

ارائه روشی مبتنی بر یادگیری عمیق جهت واژه‌یابی در گفتار فارسی

حسین سهلانی^۱، استادیار

۱- گروه فتا، دانشکده اطلاعات و آگاهی، دانشگاه جامع علوم انتظامی امین – sahlani_h@yahoo.com

چکیده: با توسعه فناوری، تحلیل بر روی گفتار انسان توسط ماشین‌های هوشمند، رشد زیادی پیدا کرده است، یکی از پیشنهادها، واژه‌یابی در گفتار است. سیستم‌های واژه‌یاب وجود و یا عدم وجود کلمات کلیدی را در گفتار بررسی می‌کند ولی به دلیل ناکافی بودن مجموعه داده، منابع پردازشی، وجود نویزها و ... توسعه یک سیستم واژه‌یاب ایده آل با دقت بالا دشوار است. در این پژوهش با ضبط ۴۲ ساعت از گفتار افراد مختلف به صورت اختصاصی در کنار سایر دادگان موجود در این حوزه و با اتکا بر روی الگوریتم‌های شبکه عصبی عمیق طراحی، آموزش و ارزیابی گردید که می‌تواند دقت واژه‌یابی و تولید یک متن متناظر با گفتار را افزایش دهد، معماری و روش پیشنهادی در دو قسمت تشکیل شده قسمت اول، طوری طراحی شده که هدفش تولید واج متناظر با گفتار، بدون داشتن اطلاعات مدل زبانی است ولی قسمت دوم با لایه‌های متعدد و داشتن مدل زبانی سعی در تولید یک متن صحیح با توجه به دامنه زبان فارسی دارد. روش پیشنهادی ترکیبی ارائه شده با دقت ۸۸.۰۱٪ می‌تواند کلیدواژه‌های موجود در گفتار را تشخیص و با دقت ۹۹.۸۰٪ عدم وجود واژه‌ها را مشخص نماید و دقت بالاتری نسبت به کارهای مشابه داشته باشد.

واژه‌های کلیدی: تبدیل گفتار به متن، واژه‌یابی، دادگان گفتار فارسی، پردازش گفتار، شبکه‌های عصبی عمیق.

A deep learning method for voice word spotting in Persian language

Sahlani, Hossein¹, assistant professor

1. FATA Department, Faculty of Information and Awareness, Amin University of Law Enforcement Sciences, Tehran, Iran

Abstract: With the development of technologies related to audio data recording and transmission, as well as the advancement of artificial intelligence science, the analysis of human speech by intelligent machines has grown greatly. One of the most important technologies in speech processing in the last decade has been word search in audio. By receiving keywords from the user, the word search system can check the presence or absence of that word in the audio file and report the result to the user. Due to the insufficient data set, it is difficult to develop an ideal software that can find all the user's words exist in the audio. In this research, by collecting 42 hours of Persian audio data along with other data available in this field and relying on deep neural network algorithms were designed, trained and evaluated. These two architectures are complementary and can increase the accuracy of word search. First architecture is designed in such a way that its goal is to produce the phoneme corresponding to the sound, without having the information of the language model, but the second architecture with multiple layers and having a language model tries to produce a correct text according to the domain of the Persian language. The presented combined proposed method can detect the keywords in the audio with 88.01% accuracy and determine the absence of words with 99.80% accuracy and has higher accuracy than similar methods.

Keywords: Audio to text conversion, audio word search, Persian audio corpus, speech processing, deep neural networks.

تاریخ ارسال مقاله: ۱۴۰۱/۱۲/۰۴

تاریخ اصلاح مقاله: ۱۴۰۲/۰۶/۱۹

تاریخ پذیرش مقاله: ۱۴۰۲/۰۷/۱۰

نام نویسنده مسئول: حسین سهلانی

نشانی نویسنده مسئول: ایران - تهران - گروه فتا، دانشکده اطلاعات و آگاهی، دانشگاه جامع علوم انتظامی امین



