

Павел Малецкий

АГТ Науково-Технічний Університет/AGH University of Krakow

ORCID: 0000-0001-5050-2111

Магдалена Піотровска

АГТ Науково-Технічний Університет/AGH University of Krakow

ORCID: 0000-0002-3974-9559

Аналіза і класифікація русинської бесіди языковим модельом штучної інтелегенції OpenIA Whisper

Abstrakt

Analiza i klasyfikacja języka rusińskiego przy użyciu modelu sztucznej sieci neuronowej ASR OpenIA Whisper

Artykuł przedstawia analizę lingwistyczną języka rusińskiego, koncentrując się na jego złożonych i zmieniających się aspektach, takich jak wymowa oraz różnice indywidualne, regionalne i historyczne. Do przeprowadzenia badania wykorzystano sztuczną sieć neuronową opartą na modelu OpenAI Whisper. Model ten, choć szkolony na danych z większości państwowych języków urzędowych, nie był bezpośrednio trenowany na bazach próbek języka rusińskiego ze względu na jego lokalny i mniejszościowy/etniczny charakter. Stąd próbki mowy tego języka klasyfikowane były przy użyciu najbardziej zbliżonych dostępnych etykiet, co pozwoliło na wyznaczenie podobieństwa języka rusińskiego do innych słowiańskich języków. Badanie objęło użytkowników zróżnicowanych pod względem płci, wieku i lokalizacji (Polska, Ukraina, Słowacja, Serbia), wykazując znaczące podobieństwa do języków dominujących w tych krajach oraz zależności między wyznaczonym podobieństwem językowym a wiekiem mówców.

Słowa kluczowe: język rusiński, fonetyka, klasyfikacja, asymilacja, IA, sztuczne sieci neuronowe, ASR

Abstract

Analysis and Classification of the Rusyn Language Using the OpenAI Whisper ASR Model

The paper presents a linguistic analysis of the Rusyn language, focusing on its complex and dynamic aspects, such as pronunciation and individual, regional, and historical variations. The study employed a neural network based on the OpenAI Whisper automatic speech recognition (ASR) model. While trained on data from a majority of official state languages, the model lacked direct training on Rusyn language samples due to its localized and minority/ethnic nature. Consequently, speech samples were classified using the closest available language labels, enabling the identification of similarities between Rusyn and other Slavic languages. The study encompassed a diverse range of speakers across gender, age, and location (Poland, Ukraine, Slovakia, Serbia), revealing significant similarities to the dominant languages in these respective countries. Furthermore, the research highlights correlations between the identified linguistic similarities and the age of the speakers.

Keywords: Rusyn language, Phonetics, Classification, Assimilation, IA, ANN, ASR

Ключовы слова: русиньскій язык, фонетыка, класифікация, асиміляция, ШІ, штучны невронны сїти, ASR

1. Вступ

Русиньскій язык то східньославянскій язык хоснуваний через русиньску етнічну групу, головні в карпатском регіоні Серединовой Європы, котрый обнимат част України, Словачії, Польщы, Мадяр і Румунії. Має місцевы одміны, што оддзеркаляють зложены історичны і культуровы вплины поєднаных регіонів. Класифікация русиньского языка в шыршым контексті східньославянських і інчых славянських языків є зложеном і дискусийном темом з огляду на його унікальне положыня і історичны вплины. Русиньскій язык має спільны приметы зо східньославянськыма языками (як російскій, українскій і білорускій), а тіж выказує вплины західньославянських языків (польскій, словацькый) і полудньовославянських (сербскій, хорватскій). Тота вельоаспектова мішанина оддзеркалят його положыня на скривуваню тых языковых груп (Kushko 2007).

Історичні русиньскій язык был залежны од языковой політыкы пануючых держав ци цисарств, медже інчыма Австро-Мадярской Імперії ци Совітского Союзу. Языковой проблем є кісно повязаны з достоменністю, а дискусії на тему того ци русиньскій то окремый язык, ци диалект українского, тырвають іщы аж і днес. Прихильниками той другой теорії сут мало не выключні дослідники і політыкы, які ідентифікуют Руснів як Українців. В Україні Русины не сут офіційно узнаны як окрема етнічна група, што іщы веце комплікує охорону языка (Nikitin і ін. 2009). Русиньскій язык был скодифікованы в XX ст., што оддзеркаляло одроджыня народовой і языковой достоменности. Кодифікацийны діяня довели до выокремліня регіональных літературных стандартів опертых на місцевых приметах – в Словачії, Польщы, Україні і в Сербії (Plišková 2008).

Русиньскій єст переважні класифікуваний як східньославянскій язык з огляду на його історичне і языкове коріня. Його класифікацию комплікує єднак значуче вплины західньо- і полудньовославянських языків (Moser 2016).

