

Объединенный институт проблем информатики
Национальной академии наук Беларуси

XVI Международная конференция

**РАЗВИТИЕ ИНФОРМАТИЗАЦИИ
И ГОСУДАРСТВЕННОЙ СИСТЕМЫ
НАУЧНО-ТЕХНИЧЕСКОЙ ИНФОРМАЦИИ
РИНТИ-2017**

16 ноября 2017 года, Минск

Доклады

Минск
ОИПИ НАН Беларуси
2017

Развитие информатизации и государственной системы научно-технической информации (РИНТИ-2017) : доклады XVI Международной конференции, Минск, 16 ноября 2017 г. – Минск : ОИПИ НАН Беларуси, 2017. – 414 с. – ISBN 978-985-6744-97-9.

Представлены доклады XVI Международной конференции «Развитие информатизации и государственной системы научно-технической информации» (РИНТИ-2017), Минск, 16 ноября 2017 г., в которых рассмотрены перспективные направления развития системы научно-технической информации в Беларуси, концепция формирования электронного государства в республике, основные направления и технологии цифровой трансформации в образовании, вопросы формирования информационного общества и цифровой трансформации системы государственного управления, информационно-правовое обеспечение нормотворческой деятельности, научно-методическое обеспечение развития информатизации и государственной системы научно-технической информации НАН Беларуси в 2017 г. Рассмотрены вопросы проектирования и внедрения автоматизированных систем научно-технической информации, информационного обеспечения научной, научно-технической и инновационной деятельности в организациях, министерствах и различных отраслях экономики, в том числе создания корпоративных автоматизированных библиотечно-информационных систем и технологий, а также психологические аспекты в информатизации.

Материалы конференции будут полезны специалистам в области информационно-коммуникационных технологий, занимающимся разработкой и внедрением автоматизированных информационных систем управления, систем научно-технической информации, автоматизированных библиотечно-информационных систем и технологий, а также развитием информационной инфраструктуры Беларуси и других стран, реализацией проектов государственных и отраслевых программ в сфере информатизации.

Одобрены программным комитетом и печатаются по решению редакционной коллегии Объединенного института проблем информатики Национальной академии наук Беларуси в виде, представленном авторами.

Научные редакторы:

член-корреспондент А.В. Тузиков;
кандидат технических наук, доцент Р.Б. Григянец;
кандидат технических наук, доцент В.Н. Венгеров

Солодков А.Т. Библиотека будущего. Некоторые прогнозы футуролога	234
Молчан Ж.М. Оптимизация технологии обслуживания в Центральной научной библиотеке Национальной академии наук Беларуси	239
Григянец Р.Б., Молчан Ж.М. Развитие системы автоматизации Центральной научной библиотеки Национальной академии наук Беларуси	244
Аксютю Е.В., Бабарико-Омельченко В.Б., Шакура Н.С. Информационно-библиотечное обслуживание в Белорусской сельскохозяйственной библиотеке на современном этапе	248
Шпаковский Ю.Ф. Разработка автоматизированных средств контроля качества учебных материалов	252
Гараничева С.Л. Особенности применения ИКТ в медицинском вузе в условиях создания электронного образования	258
Гарновская И.И. Исторический анализ развития информационных технологий в контексте медицинского образования	263
Буравкин А.Г., Липницкий С.Ф., Степура Л.В. Модель представления знаний в системе аналитической обработки научно-технической информации на основе коммуникативных фрагментов	269
Мыхлик В.В. Электронный каталог как информационно-поисковый ресурс библиотеки	275
Астапович Л.Л., Лаужель Г.О. Разработка политематического рубрикатора для подсистемы избирательного распространения информации в ЦНБ НАН Беларуси	280
Астапович Л.Л., Старовойтова О.И. Редактирование словарей электронного каталога	285
Кондратенко С.А., Григянец Р.Б., Мисякова Г.Т. Информатизация научно-технической и инновационной деятельности в сфере продовольственной безопасности Республики Беларусь	290
Борисевич Н.Я. Дистанционное консультирование как средство повышения радиоэкологической грамотности школьников	297
Казлоўская Н.Д., Шыбко М.В., Пратасеня А.І., Гецэвіч Ю.С. Распрацоўка алгарытмаў і сістэм дыктаранезалежнага распазнавання беларускага маўлення	299
Марчык М.У., Станіславенка Г.Р., Лысы С.І., Гецэвіч Ю.С. Вычытка і генерацыя тэкстаў вялікага памеру на беларускай мове	305
Коваленко Н.С., Венгеров В.Н. Организация выполнения синхронных процессов при ограниченном числе каналов обмена	311