۱- مقدمه

تشخیص گفتار به سرعت در حال تبدیل شدن به یک فناوری پرکاربرد و مؤثر است. فناوری‌های تشخیص گفتار در حال تغییر نحوه تعامل افراد با دستگاه‌ها، خانه‌ها، اتومبیل‌ها و مشاغل خود هستند و کاربرد آن در صنایعی مانند تجارت، بانکداری، بازاریابی و مراقبت‌های بهداشتی به سرعت آشکار می‌شود. آینده تعامل انسان و سیستم، کار با سامانه‌های گفتار است و نه با صفحه کلید و ... با پیشرفت فناوری، بحث کمک گرفتن از هوش مصنوعی برای تبدیل گفتار به نوشتار و واژه‌یابی به عنوان یکی از شاخه‌های مهم هوش مصنوعی مطرح شده است.

در این پژوهش با بهره‌گیری از هوش مصنوعی و یادگیری عمیق می‌بایست متن متناظر با گفتار تشخیص داده‌شده و سپس با جستجو در آن متن، حضور یا عدم حضور واژه‌ی مدنظر کاربر مشخص شود. با پیشرفت شبکه‌های عصبی، تغییرات زیادی در الگوریتم‌های هوش مصنوعی ایجاد شده به طوری که حوزه شناسایی گفتار در گذشته بیشتر با روش‌های مبتنی بر مدل مخفی مارکوف و با ابزارهایی مانند HTK صورت می‌گرفت [1]، که برای انجام این مهم نیازمند دادگان مبتنی بر برچسب برای هر یک از واژه‌ها و استخراج ویژگی‌های متداول از آنها مانند MFCC از فایل‌های گفتاری بود که این خود زمان و نیروی انسانی مختص خود را می‌طلبد؛ اما در سال‌های اخیر این کار با استفاده از الگوریتم‌های یادگیری عمیق صورت می‌گیرد. این الگوریتم‌ها برای تبدیل گفتار به متن ابتدا ویژگی‌های مشخصی را از گفتار استخراج می‌کنند، سپس با کمک مدل آموزش داده‌شده این ویژگی‌ها را به دنباله‌ای از واژه‌های احتمالی تبدیل می‌کنند و سپس به کمک مدل زبانی این واژه‌ها را به رشته‌ای متنی تبدیل می‌کنند. این رویکرد که با عنوان انتها به انتها مطرح است، سیگنال خام گفتار و متن متناظر آن را لازم دارد و نیاز به دادگان با دنباله برچسب واجی نیست. به عبارتی شبکه‌هایی مانند CNN, RNN کار استخراج ویژگی را بر عهده دارند و ویژگی‌های استخراج‌شده را به یک طبقه بند می‌دهند تا عمل یادگیری را انجام دهد [2].

پیاده‌سازی سامانه‌ای که بتواند با دقت بالایی شناسایی گفتار را انجام دهد، همواره یکی از چالش‌های این حوزه بوده است. هرچند با استفاده از روش‌های یادگیری عمیق در این بخش، دقت‌های قابل قبولی گزارش شده است [3]، اما به دلیل این که استفاده از روش‌های یادگیری عمیق نیازمند دادگان حجیم است، برای زبان‌های با داده کم این روش‌ها کارا نیستند. برای زبان فارسی دادگان فارسی‌دات وجود دارد که شامل دو نسخه کوچک و بزرگ است. اکثر مقالات این حوزه برای زبان فارسی از دادگان فارسی‌دات استفاده کرده‌اند [5] [4]. در سال‌های اخیر شرکت موزیلا پروژه‌ای را راه‌اندازی کرد تا بتواند با آن مشکل کمبود داده برای سایر زبان‌ها را تا حد زیادی حل کند. این پروژه که با کمک کاربران هر زبان شروع به جمع‌آوری داده‌ها کرده تا به امروز توانسته حجم زیادی داده گفتاری را تهیه نمایند.