Змінніст языка, што обнимат дияхронны, диястратычны, дияфазычны і диятопічны аспекти, то вызваны для обробкы натурального языка (NLP, ен. Natural Language Processing). Диалектычны ріжницї медже одмінами того самого языка вплиют на ефективніст аплікаций, таких як машынове тлумачыня ци розпознаваня бесіды (Zampieri і ін. 2020). Прото тіж росне заінтерсуваня досліджынями над обробком

споріднених мов, мовних відмін і діалектів, на що доказом є численні публікації і наукові події, такі як варшавські VarDial (Zampieri і ін. 2020).

В примірі мов з низькими ресурсами, як нп. русинської, придбати відповідні мовні корпуси є особливо важко (Scherrer, Rabus 2019). Часто мови розв'язати транскрибуючи бесіду, як в примірі корпусу ArchiMob для німецьких діалектів або хоснуючи тлумачення, як в примірі корпусу MADAR для арабських діалектів (Bouamor і ін. 2019). Ахім Rabus і Ів Scherrer описують виклики пов'язані з творінням ресурсів NLP для русинської мови і пропонують індукцію морфосинтактичного лексикону з хоснуванням ресурсів слов'янських мов (Rabus, Scherrer 2017). Досліджують вони морфосинтактичне етикетування (англ. tagging) для русинської мови, хоснуючи трансфер вчення з споріднених мов. Вказані вибрані одніє є доказом того, що тема є інтенсивно досліджувана і експлуатована. Існують спеціалізовані програми і аплікації до дослідження мовних примет, але хоснування моделю так барз зложеного і чутливого як в тій дописі є іновативним підходом.

Цілю праці є дослідити ступінь розпознавальності і класифікації русинської мови через модель автоматичного розпознавання бесіди Open IA Whisper і оцінити вплив домінуючих мов держав, де живуть Русини на результат класифікації. Обсяг праці обнімає аналіз звукових даних зібраних з русинськомовних радіо-аудіо з чотирьох держав: Польщі, України, Словаччини і Сербії, а тіж другі (додаткові) аналізи з огляду на вікові групи бесідників. В статті поміщено основні інформації на тему моделю OpenIA Whisper, в тій його архітектуру і примети, що чинят можливим вельомовну транскрипцію і дрібницьовий опис дослідничой методології, що обнімає зміст бази звукових даних, алгоритмів сегментації і аналізи звуку та параметрів хоснуваних в моделі.

Як результати представлено класифікацію русинської мови при увзгляднію держав, одкале походять бесідники і різниця між віковими групами. Результати омовлено в контексті впливу домінуючих мов на мову Русинів і проходять асиміляційних процесів. Представлено тіж квести пов'язані з обмеженнями моделю і специфіки мовної класифікації. Працю завершат підсумування, котре вказує на основні висновки і можливі напрями дальших досліджень в обшарі автоматичного розпознавання бесіди для меншомовних мов.

2. Мовна модель OpenIA Whisper

OpenIA Whisper то заавансуваный модель автоматичного розпознаваня бесіды ASR (анг. Automatic Speech Recognition) (Radford і ін. 2022). Модель і програмовый код сут одкрыто доступны на ліценції тыпу *open source*, што уможливят динамічне розвита аплікації і повязаны дослїджыня в обшыри заавансуваной переробкы бесіды. Сесый модель вытренувано на основі веце як 680 тысячы годин вельоязычных награнь з транскрипціям. Обшырна і ріжнорідна база даных дозволила осягнути велику одпорніст системы на зріжницюваны акценты, галас на фоні ці спецыялістычну термінологію, спомогла при тым транскрипцію в вельох языках.

Архітектура OpenIA Whisper єст приміром реалізації системы ASR зо схоснуваньом нейронной сіті тыпу *Transformer* в конфігурації енкодер-декодер (анг. encoder-decoder). Звуковы даны на вході моделю ділены сут на 30-секундовы часткы, пак переформлюваны на *log-Mel spectrogram* (є то дигітальна, графічна репрезентація звуку, яка бере до увагы приметы перцепції чловека), а дале перерабляны через енкодер. Декодер передвидує зміст, што одповідат звукам і вводит додатковы інформаційны токены, такы як ідентифікація языка бесідника і часовы значныкы з докладністю до рівня фраз бесіды.

Во вчыню моделю OpenIA Whisper схоснувано барз велику і зріжницювану базу даных, што ілюструє Рys. 1, на котрым вказано залежніст процента невластиві розпознаных слів од кількості даных до тренуваня. В базі была менше-веце третя част награнь в інчым як англійцкій языку. І хоц в базі єст о вельо менше тестовых награнь в інчых языках, то ефектывніст розпознаваня языка, при базі о велькості більше як сто годин, єст на рівні 80 проц. і высшым. До того зараховуют ся вшыткы значучы з перспектывы дослїджынь славянскы языкы, што тіж зазначено на Рys. 1.