РАСПРАЦОЎКА АЛГАРЫТМАЎ І СІСТЭМ ДЫКТАРАНЕЗАЛЕЖНАГА РАСПАЗНАВАННЯ БЕЛАРУСКАГА МАЎЛЕННЯ

Н.Д. Казлоўская, М.В. Шыбко, А.І. Пратасеня, Ю.С. Гецэвіч
Аб'яднаны інстытут праблем інфарматыкі НАН Беларусі, Мінск

Прадстаўлены падыход да стварэння сістэмы дыктаранезалежнага распазнавання беларускага маўлення, якая заснавана на CMU Sphinx. Прадстаўлены алгарытмы і працэсары перадапрацоўкі аўдыякорпуса беларускага маўлення і складання фанетычнага слоўніка для далейшага іх выкарыстання ў трэніроўцы акустычнай і моўнай мадэляў распазнавання маўлення. Паказаны шляхі далейшага папаўнення аўдыякорпуса і новых ітэрацый трэніроўкі мадэляў распазнавання беларускага маўлення.

Уводзіны

У цяперашні час існуюць сістэмы сінтэзу і распазнавання маўлення для шматлікіх моў, такіх як англійская, нямецкая, французская, украінская, польская, руская і інш. Што датычыцца беларускай мовы, існуе шэраг сінтэзатараў маўлення, сярод якіх стацыянарны [1, 2], мабільны [3], інтэрнэт-сінтэзатар [3, 4]. У той жа час для вырашэння задачы распазнавання маўлення было зроблена некалькі незавершаных спробаў. Была распрацавана пытална-адказная сістэма з выкарыстаннем трэніроўкі для аднаго дыктара на дзясятку гукавых фраз для інструмента Hidden Markov Model Toolkit (НТК) [5]. Таксама ёсць праца, дзе аўтаматычнае распазнаванне беларускага маўлення было натрэніравана на корпусе, у які ўваходзіла каля 1400 сказаў-загадаў для мабільнага робата, начытаных галасамі розных дыктараў. У гэтай распрацоўцы таксама выкарыстоўваўся пакет НТК, а дакладнасць распазнавання маўлення вагалася ад 57 да 92 % у залежнасці ад велічыні абранай часткі корпуса трэніроўкі [6].

У гэтым дакладзе аўтары імкнуліся перагледзець падыходы да распрацоўкі сістэмы распазнавання беларускага маўлення. Ставіліся задачы павышэння якасці распазнавання маўлення з дапамогай збору больш прадстаўнічага корпуса для трэніроўкі праз выкарыстанне зручнай сістэмы трэніроўкі маўлення і ранейшых напрацовак па сінтэзе маўлення для беларускай мовы.

Сёння распрацоўшчыкам найбольш вядомыя такія сістэмы распазнавання маўлення, як CMU Sphinx, НТК, Julius, Kaldi. Для НТК распрацоўка распазнавання беларускага маўлення працягваецца праз пабудову больш зручнай сістэмы разметкі для экспертаў-лінгвістаў, таму ў гэтым артыкуле НТК разглядацца не будзе. Што датычыць распрацоўкі сістэмы распазнавання маўлення на Kaldi, яна прыпынілася, паколькі патрабуе вялікіх вылічальных магутнасцяў (больш за 200 ГБ RAM).

У гэтым артыкуле будзе разгледжана распрацоўка сістэмы распазнавання беларускага маўлення для сістэмы CMU Sphinx, бо сёння гэта добра дапрацаваная сістэма. Адносна яе падрабязна дакументаваны працэс паляпшэння існуючых і стварэння новых мадэляў для моў, якія яшчэ не падтрымліваюцца сістэмай [7]. У рэпазіторыі натрэніраваных мадэляў распазнавання маўлення CMU Sphinx ёсць даступныя для загрузкі і азнаямлення мадэлі для нямецкай, французскай, галандскай, амерыканскай англійскай, іспанскай, італьянскай, мандарын, хіндзі, казахскай і рускай моў. Найбольш дасканалай з'яўляецца мадэль для англійскай мовы, у той час як астатнія мадэлі прыдатныя для выкарыстання, але знаходзяцца ў стане дапрацоўкі. Таму пры стварэнні ўласнай мадэлі можна звяртацца да досведу іншых распрацоўшчыкаў.