هدف مقاله پیش رو، ارائه سامانه واژه‌یاب گفتار فارسی مبتنی بر یادگیری عمیق با دقت بالا می‌باشد که حین کار در کنار دادگان رایگان موزیلا، مجموعه داده گفتار فارسی مناسب جمع‌آوری و پالایش و ایجاد گردید که بعد از پیش‌پردازش‌های صورت گرفته روی اصوات موجود در دادگان تهیه شده (استخراج ویژگی، حذف نویز و ...) بهینه‌سازی و بومی‌سازی یک معماری شبکه عصبی عمیق (تعیین تعداد پالایه‌ها، تعداد لایه‌ها و نورون‌های هر لایه و ...) آموزش و مدل‌سازی گردید و نهایتاً کارایی روش پیشنهادی با مدل‌های مختلف ارزیابی گردید که نتایج نشان دهنده بهبود عملکرد روش پیشنهادی با دقت واژه‌یابی ۸۸٫۰۱ درصد برای تشخیص وجود کلیدواژه‌ها در گفتار و دقت ۹۹٫۸۰ درصد برای تشخیص عدم وجود واژه‌ها در گفتار می‌باشد. و از دستاوردهای جانبی این پژوهش می‌توان به ساخت یک سامانه تبدیل گفتار به متن اشاره کرد که توانسته متن حاصل از گفتار را اصلاح و خروجی قابل قبولی ارائه دهد همچنین یکی از محصولات مهم این پژوهش که در آینده بیشتر قابل استفاده خواهد بود ساخت پایگاه داده اختصاصی می‌باشد که در بخش‌های بعدی به صورت کامل تشریح شده‌اند.

۲- پیشینه تحقیق

باز شناسی گفتار، گستره وسیعی از کاربردها از قبیل فهم ارتباط میان انسان و رایانه، شنود مکالمات تلفنی با اهداف امنیتی و ... را شامل می‌شود که در آن نیازی به بازتولید تمام عبارات و کلمات ادا شده توسط گوینده نمی‌باشد و تنها باز شناسی دقیق تعدادی از کلمات مهم کفایت می‌کند که به این آشکار سازی کلمات کلیدی در گفتار پیوسته ورودی، واژه‌یابی گفتار^۱ اطلاق می‌شود [6].

رویکردهایی که برای آموزش سامانه‌های واژه‌یاب گفتار مورد استفاده قرار می‌گیرند را می‌توان به دو دسته تقسیم‌بندی کرد: رویکردهای

مبتنی بر مدل مخفی مارکوف و رویکردهای مبتنی بر طبقه‌بندی [7] از مزایای روش‌های واژه‌یابی گفتار مبتنی بر مدل مخفی مارکوف سرعت مناسب، استفاده از اطلاعات وابسته به محتوا می‌باشد، در مقابل، به دلیل در اختیار نبودن تخمین‌های دقیق در آموزش مدل، الگوریتم آموزش آن در بهینه‌های محلی گرفتار شده و دقت واژه‌یابی گفتار کاهش می‌یابد. رویکردهای واژه‌یابی گفتار مبتنی بر مدل مخفی مارکوف به دو دسته تقسیم می‌شوند: واژه‌یابی گفتار مبتنی بر باز شناسی گفتار پیوسته با دادگان و واژه‌یابی گفتار مبتنی بر جستجوی واجی. در رویکرد اول، ابتدا با استفاده از یک باز شناس گفتار مبتنی بر دادگان، آر شیوهای گفتاری بزرگ به شبکه‌های واجی یا کلمه‌ای تبدیل می‌شوند. سپس با به کارگیری روش‌های جستجو در شبکه‌های واجی یا کلمه‌ای، کلمات ورودی مورد جستجو قرار می‌گیرند که نیاز به وجود حجم زیاد داده آموزشی برچسب زده‌شده و پیچیدگی محاسباتی بالا از چالش‌های اساسی آن محسوب می‌شود. در رویکرد دوم (جستجوی واجی)، زیر مدل‌های مستقل از محتوا یا وابسته به محتوا ساخته می‌شوند و از اتصال

¹ Keyword spotting

مقادیر مختلفی بگیرند و اثر متفاوتی بر روی سیگنال صدا بگذارند مانند صدای موسیقی یا صحبت فرد دیگری در پس زمینه. به صورت کلی کاهش تأثیر نویزهای غیر ثابت بسیار سخت تر از نویزهای ثابت می باشد، زیرا در نویزهای ثابت یک الگو وجود دارد که در طول زمان تکرار می شود و با شناسایی آن الگو می توان آن را شناسایی کرده و تأثیر آن را کاهش داد، منتهی در نویزهای غیر ثابت الگوی مشخصی وجود ندارد و باید با استفاده از روش های پیچیده ی غیر خطی به مانند یادگیری عمیق، تأثیر این نویز را کاهش داد. یکی از مهم ترین شبکه های عصبی مربوط به کاهش نویز، شبکه های خود رمز کننده هستند [11].