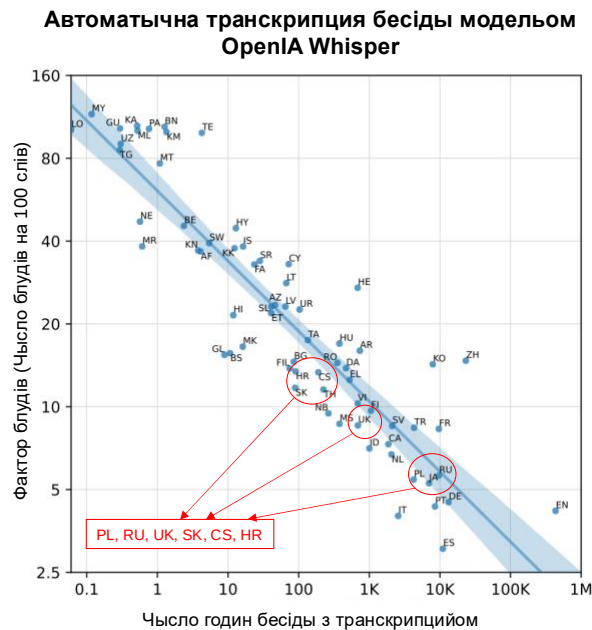


Рис. 1. Залежність ефективності моделю ASR, в залежності од велькості даних до тренування (Radford і ін. 2022). На выкресі зазначено выбраны славянськы языки (польскій, українській, словацкій і російській).

Модель OpenIA Whisper розпознає язык выповідичерез вельюетаповый процес, котрый єст зінтегрований з його архітектуром енкодер-декодер. Логарытмічны Мелы спектрограмы сут передаваны до енкодера, што выокремлят зложены звуковы приметы. Енкодер выокремлят взірці в звуковых даних на барз высокым рівни зложености, які характеризуют ріжны языки. Прогноза основана єст на высокоплощыновых приметах, што обнимяют фонетычны і фонологічны взірці особливы для каждого языка. В час вчыня модель был експонуваний на звуковы даны в барз вельох языках, ведно з відповідніма для них текстовима етикетами, в котрых сут токены до ідентифікації языка. В процесі декодування звуку декодер хоснує сесы вивчени токены, жебы окрислити язык на основі введеных даних/звуків. Завдякы так обшырному вчыню модель аналізує сигналы в комплексовый спосіб в ріжных языках і акцентах, што поправлят його здібніст до прецизийного розпознавання языка. OpenIA Whisper хоснує тіж здібности *zero-shot learning*, што означат, же на основі даних з вчыня може робити узагальніня, жебы розпознати і транскрибувати языки, на яких не был конкретні вчений. Осігат того хоснуючы вельюязычний характер даних до тренування і заавансуваны можливости выокремляня примет через енкодер.

Схоснування моделю OpenIA Whisper до класифікації русинського языка, хоч модель тот не был вчений на русинській базі, мож пояснити парома чынниками. Модель OpenIA Whisper был вытренуваний на великій, вельоязычній базі даных, в якій были ріжны славянські языки (таки як польскій, українскій, словацкій, сербскій, хорватскій ци російскій), што сут зближены до русинського языка під оглядом граматичной структури, фонетыки ци лексики. Завдяки тому модель годен розпознати деяки взірці і приметы характерны для группы славянських языків, што уможливят му зближену класифікацію русинського языка, хоч не было безпосереднього вчыня на його базі. В примірі русинського языка, котрый не єст релятивні масовый, схоснування моделю вытренованого на шыроком спектрі языків дозвелят на його приближену класифікацію на основі примет спільных з інчыма языками.

3. Методологія

а. База даных

До языковой аналізы схоснувано архівальны награня лемківской радийовой стациі Lem.fm. В базі звуковых даных є барз вельо зріжницюваных примірів русинського бесідуваного языка, што уможливило комплексове досліджыня з шырокой перспектывы. З каждой авдициі усунено фрагменты музыкы, джінглі і реклямовы споты, а з інтервю лишено голос лем єдной особы.

Так приготовлена база складала ся з 470 ріжных авдиций. Час тырвання цілою базы то даде 121,8 годин. Поділ походжыня бесідників з кількістю авдиций і часом тырвання досліджаных пробок представлено в Табелі 1.

Скорочыня	Держава	Кількіст авдиций	Час тырвання [годин]	Чысло лекторів
PL	Польша	209	45,4	веце як 100
UK	Україна	99	37,5	веце як 40
SK	Словачия	131	23,2	веце як 50
CS	Сербія	31	11,5	10

Таб. 1.

Додаткові метадані приписані до кожної аудіофайлу:

- Вікова група: дві категорії – бесідники що мають вік або менше як 70 років (лем для бесідників з Польщі).
- Ідентифікація бесідника: мена або псевдонім.
- Назва файлів згенерованих в час аналізу.
- Полный час награна (в секундах) каждой звуковой пробки.

6. Розпознавання мови бесідника за сховуванням моделлю OpenIA Whisper

До аналізу і категоризації даних опрацьовано скрипт в мові *Python*, котрий уможливає автоматичний вибір і переробку звукових файлів, мовний аналіз і реєстрацію результатів.