CMU Sphinx мае добрыя алгарытмы па аўтаматычнай разметцы падрыхтаваных гукавых і тэкставых дадзеных (аўдыякорпусу), прычым яны не патрабуюць высокіх

вылічальных магутнасцяў (8–16 ГБ RAM). Як вынік праведзенай працы аўтары прыводзяць не толькі алгарытмы стварэння сістэмы распазнавання беларускага маўлення, але і два даступныя ўсім жадаючым для тэставання эксперыментальныя прыклады: мабільную праграму і інтэрнэт-сэрвіс.

1. Пабудова мадэляў для аўтаматычнага распазнавання беларускага маўлення у CMU Sphinx

Для атрымання аўтаматычнага распазнавання беларускага маўлення з выкарыстаннем пакета CMU Sphinx былі створаны алгарытм і працэсар апрацоўкі адвольнага слова на беларускай мове ў фанемны выгляд, спецыяльна размечаны аўдыякорпус, алгарытм і працэсар перадапрацоўкі размечанага аўдыякорпуса. Далей гэтыя дадзеныя і праграмныя сродкі выкарыстоўваліся пакетам трэніроўкі CMU Sphinx для атрымання мадэлі распазнавання беларускага маўлення.

1.1. *Фанетычны працэсар апрацоўкі слова беларускай мовы ў фанемны выгляд*

Для трэніроўкі сістэмы аўтаматычнага распазнавання маўлення CMU Sphinx трэба падрыхтаваць слоўнік ці алгарытм, які ў будучым зможа выдаваць адвольнаму слову фанемны запіс. У выпадку адсутнасці слоўніка з пазнакамі фанем для слоў абранай мовы пры стварэнні такога слоўніка можа быць выкарыстаны фанетычны працэсар сінтэзатара маўлення [8].

Сістэмы сінтэзу маўлення па тэксце, пад якімі разумеюць сістэмы, здольныя пераўтвараць друкаваны тэкст у адпаведны маўленчы сігнал, зазвычай складаюцца з шэрагу працэсараў, якія паслядоўна апрацоўваюць электронны тэкст. Сярод іх можна выдзеліць тэкставы, прасадны, фанетычны, акустычны. Пад фанетычным працэсарам разумеюць сукупнасць алгарытмаў, пры дапамозе якіх адбываецца канвертаванне нармалізаванага арфаграфічнага тэксту ў фанетычны запіс. Асноўнымі двума алгарытмамі фанетычнага працэсара з'яўляюцца алгарытмы пераўтварэння «графема – фанема» (або «літара – фанема») і «фанема – алафон».

Выкарыстанне фанетычнага працэсара анлайн-сінтэзатара беларускага маўлення дазваляе канвертаваць нармалізаваны арфаграфічны тэкст у алафонны.

Каб пабудаваць машыначытаемы фанетычны слоўнік беларускай мовы, патрэбны прадстаўнічы электронны слоўнік беларускай мовы. Ён быў узяты як копія электроннага слоўніка з пазначанымі націскамі і выдаленымі дублікатамі слоў (больш за 1 000 000 слоў) з анлайн-сінтэзатара беларускага маўлення [8].

1.2. *Алгарытм і працэсар апрацоўкі адвольнага слова беларускай мовы ў фанемны выгляд*

Для стварэння мадэлі «графема – фанема» беларускай мовы быў задзейнічаны фанетычны працэсар, апісаны вышэй, і прылада g2p-seq2seq, заснаваная на рэкурэнтных нейронных сетках [9].

Абагулены алгарытм стварэння мадэлі «графема – фанема» для адвольнага слова складаецца з наступных крокаў:

1. Загрузка электроннага слоўніка S з прастаўленымі націскамі і выдаленымі дублікатамі.

2. Транскрыбаванне слоў w_i слоўніка S у фанемны запісы ph_i фанетычным працэсарам F сінтэзатара маўлення па тэксце.

3. Выдаленне для кожнага ph_i двух правых індэксаў кожнага алафона, замена сімвала мяккасці ў алафоне з $'$ на J пры наяўнасці.

