций долженствования при оформлении логических доказательств языковыми средствами. Такие показатели непрямого эмоционального воздействия практически не выражены в трудах по лингвистике, которым ввиду теоретического характера свойственно более активное обращение к авторитету иных исследователей.

4. Достигнуты приемлемые показатели качества при компьютерном распознавании элементов аргументации (тезисов, аргументов по методам убеждения, отдельных типов аргументов, связей в их основе) посредством реализованного программного комплекса из различных блоков компьютерной обработки аргументации. На этапе распознавания тезисов достигнута F-мера, равная 69%, для выявления аргументов с пафосом и этосом – 66% и 65%, для аргументов отдельных шести типов – от 54% до 68%.

5. Выявлена явная жанровая специфика составных приёмов аргументации в текстах разных жанров научной коммуникации. Научным статьям корпуса свойственна устойчивая вариативность доказательств при ограниченном числе уникальных схем, научно-популярные тексты выделяются частотой полемических приёмов, научные новости отличаются низким разнообразием приёмов ввиду непродуктивности их преобладающих схем в образовании регулярных составных конструкций. Посредством кластеризации текстов, представленных по совокупности реализованных в них приёмов, показана высокая различительная способность приёмов аргументации в передаче жанровой и тематической специфики текстов научной коммуникации.

6. На основе интеграции формальных характеристик аргументации в разных аспектах (обработанных в ходе отдельных предварительных экспериментов) реализован метод автоматической кластеризации научных статей

со сходной организацией аргументации. Продemonстрировано, что компьютерная обработка аргументационных структур по их формальным показателям обладает достаточной различительной способностью для разграничения групп текстов с наглядной передачей многоаспектной специфики доказательств внутри каждой группы. Эффективность формальных аргументационных структур в передаче содержательной специфики текстов продемонстрирована и посредством эксперимента по жанровой атрибуции. Это даёт возможность компьютерной характеристики стратегий аргументации, реализуемых на уровне полных текстов, по структурной и языковой специфике излагаемых в них аргументов.

В перспективе дальнейших исследований в области формального анализа аргументации следует отметить расширение корпуса (за счёт увеличения числа статей с аргументационной разметкой, в том числе иных тематик и жанров). В свою очередь, увеличение объёма данных обеспечит возможность реализации и улучшения методов компьютерной обработки аргументации (например, автоматического связывания смежных и позиционно разнесенных утверждений в аргумент на основе дискурсивных маркеров и тематической структуры текста, а также распознавания аргументов отдельных типов). Разработка новых методов компьютерного анализа аргументации в текстах на русском языке актуальна для решения отдельных прикладных задач (таких как автоматическое реферирование научных статей, или извлечение их ключевого содержания, с учётом их аргументационной структуры, компьютерная количественная оценка сложности аргументации в научном тексте для восприятия аудиторией, выявление автоматически сгенерированных псевдонаучных текстов, жанровая и тематическая классификация, проверка доказательств в студенческих работах).

Список литературы

1. Аристотель. Риторика / Аристотель // Античные риторики. – М., 1978. – С. 15–167.
2. Архипова Е. В. Системный подход к обучению языку и методическая система речевого развития школьников / Е. В. Архипова // Русский язык в школе. – 2005. – № 4. – С. 3–11.
3. Бакиева А. М. Исследование применимости теории риторических структур для автоматической обработки научно-технических текстов / А. М. Бакиева, Т. В. Батура // Cloud of Science. – 2017. – Т. 4, № 3. – С. 450–464.
4. Баранов А. Н. Лингвистическая теория аргументации (когнитивный подход): Автореф. дис. ... д-ра филол. наук / А. Н. Баранов. – М., 1990. – 46 с.
5. Батура Т. В. Математическая лингвистика и обработка текстов на естественном языке / Т. В. Батура. – Новосибирск, НГУ, 2016. – 166 с.
6. Белоусов К. И. Ментальная репрезентация научного события: содержание, структура, вариативность / К. И. Белоусов, В. Н. Гатаулин, Е. В. Ерофеева // Репрезентация событий: интегрированный подход с позиции когнитивных наук. Отв. ред. В. И. Заботкина. – 2-е изд. – М.: Издательский Дом ЯСК: Языки славянской культуры, 2017. – С. 131–149.
7. Бернацкая А. А. О понятии «текст» / А. А. Бернацкая // Журнал Сибирского федерального университета. Серия: гуманитарные науки. – Т. 2. – 2009. – С. 30–37.
8. Буйлова Н. Н. Лексико-синтаксические маркеры малых литературных жанров / Н. Н. Буйлова, О. Н. Ляшевская // Вестник ПСТГУ. Серия III: Филология. – Вып. 66. – 2021. – С. 11–23.
9. ван Дейк Т. А. Язык. Познание. Коммуникация. Благовещенск: БГК им. И. А. Бодуэна де Куртенэ, 2000. – 308 с.
10. Витгенштейн Л. Философские исследования / Л. Витгенштейн // Людвиг Витгенштейн, философские работы. Ч. I. – М.: Гнозис, 1994. – 612 с.

11. Ворожбитова А. А. Концептуальные основы лингвориторической парадигмы / А. А. Ворожбитова. Теория текста: антропоцентрическое направление: учеб. пособие (3-е издание). М., 2014. – С. 121–173.
12. Ильина Д. В. Аргументация и референция как не прямые способы выражения отношения автора к адресату научно-популярной статьи лингвистической тематики: Дис. ... канд. филол. наук / Д. В. Ильина. – Красноярск, СФУ, 2023. – 251 с.
13. Иоанесян Е. Р. Функционирование языковых единиц в аргументативном дискурсе (сопоставительный аспект) / Е. Р. Иоанесян. – М., 2022. – 214 с.
14. Иргашева Т. Г. Лингвистика текста как интегрирующая филологическая дисциплина / Т. Г. Иргашева // Наука и школа. – 2011. – № 3. – С. 56–61.
15. Караулов Ю. Н. Русский язык и языковая личность. Изд. 7-е / Ю. Н. Караулов. М.: Издательство ЛКИ, 2010. – 264 с.
16. Кибрик А. А. Рассказы о сновидениях: Корпусное исследование устного русского дискурса / А. А. Кибрик, В. И. Подлесская (ред.). М.: Языки славянских культур, 2009. – 736 с.
17. Ким И. Е. Языковая личность автора научно-популярного текста: аргументативный аспект / И. Е. Ким, Д. В. Ильина // Вестник НГУ. Серия: История, филология. – 2020. – Т. 19, № 9: Филология. – С. 19–30.
18. Костомаров В. Г. Наш язык в действии. Очерки современной русской стилистики / В. Г. Костомаров. – М., 2005. – 287 с.
19. Котельников Е. В. Извлечение аргументации из текстов и проблема отсутствия русскоязычных текстовых корпусов / Е. В. Котельников // Advanced Science. – 2018. – № 3 (11). – С. 44–47.
20. Красина М. Н. Дискурс, дискурс-анализ и методы их применения в междисциплинарных проектах / М. Н. Красина // Вестник ТвГУ. Серия «Филология». – 2018. – № 2. – С. 159–165.
21. Кубрякова Е. С. В поисках сущности языка. Когнитивные исследования / Е. С. Кубрякова // М.: Знак, 2012. – 208 с.

22. Кутузов А. Б. Методики определения сложности текста в рамках переводческого анализа / А. Б. Кутузов // Вестник Нижегородского государственного лингвистического университета им. Н. А. Добролюбова. Вып. 4. Лингвистика и межкультурная коммуникация. – Нижний Новгород: НГЛУ, 2009. – С. 30–36.

23. Лагутина К. В. Классификация текстов по жанрам на основе ритмических характеристик / К. В. Лагутина, Н. С. Лагутина, Е. И. Бойчук // Моделирование и анализ информационных систем. – Т. 28, № 3. – 2021. – С. 280–291.

24. Лапошина А. Н. Анализ релевантных признаков для автоматического определения сложности русского текста как иностранного / А. Н. Лапошина // Компьютерная лингвистика и интеллектуальные технологии: По материалам ежегодной международной конференции «Диалог». – М., 2017. – С. 1–7.

25. Левицкий Ю. А. Лингвистика текста / Ю. А. Левицкий. – М.: Высш. шк., 2006. – 207 с.

26. Ляшевская О. Н. Лексико-синтаксические маркеры малых литературных жанров / О. Н. Ляшевская, Н. Н. Буйлова // Вестник ПСТГУ. Серия III: Филология. – 2021. – Вып. 66. – С. 11–23.

27. Манерко Л. А. Научное открытие как событие: конструирование структур знания в специальном дискурсе / Л. А. Манерко // Репрезентация событий: интегрированный подход с позиции когнитивных наук. Отв. ред. В. И. Заботкина. – 2-е изд. – М.: Издательский Дом ЯСК: Языки славянской культуры, 2017. – С. 150–177.

28. Минина О. Г. Лингвистика текста: прошлое и будущее / О. Г. Минина, А. В. Карзунина // Альманах современной науки и образования. – Тамбов: Грамота. – № 12 (43). – 2010. – С. 220–222.

29. Мурзин Л. Н. Текст и его восприятие / Л. Н. Мурзин, А. С. Штерн. Свердловск: Изд-во Урал. ун-та, 1991. – 172 с.

30. Негрышев А. А. Новости в прессе: к моделированию макротекстовой структуры // Язык и дискурс СМИ в XXI веке / Колл. монография под ред. М. Н. Володиной. – М.: Академический проект, 2011. – С. 85–97.

31. Олейник А. Н. Надежность и достоверность в контент анализе текстов: выбор показателей / А. Н. Олейник, И. П. Попова, С. Г. Кирдина, Т. Ю. Шаталова // Психологический журнал. – 2014. – Т. 35, № 6. – С. 99–113.

32. Пименов И. С. Специфика аргументационного аннотирования научных и научно-популярных текстов / И. С. Пименов // Труды международной конференции «Корпусная лингвистика-2021». – СПб.: Скифия-принт, 2021. – С. 330–337.

33. Пименов И. С. Сочетаемость аргументов разных функциональных групп в научных текстах / И. С. Пименов // Тамбов: Грамота. – 2022, № 11. – С. 3672–3680.

34. Пименов И. С. Анализ расхождений в аргументационной разметке научных статей на русском языке / И. С. Пименов // Вестник НГУ. Серия: Лингвистика и межкультурная коммуникация. – 2023. – Т. 21, № 2. – С. 89–104.

35. Пименов И. С. Формальное выявление приемов аргументации в научных текстах / И. С. Пименов, Н. В. Саломатина, М. К. Тимофеева // Вестник НГУ. Серия: Лингвистика и межкультурная коммуникация. – 2022. – Т. 20, № 1. – С. 21–36.

36. Пименов И. С. Маркерный подход к автоматизации извлечения аргументативных структур из научных статей / И. С. Пименов, Н. В. Саломатина, А. С. Засыпкин // Труды международной конференции «Корпусная лингвистика – 2023». – СПб., 2024. – С. 184–191.

37. Пименов И. С. Компьютерный анализ приемов и стратегий аргументации в текстах научной коммуникации / И. С. Пименов, Н. В. Саломатина, М. К. Тимофеева // Вестник НГУ. Серия: Лингвистика и межкультурная коммуникация. – 2025. – Т. 23, № 1. – С. 93–109.