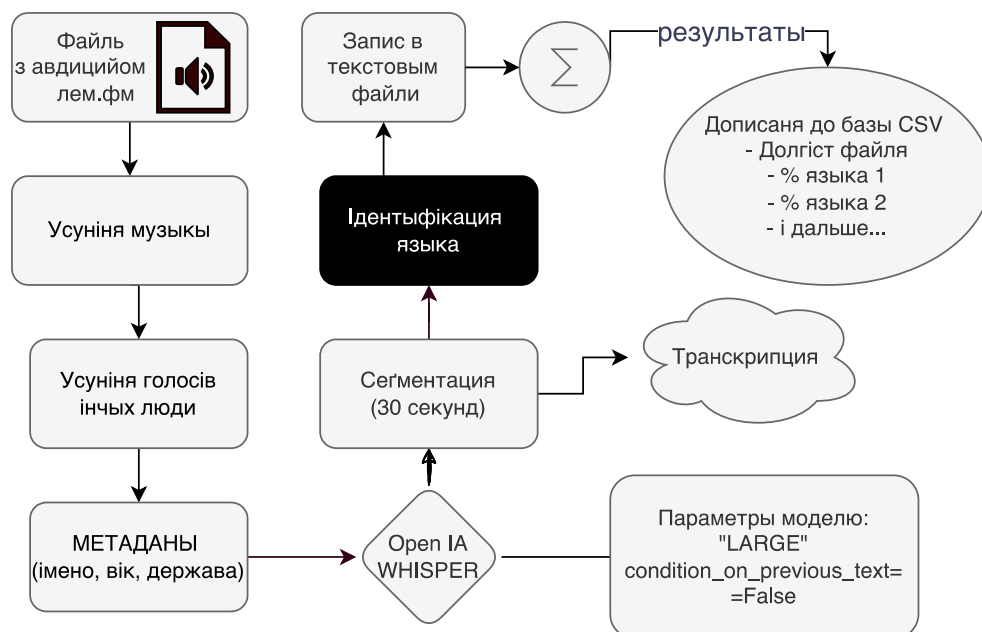


Рис. 2. Схема процесу ідентифікації мови радіоаудіофайлу.

Діяння скрипту схемово представлено на Рис. 2, а поєдні етапи його діяння є наступні:

- вибір каталогу: вибір каталогу з аудіо файлами при схованому графічному інтерфейсу (GUI),
- сегментація звуку: кожен аудіо файл є ділений на частки довготи 30 секунд,
- діяння моделю: до транскрипції бесіди і визначення мови сховану модель Whisper в варіанті «*large*». Архітектура моделю запроєктована є з 32 верств сітної широкоти 1280 нейронів і 20 голів (англ. heads), що разом дає 1,55 мільярда параметрів. Є то найкращий модель в родині Whisper, що перекладає на високу прецизію транскрипції і розпізнавання мови завдяки його багаторівній структурі. Модель робить транскрипцію бесіди для кожного сегменту аудіо, а на виході моделю є додаткові інформації, в тому визначений мову. Параметри моделю сконфігуровані є на вартості: *no_speech_threshold=0.8* і *condition_on_previous_text=False*, завдяки чому кожен наступний 30-секундний пробка аналізується є незалежний від попередньої. Широкот сітної дефініює число нейронів в кожній верстві моделю, значить розмір вектора репрезентації, що є перетворюється через модель на даній верстві, а число голів одностить до числа «голів уваги» (англ. attention heads) в кожній верстві механізму уваги моделю. Тот механізм дозволяє зосередити ся моделю на різних частях секвенції на вході в час переробки. Кожен голівка аналізує дані під іншим кутом, завдяки чому модель може краще виїмати зложені залежності в мовних секвенціях,
- реєстрування результатів: результати транскрипції і визначений мову для кожного сегменту є записувані за порядком до текстових файлів,
- експорт до CSV: результати є записувані в файли CSV, в котрих є первістні метадані і інформації о розпізнаних мовах і їх процентова участь в повному часі награния.

Скрипт додатково хоснує бібліотеку *librosa* до ладування звуку і *pydub* до сегментації даних.

4. Результати

Дані переаналізувано в тот спосіб, же в кожній аудіоці окреслено шор розпізнаних мов, зачинаючи від найчастіше розпізнаваного (#1), ведно з ідентифікатором держави бесідника. Пак порохувано процентову участь кожного з мов в ролі мови #1,

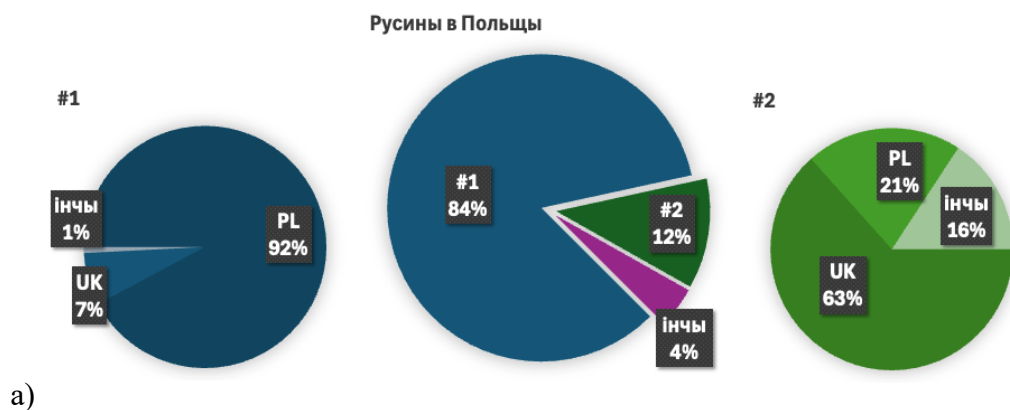
#2 ітд. Результати графічні представлено на выкресах (Рыс. 3а-г), вказуючы найчастійше розпознаваны языки в бесіді Русинів, што жыют в Польцы, Словації, Україні і Сербії.