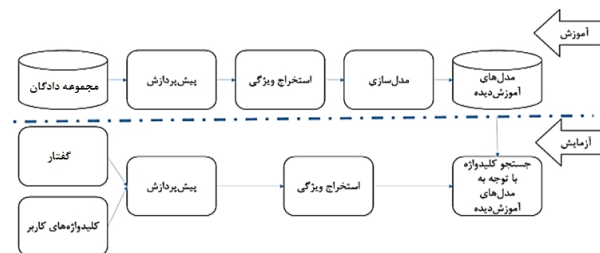
در بخش پیش پردازش، بیشترین تحقیقات بر روی بهبود کیفیت صدا و افزایش نرخ نمونه برداری در سامانه های تشخیص گفتار طراحی می شوند. روال های مرسوم شامل روش های تبدیل سیگنال آنالوگ به دیجیتال، پنجره گذاری^۴، تبدیل فوریه سریع، محاسبه ی توان اسپکترال^۵ و تولید ضرایب کپسترال فرکانس مل MFCC^۶ یا موارد دیگر به مانند پیش بینی خطی، پیش بینی خطی ادراکی^۷ و غیره می باشند. به صورت کلی دو روش برای افزایش نرخ نمونه برداری در تحقیقات پیشین استفاده شده است. روش های سنتی درون بایی مانند Spline که با متوسط گیری و یا برازش توابع چند جمله ای بر روی داده ها، سعی بر درون بایی فرکانس های از دست رفته دارند و دومین روش، روش های مبتنی بر یادگیری می باشند. روش های یادگیری، با مدل سازی از روی دادگانی با فرکانس پایین، دادگانی با نرخ وضوح بالا تولید می کنند. محققین در [12] یک مدل یادگیری عمیق خود رمز کننده مستقل برای افزایش وضوح سیگنال صدا در دامنه ی زمان پیشنهاد دادند که از خصوصیات مهم این طرح، افزایش ۴ برابری نرخ نمونه برداری فرکانس می باشد.

۲-۲- استخراج ویژگی

بعد از اعمال پیش پردازش های مختلف بر روی سیگنال گفتار سیگنال وارد قسمت استخراج ویژگی می شود. در این بخش، استفاده از فیلتر بانک ها رونق زیادی دارد و فیلتر بانک های مل و بارک از مرسوم ترین فیلتر بانک های استخراج ویژگی به حساب می آمدند. سیگنال صدا با کانوالو شدن در فیلترهای مختلف، وزن هایی از آن فیلتر دریافت می کند که با کنار هم قرار دادن آن ها می توان به یک ویژگی منحصر به فرد از سیگنال صدا رسید [13]. در ادامه این ویژگی ها برای این که مشخص شود هر واج به چه احتمالی در گفتار وجود دارد به الگوریتم های برآوردگر احتمال واج^۸ منتقل می شود. برآوردگر احتمال واج به همراه مدل HMM و مدل زبانی n-gr^۹ توسط رمزگذار^۹ برای رمزگذاری گفتار استفاده می شود. کلمات