38. Пименов И. С. Анализ приёмов аргументации в текстах научной коммуникации / И. С. Пименов // Обработка естественного языка для лингвистики, IT, бизнеса: сборник материалов Международного научно-практического семинара (отв. ред. З. И. Резанова). – Томск: Изд-во Томск. гос. ун-та, 2025. – С. 42–44.

39. Саломатина Н. В. Распознавание аргументативных связей в научно-популярных текстах / Н. В. Саломатина, И. С. Кононенко, Е. А. Сидорова, И. С. Пименов // Системная информатика. – № 16. – 2020. – С. 149–164.

40. Саломатина Н. В. Распознавание в научных текстах аргументов с этосом и пафосом / Н. В. Саломатина, И. С. Пименов, Е. А. Сидорова // Тезисы Международной научной конференции «Марчуковские научные чтения – 2022». – Новосибирск, 2022. – С. 140.

41. Саломатина Н. В. Автоматическое обнаружение главного и локальных базовых аргументов в научном тексте на русском языке / Н. В. Саломатина, И. С. Пименов // Тезисы Международной научной конференции «Марчуковские научные чтения 2024». – Новосибирск, 2024. – С. 128–129.

42. Саломатина Н. В. Применение методов машинного обучения для выявления аргументативных связей в текстах научной коммуникации / Н. В. Саломатина, Е. А. Сидорова, И. С. Пименов // Онтология проектирования. – №14 (51). – 2024. – С. 82–93.

43. Самсонов Н. Б. Разработка и апробация лингвистической методики оценки когнитивной сложности научно-учебного текста / Н. Б. Самсонов, Е. В. Чмыхова, Д. Г. Давыдов // Психологические исследования. – 2015. – № 8 (41). – С. 1–6.

44. Селищев А. М. Язык революционной эпохи. Из наблюдений над русским языком последних лет (1917–1926) / А. М. Селищев. М., 1928. – 248 с.

45. Сидорова Е. А. Платформа для исследования аргументации в научно-популярном дискурсе / Е. А. Сидорова, И. Р. Ахмадеева, Ю. А. Загорулько, А. С. Серый, В. К. Шестаков // Онтология проектирования. – 2020. – Т. 10, № 4 (38). – С. 489–502.

46. Сидорова Е. А. Комплексный подход к анализу аргументативных отношений в текстах научной коммуникации / Е. А. Сидорова, И. Р. Ахмадеева, Ю. А. Загорулько, И. С. Кононенко, А. С. Серый, П. М. Чагина, В. К. Шестаков // Онтология проектирования. – 2023. – Т. 13, №4(50). – С. 562–579.

47. Солнышкина С. И. Сложность текста: этапы изучения в отечественном прикладном языкознании / С. И. Солнышкина, А. С. Кисельников // Вестник

Томского государственного университета. Филология. – 2015. – № 6 (38). – С. 86–99.

48. Соссюр Ф. Курс общей лингвистики / Ф. де Соссюр // Екатеринбург: Изд-во Уральск. ун-та, 1999. – С. 112–123.

49. Тимофеева М. К. Теория риторической структуры как инструмент анализа стихотворных текстов / М. К. Тимофеева // Вестник Томского государственного университета. Филология. – 2020. – № 68. – С. 109–137.

50. Тимофеева М. К. Теория риторической структуры как инструмент когнитивного и стилистического анализа / М. К. Тимофеева // Язык, общество, образование. – 2020. – С. 116–120.

51. Филиппов К. А. Лингвистика текста: курс лекций / К. А. Филиппов. СПб., 2003. – 336 с.

52. Чаплина С. С. Понятие «тип текста»: принципы дефинирования и проблемы классификации / С. С. Чаплина // Вестник Бурятского госуниверситета. – 2009. – № 11. – С. 132–137.

53. Чернявская В. Е. Лингвистика текста: Поликодовость, интертекстуальность, интердискурсивность / В. Е. Чернявская. М.: Книжный дом «ЛИБРОКОМ», 2009. – 248 с.

54. Якубинский Л. П. Ленин о «революционной фразе» и смежных явлениях / Л. П. Якубинский // Печать и революция. – 1926. – № 3. – С. 5–17.

55. Якубинский Л. П. О диалогической речи / Л. П. Якубинский // Якубинский Л. П. Избранные работы. – М., 1986. – С. 17–58.

56. Accuosto P. Transferring knowledge from discourse to arguments: A case study with scientific abstracts / P. Accuosto, H. Saggion // Proceedings of the 6th Workshop on Argument Mining. – Florence, Italy, 2019. – Pp. 41–51.

57. Ajjour Y. Unit segmentation of argumentative texts / Y. Ajjour, W.-F. Chen., J. Kiesel, H. Wachsmuth, B. Stein // Proceedings of the 4th Workshop on Argument Mining. – Copenhagen, 2017. – Pp. 118–128.

58. Al-Khatib K. Patterns of Argumentation Strategies across Topics / K. Al-Khatib, H. Wachsmuth, M. Hagen, B. Stein // Proc. of the 2017 Conference on Empirical

Methods in Natural Language Processing. – Copenhagen, Denmark. – September 7–11, 2017. – Pp. 1351–1357.

59. Amossy R. Argumentation in discourse. A socio-discursive approach to arguments / R. Amossy // *Informal Logic*. – 29 (3). – 2009. – Pp. 252–267.

60. Anand P. Believe Me – We Can Do This! Annotating Persuasive Acts in Blog Text / P. Anand, J. Eichenbaum, J. Boyd-Graber, E. Wagner // *Computational Models of Natural Argument*. – San Francisco, California, USA, 2011. – Pp. 1–5.

61. Anscombe J.-C. Argumentativity and informativity / J.-C. Anscombe, O. Ducrot // M. Meyer (Ed.). *From metaphysics to rhetoric*. – Dordrecht: Kluwer, 1989. – Pp. 71–87.

62. Argamon S. Routing Documents According to Style / S. Argamon, M. Koppel, G. Avneri // *First International Workshop on Innovative Information Systems* – Pisa, Italy, 1998. – Pp. 85–92.

63. Arthur D. K-Means++: the Advantages of Careful Seeding / D. Arthur, S. Vassilvitskii // *Proceedings of the Eighteenth Annual ACM-SIAM Symposium on Discrete Algorithms, SODA*. – 2007. – New Orleans, Louisiana. – Pp. 1–9.

64. Artstein R. Inter-coder agreement for computational linguistics / R. Artstein, M. Poesio // *Computational Linguistics*. – 34 (4). – 2008. – Pp. 555–596.

65. Buhmann M. D. Radial Basis Functions: Theory and Implementations / M. D. Buhmann // *Cambridge University Press*. – 2003. – 272 p.

66. Byjlova N. Amateur Prose on the Web: Verb Construction as a Feature of Genre Classification / N. Byjlova // *ARANEIA 2018: Web Corpora as a Language Training Tool: Proceedings*. – A Butašová, V. Benko, Z. Puchovská (eds.). Bratislava: Univerzita Komenského v Bratislave, 2018. – Pp. 25–30.

67. Carel M. Pourtant: Argumentation by exception / M. Carel // *Journal of Pragmatics*. – 24. – 1995. – Pp. 167–188.

68. Castro S. Fast Krippendorff: Fast computation of Krippendorff's alpha agreement measure / C. Castro – 2017. – URL: <https://github.com/pln-fing-udelar/fast-krippendorff> [Дата обращения: 25.09.03].

69. Chen G. Exploring the Potential of Large Language Models in Computational Argumentation / G. Chen, L. Cheng, L. A. Tuan, L. Bing // Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics. – Volume 1. – 2024. – Pp. 2309–2330.

70. Chistova E. Towards the Data-Driven System for Rhetorical Parsing of Russian Texts / E. Chistova, M. Kobozeva, D. Pisarevskaya, A. Shelmanov, I. Smirnov, S. Toldova // Proceedings of Discourse Relation Parsing and Treebanking. DISRPT2019. – Minneapolis, MN. – 2019. – Pp. 82–87.

71. Cohen J. A coefficient of agreement for nominal scales / J. Cohen // Educational and Psychological Measurement. – 20 (1). – 1960. – Pp. 37–46.

72. Cordella L. P. (Sub)Graph Isomorphism Algorithm for Matching Large Graphs / L. P. Cordella, P. Foggia, C. Sansone, M. Vento // IEEE Transactions on Pattern Analysis and Machine Intelligence. – Vol. 26, no. 10. – Oct., 2004. – Pp. 1367–1372.

73. Daxenberger J. What is the Essence of a Claim? Cross-Domain Claim Identification / J. Daxenberger, S. Eger, I. Habernal, C. Stab, I. Gurevych // Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing. – 2017. – Pp. 2055–2066.

74. Delgado R. Why Cohen's Kappa should be avoided as performance measure in classification / R. Delgado, X.-A. Tibau // PLOS One. – Vol. 14 (9). – 2019. – Pp. 1–26.

75. Demberg V. How compatible are our discourse annotation frameworks? Insights from mapping RST-DT and PDTB annotations / V. Demberg, M. Schoman, F. T. Asr // Dialogue & Discourse. – 10 (1). – 2019. – Pp. 87–135.

76. Duthie R. Mining ethos in political debate / R. Duthie, K. Budzynska, C. Reed // 6th International Conference on Computational Models of Argument (COMMA 16). – 2016. – Pp. 299–310.

77. El Baff R. Computational Argumentation Synthesis as a Language Modeling Task / R. El Baff, H. Wachsmuth, K. Al Khatib, M. Stede, B. Stein // Proceedings of the 12th International Conference on Natural Language Generation. – Tokyo, Japan, 2019. – Pp. 54–64.

78. Fishcheva I. Cross-Lingual Argumentation Mining for Russian Texts / I. Fishscheva, E. Kotelnikov // Proceedings Of the 8th International Conference “Analysis of Images, Social Networks and Texts”. Kazan, 2019. – Pp. 134–144.
79. Fleiss J. L. Measuring nominal scale agreement among many raters / J. L. Fleiss // Psychological Bulletin. – Vol. 76, no. 5. – 1971. – Pp. 378–382.
80. Green N. Identifying Argumentation Schemes in Genetics Research Articles / N. Green // Proceedings of the 2nd Workshop on Argumentation Mining. – Denver, 2015. – Pp. 12–21.
81. Hagberg A. A. Exploring network structure, dynamics, and function using NetworkX / A. A. Hagberg, D. A. Schult, P. J. Swart // Proc. of the 7th Python in Science Conference (SciPy2008). – Pasadena, CA USA, 2008. – Pp. 11–15.
82. Hamblin C. L. Fallacies / C. L. Hamblin // London: Vale Press, 1970. – Pp. 224–252.
83. Hidey C. Persuasive Influence Detection: The Role of Argument Sequencing / C. Hidey, K. McKeown // Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence. – New Orleans, LA, 2018. – Pp. 5173–5180.
84. Hintikka J. The Role of Logic in Argumentation / J. Hintikka // The Monist. – Vol. 72, No. 1. – 1989. – Pp. 3–24.
85. Hua X. Understanding and Detecting Supporting Arguments of Diverse Types / X. Hua, L. Wang // Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Short Papers). – 2017. – Pp. 203–208.
86. Ilynska L. Rhetorical strategies in the context of professional communication / L. Ilynska, M. Platonova, T. Smirnova // Linguistics & Polyglot Studies. – No. 5. – 2016. – Pp. 32–40.
87. Jiang C. A Survey of Frequent Subgraph Mining Algorithms / C. Jiang, F. Coenen, M. Zito // The Knowledge Engineering Review. – 000(1). – 2004. – Pp. 1–31.
88. Kinneavy J. L. Kairos in Aristotle’s Rhetoric / J. L. Kinneavy, C. R. Eskin // Written Communication. – 17 (3). – 2000. – Pp. 432–444.

