



АНАЛИЗ ВОЗМОЖНОСТИ ПРИМЕНЕНИЯ МОДЕЛЕЙ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА ДЛЯ РАСПОЗНАВАНИЯ РЕЧИ ПРИ ПРОВЕДЕНИИ СОЦИОЛОГИЧЕСКИХ ОПРОСОВ

Представлены результаты сравнительного анализа открытых моделей искусственного интеллекта для автоматического распознавания речи в условиях полевых социологических опросов. Целью исследования было определение наиболее эффективной модели семейства Whisper для транскрибации интервью при наличии акустических помех, характерных для уличных условий проведения опросов. Методология включает экспериментальную оценку шести версий модели Whisper (Tiny, Base, Small, Medium, Large-v3, Large-v3-turbo) с применением метрики IWER (Inflectional Word Error Rate). Моделирование проводилось на аудиофайлах с искусственно внесёнными шумами различной интенсивности (соотношение сигнал/шум от 9,9 до 22 LUFs), имитирующими пять типичных сценариев проведения опросов. Установлено, что модель Large-v3-turbo демонстрирует наилучшую точность распознавания: показатель метрики IWER составляет от 0,79% (условия жилого сектора) до 4,86% (непосредственно у дороги), что на 20% ниже по сравнению с базовой моделью Tiny. Новизна исследования заключается в использовании сравнительной оценки метрики IWER для моделей Whisper в условиях, специфичных для социологических полевых исследований. Полученные данные позволяют обосновать выбор модели для создания специализированных информационных систем автоматизированной обработки речевых ответов респондентов.

Ключевые слова: социологический опрос, распознавание речи, искусственный интеллект, Whisper, метрика IWER, акустические помехи, транскрибация, полевые исследования, нейронные сети, автоматизация обработки данных.

ANALYSIS OF THE POSSIBILITY OF USING ARTIFICIAL INTELLIGENCE MODELS FOR SPEECH RECOGNITION IN SOCIOLOGICAL SURVEYS

The results of a comparative analysis of open-source artificial intelligence models for automatic speech recognition under field sociological survey conditions are presented. The aim of the study was to determine the most effective model of the Whisper family for interview transcription in the presence of acoustic noise characteristic of outdoor survey conditions. The methodology includes an experimental evaluation of six versions of the Whisper model (Tiny, Base, Small, Medium, Large-v3, Large-v3-turbo) using the IWER (Inflectional Word Error Rate) metric. Modeling was performed on audio files with artificially added noise of varying intensity (signal-to-noise ratio ranging from 9.9 to 22 LUFs), simulating five typical survey scenarios. It was established that the Large-v3-turbo model demonstrates the highest recognition accuracy: the IWER metric ranges from 0.79% (residential area conditions) to 4.86% (directly near the road), which is 20% lower compared to the baseline Tiny model. The novelty of the research lies in the comparative evaluation of the IWER metric for Whisper models under conditions specific to sociological field research. The obtained data enable substantiating the model choice for creating specialized information systems for automated processing of respondents' speech responses.

Keywords: sociological survey, speech recognition, artificial intelligence, Whisper, IWER metric, acoustic noise, transcription, field research, neural networks, data processing automation.

В настоящее время известны различные области применения систем автоматического распознавания речи [1]: построение систем управления устройствами голосовыми командами, автоматизация обработки звонков и улучшения качества обслуживания клиентов, транскрибация аудиозаписей в текст, в т.ч. в условиях помех и др. Среди вышеназванных, отдельный интерес представляет обработка социологических данных, получаемых в ходе проведения социологических опросов в форме интервьюирования, где также могут применяться подобные системы [2].

Известны различные средства распознавания речи, среди которых можно выделить коммерческие (Google Speech-to-Text, Microsoft Azure Speech Service, Amazon Transcribe, IBM Watson Speech to Text и т.д.) [3, 4] и открытые (Whisper (OpenAI), Vosk, SpeechBrain и т.д.) [5]. Среди них наиболее точными для решения

прикладных задач являются Google Speech-to-Text (работает на основе SpeechRecognition), Whisper, Vosk и SpeechBrain [5-8]. Сравнительная характеристика этих средств представлена в табл. 1.