زیر مدل ها، دنباله واجی کلمه کلیدی حاصل می شود. ضعف اساسی این بالا بودن خطاهای درج، حذف و جایگزینی می باشد [8]. راهکار دیگر استفاده از مدل های مبتنی بر طبقه بندی می باشد که در این بین شبکه های عصبی، ویژه ماشین بردار پشتیبان، شبکه های عصبی عمیق از اقبال بیشتری برخوردار بوده اند [9]. پیچیدگی های محاسباتی و زمانی بالا، عدم امکان استفاده از اطلاعات وابسته به محتوا و عدم امکان استفاده از مدل های زبانی از نقاط ضعف مدل های مبتنی بر طبقه بندی بوده که اقدامات مختلفی برای آن ها انجام شده است. چراکه در این مدل واژه یاب گفتار به عنوان یک دسته بند دودویی در نظر گرفته شده که کلاس جملات حاوی کلمات کلیدی را از کلاس جملات فاقد کلمات کلیدی تفکیک می کند. رویکرد واژه یابی گفتار مبتنی بر طبقه بندی که در ادامه این مقاله هم تنها به آن اشاره شده از دو بخش اصلی تشکیل شده است: استخراج ویژگی و الگوریتم یادگیری [10].

چارچوب کلی اکثر سامانه های واژه یاب در تحقیقات پیشین در شکل (۱) آورده شده است، این ساختار به دو بخش آموزش^۱ و آزمایش^۲ تقسیم می شود مرحله ی آموزش شامل فراهم آوری مجموعه داده گفتار، پیش پردازش صدا، استخراج ویژگی و در نهایت مدل سازی می باشد. مرحله ی آزمایش به مانند مرحله ی آموزش است با این تفاوت که به جای بخش مدل سازی، به جستجوی کلمات و یا واج ها بر روی مدل آموزش دیده پرداخته خواهد شد تا وجود و یا عدم وجود یک کلیدواژه در گفتار مشخص شود. در ادامه هر یک از مراحل تشریح شده است.



شکل ۱: چارچوب کلی سامانه های واژه یاب

۲-۱- پیش پردازش

در محیط های واقعی گفتار انسان با عوامل تخریب کننده صدا ترکیب شده و سرانجام یک گفتار ناخالص تولید می شود. عوامل تخریبی صدا باعث کاهش کارایی نهایی سیستم تشخیص گفتار می شوند، از این رو باید به نحوی مؤثر، تأثیر این عوامل را در صدا کاهش داد. عوامل تخریبی صدا به دودسته ی ثابت و غیر ثابت تقسیم می شوند. دسته ی ثابت نویزهایی هستند که در طول زمان در یک روند متناوب تکرار می شوند مثل نویز سفید و دسته ی غیر ثابت نویزهایی هستند که در طول زمان می توانند

⁶ Mel frequency cepstral coefficients

⁷ Perceptual linear prediction

⁸ Phone likelihood estimator

⁹ Decoder

¹ Train

² Test

³ Preprocessing

⁴ Windowing

⁵ Computing power spectrum

خروجی سپس به رمزگشا^۱ فرستاده می‌شوند تا آن را به شکل قابل خواندن برای انسان تبدیل کند [14].

۲-۳- مدل سازی و فناوری یادگیری عمیق

جهت مدل سازی واج ها یا کلمات، محققان در گذشته از روش های آوایی و زبانی استفاده می کردند. روش های آوایی تنها بر روی تشخیص آوا و واج ها به صورت مستقل کاربرد دارد در این رویکرد گفتار به کمک تعدادی واحد آوایی (مانند کلمه، هجا، سه واجی یا واج) مدل می شود و برای بازشناسی نیز از تشخیص این واحدها و کنار هم قرار دادن آن ها، متن متناسب با گفتار تشخیص داده می شود. روش های زبانی به گرامر زبان، هم خوانی واج ها و کلمات با یکدیگر و یا احتمال وجود کلمه ها در آن زبان تمرکز دارد. مدل های زبانی معمول مورداستفاده در سامانه های تشخیص گفتار امروزی شامل روش های گرامری و آماری هستند. در روش گرامری سعی می شود که به جملات خروجی، ساختار گرامری آن زبان اعمال شود و در روش آماری احتمال پشت سرهم آمدن کلمات (مثل مونوگرام یا احتمال وقوع کلمات در زبان، باایگرام یا آمار وقوع دو کلمه پشت سر هم در زبان و ...) مورداستفاده قرار می گیرند. خروجی شامل لیست واژگانی است که توسط سیستم بازشناسی می گردند، در سامانه های بازشناسی گفتار پیوسته با تعداد واژگان زیاد، علاوه بر لیست خود کلمات، اطلاعات مختلفی در مورد هر کلمه مانند احتمال وقوع آن در زبان، احتمال وقوع آن بعد از سایر کلمات، نقش (های) گرامری در جمله و ... را نیز شامل می شود.