В аналізованих авдициях з Польцы найчастійше розпознаваный язык (#1) становил 84 проц. полной долгости авдиции, а з того аж 92 проц. становил польскій язык, а 7 проц. українскій язык. Другій найчастійше розпознаваный язык (#2) становил 12 проц., з чого 63 проц. приписано українському языкови, а 21 проц. польському (Рыс. 3а).

В авдициях бесідників з України (Рыс 3б), язык розпознаваный як першый (#1) становил 91 проц. цілой базы, з чого аж 99 проц. тых сегментів приписано українському языкови .

Для Русинів на Словації найчастійше розпознацаный язык становил 66 проц., з чого 77 проц. приписано словацкому языкови. Другій найчастійше розпознаваный язык (#2) становил 19 проц., де польскій і словацкій становили 27 проц., а українскій язык 16 проц. (Рыс. 3в).

В примірі авдиций з Сербії (Рыс. 3г), першый найчастійше розпознаваный язык (#1) становил 71 проц. награнь, з чого 87 проц. тых авдиций приписано хорватському языкови. Другій найчастійше розпознаваный язык (#2) становил 22 проц. часу, з чого 78 проц. становил словенскій язык.



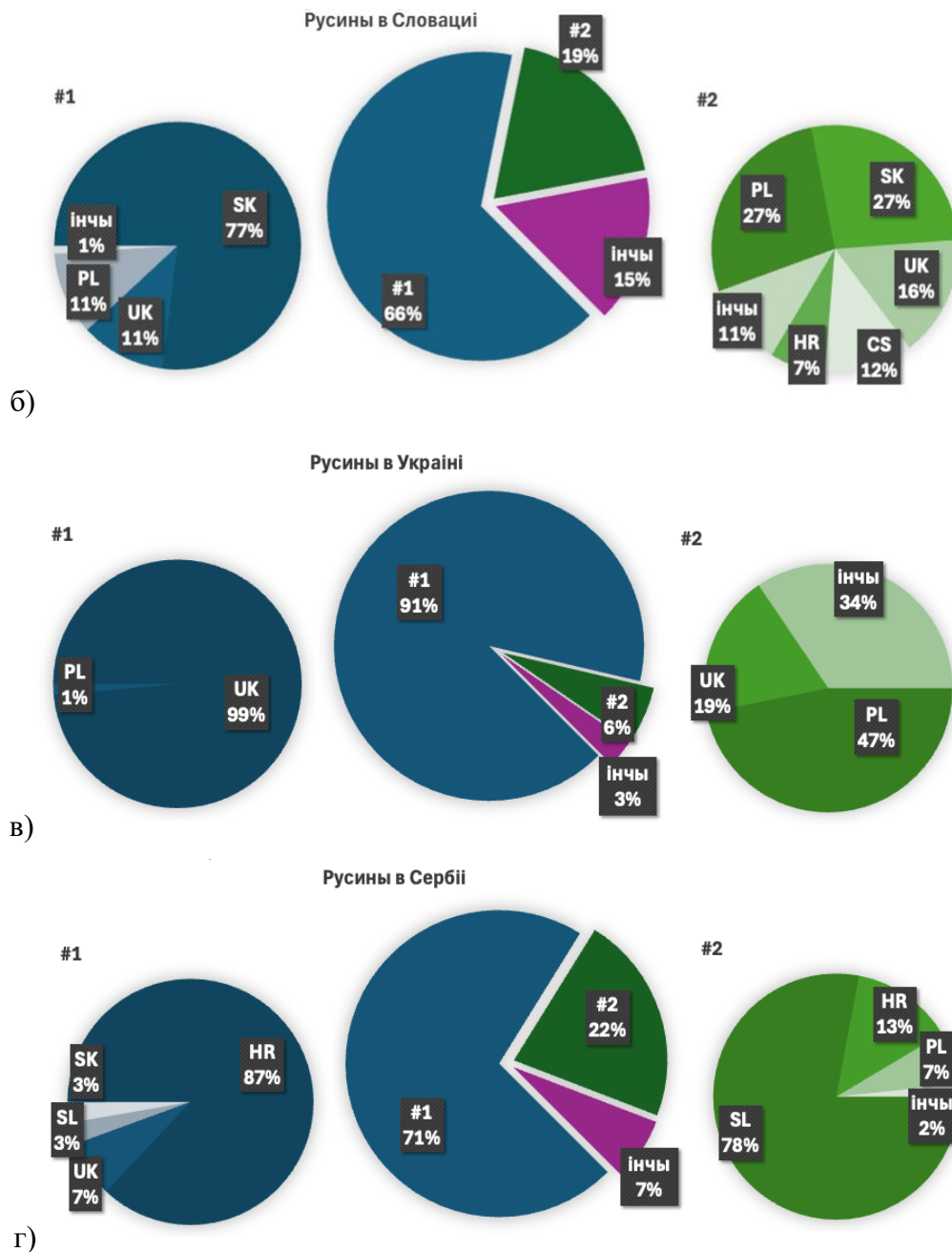


Рис. 3 а–д. Процентове сопоставління розпознаних языків Русинів з різних держав.