4. Стварэнне фанетычнага слоўніка S_{ph} як трэніровачнага мноства з апрацаваных фанетычным працэсарам слоў выгляду (w_i, ph_i) для прылады g2p-seq2seq з параметрамі 4 layers, 512 units.

5. Атрыманне мадэлі F_{ph} «графема – фанема» з дапамогай прылады g2p-seq2seq.

Атрыманая мадэль «графема – фанема» паказала дакладнасць 99,91 % на тэставым мностве, якое было адлучана ад мноства, падрыхтаванага для трэніроўкі. Дадзеная дакладнасць была адзначана як прымальная, і атрыманая мадэль «графема – фанема» была выкарыстана ў далейшым для падрыхтоўкі файлаў трэніроўкі сістэмы аўтаматычнага распазнавання маўлення CMU Sphinx.

1.3 Стварэнне размечанага аўдыякорпуса

Праблема стварэння сістэмы распазнавання дыктаранезалежнага беларускага маўлення на CMU Sphinx патрабавала распрацоўкі спецыяльна размечанага аўдыякорпуса, які мае прадстаўнічы аб'ём і грунтуецца на тэкстах, што пакрываюць усе алафоны беларускай мовы. Будзем называць такі аўдыякорпус фанетыка-акустычнай базай дадзеных. Галоўным элементам такой базы будзе сукупнасць алічбаванай гукавой хвалі фрагменту прамоўленага чалавекам тэксту на беларускай мове і дадатковай інфармацыі пра гэтую хвалю. Дадатковая інфармацыя змяшчае ў сабе звесткі пра дыктара, які прамовіў тэкст, і пра пакрыццё ўсіх алафонаў беларускай мовы.

Стварэнне фанетыка-акустычных баз дадзеных уключае ў сябе наступныя этапы:

1) збор акустычных і тэкставых дадзеных – сюды ўваходзяць пошук запісаў і транскрыпцый да іх (аўдыякнігі, запісы праграм тэлебачання/радыё, агучка фільмаў і г.д.), падрыхтоўка тэкстаў для начыткі, начытка падрыхтаваных тэкстаў рознымі дыктарамі;

2) апрацоўка сабраных дадзеных.

3 пункту гледжання тэкставага матэрыялу пабудаваная база дадзеных складаецца з двух частак. Першая частка ўтрымлівае ў сабе тэксты мастацкага стылю мовы. Пры абранні гэтай часткі тэкстаў аўтары зыходзілі з гіпотэзы пра тое, што любы досыць вялікі ўрывац літаратурнага тэксту з'яўляецца фанетычна рэпрэзентатыўным. Мастацкія тэксты былі начытаны прафесійнымі дыктарамі або актарамі (напрыклад, «Каласы пад сярпом тваім» У. Караткевіча, «Бездакорная пакаёўка» А. Крысці, «Вераснёвыя ночы» К. Чорнага). У іх склад уваходзяць сказы розных інтанацыйных тыпаў (апавядальныя, пытальныя і клічныя), а таксама сказы, якія змяшчаюць простую мову. Агульная колькасць начытаных тэкстаў гэтай групы – 16 гадзін.

Другая частка змяшчае тэксты публіцыстычнага стылю – артыкулы з часопісаў і газет на беларускай мове. Агульная колькасць начытаных тэкстаў такой групы – 2 гадзіны 30 хвілін.

Абедзве часткі тэкставага матэрыялу з'яўляюцца фанетычна збалансаванымі, яны былі правераны на наяўнасць усіх алафонаў беларускай мовы з дапамогай сэрвісу «Падлік частотнасці алафонаў» інтэрнэт-пляцоўкі www.corpus.by [10].

Для стварэння маўленчых матэрыялаў на базе тэкставых запісаў 23 дыктары (12 мужчын, 11 жанчын). Толькі пяцёра з іх былі прафесійнымі дыктарамі або актарамі (3 мужчыны і 2 жанчыны), астатнія нават не мелі вопыту ў мастацтве чытання. Дыктары не выбіраліся з пункту гледжання разнастайнасці дыялектаў. Усе тэксты прамаўляліся з захаваннем інтанацыйных асаблівасцей. Усе тэксты былі начытаны ў асобным кабінце, спецыяльна абсталяваным для запісу голасу з адсутнасцю знешніх шумоў, але з прысутнасцю натуральнага шумавога фону.