Анализ данных табл. 1 свидетельствует о том, что наибольший научный интерес на текущем этапе представляют модели системы распознавания речи Whisper. К их ключевым преимуществам перед аналогами относятся поддержка русского языка, открытый доступ к архитектурам моделей и возможность автономной работы.

Особенности измерения и представления громкости звука в средствах вычислительной техники

При распознавании речи ключевой характеристикой в условиях четкой дикции интервьюера и респондента является громкость.

Сравнительная характеристика средств распознавания

Параметры \ Средства распознавания	Whisper	Speech Recognition	SpeechBrain	Vosk
Режим работы	Офлайн/Онлайн	Онлайн	Онлайн	Офлайн/Онлайн
Доступ к архитектурам моделей	Есть	Есть	Есть	Отсутствует
Распознавание русского языка	+	+	–	+

Таблица 2

Сравнительная таблица моделей Whisper

Модель	Количество параметров	Область применения
Tiny	39 млн	Быстрое распознавание речи с низкими требованиями к вычислительным ресурсам и точности
Base	74 млн	Задачи, требующие более качественного распознавания при ограниченных вычислительных ресурсах
Small	244 млн	Более сложные задачи распознавания речи, где требуется баланс между качеством и скоростью
Medium	769 млн	Профессиональные задачи, где важна высокая точность и качество распознавания с ограниченными вычислительными ресурсами
Large-v3	1,55 млрд	Крупномасштабные и профессиональные проекты, где важна максимальная точность и поддержка множества языков, при этом задействуются большие вычислительные ресурсы
Large-v3-turbo	809 млн	Задачи, где требуется быстрое и точное распознавание речи, с ограниченными вычислительными ресурсами

Громкость звука представляет собой субъективное восприятие силы звукового сигнала, которое коррелирует с интенсивностью звуковых колебаний и уровнем звукового давления. В окружающей среде громкость измеряется с помощью приборов, фиксирующих интенсивность звука непосредственно, например, при разговоре человека, когда его голос регистрируется специальными устройствами, и результат выражается в децибелах (дБ) [9]. В цифровых вычислительных системах, когда речь представлена в виде звукового файла, прямое измерение громкости в дБ невозможно, поскольку вычислительная система не осуществляет прямое измерение акустического сигнала. В таких условиях для оценки громкости применяется величина LUFS (Loudness Units relative to Full Scale), представляющая собой стандартизированный способ измерения средней громкости аудиосигнала с учетом чувствительности человеческого уха к различным частотам. Она предназначена для выравнивания уровня

громкости звукового файла при воспроизведении на различных платформах с целью устранения его внезапного изменения при смене аудиофайлов, что обеспечивает комфортное восприятие звука [10].

Обзор моделей средства распознавания речи Whisper

Whisper – модель преобразования речи в текст, являющаяся открытой и основанной на применении нейронных сетей – трансформеров. В ней реализованы две ключевые составляющие: кодер и декодер. Кодер обрабатывает акустический сигнал, разбивая его на фрагменты и преобразуя в спектрограммы. Декодер на основе этих преобразований генерирует текстовую расшифровку (наиболее вероятную последовательность символов или слов, соответствующих речи) в зависимости от языка исходного речевого сообщения [11].

Наиболее распространенные варианты моделей Whisper, отличающиеся друг от друга числом параметров и структурой, влияю-

щих на их производительность, точность и требования к вычислительным ресурсам, представлены в табл. 2.

Модели Tiny, Base, Small, Medium обучены на 680 тыс. часах аудиозаписей, представленных сообщениями на 97 языках с посторонними шумами, смехом и другими звуками. Модель Large-v3 была обучена на смеси примерно 1 миллиона часов размеченного и 4 миллионов часов псевдоразмеченных аудиоданных. Входные данные включали аудиозаписи на множестве языков и в различных условиях записи. Модель Large-v3-turbo также обучалась на 5 миллионах часов размеченного аудио, что обеспечивает её способность точно распознавать речь на более чем 99 языках, включая шумные и сложные аудиозаписи. За счёт архитектурных изменений (сокращение числа слоев декодера с 32 до 4) модель стала легче и быстрее, сохранив при этом качество.

Выбор конкретной реализации модели Whisper определяется преимущественно типом решаемой задачи и доступными вычислительными ресурсами. Следовательно, при выборе модели важно учитывать баланс требований по точности распознавания и вычислительных ресурсов, а также условия формирования речевых сообщений (наличие посторонних шумов).