Додатковы аналізы зроблено для Русинів з Польщы. Выкресы (Рис. 4 а, б) представляють аналізу розпознаних авдиций Русинів з Польщы, котрых поділено на віковы группы: особы старшы як 70 років «70+» і особы молодшы як 70 років «70-». Треба додати, што база награнь лекторів «70+» становит понад 60% лекторів з Польщы.

Ограничыня досліджыня з поділом на віковы группы выключні для лекторів з Польщы, выникат з двох причин. По перше, кількіст дальшых груп была за мала, жебы отримати одповідньо великы пробы. По друге, проблемом было рішыти вік лекторів з інчых держав, бо част редакторів, што робили награня, не робит уж в радію.

В групі «70+» (Рис. 4а) найчастіше розпознаваним языком (#1) был польскій, котрый становил 89 проц. авдиций, обнимаючи 77 проц. цілковитого аналізуваного часу. Як другій найчастіше розпознаваний язык (#2) выкрыто українскій (75 проц. в тій категорії), што становит 17 проц. вшытких аналізованих авдиций. Сут гев тіж присутны інчы языки, яки появляють ся в шестьох процентах розпознань.

В молодшій віковій групі «70-» найчастіше розпознаваним языком (#1) был польскій, обнимаючи веце як 99 проц. часу в тій категорії, з 98 проц. вшытких аналізованих авдиций (Рис. 4б). Інчы языки незначні фігурують в тій групі.

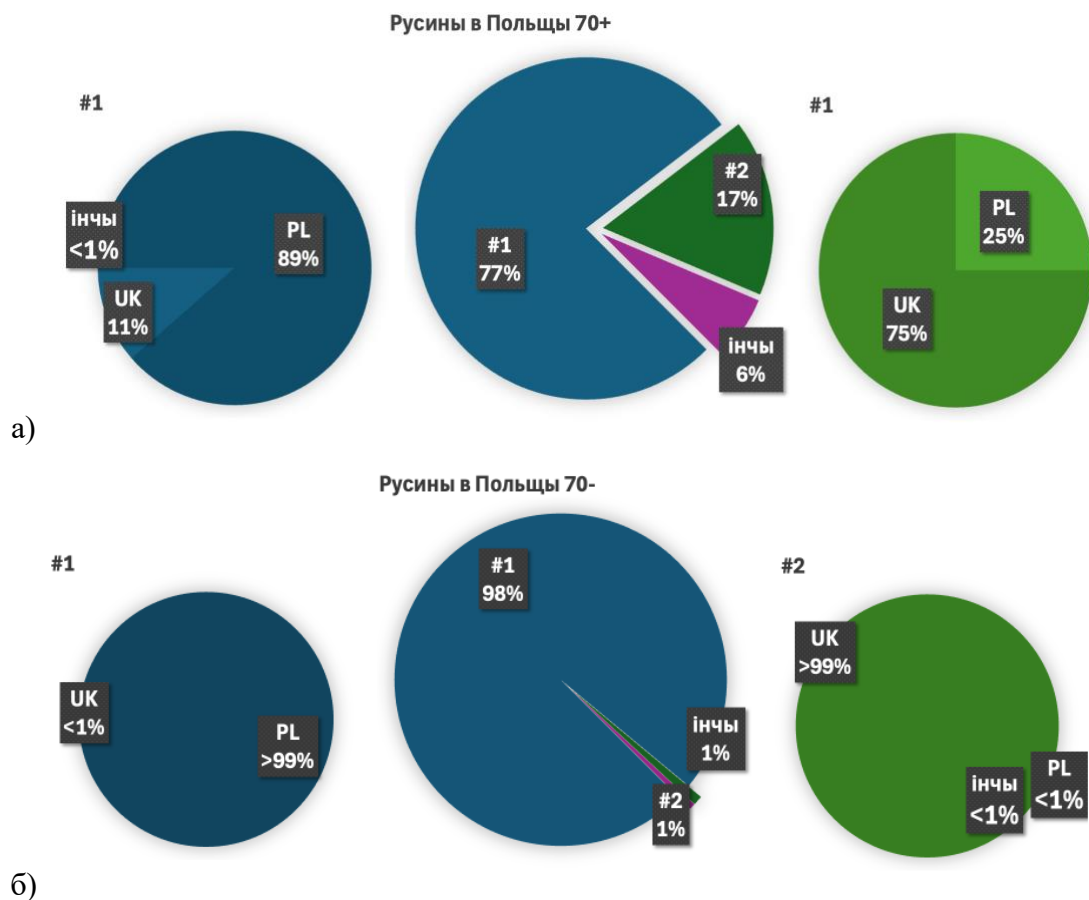


Рис. 4 а, б. Процентове сопоставління розпознаваних языків Русинів з Польщы з поділом на віковы группы (веце або менше як 70 років).