89. Kohavi R. A Study of Cross-Validation and Bootstrap for Accuracy Estimation and Model Selection / R. Kohavi // Proceedings of the International Joint Conference on Artificial Intelligence (IJCAI) – 1995. – Pp. 1137–1145.

90. Kononenko I. Comparative Analysis of Rhetorical and Argumentative Structures in the Study of Popular Science Discourse / I. Kononenko, E. Sidorova, I. Akhmadeeva // International Conference on Computational Linguistics and Intellectual Technologies “Dialogue”. – 2020. – Pp. 432–444.

91. Korobov M. Morphological Analyzer and Generator for Russian and Ukrainian Languages / M. Korobov // Analysis of Images, Social Networks and Texts. – 2015. – Pp. 320–332.

92. Kotelnikov E. RuArg-2022: Argument Mining Evaluation / E. Kotelnikov, N. Loukashevitch, I. Nikishina, A. Panchenko // Computational Linguistics and Intellectual Technologies: Proceedings of the International Conference “Dialogue 2022”. – 2022. – Pp. 1–16.

93. Krippendorff K. Content analysis: An introduction to its methodology (2nd edition) / K. Krippendorff // Sage Publications, Beverly Hills, California, 2005. – Pp. 211–257.

94. Kuzman T. Exploring the Impact of Lexical and Grammatical Features on Automatic Genre Identification / T. Kuzman, N. Ljubešić // SIKDD 2022. – D. Mladenović & M. Grobelnik (Eds.), Odkrivanje znanja in podatkovna skladišča. – 2022. – Pp. 1–4.

95. Landis J. R. The measurement of observer agreement for categorical data / J. R. Landis, G. G. Koch // Biometrics. – Vol. 33, no. 1. – 1977. – Pp. 159–174.

96. Lawrence J. AIFdb corpora / J. Lawrence, C. Reed // Proceedings of the Fifth International Conference on Computational Models of Argument (COMMA 2014). – 2014. – Pp. 465–466.

97. Lawrence J. Argument Mining: A Survey / J. Lawrence, C. Reed // Computational Linguistics. – Vol. 45 (4). – 2019. – Pp. 765–818.

98. Lawrence J. Mining Arguments From 19th Century Philosophical Texts Using Topic Based Modelling / J. Lawrence, C. Reed, C. Allen, S. McAlister, A. Ravenscroft,

D. Bourget // *Proceedings of the First Workshop on Argumentation Mining*. – Baltimore, Maryland, 2014. – Pp. 79–87.

99. Levy R. Context Dependent Claim Detection / R. Levy, Y. Bilu, D. Hershcovich, E. Aharoni // *Proceedings of the 25th International Conference on Computational Linguistics: Technical Papers*. – Dublin, Ireland, 2014. – Pp. 1489–1500.

100. Lidstone G. J. Note on the general case of the Bayes-Laplace formula for inductive or a posteriori probabilities / G. J. Lidstone // *Transactions of the Faculty of Actuaries* – vol. 8. – 1920. – Pp. 182–192.

101. Lindahl A. Towards Assessing Argumentation Annotation – A First Step / A. Lindahl, L. Borin, J. Rouces // *Proceedings of the 6th Workshop on Argument Mining*. – Florency, Italy, 2019. – Pp. 177–186.

102. Lukin S. Argument Strength is in the Eye of the Beholder: Audience Effects in Persuasion / S. Lukin, P. Anand, M. Walker, S. Whittaker // *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics*. – Valencia, 2017. – Pp. 742–753.

103. Luxburg U. von. A Tutorial on Spectral Clustering / von. U. Luxburg // *Statistics and Computing*. – Vol. 17. – 2007. – Pp. 395–416.

104. Madnani N. Identifying high-level organizational elements in argumentative discourse / N. Madnani, M. Heilman, J. Tetreault, M. Chodorow // *Proceedings of the 2012 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. – Montreal, 2012. – Pp. 20–28.

105. Mann W. Rhetorical Structure Theory: Toward a functional theory of text organization / W. Mann, S. Thompson // *Text – Interdisciplinary Journal for the Study of Discourse*. – 1988. – Vol. 8 (3). – Pp. 243–281.

106. Manning C. *Foundation of Statistical Natural Language Processing* / C. Manning, H. Schutze // The MIT Press, Cambridge, Massachusetts. – 2000. – P. 680.

107. Matthews B. W. Comparison of the predicted and observed secondary structure of T4 phage lysozyme / B. W. Matthews // *Biochimica et Biophysica Acta*. – Vol. 405, no. 2. – 1975. – Pp. 442–451.

108. Mochales Palau R. Argumentation mining: The detection, classification and structure of arguments in text / R. Mochales Palau, M.-F. Moens // Proceedings of the 12th International Conference on Artificial Intelligence and Law. – Barcelona, ACM NY, 2009. – Pp. 98–109.

109. Moens M.-F. Automatic Detection of Arguments in Legal Texts / M.-F. Moens, E. Boiy, R. Mochales Palau, C. Reed // Proceedings of the 11th International Conference on Artificial Intelligence and Law. (NY, USA). – ACM Press, 2007. – Pp. 225–230.

110. Park J. Identifying Appropriate Support for Propositions in Online User Comments / J. Park, C. Cardie // Proceedings of the First Workshop on Argumentation Mining (Baltimore, Maryland USA). – 2014. – Pp. 29–38.

111. Passon M. Predicting the Usefulness of Amazon Reviews Using Off-The-Shelf Argumentation Mining / M. Passon, M. Lippi, G. Serra, C. Tasso // Proc. of the 5th Workshop on Argument Mining (Brussels, Belgium). – 2018. – Pp. 35–39.

112. Peldzus A. From Argument Diagrams to Argumentation Mining in Texts: A Survey / A. Peldzus, M. Stede // International Journal of Cognitive Informatics and Natural Intelligence. – Vol. 7 (1). – 2003. – Pp. 1-13.

113. Perelman C. The new rhetoric. A treatise on argumentation / C. Perelman, L. Olbrechts-Tyteca // Notre Dame: University of Notre Dame Press, 1969. – 566 p.

114. Plantin C. Argumentation Studies in France: A New Legitimacy / C. Plantin // Anyone who has a view. Theoretical contributions to the study of argumentation. – Dordrecht: Kluwer, 2003. – Pp. 173–187.

115. Pimenov I. S. The Quantitative Evaluation of the Pathos to Ethos Ratio in Scientific Texts / I. S. Pimenov, N. V. Salomatina, M. K. Timofeeva // 2022 IEEE 23rd International Conference of Young Professionals in Electron Devices and Materials (EDM). – Russia, Altai, 2022. – Pp. 312–317.

116. Pimenov I. S. Identification of Scientific Texts with Similar Argumentation Complexity / I. S. Pimenov, N. V. Salomatina // 2022 IEEE International Multi-Conference on Engineering, Computer and Information Sciences (SIBIRCON). – 2022. – Pp. 870–875.

117. Pimenov I. S. Evaluating the Influence of Argumentation Markers on the Identification of Reasoning Models / I. S. Pimenov, N. V. Salomatina // Communications in Computer and Information Science. – 2024 (a). – Vol. 2086. – Pp. 282–297.

118. Pimenov I. S. An Automatic Method for Standartizing Argumentative Annotations across Annotators Genres / I. S. Pimenov, N. V. Salomatina // 2024 IEEE 25th International Conference of Young Professionals in Electron Devices and Materials (EDM). – 2024 (b). – Pp. 2260–2265.

119. Pimenov I. S. Productivity-Based Analysis of Argumentation Patterns across Texts of Different Genres / I. S. Pimenov, N. V. Salomatina // 2024 IEEE 25th International Conference of Young Professionals in Electron Devices and Materials (EDM). – 2024 (c). – Pp. 2270–2275.

120. Pimenov I., Salomatina N. Employing Argumentation Patterns for Genre Classification of Scientific Communication Texts / I. Pimenov, N. Salomatina // 2025 IEEE 26rd International Conference of Young Professionals in Electron Devices and Materials (EDM). – 2025. – pp. 1780–1785.

121. Potter A. Inferring Inferences: Relational Propositions for Argument Mining / A. Potter // Proceedings of the Society for Computation in Linguistics. – 2022. – Pp. 89–100.

122. Puig L. Doxa and persuasion in lexis / L. Puig // Argumentation. – 26 (1). – 2012. – Pp. 127–142.

123. Pungeršek T. K. Automatic genre identification: a survey / T. K. Pungeršek, N. Ljubešić // Language Resources and Evaluation. – 59 (1). – 2023. – Pp. 537–570.

124. Rahwan I. The argument interchange format / I. Rahwan, C. Reed // Argumentation in artificial intelligence. Rahwan I. and Simari G. (Eds.). – Springer, 2009. – Pp. 383–402.

125. Reed C. Araucaria: Software for argument analysis / C. Reed, G. Rowe // International Journal of AI Tools. – 14 (3–4). – 2004. – Pp. 961–980.

126. Rigotti E. From argument analysis to cultural keywords (and back again) / E. Rigotti, A. Rocci // F. H. van Eemeren & P. Houtlosser (Eds.), *Argumentation in practice*. – Amsterdam-Philadelphia: John Benjamins, 2005. – Pp. 125–142.

127. Ryabko B. Y. Information-Theoretic method for classification of texts / B. Y. Ryabko, A. E. Gus'kov, I. V. Selivanova // *Problems of Information Transmission*. – Vol. 53. – 2017. – Pp. 294–304.

128. Salomatina N. V. Identification of connected arguments based on reasoning schemes “from expert opinion” / N. V. Salomatina, I. S. Kononenko, E. A. Sidorova, I. S. Pimenov // *Journal of Physics: Conference Series*. Vol. 1715, International Conference «Marchuk Scientific Readings 2020» (MSR-2020). – Novosibirsk, 2020. – Pp. 1–10.

129. Salomatina N. V. Identification of argumentative sentences in Russian scientific and popular science texts / N. V. Salomatina, E. A. Sidorova, I. S. Pimenov // *Journal of Physics: Conference Series*. – Vol. 2099. – International Conference «Marchuk Scientific Readings 2021» (MSR-2021). Russia, Novosibirsk, 2021. – Pp. 1–8.

130. Sardianos C. Argument Extraction from News / C. Sardianos, M. Katakis, G. Petasis, V. Karkaletis // *Proceedings of the 2nd Workshop on Argumentation Mining*. – Denver, Colorado, 2015. – Pp. 56–66.

131. Schaeffer R. Are emergent abilities of large language models a mirage? / R. Schaeffer, B. Miranda, S. Koyejo // *Advances in Neural Information Processing Systems*. – 36. – 2023. – Pp. 55565–55581.

132. Sheng L. Are Emergent Abilities in Large Language Models Just In-Context Learning? / L. Sheng, I. Bigoulaeva, R. Sachdeva, H. T. Madabushi, I. Gurevych // 2023. – Pp. 1–42. – DOI: 10.48550/arXiv.2309.01809.