Экспериментальная оценка точности распознавания речи моделью искусственного интеллекта Whisper при проведении социологических опросов (результаты исследования)

В исследовании использовалась социологическая анкета из 30 вопросов, в т.ч. один вопрос предусматривал множественный выбор ответов. Анализируемый аудиофайл представлял собой запись ответов респондентов на вопросы анкеты, выполненную микрофоном устройства Samsung S20. Запись производилась с расстояния 3 – 4 см от микрофона обычной разговорной громкостью голоса. Среднее значение громкости речи при измерении составило –18 LUFS (для измерения используется программа Adobe Audition 2023 с установленным стандартом EBU R 128 и глубиной кодирования файла 16 бит), что соответствует громкости речи 50 дБ.

Для оценки качества распознавания речи в начале распознавание происходит на аудиоданных, характеризующихся отсутствием фонового шума или его уровнем, существен-

но ниже уровня громкости речи, а в конце – на аудиоданных, в которых искусственно внесены акустические шумы различной громкости.

Данные средства распознавания речи будем сравнивать с помощью показателя флексивной частоты ошибки слов – IWER (Inflectional Word Error Rate).

$$IWER = \frac{S_{hard} + C_{hard} + S_{soft} + C_{soft} + D + I}{N},$$

где C_{hard} – вес «грубой» ошибки, приведшей к замене изначального слова другим отличающимся по смыслу словом (например, диаграмма – дейтаграмма); S_{hard} – количество «грубых» ошибок; C_{soft} – вес «негрубой» ошибки, не повлиявшей на основу слова (например, сделано – сделана); S_{soft} – количество «негрубых» ошибок; D – количество удаленных слов; I – количество излишне добавленных слов; N – общее количество слов в эталоне (речи) [12].

По определению, $C_{soft} < C_{hard}$, $C_{soft} < 1$, $C_{hard} \geq 1$, при этом в данном исследовании принято $C_{soft} = 0,5$, $C_{hard} = 1$.

Исследование предполагает распознавание идентичного аудиофайла с использованием каждой модели, после чего вычисляется флексивная частота ошибок слов. На основании полученных показателей выполняется выбор модели с наибольшей точностью распознавания.

Количество слов в эталонной фразе ответов респондентов $N = 257$ слов. Например, расчет IWER для Large-v3-turbo согласно (1):

$$IWER = (6 \cdot 0,5 + 3) / 257 = 0,0233 \text{ (2,33\%)}.$$

Результаты эксперимента представлены в табл. 3.

После произведенной оценки распознавания аудиофайла без фонового шума каждой моделью Whisper, можно прийти к выводу, что large-v3-turbo лучше всех распознала данный аудиофайл.

В дальнейшем необходимо проверить каждую модель на основе зашумленной речи, т.е. в изначальной речи содержится шум улицы.

Громкость речи: –18 LUFS

1) У дороги: –27,9 LUFS

2) 10 метров от дороги: –29,6 LUFS

3) 15 метров от дороги: –30,7 LUFS

4) 20 метров от дороги: –36 LUFS

5) Во дворе жилого сектора: –40 LUFS

Шум добавлялся методом линейного микширования с коэффициентом β , рассчитанным по формуле $\beta = 10^{(L_{шум} - L_{речь})/20}$, где

Таблица 3

Точность распознавания речи при отсутствии шумов открытого пространства

	Удалено слов	Грубые ошибки	Негрубые ошибки	Излишне добавлено слов	IWER, %
Large-v3-turbo	0	0	6	3	2,33
Large-v3	1	4	11	2	4,86
Medium	0	8	9	1	5,25
Small	1	21	16	3	12,84
Base	19	37	26	5	28,79
Tiny	1	41	25	14	26,65

L_речь = –18 LUFS, L_шум варьировалось от –27,9 до –40 LUFS в зависимости от сценария. После микширования сигнал нормировался к –18 LUFS по стандарту EBU R 128 [10]. Все операции выполнены в Adobe Audition 2023.

Значения параметра IWER в зависимости от условий для каждого средства распознавания представлены в табл. 4 – 8.

Значения параметра IWER в зависимости от условий для различных моделей представлены в табл. 9.