مدل های آوایی نیز در تحقیقات گذشته بیشتر به دو صورت ۱- مدل های بدون نظارت همچون مدل مخفی مارکوف یا مدل مخلوط گاوسی و ۲- مدل های با نظارت مانند ماشین بردار پشتیبان^۲ (SVM) بودند که با مرزبندی بین واج ها و همچنین اختصاص وزن های مختلف به داده های گفتار مختلف، می توانستند با دقت مناسبی واج ها را شناسایی و تشخیص دهند. در مدل سازی بدون نظارت با اینکه دقت مناسبی در داده های مختلف مشاهده می شد، اما هزینه ای آموزش و محاسبات بالا بود. در سال های گذشته روش های بانظارت شبکه عصبی بازگشتی توسعه پیدا کردند که از نظر دقت، نسبت به روش های بدون نظارت برتری محسوسی داشتند [15]. محققان توانستند با ارائه ی شبکه ی عصبی عمیق^۳ (DNN) تمام متصل از نظر دقت و سرعت از روش های بدون نظارت سبقت بگیرند. با گذشت زمان شبکه ی عصبی کانولوشنی (CNN)^۴ با استفاده از اسپکتوگرام صدا و تکیه بر شبکه های عصبی کانولوشنی توانستند دقت بیشتری نسبت به DNN ها کسب کنند. باوجود کارایی مناسب CNN ها، در نظر نگرفتن وابستگی بلندمدت داده ها، اشکال اصلی این روش بود. از این رو با تولید روش های ترکیبی از شبکه ی عصبی بازگشتی RNN^۵ و CNN مدلی هایی به نام شبکه های عصبی کانولوشنی بازگشتی طراحی شد که در محیط های نوینی پایداری بیشتری نسبت به روش CNN از خود

نشان می دادند. مدل های بازگشتی واحد بازگشتی دروازه GRU^۶ و حافظه ی کوتاه مدت بلند LSTM^۷ دو مورد از مدل های متداول و پرکاربرد از شبکه های بازگشتی هستند [16].

برای بررسی عملکرد روش های موجود، یکی از محبوب ترین تابع محاسبه ی خطا روش مدل زمانی پیوسته^۸ (CTC) است. در CTC اختلاف متن به دست آمده حاصل از پیمایش بر روی شبکه ی عصبی با متن اصلی محاسبه شده و سپس به مقدار خطای تشخیص حاصل شده، وزن های شبکه ی عصبی به هنگام می شوند، این روند در پروسه ی آموزش و توسعه ی سیستم تا هنگامی که وزن ها دیگر به هنگام نشوند و یا دقت موردنظر حاصل شود ادامه پیدا خواهد کرد. مدل زمانی پیوسته با در نظر گرفتن حالت های مختلفی از خروجی حقیقی برنامه، اختلاف کاراکترهای پیش بینی شده و کاراکترهای اصلی را به صورت میانگین محاسبه می کند. مدل های زبانی، احتمال رخداد واژه در یک جمله را با توجه به کلمات پیشین و پسین به دست می آورند و در صورتی که کلمه ای به اشتباه نمایش داده شود با توجه به دیکشنری از پیش ساخته شده آن را اصلاح می کند.

در مرحله جستجو در قسمت آزمایش نیز عمدتاً به صورت سلسله مراتبی انجام می شود. تحقیقات پیشین جهت جستجو در یک گفتار، ابتدا گفتار اصلی را به قاب هایی در سطح کلمات تقسیم بندی می کردند، در ادامه اولین واج از کلیدواژه ی موردنظر بر روی قاب های تقسیم بندی شده به ترتیب جستجو می کردند، اگر قاب منتخب دارای این واج بود، واج بعدی در همین قاب آزمایش می شد، این روند به همین صورت ادامه داشت تا اینکه با پیمایش بر روی کل گفتار، کلمه ی مورد نظر شناسایی شود، در صورت عدم وجود واج های پشت سر هم در گفتار، نتیجه گرفته می شود کلیدواژه در این بخش حضور ندارد و سیستم جستجو به قاب بعدی می رود. جستجوی درختی ساده، جستجو ویتربی، جستجوی مبتنی بر پشته و جستجوی A*، سامانه های جستجوی مطرح در این حوزه می باشند [15].