133. Shum S. B. Modelling naturalistic argumentation in research literatures: Representation and interaction design issues / S. B. Shum, V. Uren, G. Li, B. Sereno, C. Mancini // *International Journal of Intelligent Systems, Special Issue on Computational Modelling of Naturalistic Argumentation*. – 22 (1). – 2007. – Pp. 17–47.

134. Stab C. Identifying argumentative discourse structures in persuasive essays / C. Stab, I. Gurrevych // *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (Doha)*. – 2014. – Pp. 46–56.
135. Stubbs M. *Text and Corpus Analysis* / M. Stubbs // *Computer-assisted Studies of Language and Culture*. – London: Blackwell, 1996. – 288 p.
136. Sun Y. PITA: Prompting Task Interaction for Argumentation Mining / Y. San, M. Wang, J. Bao, B. Liang, Z. Zhao, C. Yang, M. Yang, R. Xu // *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics*. – Vol. 1. – 2024. – Pp. 5036–5049.
137. Taboada M. *Applications of Rhetorical Structure Theory* / M. Taboada, W. Mann // *Discourse Studies*. SAGE Publications. – 2006 (a). – Vol. 8, № 4. – P. 567–588.
138. Taboada M. *Rhetorical Structure Theory: Looking back and moving ahead* / M. Taboada, W. Mann // *Discourse Studies*. SAGE Publications. – 2006 (b). – Vol. 8, no. 3. – Pp. 423–459.
139. Teruel M. Increasing Argument Annotation Reproducibility by Using Inter-annotator Agreement to Improve Guidelines / M. Teruel, C. Cardellino, F. Cardellino, L. A. Alemany, S. Villata // *Proceedings of the Eleventh International Conference on Language Resources and Evaluation*. – 2018. – Pp. 4061–4064.
140. Timofeeva M. Comparative Analysis of Reasoning in Russian Classic Poetry / M. Timofeeva // *Applied Sciences*. – 2021. – 11, 8665. – Pp. 1–12.
141. Toulmin S. *The Uses of Argument* / S. Toulmin // New York: Cambridge University Press, 2003. – 259 p.
142. Ullmann J. R. An Algorithm for Subgraph Isomorphism / J. R. Ullmann // *Journal of the ACM*. – 23 (1). – 1976. – Pp. 31–42.
143. van Eemeren F. H. *A systematic theory of argumentation. The pragma-dialectical approach* / F. H. van Eemeren, R. Grootendorst // Cambridge: Cambridge University Press, 2004. – 216 p.

144. van Eemeren F. H. Handbook of Argumentation Theory / F. H. van Eemeren, B. Garssen, E. Krabbe, F. Henkemans, B. Verheij, J. Wagemans. – Springer, 2014. – 988 p.
145. Wachsmuth H. Argumentation Synthesis following Rhetorical Strategies / H. Wachsmuth, M. Stede, R. El Baff, K. All-Khatib, M. Skeppstedt, B. Stein // Proceedings of the 27th International Conference on Computational Linguistics. – Santa Fe, New Mexico, USA, 2018. – Pp. 3753–3765.
146. Walker W. Annotating Patterns of Reasoning about Medical Theories of Causation in Vaccine Cases: Toward a Type System for Arguments / W. Walker, K. Vazirova, C. Stanford // Proceedings of the First Workshop on Argumentation Mining. – Baltimore, MD, 2014. – Pp. 1–10.
147. Walton D. Argumentation schemes / D. Walton, C. Reed, F. Macagno // New York: Cambridge University Press, 2008. – 443 p.
148. Ward J. H., Jr. Hierarchical Grouping to Optimize an Objective Function / J. H. Ward, Jr. // Journal of the American Statistical Association. – 58. – 1963. – Pp. 236–244.
149. Wei J. Emergent Abilities of Large Language Models / J. Wei, Y. Tay, R. Bommasani, C. Raffel, B. Zoph, S. Borgeaud, D. Yogotama, M. Bosma, D. Zhou, D. Metzler, E. H. Chi, T. Hashimoto, O. Vinyals, P. Liang, J. Dean, W. Fedus // Transactions on Machine Learning Research. – 2022. – Pp. 1–30.
150. Wierzbicka A. Understanding Cultures through Their Key Words / A. Wierzbicka // Oxford: Oxford University Press, 1997. – 288 p.
151. Zagorulko Yu. A. Analysis of the persuasiveness of argumentation in popular science texts / Yu. A. Zagorulko, O. A. Domanov, A. S. Sery, E. A. Sidorova, O. I. Borovikova // Artificial Intelligence, Proceedings of the 18th Russian Conference RCAI. – 2020. – Pp. 351–367.
152. Zasytkin A. S. The Combined Approach to Identifying Argumentation Structures in Short Scientific Papers / A. S. Zasytkin, I. S. Pimenov, N. V. Salomatina // 2023 IEEE 24rd International Conference of Young Professionals in Electron Devices and Materials (EDM). – Russia, Altai, 2023. – Pp. 1800–1805.

153. Ziegenbein T. LLM-based Rewriting of Inappropriate Argumentation using Reinforcement Learning from Machine Feedback / T. Ziegenbein, G. Skitalinskaya, A. B. Makou, H. Wachsmuth // Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics. – Vol. 1. – 2024. – Pp. 4455–4476.

Приложение А

Маркеры главного тезиса

Таблица А.1 – Маркеры главного тезиса, выявленные в статьях корпуса

№	Маркер	Позиция в тексте
1	Таким образом, [можно сделать вывод / ...] ...	последняя фраза в тексте; первая фраза в последнем, предпоследнем абзаце
2	В итоге, ...	последний абзац, первое предложение
3	Подводя итог [вышесказанному], можно сделать вывод / можно утверждать / ...	последний абзац
4	Как итог, ...	последняя фраза в тексте
5	Делая вывод, можно... [заявить / сказать / утверждать...], что...	последний абзац
6	Вывод(ы) [В [настоящей, ...] статье были рассмотрены...]...	первая фраза в разделе
7	Все это говорит... о том, что...	последняя фраза в тексте
8	В результате [проведенных [исследований / проделанной работы, ...] [было] выявлено / установлено, что...	последняя фраза, предпоследняя фраза текста
9	Из / На основании [всего] вышесказанного следует\ можно заключить, что...	последняя фраза текста; предпоследний абзац, первая фраза
10	Итак, ...	первая фраза абзаца перед иллюстративным материалом (при разборе конкретной задачи); последняя фраза последнего абзаца
11	В статье рассмотрены / исследованы...	первое предложение в последнем абзаце
12	Как видим, ...	первое предложение в последнем абзаце
13	Исследование показало, что...	последняя фраза текста
14	Поэтому...	последняя фраза текста
15	В заключение можно сказать, что...	последний абзац, первое предложение

Окончание таблицы А.1

16	Заключение В [данной / настоящей / ...] статье / работе предложен / представлен...	первое предложение раздела
17	В [данной, ...] работе / статье ...	в начальной части текста
18	Задачей данной работы был(а)...	в начальной части текста
19	Из [проведенного / данного / настоящего / ...] исследования следует, что...	последняя фраза текста
20	Изучение [...] показало, что...	предпоследнее предложение последнего абзаца
21	Как видно из [проведенного] анализа / ...	последняя фраза текста
22	В заключение можно / следует / хотелось бы ... отметить	первое предложение в последнем абзаце
23	Из [проведенного] исследования следует, что...	последняя фраза текста
24	...на основе [имеющихся / проведенных, ...] исследований, ..., можно сделать вывод, что...	последняя фраза текста
25	Несомненно, [что]...	последняя фраза текста
26	Изучение [...] показало, что...	предпоследняя фраза текста

Приложение Б

Маркеры аспектов содержания

Таблица Б.1 – Маркеры аспектов содержания, выявленные в статьях корпуса

Аспект	Типовая(ые) схема(ы)	Пример маркеров
1. Цель работы	<i>Practical Reasoning</i>	«(настоящее) исследование посвящено», «данная работа посвящена», «целью (работы / данной статьи) является», «нашей целью является», «цель (представленного/данного) исследования →», «в рамках статьи будет рассмотрен», «в данной работе нас интересует»
2. Актуальность	<i>Practical Reasoning, Popular Practice</i>	«не раз становился объектом исследований», «проблема начала разрабатываться ещё в», «актуально, так как», «актуально для многих практических применений»
3. Методы	<i>Applied Method</i>	Х избран в качестве метода», «при помощи набора методов», «использование такого метода, как», «одним из методов является», «для решения задачи выбран метод»
4. Данные	<i>Applied Method</i>	«источником послужил корпус», «в качестве материала использован», «материалом послужил», «материалом для нашего исследования выступают», «исследование проводилось на материале корпуса»

Окончание таблицы Б.1

5. Представление выводов	(Как правило, соответствует главному тезису текста, либо вводит его содержащий абзац)	«таким образом», «можно утверждать, что», «исследование показало, что», «проведённый эксперимент показал, что», «позволил сделать следующие выводы», «результаты позволяют выделить», «подводя итоги, отметим, что», «в результате проведённого исследования можно сделать вывод, что», «итак», «в заключение хотелось бы отметить, что», «в работе показано, что», «в заключение —», «в результате было выявлено, что»
6. Перспективы дальнейшей работы	<i>Practical Reasoning, Positive Consequences</i>	«в перспективе (настоящего) исследования», «результаты (текущего) исследования могут быть использованы», «в дальнейшем планируется», «результаты могут быть полезны в», «в перспективе можно»

Приложение В

Актуальные схемы аргументации в научных статьях

Таблица В.1 – Актуальные схемы аргументации в научных статьях корпуса

Схема, формализованное описание	Неформальное описание	Пример реализации
1. <i>Analogy</i> Premise: <u>A</u> истинно (ложно) в случае <u>C1</u>. Premise: В общем, случай <u>C1</u> похож на случай <u>C2</u>. Conclusion: <u>A</u> истинно (ложно) в случае <u>C2</u>.	Вводит тезис, корректность которого обосновывается через его сходство с некоторым другим заявленным тезисом.	Pr: Подобный принцип моделирования был использован при <u>построении лексико семантической карты славянских языков (по коэффициенту близости больших параметрических ядер лексики)</u> в работе А.А. Кретьева и И.А. Меркуловой [16]. Con: При помощи корреляционных коэффициентов была <u>построена карта имени coronavirus по данным корпуса NOW (рис. 3), где длина линии, соединяющей все идиомы, указывает на тесноту их связи: чем короче линия, тем выше коэффициент корреляции между ними и теснее связь</u> [12], [13].
2. <i>Applied Method</i> Premise: метод <u>E</u> применяется в области <u>D</u>, содержащей пропозицию <u>A</u>. Premise: <u>A</u> получено методом <u>E</u>. Conclusion: <u>A</u> истинно (ложно).	Обосновывает корректность описываемых результатов через указание особенностей метода, с помощью которого эти результаты получены. Под методом в данном случае могут	Pr: <u>Пропозициональный анализ</u> избран нами в качестве продуктивного метода выявления <u>семантики 50 текстов газетных спортивных новостей</u>, отобранных из британского национального корпуса. Pr: <u>Качественный и количественный анализ</u> исследуемых текстов установил <u>ситуации, которые могут являться экстралингвистической основой ядерного элемента суперструктуры текстов</u>