Сабраныя акустычныя і тэкставыя дадзеныя перад тым, як былі адпраўлены на трэніроўку распазнавання маўлення, праходзілі спецыяльную апрацоўку. Акустычныя дадзеныя былі падзеленыя на часткі A_i працягласцю ад 5 да 30 с, для кожнага

атрыманага гукавога файла скапіраваны адпаведныя тэкставыя фрагменты T_i (транскрыпты). Далей гукавыя часткі былі пераведзеныя ў гукавы wav-фармат.

У выніку праведзенай працы маем фанетыка-акустычную базу $SC = \langle A_i, T_i \rangle$ агульнай працягласцю каля васьмі гадзін агучнага рознымі дыктарамі і спецыяльна апрацаванага тэксту. Прычым зразумелыя шляхі яе далейшага павелічэння.

1.4. Алгарытмы апрацоўкі размечанага аўдыякорпусу для трэніроўкі мадэляў сістэмы аўтаматычнага распазнавання маўлення CMU Sphinx

Атрыманая фанетыка-акустычная база выкарыстоўваецца спачатку для пабудовы моўнай мадэлі для сістэмы аўтаматычнага распазнавання беларускага маўлення. Наступны алгарытм на выніку выдае неабходную моўную мадэль:

1. Валідацыя дадзеных. Прыводзяцца ўсе тэксты T_i да ніжняга рэгістра, выдаляюцца знакі прыпынку, шукаюцца абдрукоўкі (лацінскія літары замест кірылічных) і, калі магчыма, выпраўляюцца ў аўтаматычным рэжыме, калі не – вяртаюцца на ручную апрацоўку.

2. Вылучэнне з тэкстаў асобных слоў і сказаў, захаванне іх у адпаведныя спісы – S_w, S_s ;

3. Стварэнне слоўніка «графема – фанема» S_{wph} для слоў S_w пры дапамозе падрыхтаванай раней F_{ph} мадэлі «графема – фанема» і прылады g2p-seq2seq.

4. Пабудова моўнай мадэлі LM праз перадачу спіса сказаў S_s у прыладу CMU (Cambridge Statistical Language Modeling Toolkit) [11].

Атрыманая моўная мадэль LM будзе задзейнічана у аўтаматычным тэставанні натрэніраванай мадэлі і распазнаванні беларускага спантаннага маўлення.

Пасля размечаны аўдыякорпус SC і слоўнік S_{wph} выкарыстоўваецца для трэніроўкі акустычнай мадэлі сістэмы аўтаматычнага распазнавання маўлення CMU Sphinx. Абазначым крокі алгарытму:

1. Аўдыяфайлы A_i правяраюцца на адпаведнасць параметрам, патрэбным сістэме аўтаматычнага распазнавання маўлення (гэта частата 16000 КГц, шырыня 16 біт, 1 канал). Калі запіс не адпавядае гэтым параметрам, то аўтаматычна канвертуем яго ў неабходны. Атрымліваем мноства S_a .

2. У спецыяльны файл f_x заносім адпаведна шляхі да файлаў з аўдыязапісамі маўлення S_a і тэксты транскрыпцый маўлення S_s .

3. Запускаецца прылада sphinxtrain па файле f_x , якая накладвае тэкставыя дадзеныя на гукавыя дадзеныя аўдыякорпуса, ствараецца акустычная мадэль AM .