۳- روش پیشنهادی

یک نرم افزار واژه یاب باید تشخیص دهد که آیا کلیدواژه در گفتار وجود دارد یا خیر. لذا ساختار نرم افزار واژه یاب دارای ۳ مؤلفه ی اصلی می باشد که مهم ترین آن تبدیل گفتار به متن می باشد. بعد از آن، متن به دست آمده باید طبق ساختار مدل زبانی فارسی، غلطیابی شده و سپس کلیدواژه ها در متن جستجو شوند.

در مرحله اول باید مجموعه داده گفتار فارسی مناسب ساخته و توسعه داده شود سپس از این داده ها جهت توسعه مدل یادگیرنده استفاده شود. در این بخش، از روش های مبتنی بر یادگیری عمیق استفاده شده به طوریکه در صورت بهره گیری از معماری صحیح و وجود دادگان مناسب

⁵ Recurrent neural network

⁶ Gated recurrent network

⁷ Long short term memory

⁸ Connectionist temporal classification

¹ Parser

² Support vector machine

³ Deep neural network

⁴ Convolutional neural network

صحیح و دقیق به شبکه‌ی طراحی شده اختصاص یابد، باید به کلمه‌ی "ketab" تبدیل شود تا خوانش واقعی آن لحاظ شود. لذا باید یک لغت‌نامه از صورت فارسی و خوانش واقعی وجود داشته باشد که بعد از اختصاص متن به گفتار، این متن به خوانش واقعی تبدیل شود. همان‌طور که مشخص است، ساختن مجموعه داده بسیار پرهزینه و زمان‌بر هست، ولی برای معماری B نیازی به خوانش واقعی نیست و به همان صورت که نوشته می‌شوند به فینگلیش تبدیل می‌گردند زیرا از مدل زبانی برای اصلاح متن خروجی استفاده می‌شود.

با توجه به موارد عنوان شده ۳۲ ساعت گفتار به همراه معادل آوایی متناظر برای هر فایل تهیه شد. این مجموعه داده با دقت بالا و در حد تمامی کسره‌های آخر و سایر حرکت‌گذاری‌ها هر کلمه مشخص شد، تمامی فایل‌های گفتاری ابتدا به زبان فارسی نوشته شدند و سپس با استفاده از واژه‌نامه‌ی آوایی متناظر به صورت معادل آوایی درآمدند. این مجموعه داده شامل ۱۸۳۶۳ عدد کلمه‌ی منحصر به فرد است و تعداد کل کلمات ۲۵۱۲۴۱ است. همراه با این مجموعه داده یک واژه‌نامه‌ی آوایی برای کلمات فارسی تولید شد که برای تبدیل کلمه‌های فارسی به کلمه ی آوایی بسیار مهم است.

۳-۱-۲- دادگان موزیلا

در سال‌های اخیر شرکت موزیلا مشکل کمبود داده برای سایر زبان‌ها را تا حد زیادی حل کرده است، این کار با کمک کاربران هر زبان انجام شده و تا به امروز توانسته حجم زیادی از داده‌های گفتاری را تهیه نماید. در این مقاله از مجموعه داده‌ی فارسی رایگان ۱۳۲ ساعته شرکت موزیلا^۱ نیز استفاده شده است؛ که با پالایش و منظم سازی فایل‌های این مجموعه داده، در مجموع حدود ۱۶۵ ساعت گفتار به همراه متن گردآوری شد.

