

Ученые СПб ФИЦ РАН разработали цифровой инструмент для повышения качества работы систем синтеза чеченской речи

14.07.2026



Специалисты Санкт-Петербургского Федерального исследовательского центра РАН (СПб ФИЦ РАН) и Академии наук Чеченской Республики создали программный модуль, который позволяет с высокой точностью проводить распознавание омографов (пишутся одинаково, произносятся по-разному) в текстах на чеченском языке по аудиозаписям без использования нейросетей. Разработка встроена в систему синтеза чеченской речи. Результаты исследования опубликованы в научном журнале [Big Data and Cognitive Computing](#) (Q1).

Цифровые системы синтеза речи анализируют входной текст, разбивают его на составляющие части, такие как слова, слоги, ударения и паузы, а затем используют алгоритмы и базы данных для формирования звуковых волн, имитирующих человеческую речь (в первую очередь используются большие нейросетевые модели). В результате получается аудиозапись, которая звучит как голос человека, но полностью создана искусственно. Современные системы синтеза

речи делают взаимодействие человека с компьютером более естественным и удобным, и их применение продолжается расширяться с развитием технологий. Так, они широко используются в различных областях: в голосовых помощниках и интеллектуальных ассистентах, системах навигации, чат-ботах и телефонных голосовых меню для взаимодействия с пользователями цифровых систем.

Одним из направлений для исследований и разработок в сфере систем синтеза речи является создание подходов и алгоритмов для языков с острой нехваткой аудио- и текстовых данных, поскольку применение больших нейросетевых моделей в таком случае ограничено. Как правило, это касается языков малых народов, число носителей которых сравнительно невелико, а от количества и качества этих данных зависит точность синтеза речи.

«Мы разработали подход для распознавания омографов в процессе синтеза чеченской речи без использования больших нейросетевых моделей, которые неприменимы к этому языку из-за нехватки текстовых и аудиоданных на чеченском языке. Созданные нами алгоритмы продемонстрировали высокую точность распознавания – около 80%. Разработанный на их основе программный модуль интегрирован в систему синтеза чеченской речи, он сделал ее более точной и естественной. В перспективе наши разработки могут быть адаптированы и использоваться для распознавания омографов других языков малых народов России», – рассказывает соискатель лаборатории речевых и многомодальных интерфейсов СПб ФИЦ РАН, научный сотрудник лаборатории прикладной математики федерального государственного бюджетного учреждения науки «Комплексный научно-исследовательский институт им. Х.И. Ибрагимова Российской академии наук» **Элиса Израилова**.

Омографы – это слова, которые пишутся одинаково, но звучат по-разному (например, ударение ставится по-разному) и имеют разные значения. Для разработки алгоритмов распознавания команда ученых собрала базу данных чеченских текстов, которая содержит более 15 тысяч предложений, отобранных вручную из 5 миллионов аннотированных слов и отражающих естественную частоту многозначности омографов по грамматическим категориям в различных контекстах. База данных была размечена лингвистами Академии наук Чеченской Республики в соответствии со значениями слов, долготой гласных, наличием дифтонгов и редукции.

Ученые разработали три разных алгоритма (WEN, WAW и WATCH) для разрешения неоднозначности в синтезе речи: например, как правильно цифровой системе ставить ударение в зависимости от значения в конкретного омографа.

Эксперименты по распознаванию омографов показали, что один из предложенных алгоритмов – WATCH – продемонстрировал точность 85%. Сравнительный анализ с существующими, находящимися в свободном доступе, алгоритмами для синтеза речи языков с ограниченным набором данных (вьетнамский, урду, ассамский, кашмирский, хауса) показал, что разработка команды ученых из Санкт-Петербурга и Чеченской республики демонстрирует второй по точности в мире результат.

Разработанные база данных, алгоритмы и программа распознавания омографии были созданы в рамках Года единства народов России, что подчеркивает вклад в сохранение и развитие языкового многообразия народов страны.

Материал: пресс-служба СПб ФИЦ РАН