Продолжение таблицы В.1

	пониматься как конкретные алгоритмы, так и наборы техник из анализируемой предметной области, так и данные, на которых получен заявленный результат.	<u>спортивных новостей – основного спортивного события.</u> Соп: Наиболее частотным типом ситуации, актуализируемой основным событием в исследуемых текстах, является <u>итог соревновательной деятельности, отражающий победу, поражение, ничью, получение медалей и т.п.</u>
3. <i>Cause to Effect</i> Premise: <u>А</u> является причиной <u>В</u>. Premise: В этом случае <u>А</u> имеет место. Conclusion: В этом случае <u>В</u> будет иметь место.	Самая общая причинно-следственная связь, применяемая в следующих целях: – для указания на последовательность процедурных этапов; для соединения выполняемых процедур и технического описания результатов (без оценочного толкования); – для указания на последовательность событий во времени; – для организации абстрактного рассуждения; – для развития оценочных блоков или блоков практического рассуждения.	Pr: <u>Стремительное развитие информационных технологий в корне изменило отношение к традиционному процессу перевода.</u> Pr: Сама же <u>история исследований и разработок в данной сфере насчитывает уже более пятидесяти лет.</u> Соп: С середины семидесятых годов двадцатого века наблюдается <u>устойчивый рост интереса к машинному переводу.</u>

Продолжение таблицы В.1

4. <i>Effect To Cause</i> Premise: <u>A</u> является причиной <u>B</u>. Premise: В этом случае <u>B</u> будет иметь место. Conclusion: В этом случае <u>A</u> имеет место.	– выражает причинно-следственную связь в её обратном рассмотрении (от следствий к причинам); – посылка обозначает явление, причина которого выражается в выводе; – схема обозначает переход от более поверхностной идеи к глубинной.	Pr: Для <u>грамматического оформления системы</u> понадобились <u>специальные маркеры (классные показатели)</u>. Pr: В даргинском языке <u>1 класс обслуживает маркер в, 2 класс обслуживает маркер p, 3 класс обслуживают маркеры б, д</u>. Соп: В современном даргинском языке <u>представлены 3 класса</u>: 1 класс – класс мужчин (лица, обозначающие мужской пол), 2 класс – класс женщин (лица, обозначающие женский пол), 3 класс – всех остальных.
5. <i>Example</i> Premise: В этом случае индивид <u>A</u> имеет свойство <u>F</u>, а также свойство <u>G</u>. Premise: а типично для вещей, имеющих <u>F</u> и имеющих или не имеющих <u>G</u>. Conclusion: Как правило, если x имеет свойство <u>F</u>, то (обычно, вероятно, типично) x также имеет свойство <u>G</u>.	Схема объединяет конкретный пример и иллюстрируемый им общий тезис. Посылка может включать следующие элементы: – указание на частный случай некоего общего явления; – конкретное описание алгоритма; – представление используемой формулы; – указание на иллюстративный материал без явного комментария.	Pr: Так, например, всем известны такие «магические» <u>числа</u>, как «три», «семь», «девять», «двенадцать», «тринадцать». Pr: Эти <u>числа</u> отражают <u>особенности мышления, мироощущения, религиозные суеверия</u> людей с давних времен. Соп: Издавна <u>числа</u> наделялись людьми некоей <u>магией</u>.

Продолжение таблицы В.1

6. <i>Expert Opinion</i> Premise: Источник <u>Е</u> это эксперт в области <u>Д</u>, к которой относится пропозиция <u>А</u>. Premise: <u>Е</u> утверждает, что <u>А</u> истинно (ложно). Conclusion: <u>А</u> истинно (ложно).	– вводит тезис за авторством явно указываемого авторитета; – вводит тезис с явным указанием на авторитетность его источника; – вводит тезис в формулировке из цитируемого источника.	Pr: <u>Наблюдение относительно воздействующих особенностей рекламы подтверждается мнением <u>авторитетных специалистов</u>.</u> Pr: <u>«Рекламное воздействие по определению является манипулятивным, так как <u>изменение структуры сознания покупателя происходит [...] в результате принятия на веру эмоционально, экспрессивно вводимой не прямой и не исчерпывающей информации</u>» [1, с.67].</u> Соп: Актуализация тех или иных мотивов в рекламе относится к приемам подсознательного стимулирования, когда <u>отношение аудитории к рекламируемому объекту формируется с помощью различных представлений (стереотипов, мифов, имиджей), автоматически вызывающих в массовом сознании положительную реакцию.</u>
7. <i>Logical</i> Схема поддерживает две роли компонентов: Conflicted statement – опровергаемый тезис. Conflicting statement – опровергающий тезис.	– общая схема для случаев опровержения одного тезиса другим. – вводит тезис через указание противоречия между этим тезисом и иным, введённым отдельно.	Conflicted statement: Нельзя отрицать, что нынешние относительно устойчивые часто употребляемые словосочетания можно назвать современными фразеологизмами. Conflicting statement: Одним из трёх важных параметров, характеризующих рассматриваемые выражения, является устойчивость (а это прямая ассоциативная параллель с историей).

Продолжение таблицы В.1

8. <i>Negative Consequences</i> Premise: если осуществится <u>A</u>, то, вероятно, возникнут плохие последствия. Conclusion: <u>A</u> не следует осуществлять.	Обозначает причинно следственную связь, где – вывод либо посылка явно содержит отрицательный оценочный компонент; – вывод описывает отрицательный результат действия / ситуации в посылке.	Pr: Но при этом, <u>все вышеперечисленные методики</u> имеют один общий недостаток: качественную оценку (экспертную оценку) определения угроз безопасности информации. Con: На основании вышеизложенного, можно сделать вывод о том, что <u>применение только качественных методик оценки угроз безопасности информации недостаточно.</u>
9. <i>Part to Whole</i> Premise: <u>B</u> подвид (часть) <u>A</u>. Premise: <u>B</u> имеет свойство <u>G</u>. Conclusion: <u>A</u> имеет свойство <u>G</u>.	– Обозначает меронимическую связь (от части к целому). – Реализуется при анализе составных сущностей через анализ их составных частей.	Pr: <u>Чалдини</u> дает читателям <u>уникальный взгляд на психологические процессы, лежащие в основе наших решений и способность убеждать.</u> Pr: В <u>работе</u> выделены <u>шесть фундаментальных принципов, которые оформляют наши повседневные решения.</u> Con: <u>Глубокое погружение в мир психологии убеждения</u> представлено в <u>работе Роберта Чалдини «Influence: The Psychology of Persuasion».</u>
10. <i>Popular Opinion</i> Premise: общепринято, что <u>A</u> истинно. Conclusion: <u>A</u> истинно.	– Вводит тезис с указанием (возможно неявным) на его распространённость. – Вводит тезис за авторством неопределённого источника.	Pr: Совершенно очевидно, что в настоящее время <u>когнитивное направление в лингвистике заявило себя как новая научная парадигма, имеющая свой объект и инструментарий анализа.</u> Con: Это направление предполагает <u>исследование языковых явлений в совершенно новом аспекте.</u>

Продолжение таблицы В.1

11. <i>PopularPractice</i> Premise: Если <u>A</u> – общепринятая практика среди знакомых с тем, что допустимо или нет в отношении <u>A</u>, то это даёт повод думать, что <u>A</u> допустимо. Premise: <u>A</u> это общепринятая практика среди тех, кто знаком с тем, что приемлемо в отношении <u>A</u>. Conclusion: <u>A</u> – допустимое действие.	– вводит тезис с предписанием некоего действия (выбором конкретного алгоритма, исследования определённой проблемы) ввиду его частого совершения иными агентами. – посылка первого типа практически не выражается явно в текстах выбранного жанра.	Pr: <u>Модель «программное обеспечение как услуга» (SaaS), подразумевающая использование программного обеспечения, установленного на сервере, в настоящее время наиболее популярна среди пользователей.</u> Con: Для организации образовательного процесса по информатике и другим дисциплинам хорошо подходят <u>публичные, общественные или гибридные сервисы облачных технологий.</u>
12. <i>Position to Know</i> Premise: <u>E</u> утверждает, что <u>A</u> истинно (ложно). Premise: <u>E</u> известно об истинности <u>A</u>. Conclusion: <u>A</u> истинно (ложно).	– вводит тезис за авторством явно указанного источника без указания на его экспертность; – вводит собственный тезис автора статьи, который подаётся как предположение.	Pr: <u>На наш взгляд, для получения прозрачной картины необходимо применять тексты разных переводчиков.</u> Con: Важным условием создания корпусов для переводоведческих задач является <u>использование переводов разных переводчиков.</u>
13. <i>Positive Consequences</i> Premise: Если <u>A</u> осуществится, то это вероятно, приведет к хорошим последствиям.	Обозначает причинно-следственную связь, где – вывод либо посылка явно содержит положительный оценочный компонент;	Pr: <u>Подобный анализ</u> позволит учитывать изменения в ответах пользователя. Con: Таким образом, необходимо проводить <u>не только нейросетевой анализ тональности каждого сообщения, но и анализ групп сообщений.</u>

Продолжение таблицы В.1

Conclusion: <u>A</u> следует осуществить.	– вывод описывает положительный результат действия / ситуации в посылке.	
14. <i>Practical Reasoning</i> Premise: Осуществление <u>B</u> есть способ осуществить <u>A</u> . Premise: Цель состоит в осуществлении <u>A</u> . Conclusion: <u>B</u> следует осуществить.	Обозначает причинно-следственную связь, где – вывод обозначает некое действие с прагматическим компонентом его предпочтительности; – посылка указывает цель или способ достижения цели без включения модального компонента практичности.	Pr: Читателю требуются <u>определенные навыки прочтения публицистического текста, включая умение согласовать содержание с названием для получения подтекста</u> [6]. Pr: Для <u>формирования у школьников и студентов такого навыка</u> следует <u>типологизировать модели экспликации в заголовке смыслов для рассмотрения на уроках анализа текста</u> [7]. Con: Предпринятое исследование преследовало цель <u>выявить продуктивные средства иллюстрации в заголовках</u> .
15. <i>Sign</i> Premise: <u>A</u> истинно в данной ситуации. Premise: как правило, истинность знака <u>A</u> , обозначающего <u>B</u> , указывает на истинность <u>B</u> . Conclusion: <u>B</u> истинно в этой ситуации.	– доказываемый тезис и свидетельство его истинности оперируют сущностями разного уровня, дополняют друг друга как означающее и означаемое. – обосновывает корректность тезиса через приведение фактов в его поддержку.	Pr: Одно из самых весомых событий мирового масштаба произошло в 1997 году, когда <u>компьютер "IBM Deep Blue" победил действующего чемпиона мира по шахматам Гарри Каспарова</u> . Pr; <u>Победа машины над человеком ознаменовала важный поворотный момент для ИИ</u> . Con: Текущий период можно охарактеризовать как <u>время "зарождающихся технологических прорывов в самом широком спектре областей"</u> , включая ИИ.