Для удобства визуального анализа таблицы 9 представим полученные данные в виде графиков (рис. 1), из которых следует, что наилучшие результаты продемонстрировала модель Large-v3-turbo.

Таблица 4

Точность распознавания речи при нахождении непосредственно около источника посторонних шумов (дороги)

	Разность громкостей речи и шума, LUFS	Удалено слов	Грубые ошибки	Негрубые ошибки	Излишне добавлено слов	IWER, %
Large-v3-turbo	9,9	2	7	7	0	4,86
Large-v3	9,9	2	6	16	5	8,17
Medium	9,9	12	8	25	2	13,42
Small	9,9	3	29	18	3	17,12
Base	9,9	9	60	38	8	37,35
Tiny	9,9	15	71	25	11	42,61

Таблица 5

Точность распознавания речи при нахождении на расстоянии 10 м от источника посторонних шумов (дороги)

	Разность громкостей речи и шума, LUFS	Удалено слов	Грубые ошибки	Негрубые ошибки	Излишне добавлено слов	IWER, %
Large-v3-turbo	11,6	1	2	5	1	2,53
Large-v3	11,6	2	4	16	1	5,84
Medium	11,6	3	1	15	0	4,47
Small	11,6	3	14	13	1	9,53
Base	11,6	1	48	22	2	24,12
Tiny	11,6	5	56	21	4	29,38

Таблица 6

**Точность распознавания речи при нахождении на расстоянии 15 м от источника
посторонних шумов (дороги)**

	Разность громкостей речи и шума, LUFS	Удалено слов	Грубые ошибки	Негрубые ошибка	Излишне добавлено слов	IWER, %
Large-v3-turbo	12,7	1	2	4	0	1,95
Large-v3	12,7	1	4	10	1	4,28
Medium	12,7	0	13	5	0	6,03
Small	12,7	2	16	17	2	11,09
Base	12,7	7	43	15	6	24,71
Tiny	12,7	10	57	23	3	31,71

Таблица 7

**Точность распознавания речи при нахождении на расстоянии 20 м от источника
посторонних шумов (дороги)**

	Разность громкостей речи и шума, LUFS	Удалено слов	Грубые ошибки	Негрубые ошибка	Излишне добавлено слов	IWER, %
Large-v3-turbo	18	1	4	6	0	2,14
Large-v3	18	1	2	11	1	3,7
Medium	18	1	11	3	0	5,25
Small	18	4	11	20	0	9,73
Base	18	2	30	22	2	17,51
Tiny	18	4	50	27	3	27,43

Таблица 8

**Точность распознавания речи при нахождении в жилом секторе при отсутствии
движения транспортных средств**

	Разность громкостей речи и шума, LUFS	Удалено слов	Грубые ошибки	Негрубые ошибка	Излишне добавлено слов	IWER, %
Large-v3-turbo	22	1	0	2	0	0,79
Large-v3	22	4	0	1	0	1,75
Medium	22	1	2	9	1	3,31
Small	22	3	16	11	0	9,53
Base	22	0	23	19	1	13,06
Tiny	22	7	40	19	5	23,93

Таблица 9

Сводная таблица значений параметра IWER для моделей, %

	У дороги	10 метров от дороги	15 метров от дороги	20 метров от дороги	Во дворе
Large-v3-turbo	4,86	2,53	1,95	2,14	0,79
Large-v3	8,17	5,84	4,28	3,7	1,75
Medium	13,42	4,47	6,03	5,25	3,31
Small	17,12	9,53	11,09	9,73	9,53
Base	37,35	24,12	24,71	17,51	13,06
Tiny	42,61	29,38	31,71	27,43	23,93