5. Дискусия

На основі досягнених результатів явні видно впливля домінуючого языка (урядового) даной державы на класифікацію русинського языка через языковий модель. З вынятком

Русинів з Войводіны, бесіда Русинів з України, Польщі і Словачії была класифікувана через алгоритм як найбарже зближена до домінуючого в даній державі языка. Хоц флексійны, лексикальны і синтактычны ріжніці медже варіянтами русиньского языка сут з перспектывы його хоснователи або языкознавців рішучо меншы як ріжніці медже русиньскым а домінуючым языком даной державы, алгоритм, што комплексово аналізує языкову подібніст, выказує значучо векше зближыня до домінуючых языків. Особливо українській язык розпознаваний єст релятивні часто серед Русинів з Польщі і Словачії, але в Сербії є інакше. В примірі Русинів з Войводіны найвекшу подібніст выказуют хорватській і словенській, при незначній участі сербского языка. Треба єднак зазначыти, же база пробок з Сербії, на якій оперто аналізу, была найменше чысленна з увагы на ограниченный доступ до материялів.

Не сут знаны докладны механізмы і критерії, на основі яких модель рішат о класифікації языка. Діяня сіти тыпу *Transformer* єст барз зложеным процесом статистычної аналізу, опертым на ідентыфікації взірців в даных языках, што дозволят вызначати подібніст без фактычного розумліня змісту.

Традиційны моделі ASR, основаны на вкрытых моделях Маркова (HMM, англ. Hidden Markov Models), характеризували ся деяком міром детермінізму, што злегшало аналізу влияния акустычных примет на результат транскрипції/класифікації языка. Сучасны моделі ASR, основаны на глубоких нейронных сітях, в тым рекурентных нейронных сітях (RNN, англ. Recurrent Neural Networks), конволюційных нейронных сітях (CNN, англ. Convolutional Neural Networks), і архітектурі тыпу *Transformer*, сут неявны і недетерміністычны. Означат то, же тоты самы входовы даны (сигнал бесіды) можут в ріжних примірах вести до кус інчых транскрипцій. Тота недетерміністычність выникат з великой кількосты параметрів моделю і зложеных, нелінійных реляцій медже нима. Утруднят то безпосередню аналізу і выокремліня особливых примет акустычного сигналу, які рішают о конкретным результаті транскрипції. Інакше бесідуючы, тяжко є докладні окрислити, котры фрагменты сигналу бесіды і яким способом мали влияния на класифікації поєдных фонемів ци слів. Істніють єднак методы, што дозволяють на досліджыня і інтерпретатицію моделей ASR, хоц сут зложены і вымагають заавансуваного інформатычного знаня. Належат до них м.ін. (Rahate і ін. 2022):

- аналіза верств увагы (attention mechanism) дозволят візуалізувати звязок медже фрагментами сигналу бесіды а елементами транскрипції, вказуючы, на які части звуку модель «зверат увагу» в час розпознаваня,

- методи основаны на градієнтовим пошпирюванню, дозволяють ідентифікувати фрагменти сигналу на вході, які мають найвекше впливання на активацію окремих нейронів і в тот спосіб на кінцевый результат транскрипції,
- методи пертурбації входу основаны сут на тым, же вводит ся невеликы зміны в сигналі бесіды і обсервує, як вплиають на результат транскрипції, што дозволять ідентифікувати основны акустычны приметы.

Результаты вказуют праві зерове визначання російского языка, што дакус заскакує, беручы до увагы його сутьове вплиання і хоснування в Україні, особливо на сході державы, одкале походила част бесідників.

Аналіза лемківского языка выказала тіж, што результаты ріжняты ся в залежности од віку бесідників. В звязку з тым проведено другу аналізу. Поділено пробкы на дві віковы групы: бесідників маючых менше і веце як 70 років. Подібніст лемківского языка до польского серед старшых осіб была значучо менша як серед молодшых осіб, што сугерує поступуючу языкову асиміляцію, особливо в обшыри фонетыкы. В лексикальным аспекті старшы бесідники хоснували подібне, аж і векше чысло польонізмів або запозычынь, што може вказувати на зложены механізмы языковой асиміляції медже поколіннями.