Окончание таблицы В.1

16. <i>Verbal Classification</i> Premise: Для всех <u>x</u>, если <u>x</u> имеет свойство <u>F</u>, то <u>x</u> можно отнести к имеющим свойство <u>G</u>. Premise: <u>A</u> имеет свойство <u>F</u>. Conclusion: <u>A</u> имеет свойство <u>G</u>.	– Посылки содержат общие тезисы, основание классификации и/или свойство принадлежности к классу. – Вывод – обозначение общего класса. – Часто используется при определении нескольких классов.	Pr: <u>Каждый уровень</u> напрямую зависит от <u>элементов содержания, обеспечивающих</u> <u>достижение эквивалентности</u>. Pr: Он выделил <u>пять уровней (типов)</u> <u>эквивалентности</u>: <u>цели коммуникации,</u> <u>описания ситуации, высказывания,</u> <u>сообщения и языковых знаков</u>. Con: Только при <u>достижении идентичности</u> <u>на всех уровнях</u> содержания текста оригинала и текста перевода, <u>перевод</u> <u>можно считать эквивалентным</u>.
---	---	--

Приложение Г

Пример аргументационной разметки короткой научной статьи

Аргументационная разметка научной статьи, представленная в примере, доступна в корпусе на платформе ArgNetBank Studio по следующему адресу:

<https://uniserv.iis.nsk.su/arg/graph/di?openText=6746fc36db03e7d38563a2f5&graphId=6746fccc06dcea821b01381b>.

Исходный текст размеченной статьи взят из открытой электронной научной библиотеки «КиберЛенинка»:

<https://cyberleninka.ru/article/n/obschiy-podhod-k-snizheniyu-riskov-pri-vnedrenii-kvantovyh-tehnologiy>.

Источник – Снежко В. К. Общий подход к снижению рисков при внедрении квантовых технологий / В. К. Снежко, С. А. Якушенко, А. В. Лукашев, В. Е. Егрушев, С. С. Веркин, В. В. Антонов, Е. В. Чеканова // International Journal of Humanities and Natural Sciences. – Vol. 6–3 (81). – 2023. – С. 151–155. – DOI: 10.24412/2500-1000-2023-6-3-151-155.

Текст размечаемой статьи выдан аннотатором в предобработанном виде, выданном разметчиком. Предобработка включала удаление сопроводительных данных об авторах, аннотации и списка ключевых слов, табличного и иллюстративного материала, библиографического списка. В таблице Г.1 представлен список аргументов, выявленных разметчиком, где для каждого аргумента указаны утверждения в его составе (с их номерами) и схема в основе их связи. В таблице Г.2 указаны сегменты текста, помеченные разметчиком как неаргументативные, с интерпретацией такого решения в комментариях. На Рис. Г.1 изображена визуализированная структура построенного аргументационного графа, отмечены номера вершин в соответствии с таблицей Г.1 (прямоугольные блоки соответствуют утверждениям, овальные между ними – аргументационным схемам). Единственный случай противопоставления двух утверждений через схему конфликта отмечен в таблице Г.1 индексом «C1», на Рис. 1 – красной вершиной.

Таблица Г.1 – Аргументы в разметке текста

№	Посылка(и)	Заключение	Схема
1	(1) Бурное развитие и внедрение квантовых технологий определяет основное содержание современного научнотехнического прогресса в сфере информационных технологий	(2) При этом между участниками этого процесса сформировались отношения не столько международной кооперации, сколько соревнования за достижения квантового превосходства или, как недавно принято, квантового преимущества	<i>Cause to Effect</i>
2	(2) При этом между участниками этого процесса сформировались отношения не столько международной кооперации, сколько соревнования за достижения квантового превосходства или, как недавно принято, квантового преимущества	(3) С этой целью группой западных стран 31 января 2023 года учреждён Международный совет ассоциаций квантовой промышленности, объединивший передовые национальные ассоциации Канады – Quantum Industry Canada, QIC, США – Консорциум квантового экономического развития QED-C, Японии Quantum Strategic Industry Alliance for Revolution, Q-STAR и Европейского Союза – Европейский консорциум квантовой промышленности, QuIC	<i>Practical Reasoning</i>
3	(3) С этой целью группой западных стран 31 января 2023 года учреждён Международный совет ассоциаций квантовой промышленности, объединивший передовые национальные ассоциации Канады – Quantum Industry Canada, QIC, США – Консорциум квантового экономического развития QED-C, Японии Quantum Strategic Industry	(9) Таким образом, переход на новый технологический уровень в сфере развития информационных технологий и систем коммуникаций постепенно набирает обороты	<i>Effect to Cause</i>

Продолжение таблицы Г.1

	Alliance for Revolution, Q-STAR и Европейского Союза – Европейский консорциум квантовой промышленности, QuIC		
4	(5) Признанным лидером квантовой гонки является Китай, вложивший в 2021 году в эту сферу 10 млрд. долларов США, что составило около 40 процентов мировых инвестиций [2]	(4) Большинство развитых стран имеют национальные программы создания и внедрения квантовых технологий в сферы информационного обеспечения экономической и оборонной деятельности	<i>Example</i>
5	(6) Американская программа предусматривает создание и использование квантового интернета в интересах силовых ведомств [3]	(4) Большинство развитых стран имеют национальные программы создания и внедрения квантовых технологий в сферы информационного обеспечения экономической и оборонной деятельности	<i>Example</i>
6	(7) Дорожная карта создания и развития интегрированной квантовой информационной сети с космическим сегментом Евросоюза рассчитана на период с 2020 по 2036 год [4]	(4) Большинство развитых стран имеют национальные программы создания и внедрения квантовых технологий в сферы информационного обеспечения экономической и оборонной деятельности	<i>Example</i>
7	(8) Подобная дорожная карта существует и в России [5]	(4) Большинство развитых стран имеют национальные программы создания и внедрения квантовых технологий в сферы информационного обеспечения экономической и оборонной деятельности	<i>Example</i>

Продолжение таблицы Г.1

8	(4) Большинство развитых стран имеют национальные программы создания и внедрения квантовых технологий в сферы информационного обеспечения экономической и оборонной деятельности	(9) Таким образом, переход на новый технологический уровень в сфере развития информационных технологий и систем коммуникаций постепенно набирает обороты	<i>Effect to Cause</i>
9	(9) Таким образом, переход на новый технологический уровень в сфере развития информационных технологий и систем коммуникаций постепенно набирает обороты	(10) Следует отметить, что переход на новый технологический уклад в данной сфере попутно создает угрозы существующим системам криптографии, а следовательно, и национальной информационной безопасности	<i>Cause to Effect</i>
C1	(11) До недавнего времени считалось, что квантовые компьютеры, способные взламывать современные криптографические системы появятся через 5-10 лет [6]	(12) Но действительность начала опровергать эти расчеты	<i>Logical</i>
10	(13) Об этом свидетельствует целый ряд свершившихся успешных экспериментов	(12) Но действительность начала опровергать эти расчеты	<i>Sign</i>
11	(14) 4 мая 2023 года российское электронное периодическое издание «3DNews», со ссылкой на американское агентство Semiconductor Digest, сообщило о начале проектирования квантовых процессоров для квантовых компьютеров емкостью миллион кубитов [7]	(13) Об этом свидетельствует целый ряд свершившихся успешных экспериментов	<i>Example</i>

Продолжение таблицы Г.1

12	(15) Ранее, 4 января 2023 года, китайские исследователи опубликовали [8] результат успешного эксперимента по взлому криптографического ключа RSA-48 с помощью разработанного ими протокола QAOA (Quantum Approximation Optimization Algorithm) на 10-ти кубитном квантовом компьютере применительно к алгоритму Шнорра [9] (16) Авторы эксперимента предоставили аналитические расчеты по характеристикам квантового компьютера, способного взломать криптографические ключи различной длины одной из основных систем шифрования (табл. 1)	(13) Об этом свидетельствует целый ряд свершившихся успешных экспериментов	<i>Example</i>
13	(19) К этому следует добавить, что еще в ноябре прошлого года компания IBM ввела в эксплуатацию квантовый компьютер Orsey емкостью 433 кубита, а к концу текущего года запустит систему Quantum System Two емкостью 1127 кубитов	(13) Об этом свидетельствует целый ряд свершившихся успешных экспериментов	<i>Example</i>

Продолжение таблицы Г.1

14	(12) Но действительность начала опровергать эти расчеты	(20) Сопоставляя эти сообщения, приходится сделать вывод о наличии реальных угроз информационной безопасности в критически важных областях, в том числе, в системах управления войсками и оружием	<i>Effect to Cause</i>
15	(18) Ценность опубликованного эксперимента состоит в декларации угрозы со стороны квантовых компьютеров в ближайшем будущем	(15) Ранее, 4 января 2023 года, китайские исследователи опубликовали [8] результат успешного эксперимента по взлому криптографического ключа RSA-48 с помощью разработанного ими протокола QAOA (Quantum Approximation Optimization Algorithm) на 10-ти кубитном квантовом компьютере применительно к алгоритму Шнорра [9]	<i>Positive Consequences</i>
16	(17) В описании эксперимента авторы представили экспериментальную установку и разработанный ими алгоритм взлома ключа шифрования	(16) Авторы эксперимента предоставили аналитические расчеты по характеристикам квантового компьютера, способного взломать криптографические ключи различной длины одной из основных систем шифрования	<i>Applied Method</i>
17	(21) И уже не в отдаленном, как считалось до недавнего времени, а в ближайшем будущем	(20) Сопоставляя эти сообщения, приходится сделать вывод о наличии реальных угроз информационной безопасности в критически важных областях, в том числе, в системах управления войсками и оружием	<i>Effect to Cause</i>

Продолжение таблицы Г.1

18	(20) Сопоставляя эти сообщения, приходится сделать вывод о наличии реальных угроз информационной безопасности в критически важных областях, в том числе, в системах управления войсками и оружием (22) Существуют различные подходы к решению возникшей проблемы безопасности	(23) Если отбросить позицию, что ничего не надо делать, так как угроза мнимая, то есть три основных варианта	<i>Practical Reasoning</i>
19	(24) Один из них – разработка и внедрение постквантовых (квантовобезопасных) алгоритмов шифрования	(23) Если отбросить позицию, что ничего не надо делать, так как угроза мнимая, то есть три основных варианта	<i>Verbal Classification</i>
20	(25) Другой – создание и освоение квантовых технологий, в первую очередь, квантового распределения ключей [10]	(23) Если отбросить позицию, что ничего не надо делать, так как угроза мнимая, то есть три основных варианта	<i>Verbal Classification</i>
21	(23) Если отбросить позицию, что ничего не надо делать, так как угроза мнимая, то есть три основных варианта	(26) Разумный подход заключается в оптимизации соотношения этих вариантов, когда на наиболее важных информационных направлениях используются квантовые, а на менее важных – постквантовые системы криптографии	<i>Cause to Effect</i>
22	(10) Следует отметить, что переход на новый технологический уклад в данной сфере попутно создает угрозы существующим системам криптографии, а следовательно, и национальной информационной	(39) Таким образом, в статье в общем виде предложен подход к обоснованию рациональной организации перехода на новый технологический уровень путем снижения рисков в различных	<i>Practical Reasoning</i>

Продолжение таблицы Г.1

	безопасности	сферах управления, в первую очередь, в критически важных отраслях	
23	(27) - квантовые технологии создают риски взлома существующим системам криптографии, однако нет гарантий того, что в будущем такой же угрозе не могут подвергнуться и постквантовые системы;	(28) - постквантовые системы могут и должны модернизироваться с целью повышения устойчивости к взлому;	<i>Cause to Effect</i>
24	(30) При этом необходимо уделять внимание важности исследований в создании систем связи на принципах нелокальности [11];	(29) - на первом этапе внедрения квантовых технологий в системы связи специального назначения целесообразно применять квантовое распределение ключей на наиболее важных информационных направлениях	<i>Practical Reasoning</i>
25	(28) - постквантовые системы могут и должны модернизироваться с целью повышения устойчивости к взлому;	(32) - по мере роста угроз со стороны квантовых компьютеров, следует постепенно переводить направления с постквантовых на квантовые системы.	<i>Practical Reasoning</i>
26	(26) Разумный подход заключается в оптимизации соотношения этих вариантов, когда на наиболее важных информационных направлениях используются квантовые, а на менее важных – постквантовые системы криптографии	(35) - с учетом аналитических оценок рисков для существующих и постквантовых систем криптографии, целесообразно создавать планы перехода на квантовые системы коммуникаций информационных направлений в соответствии с их ранжированием по важности	<i>Practical Reasoning</i>