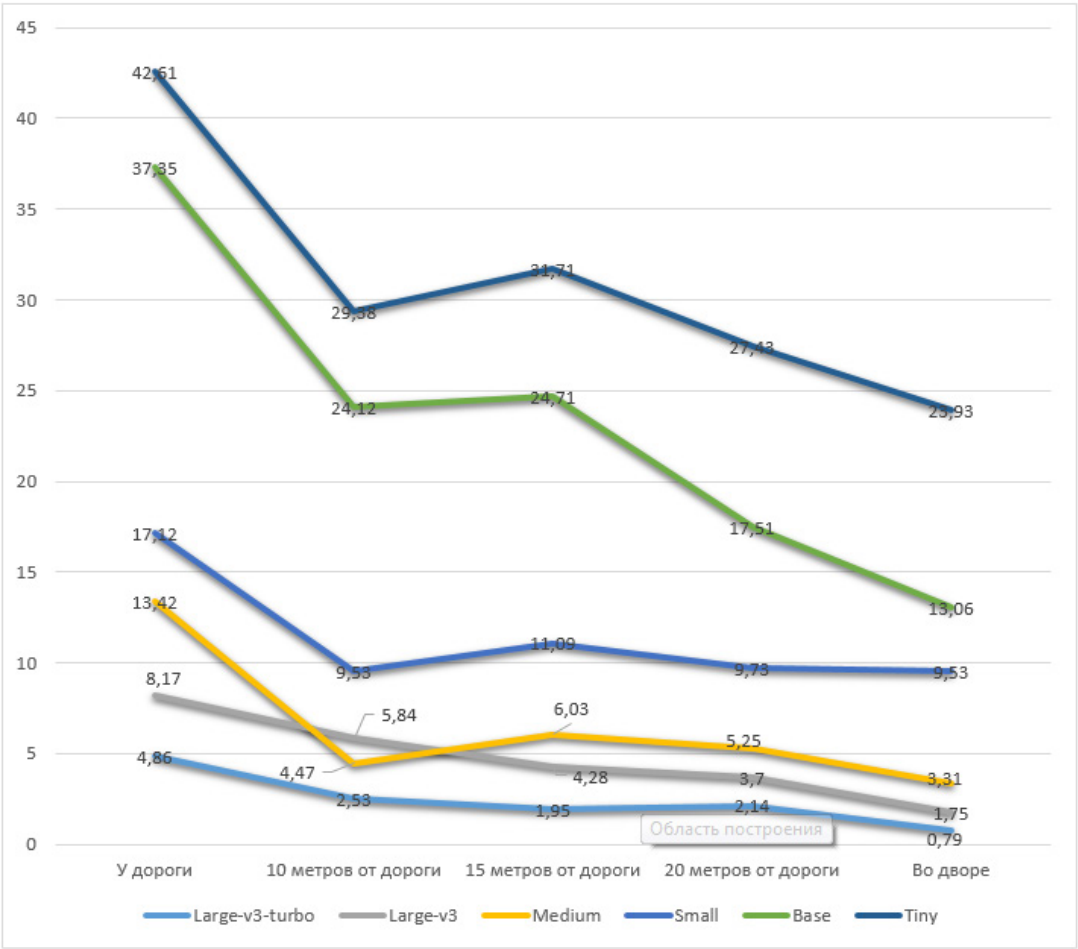


Рис. 1. Зависимости обобщенного показателя качества распознавания IWER от места проведения социологического опроса

Особенность распознавания речи у моделей Medium, Small, Base и Tiny проявляется в следующем: количество флексивных ошибок при расстоянии 10 метров от дороги меньше, чем при расстоянии 15 метров. Предположительно, данное явление обусловлено особенностями обучения этих моделей.

Кроме того, установлено, что показатели качества распознавания в условиях жилого сектора (при отсутствии интенсивных акустических помех) превышают значения, полученные при тестировании на аудиофайлах без фонового шума. Данное явление может свидетельствовать о том, что обучающая вы-

борка моделей включала аудиоданные с низким уровнем фоновых шумов, что способствовало формированию устойчивости алгоритмов к подобным акустическим условиям.

Заключение

Экспериментальная оценка подтвердила, что для автоматизации транскрибации интервью в условиях нестационарного шума оптимальным выбором является модель Large-v3-turbo. Она обеспечивает снижение показателя IWER до 0,79% в благоприятных условиях и сохраняет работоспособность при соотношении сигнал/шум около 10 дБ. В отличие от облегченных версий, данная модель реже допускает грубые смысловые ошибки, что критично для последующего анализа социологических данных. Таким образом, внедрение указанной модели в специализированные информационные системы позволит сократить трудозатраты на обработку речевых ответов

без существенной потери качества.