6. Підсумування

В статі проанаізувано здібніст моделю OpenIA Whisper до розпознавання і клясифікації русиньского языка, хоц модель не был безпосередньо вчений на русиньскых материялах. Проведено досліджыня на основі радийовых авдиций Русинів з Польщы, України, Словаччі і Сербії, при увзглядніню ріжнищ медже віковыма групамы. Результаты всказуют, же модель OpenIA Whisper клясифікує русиньскій язык головні як урядовый язык даной державы, што насуват выснoveк о сутьовым вплианню домінуючых языків на автoматычну клясифікацію меншынового языка. Достережено тіж значучу языкову асиміляцію в молодшій віковій групі, што свідчыт о поступуючым вплианню польского языка в Польщы і інчых урядовых языків в інчых державах.

В будучых досліджынях плянує ся змодифікувати модель OpenIA Whisper так, жебы выеліминувати спомагання для домінуючого языка в механізмі розпознавання, што дозволит ліпше зрозуміти подібніст русиньского языка до інчых славянських языків без сильного вплиання офіційного языка даной державы. Додатково плянує ся проаналізувати русиньскы пісні і співанкы а тіж інчы славянскы пісні. В співі выступуют інчы

механізми емісії звуку як в бесіді, што може вказати шырший контекст. Сесы досліджыня можут причыннити ся до барже прецизийной і актуальной клясифікації меншыновых языків і ліпшого зрозумліня механізмів асиміляції і языковой глобалізації.

Вшыткы первістны даны і дрібницьовы результаты вызначены проведенным експериментом сут доступны для заінтересуваных через майльовый контакт з автором.

Бібліографія

- Bouamor, Houda, Hassan, Sabit, Habash, Nizar. 2019. «The MADAR Shared Task on Arabic Fine-Grained Dialect Identification». B: *Proceedings of the Fourth Arabic Natural Language Processing Workshop*. Ред. Wassim El-Hajj, Lamia Hadrich Belguith, Fethi Bougares, Walid Magdy, Imed Zitouni, Nadi Tomeh, Mahmoud El-Haj, Wajdi Zaghouani, 199–207. Florence: Association for Computational Linguistics. <https://doi.org/10.18653/v1/W19-4622>.
- Kushko, Nadiya. 2007. «Literary Standards of the Rusyn Language: The Historical Context and Contemporary Situation». *The Slavic and East European Journal* 51, ч. 1: 111–132.
- Moser, Michael. 2016. «Rusyn: A New-Old Language In-between Nations and States». B: *The Palgrave Handbook of Slavic Languages, Identities and Borders*. Ред. Tomasz Kamusella, Motoki Nomachi, Catherine Gibson, 124–139. London: Palgrave Macmillan. https://doi.org/10.1007/978-1-137-34839-5_7.
- Nikitin, Alexey G., Kochkin, Igor T., June, Cynthia M., Willis, Catherine M., Mcbain, Ian, Videiko, Mykhailo Y. 2009. «Mitochondrial DNA Sequence Variation in the Boyko, Hutsul, and Lemko Populations of the Carpathian Highlands». *Human Biology* 81, ч. 1: 43–58. <https://doi.org/10.3378/027.081.0104>.
- Plišková, Anna. 2008. «Practical Spheres of the Rusyn Language in Slovakia». *Studia Slavica Academiae Scientiarum Hungaricae* 53, ч. 1: 95–115. <https://doi.org/10.1556/SSlav.53.2008.1.6>.
- Rabus, Achim, Scherrer, Yves. 2017. «Lexicon Induction for Spoken Rusyn – Challenges and Results». B: *Proceedings of the 6th Workshop on Balto-Slavic Natural Language Processing*. Ред. Tomaž Erjavec, Jakub Piskorski, Lidia Pivovarova, Jan Šnajder, Josef

Steinberger, Roman Yangarber, 27–32. Valencia: Association for Computational Linguistics. <https://doi.org/10.18653/v1/W17-1405>.

Radford, Alec, Kim, Jong Wook, Xu, Tao, Brockman, Greg, McLeavey, Christine, Sutskever, Ilya. 2023. «Robust Speech Recognition via Large-Scale Weak Supervision». B: *Proceedings of the 40th International Conference on Machine Learning (ICML'23)*. Ред. Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, Jonathan, 1–28 (28492–28518). Honolulu: JMLR.org.

Rahate, Anil, Walambe, Rahee, Ramanna, Sheela, Kotecha, Ketan. 2022. «Multimodal Co-learning: Challenges, Applications with Datasets, Recent Advances and Future Directions». *Information Fusion* 81: 203–239. <https://doi.org/10.1016/j.inffus.2021.12.003>.

Scherrer, Yves, Rabus, Achim. 2019. «Neural Morphosyntactic Tagging for Rusyn». *Natural Language Engineering* 25, ч. 5: 633–650. <https://doi.org/10.1017/S1351324919000287>.

Zampieri, Marcos, Nakov, Preslav, Scherrer, Yves. 2020. «Natural Language Processing for Similar Languages, Varieties, and Dialects: A Survey». *Natural Language Engineering* 26, ч. 6: 595–612. <https://doi.org/10.1017/S1351324920000492>.