Продолжение таблицы Г.1

27	(29) - на первом этапе внедрения квантовых технологий в системы связи специального назначения целесообразно применять квантовое распределение ключей на наиболее важных информационных направлениях	(32) - по мере роста угроз со стороны квантовых компьютеров, следует постепенно переводить направления с постквантовых на квантовые системы	<i>Practical Reasoning</i>
28	(31) - на менее важных направлениях могут использоваться постквантовые (квантовоустойчивые) алгоритмы шифрования;	(32) - по мере роста угроз со стороны квантовых компьютеров, следует постепенно переводить направления с постквантовых на квантовые системы	<i>Practical Reasoning</i>
29	(32) - по мере роста угроз со стороны квантовых компьютеров, следует постепенно переводить направления с постквантовых на квантовые системы (33) Для упорядоченной организации такого перехода необходимо всесторонне анализировать зарубежные эксперименты и отечественные достижения в сфере развития квантовых систем вычислений с целью обоснованной оценки рисков, а также планового создания резервов квантовых систем и оборудования, подготовки кадров для их эксплуатации;	(35) - с учетом аналитических оценок рисков для существующих и постквантовых систем криптографии, целесообразно создавать планы перехода на квантовые системы коммуникаций информационных направлений в соответствии с их ранжированием по важности	<i>Practical Reasoning</i>

Продолжение таблицы Г.1

30	(34) - поскольку сами квантовые системы коммуникации подвержены квантовым атакам, следует развивать методы защиты от них, в том числе создание гибридных систем, например, развитием стеганографии [12] внутри квантовых систем коммуникаций;	(35) - с учетом аналитических оценок рисков для существующих и постквантовых систем криптографии, целесообразно создавать планы перехода на квантовые системы коммуникаций информационных направлений в соответствии с их ранжированием по важности	<i>Practical Reasoning</i>
31	(36) В свою очередь, такие планы позволяют своевременно организовать закупки оборудования для квантовых систем в промышленности	(35) - с учетом аналитических оценок рисков для существующих и постквантовых систем криптографии, целесообразно создавать планы перехода на квантовые системы коммуникаций информационных направлений в соответствии с их ранжированием по важности	<i>Positive Consequences</i>
32	(37) В этом аспекте целесообразно установить требование к защите диссертационных и дипломных работ по темам квантовых коммуникаций по результатам успешного эксперимента	(35) - с учетом аналитических оценок рисков для существующих и постквантовых систем криптографии, целесообразно создавать планы перехода на квантовые системы коммуникаций информационных направлений в соответствии с их ранжированием по важности	<i>Practical Reasoning</i>
33	(35) - с учетом аналитических оценок рисков для существующих и постквантовых систем криптографии, целесообразно создавать планы перехода на квантовые системы коммуникаций	(39) Таким образом, в статье в общем виде предложен подход к обоснованию рациональной организации перехода на новый технологический уровень путем снижения рисков в различных	<i>Applied Method</i>

Окончание таблицы Г.1

	информационных направлений в соответствии с их ранжированием по важности	сферах управления, в первую очередь, в критически важных отраслях	
34	(38) Такое требование создаст условия для более оперативного решения задач, развития и укрепления информационной безопасности и в полной мере обеспечит загрузку развернутых экспериментальных полигонов	(37) В этом аспекте целесообразно установить требование к защите диссертационных и дипломных работ по темам квантовых коммуникаций по результатам успешного эксперимента	<i>Positive Consequences</i>
35	(39) Таким образом, в статье в общем виде предложен подход к обоснованию рациональной организации перехода на новый технологический уровень путем снижения рисков в различных сферах управления, в первую очередь, в критически важных отраслях	(40) Его суть сводится к объективной оценке угроз со стороны квантовых компьютеров и, на этой основе, планирования процесса перехода, управления финансовыми ресурсами и промышленными возможностями, начиная с научно-исследовательских работ и заканчивая организацией государственных закупок	<i>Practical Reasoning</i>

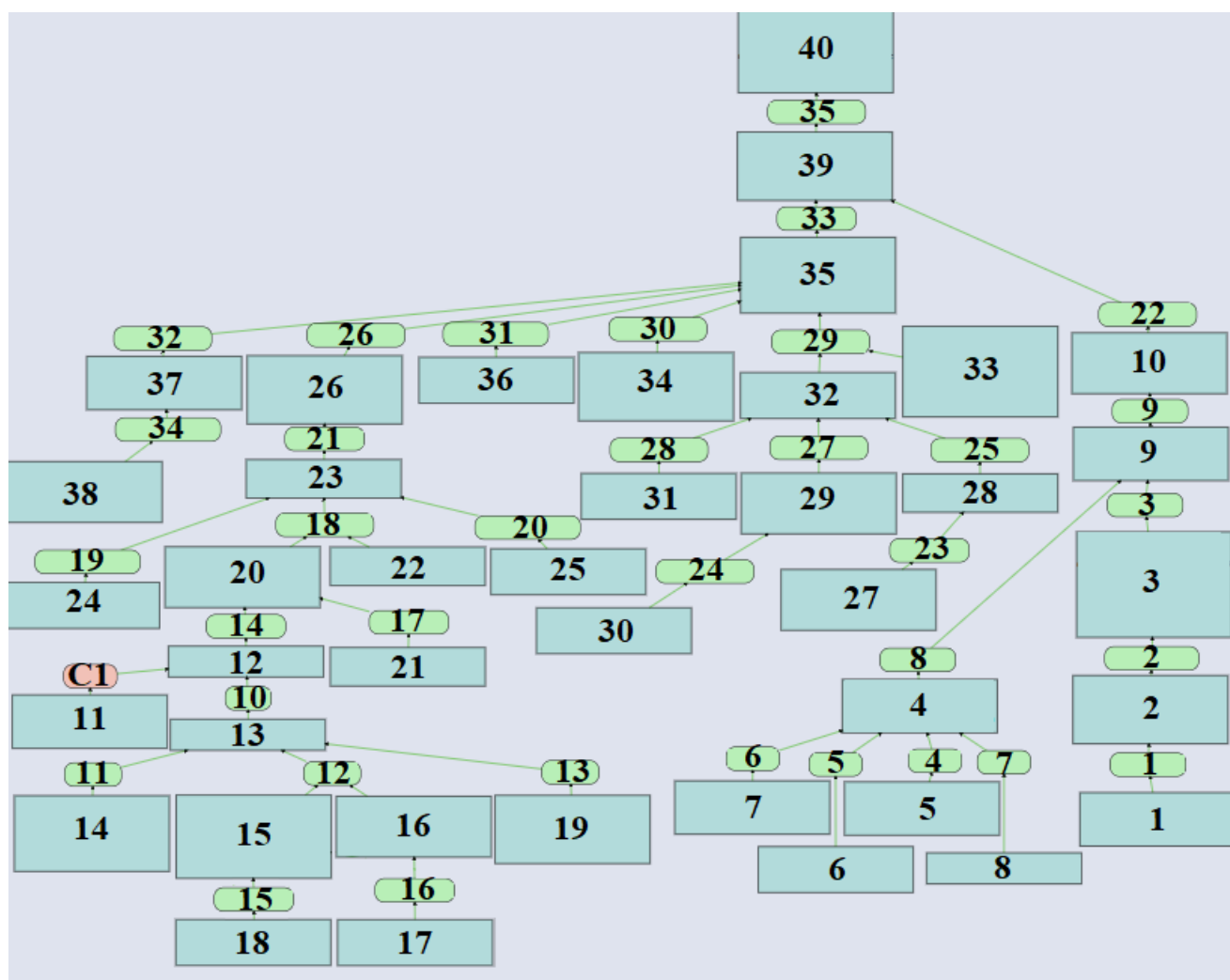


Рисунок Г.1 – Общая структура аргументационного графа для текста

Таблица Г.2 – Сегменты, отнесённые к неаргументативному содержанию

№	Сегмент	Комментарий
1	Общий подход к снижению рисков при внедрении квантовых технологий	Заголовок статьи не включается в аргументационный граф.
2	Выступая на заключительном заседании этого совета, его председатель, в частности, заявил о начале глобальной технологической революции [1]	Комментарий частного характера, о чём дополнительно свидетельствует авторский маркер «в частности». Формально может быть включен как утверждение, поддерживающее тезис (9) по схеме «Expert Opinion».
3	Поэтому ниже такие угрозы рассмотрены в аспекте обоснования подходов к их снижению	Сегмент риторической структуры текста, уточняющий организацию статьи через обозначение далее рассматриваемой темы.

Продолжение таблицы Г.2

4	Вероятные угрозы информационной безопасности со стороны квантовых технологий	Заголовок раздела статьи не включается в аргументационный граф.
5	Для оценки реальности угроз информационной безопасности рассмотрим некоторые факты, имеющиеся в доступных источниках	Сегмент риторической структуры текста, уточняющий организацию статьи (в этом сегменте пока не представляются ни факты, ни источники).
6	Таблица 1. Расчетные характеристики квантового компьютера для взлома криптосистем RSA	Название таблицы не включается в аргументационный граф.
7	Общий подход к безопасному переходу на новый технологический уровень	Заголовок раздела статьи не включается в аргументационный граф.
8	Здесь мы формулируем подход к детализации сосуществования двух различных подходов	Сегмент риторической структуры текста, уточняющий организацию статьи через обозначение её дальнейшего фокуса.
9	На рисунке 1 этот подход представлен в динамике	Общий комментарий о содержании рисунка.
10	Рис. 1. Плановый переход на новый технологический уровень при внедрении квантовых технологий	Название рисунка не включается в аргументационный граф.
11	Суть нашего подхода состоит в следующих основных положениях:	Обобщающее утверждение абстрактной семантики (положения включены в аргументационный граф напрямую).
12	На рисунке дополнительно показаны блоки доподготовки научных и преподавательских кадров, в том числе с использованием экспериментальных полигонов, районов и лабораторий для симуляции, либо натуральных систем для квантовых атак и разработки методов защиты от них	Общий комментарий о содержании рисунка.
13	Заключение	Заголовок не включается в граф.

Приложение Д

Наиболее сложный приём аргументации в научных статьях

Структура приёма: *AppliedMethod* ← [*VerbalClassification*, *VerbalClassification*, *VerbalClassification* ← {*Example*}, *PracticalReasoning* ← {*CauseToEffect* ← (*ExpertOpinion*, *ExpertOpinion*)}].