Перспективы исследований

Развитие этого направления исследований находится в адаптации алгоритмов под предметную область. Стандартные обучающие выборки Whisper не учитывают специфику социологической лексики, поэтому логичным шагом станет дообучение (fine-tuning) на профильных данных. Также планируется оценить эффективность гибридных конвейеров, где распознавание дополняется семантической коррекцией через языковые модели типа ruGPT-3. Нельзя исключать и необходимость тестирования на языках народов РФ для расширения географии исследований. Наконец, для полевых работ актуальна задача оптимизации моделей под мобильные платформы – сравнение квантованных версий поможет найти баланс между скоростью работы и точностью распознавания.

Литература

1. Садыкова А. А., Амиргалиев Е. Н. Изучение применения автоматического распознавания речи // Colloquium-journal. 2020. № 11 (63). С. 44–47. – DOI: 10.24411/2520-6990-2020-11728.
2. Ворошилова А. И. Искусственный интеллект в социологическом исследовании: опыт использования для обработки и анализа интервью // Социодинамика. 2025. № 2. – Электрон. ресурс. – URL: https://nbpublish.com/library_read_article.php?id=73330 (дата обращения: 26.04.2026).
3. Гапочкин А. В. Нейронные сети в системах распознавания речи // Science Time. 2014. № 11. С. 483–492.
4. Мырадов М.Т., Назарова С.Н. Современные технологии распознавания речи // Наука и мировоззрение. 2024. № 3. С. 45–52.
5. Рахманов М. А. Сравнительный обзор современных моделей распознавания речи // Международная научно-практическая конференция «Smart Technologies and Innovative Approaches in Research: Building a Digital Future»: сборник материалов. 2025. 22 февр. С. 572–577. – DOI: 10.5281/zenodo.14918743.
6. Graham C., Roll N. Evaluating OpenAI's Whisper ASR: Performance analysis across diverse accents and speaker traits // Journal of the Acoustical Society of America. 2024. Vol. 155, Iss. 2. P. 1123–1135. PMID: 38391582. – DOI: 10.1121/10.0024876.
7. Andreyev A. Quantization for OpenAI's Whisper Models: A Comparative Analysis // arXiv preprint arXiv:2503.09905 [cs.SD]. 2025. – DOI: 10.48550/arXiv.2503.09905.
8. Ковригина Е. А. Компьютер и распознавание речи // Вестник компьютерных и информационных технологий. 2021. № 4. С. 22–29.
9. Физика: акустика. Громкость и интенсивность звука: учеб. пособие / под ред. В.Н. Петрова. М.: Физматлит, 2019. 184 с.
10. EBU Tech Doc 3343-2011v2. Practical Guidelines for Production & Implementation of EBU R 128. Geneva: European Broadcasting Union, 2011. 48 p. – URL: <https://tech.ebu.ch/docs/tech/tech3343.pdf> (дата обращения: 26.04.2026).
11. OpenAI. Whisper: Robust Speech Recognition via Large-Scale Weak Supervision // arXiv preprint arXiv:2212.04356 [cs.LG]. 2022. – DOI: 10.48550/arXiv.2212.04356.
12. Бовбель Е. И., Паршин В. В. Нейронные сети в системах автоматического распознавания речи // Зарубежная радиоэлектроника. 1998. № 4. С. 49–65.

References

1. Sadykova A. A., Amirgaliev E. N. Izuchenie primeneniya avtomaticheskogo raspoznaniya rechi // Colloquium-journal. 2020. № 11 (63). P. 44–47. – DOI: 10.24411/2520-6990-2020-11728.
2. Voroshilova A. I. Iskusstvennyy intellekt v sotsiologicheskom issledovanii: opyt ispol'zovaniya dlya

обработki i analiza interv'yu // Sotsiodinamika. 2025. № 2. – Electronic resource. – URL: https://nbpublish.com/library_read_article.php?id=73330 (date of access: 26.04.2026).

3. Gapochkin A. V. Neyronnye seti v sistemakh raspoznavaniya rechi // Science Time. 2014. № 11. P. 483–492.

4. Myradov M. T., Nazarova S. N. Sovremennye tekhnologii raspoznavaniya rechi // Nauka i mirovozzrenie. 2024. № 3. P. 45–52.