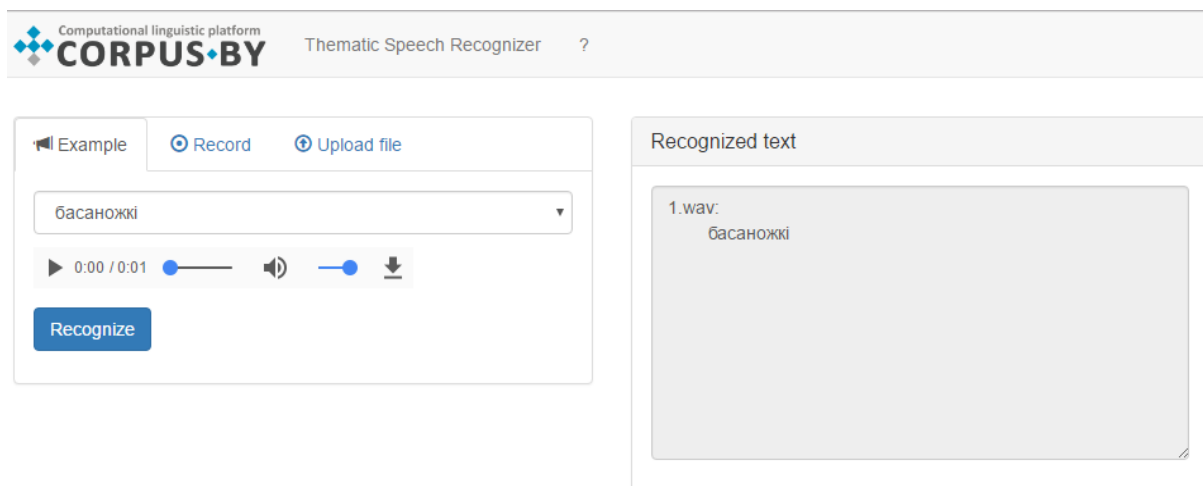
Далей атрыманая моўная мадэль LM , акустычная мадэль AM былі выкарыстаныя пакетам CMU Sphinx для пабудовы сістэмы аўтаматычнага распазнавання беларускага маўлення.

2. Вынікі трэніроўкі

Такім чынам, для трэніроўкі сістэмы аўтаматычнага распазнавання спатрэбілася восем гадзін запісаў маўлення, у якіх было каля 15 тысяч слоў, 85 фанем. Час трэніроўкі мадэлі на персанальным камп'ютары ў 16 ГБ RAM заняў сем-восем гадзін. Атрыманая сістэма распазнавання маўлення была пратэставаная на 10 % дадзеных з аўдыякорпусу для трэніроўкі сістэмы. Падлікі памылак былі наступнымі: для распазнавання сказаў – 18,6 % (49/264), асобных слоў – 5,8 % (170/2917). Далей былі распрацаваны два эксперыментальныя дэма-абразцы на аснове натрэніраванай мадэлі распазнавання беларускага маўлення для мабільнай і інтэрнэт-платформаў.

Дэма-версія інтэрнэт-сэрвіса «Тэматычнае распазнаванне маўлення» даступная анлайн: <http://ssrlab.grid.by/ThematicSpeechRecognizer/> [12]. Сэрвіс «Тэматычнае распазнаванне маўлення» (малюнак 1) дазваляе карыстальніку пераўтвараць маўленне ў

электронны тэкст анлайн. На ўваход сэрвісу можа падавацца фанаграма маўленчых слоў тэматычных даменаў памерам не больш за 20 МБ, на выхадзе сэрвіс дае распазнаны электронны тэкст фанаграмы. Фанаграма можа быць загружана на сэрвіс з цвёрдага дыску камп'ютара ў фармаце .wav ці запісана праз магчымасці аўдыязапісу сэрвісу. На дадзены момант сэрвіс распознае беларускамоўнае маўленне наступных тэматычных даменаў: вопратка, гарады, лікі, спантаннае маўленне. Спіс даменаў будзе папаўняцца.



Графічны інтэрфейс сэрвісу «Тэматычнае распазнаванне маўлення»

Дэма-версіяй мабільнай праграмы з натрэніраванай мадэллю распазнавання маўлення знаходзіцца на пляцоўцы Google Play Market [13]. Адразу пасля ўключэння праграма запуская модуль распазнавання і апрацоўвае ўваходнае маўленне. Атрыманыя вынікі выводзяцца ў рэальным часе ў выглядзе тэксту распазнаных слоў. Пры спыненні маўлення праграма спыняе распазнаванне і чакае аднаўлення гукавых сігналаў. Для распазнавання можна выбраць адзін з трох даменаў: «адзенне», «лічбы» ці «іншае».

Заклучэнне

Такім чынам, у гэтым артыкуле прадстаўлена распрацоўка сістэмы распазнавання беларускага маўлення для CMU Sphinx. Для яе распрацаваны: алгарытм і працэсар апрацоўкі адвольнага слова па-беларуску ў фанемны выгляд з выкарыстаннем даступнай сістэмы сінтэзу маўлення; спецыяльна размечаны аўдыякорпус на 18,5 гадзін з магчымасцю зручнага папаўнення; алгарытм і працэсар перадапрацоўкі размечанага аўдыякорпуса. Вынікі тэставання натрэніраванай мадэлі для беларускай мовы паказалі дакладнасць пры распазнанні сказаў 67,3 %, а пры распазнанні асобных слоў – 14,7 %. Аўтары прапануюць азнаёміцца з дэма-версіямі праграм з натрэніраванай мадэллю распазнавання беларускага маўлення для мабільнай і інтэрнэт-платформаў.