Текст 1

(10) «Таким образом, выявленные ассоциации языкового сознания педагогов Красноярского края позволяют утверждать, что любовь является одной из ценностных доминант данной профессиональной группы и содержит концептуальное ядро, тесно связанное с понятиями чувство, семья, дети» ← **AppliedMethod** ← (5) «Основываясь на данных массового свободного ассоциативного эксперимента (САЭ) среди учителей и выпускников педагогических учебных заведений Красноярского края, объединенных в электронную базу данных первого профессионального ассоциативного словаря педагога «Ассоциативный словарь профессиональной группы “педагог” Красноярского края», рассмотрим представление о любви в языковой картине мира педагогов Приенисейской Сибири» ← [**Verbal Classification** ← (6) «Ядерная зона ассоциативного поля «Любовь» представлена небольшой группой реакций с частотой от 12 до 7 и включает следующие лексемы: чувство 12, семья 8, счастье 7, нет реакции 7, показывающие в большинстве отношение педагогов к любви как ценности семейной, как к чувству, которое приносит счастье»; **Verbal Classification** ← (7) «Здесь встречаем также довольно небольшое количество реакций, связанных с прецедентными текстами (зла 6, морковь 3), семейными ценностями (муж 3, дети 3, верность 3), степенью любовного накала как движущей силы (сила 3, страсть 3, нежность 4), эмоциональным мироощущением (романтика 3), отношением к жизни (жизнь 5)»; **Verbal Classification** ← (8) «Периферийная зона обширна и показательна (частота от 2 до 1): ассоциации профессиональной группы связаны как с положительными

эмоциями, так и с отрицательными (боль, зло, вера, брак, поддержка, надежда, презрение, беспокойное сердце, разлука, радость, ненависть, болезнь, сердце, моя, благо и др.), реакции, указывающие на любовь внутри семьи, ее длительность, характеристику (к мужу, мужа, мама, муж, дети, сын, к родителям, жена, супруг, родителей, на всю жизнь, навеки, на век, вечность, навечно, не вечна, бесконечная, до гроба, неземная, взаимная, вечная, первая, большая, навсегда, крепкая, безответная, уважение, с первого взгляда, светлое чувство, взаимность, завидное, великое чувство, чистота и др.), связанные с популярным советским фильмом «Любовь и голуби» (1984) (и голуби)» ← **{Example** ← (9) «Лексема бабочка, на наш взгляд, может быть отнесена к нескольким группам ассоциаций: недолговечности чувства, ее легкости, красоте, степени положительных эмоций»}; **Practical Reasoning** ← (4) «Учитывая современное отношение мирового сообщества «ко всему русскому», когда идет постепенное овладение «русским народом через малозаметную инфильтрацию его души и воли» [Ильин, 1993, с. 169], вопрос о национальной идентичности, сохранении образа мира, в том числе через прививаемые педагогом традиционные духовные ценности, остается актуальным» ← **{Cause to Effect** ← (2) «Ценностные представления профессиональной группы «педагог» интересны нам не только с точки зрения принадлежности к данной группе, но и как необходимость изучения внутреннего мира современного учителя как важнейшего транслятора насущно необходимых знаний, «заведомо обладающего как специальными знаниями, так и четким представлением об аксиологической шкале, желаемой или доминирующей в обществе» ← (**Expert Opinion** ← (1) ««Мы живем в “смутное время”, время крушения социальных идей, скреплявших наше общество, когда в качестве реакции на социальную неопределенность возрастает роль этничности как наиболее древней и устойчивой формы информационного структурирования мира» [Уфимцева, 1998, с. 137]»; **Expert Opinion** ← (3) «Названные качества присутствуют у проводника знаний независимо от его узкой специализации, поэтому, как известно, “преподавание – одно из средств воспитания” [Цит. по: Васильева и др., 2022, с. 6]»}}].

Текст 2

(10) *«Итак, модулирующий метод в когнитивном переводе может применяться по нескольким сценариям, то есть соответствовать нескольким типам когнитивных связей, протекающих в той или логической цепочке: тип контекстуальной импликации, лингвоэкологической импликации, конъюнкции, бинарной конъюнкции и дизъюнкции»* ← **AppliedMethod** ← (5) *«В процессе изучения исходных текстов и их переводов было обнаружено несколько сценариев, по которым переводчик выстраивает свои мыслительные операции»* ← [**Verbal Classification** ← (8) *«Данный пример демонстрирует тип связи лингвоэкологической импликации, направленной на преодоление сбоя ориентации в когнитивном пространстве мышления, то есть избегания интерференции и выработку релевантных логических операций в поиске наиболее эквивалентного варианта перевода, согласно нормам переводного языка»*; **Verbal Classification** ← (9) *«Данный пример демонстрирует тип связи конъюнкции, характеризующийся приемом логической синонимии при подборе эквивалента»*; **Verbal Classification** ← (7) *«В данном примере используется тип связи контекстуальной импликации, характеризующийся подбором соответствующего эквивалента перевода в зависимости от его контекстуального окружения»* ← {**Example** ← (6) *«Рассмотрим исходное предложение и его перевод: The international exchange of evidence in criminal matters through formal mutual legal assistance arrangements is a fairly recent phenomenon [4, с. 232]. – Международный обмен доказательствами по уголовным делам, осуществляемый на основании официальных соглашений о взаимной правовой помощи, является новаторством»*}; **Practical Reasoning** ← (4) *«Таким образом, целью работы является представление модулирующего метода путем выявления когнитивных корреляций между контактными точками в переводческом процессе»* ← {**Cause to Effect** ← (1) *«Данная интеграция предназначена для решения переводческих задач посредством методологии когнитивной лингвистики, а также с целью разработки новых комплементарных методов, удовлетворяющих потребностям теории перевода, верифицирующих на разнообразном переводческом материале*

уже имеющиеся когнитивные техники и способствующих обогащению потенциала когнитивных наук» ← (**Expert Opinion** ← (2) «Когнитивный подход к переводческим проблемам начали применять с 1982 года, когда Гидеон Тури сравнил загадочную работу мозга переводчика с «черным ящиком» [11]»; **Expert Opinion** ← (3) «С тех пор ученые фокусировались на различных стадиях и этапах переводческого процесса: рассматривали когнитивные стратегии сегментации исходного текста [1; 2], обращались к ментальным моделям восприятия исходного текста [8, с. 51], исследовали явление «когнитивного наложения» (overlapping) ментальных операций восприятия, обработки и продуцирования информации [6; 7; 9], изучали когнитивные микро- и макропроцессы переводческой деятельности [3; 9] и т. д.»}}].

Приложение Е

Лексико-грамматические шаблоны с маркерами этоса и пафоса

Таблица Д.1 – представление шаблонов этоса для программной обработки

№	Шаблон
1	{в // PREP} {литература // loct}
2	{(автор учёный исследователь специалист эксперт профессор другой иной третий) // plur} {...} {(считать полагать обнаружить сомневаться выяснить утверждать предлагать отметить говорить рассказывать провести объяснять пояснять рассматривать) // plur}
3	{в // PREP} {(работа концепция статья труд) // loct} {- // NOUN,anim}
4	{как // CONJ} {известно // PRED}
5	{известно // PRED} {(что о) // (CONJ PREP)}
6	{по // PREP} {(мнение выражение слово наблюдение утверждение сообщение) // datv} {- // NOUN,anim}
7	{(считаться полагаться утверждаться говориться отмечаться) // -} {что // CONJ}
8	{принять // PRTS} {- // INFN}
9	{- // anim...nomn} {(считать полагать обнаружить сомневаться выяснить утверждать предлагать отметить говорить писать рассказывать подсчитать предсказать) // -}
10	{(считать полагать сомневаться утверждать предлагать отметить говорить писать рассказывать подсчитать предсказать) // -} {- // anim...nomn}
11	{(мнение тезис) // -} {что // CONJ}
12	{по // PREP} {...} {данные // datv}
13	{согласно // PREP} {- // NOUN...(gent datv loct)}
14	{(квалифицированный профессиональный ведущий авторитетный знаменитый почетный именитый) // masc}
15	{- // anim...nomn} {(изучить протестировать проверить испытать проанализировать) // -}
16	{(изучить протестировать проверить испытать проанализировать) // -} {- // anim...nomn}
17	{(группа коллектив сообщество) // -} {(автор учёный исследователь специалист эксперт профессор) // -}
18	{- // anim...plur} {из // -} {группа // -}
19	{(говориться утверждаться заявляться) // -} {в // PREP} {(статья работа публикация) // (loct datv)}
20	{как // CONJ} {(показать демонстрировать свидетельствовать указывать) // -} {- // nomn}

Окончание таблицы Д.1

21	{как // CONJ} {(считать полагать обнаружить сомневаться выяснить утверждать предлагать отметить говорить рассказывать провести объяснять пояснять) // plur} {(автор учёный исследователь специалист эксперт профессор) // plur}
22	{например // -}
23	{к // PREP} {пример // -}
24	{согласно // PREP} {- // PRTF...(gent datv loct)} {- // NOUN...(gent datv loct)}
25	{по // PREP} {(наш их его её) // -} {(мнение выражение слово наблюдение утверждение сообщение) // datv}

Таблица Д.2 – представление шаблонов пафоса для программной обработки

№	Шаблон
1	{\! // PNCT}
2	{- // ADJF, Supr}
3	{(лишь весьма вот ведь заведомо якобы даже столь именно слишком гораздо запросто вдруг вдобавок вполне вообще чересчур фактически попросту увы) // -}
4	{(никто нигде никогда ничто никакой незачем всегда) // -}
5	{(конечно безусловно) // CONJ}
6	{очевидный // Qual neut}
7	{(несомненно естественно определённо ясно значительно действительно) // -}
8	{(нуждаться требовать заставлять принуждать приходиться) // 3per}
9	{(необходимо надо нужно) // PRED}
10	{(важно обязательно абсолютно принципиально однозначно наверняка невозможно неизбежно) // ADVB}
11	{коренной // -} {образ // -}
12	{(совершенно явно очень заметно крайне) // -} {- // (ADJF ADJS ADVB)}
13	{интересно // -}
14	{(интересный любопытный примечательный удивительный значительный легендарный недюжинный причудливый принципиальный выгодный) // -}
15	{к // PREP} {счастье // -}
16	{благодаря // PREP}
17	{(единица десятка сотня тысяча миллион) // plur} {- // gent}

Окончание таблицы Д.2

18	{(замечательный приличный банальный ничтожный опасный фантастический губительный изысканный исключительный внушительный громадный ценный отличный серьёзный ключевой убедительный нелепый многочисленный мощный страстный целый беспрецедентный колоссальный несметный лидирующий) // -}
19	{(прорыв сенсация кладезь успех богатство ценность скачок свалка) // -}
20	{(завоевание борьба беспорядок) // -}
21	{(разгромить перечеркивать опровергать пропагандировать уничтожить покоробить ужаснуть обязать впечатлять кипеть разгадать удивить избавляться гарантировать) // -}
22	{(фундаментальный существенный огромный) // -} {(значение роль) // -}
23	{(совсем вовсе отнюдь точно далеко) // -} {не // -}
24	{должный // ADJS}
25	{всё // -} {равно // -}
26	{не // -} {(только просто) // -} {...} {(а но) // CONJ}
27	{тот // -} {самый // -}
28	{известный // -} {случай // plur}
29	{(удаться удаваться) // -} {- // INFN}
30	{(не ни) // -} {один // -}
31	{самый // -} {- // (ADJF ADJS)}
32	{не // -} {мочь // -} {не // -}
33	{без // PREP} {сомнение // -}
34	{(всё все ещё) // -} {- // COMP}
35	{(легко просто сильно) // -} {- // INFN}
36	{более // -} {тот // gent}
37	{(мягко честно строго) // -} {говорить // -}
38	{(давно хорошо прекрасно все каждый) // -} {известно // PRED}
39	{(сразу тотчас) // -} {же // PRCL}