5. Rakhmanov M. A. Sravnitel'nyy obzor sovremennykh modeley raspoznavaniya rechi // International Scientific and Practical Conference "Smart Technologies and Innovative Approaches in Research: Building a Digital Future": proceedings. 2025. 22 Feb. P. 572–577. – DOI: 10.5281/zenodo.14918743.

6. Graham C., Roll N. Evaluating OpenAI's Whisper ASR: Performance analysis across diverse accents and speaker traits // Journal of the Acoustical Society of America. 2024. Vol. 155, Iss. 2. P. 1123–1135. PMID: 38391582. – DOI: 10.1121/10.0024876.

7. Andreyev A. Quantization for OpenAI's Whisper Models: A Comparative Analysis // arXiv preprint arXiv:2503.09905 [cs.SD]. 2025. – DOI: 10.48550/arXiv.2503.09905.

8. Kovrigina E. A. Komp'yuter i raspoznavanie rechi // Vestnik komp'yuternykh i informatsionnykh tekhnologiy. 2021. № 4. P. 22–29.

9. Fizika: akustika. Gromkost' i intensivnost' zvuka: ucheb. posobie / pod red. V. N. Petrova. M.: Fizmatlit, 2019. 184 p.

10. EBU Tech Doc 3343-2011v2. Practical Guidelines for Production & Implementation of EBU R 128. Geneva: European Broadcasting Union, 2011. 48 p. – URL: <https://tech.ebu.ch/docs/tech/tech3343.pdf> (date of access: 26.04.2026).

11. OpenAI. Whisper: Robust Speech Recognition via Large-Scale Weak Supervision // arXiv preprint arXiv:2212.04356 [cs.LG]. 2022. – DOI: 10.48550/arXiv.2212.04356.

12. Bovbel' E. I., Parshin V. V. Neyronnye seti v sistemakh avtomaticheskogo raspoznavaniya rechi // Zarubezhnaya radioelektronika. 1998. № 4. P. 49–65.

АРАКЧАА Айдаш Мергенович, сотрудник, Академия ФСО России, г. Орёл. E-mail: arakchaaaydash003@gmail.com

ДЕМИН Дмитрий Михайлович, сотрудник, Академия ФСО России, г. Орёл. E-mail: d1meplg@yandex.ru

ЖУСОВ Дмитрий Леонидович, кандидат технических наук, доцент, сотрудник, Академия ФСО России, г. Орёл. E-mail: d.zhusov@mail.ru

МАКЕЕВ Сергей Михайлович, кандидат технических наук, доцент кафедры компьютерной безопасности и прикладной алгебры, Челябинский государственный университет, 454001, г. Челябинск, ул. Братьев Кашириных, 129.; доцент кафедры защиты информации, Южно-Уральский государственный университет (национальный исследовательский университет), 454080, г. Челябинск, пр. Ленина, 76. E-mail: maksm57@yandex.ru

СОКОЛОВ Александр Николаевич, кандидат технических наук, доцент, заведующий кафедрой защиты информации, Южно-Уральский государственный университет (национальный исследовательский университет), 454080, г. Челябинск, пр. Ленина, д. 76. E-mail: sokolovan@susu.ru

АРАКЧАА Aidash Mergenovich, Staff Member, Academy of the Federal Guard Service of Russia, Orel. E-mail: arakchaaaydash003@gmail.com

DEMINS Dmitry Mikhailovich, Staff Member, Academy of the Federal Guard Service of Russia, Orel. E-mail: d1meplg@yandex.ru

ZHUSOV Dmitry Leonidovich, PhD in Technical Sciences, Associate Professor, Staff Member, Academy of the Federal Guard Service of Russia, Orel. E-mail: d.zhusov@mail.ru

MAKEEV Sergey Mikhailovich, PhD in Technical Sciences, Associate Professor at the Department of Computer Security and Applied Algebra, Chelyabinsk State University, 129 Bratiev Kashirinykh St., Chelyabinsk 454001, Russia; Associate Professor at the Department of Information Security, South Ural State University (National Research University), 76 Lenin Ave., Chelyabinsk 454080, Russia. E-mail: maksm57@yandex.ru

SOKOLOV Alexander Nikolaevich, PhD in Technical Sciences, Associate Professor, Head of the Department of Information Security, South Ural State University (National Research University), 76 Lenin Ave., Chelyabinsk 454080, Russia. E-mail: sokolovan@susu.ru