Планы па дапрацоўцы сістэмы распазнавання беларускага маўлення наступныя: па-першае, стварыць сістэму аўтаматычнагабору акустычных і тэкставых дадзеных для атрымання вялікага аўдыякорпусу беларускага маўлення; па-другое, пасля таго як у акустычнай базе будзе больш за 50 гадзін маўлення, пачаць працу над Speech API для беларускай мовы; па-трэцяе, распрацаваць прыкладныя прыстасаванні для выкарыстання натрэніраванай мадэлі распазнавання беларускага маўлення ў кліентскіх, кліент-серверных і мабільных праграмах, а таксама ва ўбудаваных сістэмах (інтэрнэт рэчаў).

Спіс літаратуры

1. Гецевич, Ю.С. Система синтеза белорусской речи по тексту / Ю.С. Гецевич, Б.М. Лобанов // Речевые технологии. – 2010. – № 1. – С. 91–100.
2. Rapid prototyping of speech technologies for Belarusian as a low-resource language: Master Thesis / A. Vasileuskaya; supervisors: B. Mobius, T. Avgustinova ; Saarland University, department of computational linguistics . – Saarland, 2014. – 94 p.
3. Гецэвіч, Ю.С. Распрацоўка сінтэзатара беларускага і рускага маўленняў па тэксце для мабільных і інтэрнэт-платформаў / Ю.С. Гецэвіч, Д.А. Пакладок, Д.В. Брэж // Развитие информатизации и государственной системы научно-технической информации (РИНТИ-2012) : доклады XI Междунар. конф. (Минск, 15 нояб. 2012 г.). – Минск : ОИПИ НАН Беларуси, 2012. – С. 254–259.
4. Гецэвіч, Ю.С. Сінтэзатар маўлення па тэксце для шырокага кола карыстальнікаў Інтэрнэта / Ю.С. Гецэвіч, Д.А. Пакладок, Д.В. Брэж // Развитие информатизации и государственной системы научно-технической информации (РИНТИ-2013) : доклады XII Междунар. конф. (Минск, 20 нояб. 2013 г.). – Минск : ОИПИ НАН Беларуси, 2013. – С. 277–281.
5. Естественнo-языковые интерфейсы интеллектуальных вопросно-ответных систем / В.М. Вяльцев [и др.] // Доклады БГУИР. – 2011. – № 8 (62). – С. 80–86.
6. Nikalaenka, K. Training algorithm for speaker-independent voice recognition systems using HTK / K. Nikalaenka, Y. Hetsevich // PRIP ‘2016: Pattern recognition and information processing : proceedings of the 13th International Conference, (3-5 October 2016, Minsk, Belarus). – Minsk : Publishing Center of BSU, 2016. – P. 126-129.
7. CMU Sphinx [Electronic resource]. – Mode of access : <http://cmusphinx.sourceforge.net>. – Date of access : 20.03.2017.
8. Text-to-Speech Synthesizer [Electronic resource]. – Mode of access : <http://corpus.by/TextToSpeechSynthesizer>. – Date of access : 20.03.2017.
9. Sequence-to-Sequence G2P toolkit [Electronic resource]. – Mode of access : <https://github.com/cmusphinx/g2p-seq2seq>. – Date of access : 20.03.2017.
10. Character Frequency Counter [Electronic resource]. – Mode of access : <http://www.corpus.by/CharacterFrequencyCounter/>. – Date of access : 28.03.2017.
11. The CMU Statistical Language Modeling (SLM) Toolkit [Electronic resource]. – Mode of access : http://www.speech.cs.cmu.edu/SLM_info.html – Date of access : 28.03.2017.
12. Thematic Speech Recognizer [Electronic resource]. – Mode of access : <http://ssrlab.grid.by/ThematicSpeechRecognizer/>. – Date of access : 28.03.2017.
13. Thematic Speech Recognition [Electronic resource]. – Mode of access : <https://play.google.com/store/apps/details?id=by.ssrlab.recognition.speech.thematic.thematicspeechrecognition>. – Date of access : 28.03.2017.