در پروژه موزیلا کاربران با ضبط صدای خود و همچنین با تأیید یا عدم تأیید صدای سایر کاربران، به تهیه دادگان کمک می‌کنند. اگر حداقل دو کاربر مختلف یک داده را تأیید کنند، آن داده به عنوان داده معتبر و اگر حداقل دو کاربر مختلف یک داده را عدم تأیید کنند، آن داده به عنوان داده نامعتبر شناخته خواهد شد. در جدول (۱) سن گویندگان نشان داده شده است. همچنین ۷ درصد از گویندگان زن، ۷۵ درصد مرد و ۱۸ درصد از گویندگان جنسیت خود را مشخص نکرده‌اند.

سن گوینده	درصد دادگان
کمتر از ۱۹	۱
بین ۱۹ تا ۲۹	۳۳
بین ۳۰ تا ۳۹	۳۹
بین ۴۰ تا ۴۹	۳
بیشتر از ۵۰	۲
نامشخص	۲۲

جدول ۱: فراوانی سن گویندگان در دادگان موزیلا

می‌توان به استخراج ویژگی‌های متمایزکننده و تعمیم‌پذیری بالای این روش‌ها امید داشت. لذا در این پژوهش از دو معماری جهت ساخت مدل یادگیرنده مبتنی بر شبکه عصبی عمیق استفاده شده است. این دو معماری در کنار یکدیگر مکمل هم بوده و می‌توانند دقت واژه‌یابی و تولید یک متن متناظر با گفتار را افزایش دهند. در ادامه بخش‌های مختلف روش پیشنهادی به صورت کامل تشریح شده است.

۳-۱-۱- پایگاه داده مورد استفاده

اولین قدم در توسعه‌ی مجموعه داده گفتار، جمع‌آوری دادگان موجود می‌باشد که در این تحقیق از دو مجموعه داده استفاده شده است پایگاه داده اول مجموعه دادگان فارسی است که می‌توان به عنوان نوآوری این تحقیق در نظر گرفته شود و مجموعه دوم برای شرکت موزیلاست که در ادامه هر یک تشریح می‌شود.

۳-۱-۱- مجموعه دادگان گفتار فارسی

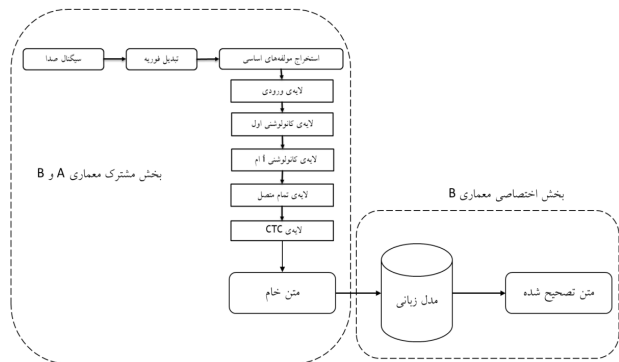
در این مرحله با توجه به وبسایت‌های ارائه‌دهنده‌ی کتاب‌های صوتی، کلیپ‌های تصویری و سایر رسانه‌های صوتی به‌مانند پادکست، بیش از ۲۰ گیگابایت فایل گفتار تهیه شد. در تهیه‌ی این اصوات سعی شد از یک گوینده بیش از یک و نیم ساعت فایل وجود داشته باشد تا اطمینان حاصل شود در توسعه‌ی سیستم، صدای گویندگان مختلف در آموزش سیستم نقش دارند. همچنین سعی شد در حالات مختلفی از صحبت، این فایل‌ها جمع‌آوری شوند. به‌طور مثال در این مجموعه داده، حالت‌های محاوره، رسمی، صحبت آرام و تند به زبان فارسی وجود دارد. متنوع بودن مجموعه داده باعث تعمیم‌پذیری مدل نهایی شده و کارایی سیستم را بیشتر خواهد کرد. برای آموزش داده‌ها باید پیش‌پردازش صورت گیرد، زیرا در فایل‌های گفتار بخش‌هایی غیر اطلاعاتی، مانند سکوت زیاد و یا آهنگ در صدا وجود دارند. لذا باید بر روی فایل‌ها پالایشی صورت گیرد تا آماده‌ی تزریق به شبکه‌ی عصبی طراحی شده باشند. از این‌رو هر یک از فایل‌های گفتار به فایل‌های زیر ۱۰ ثانیه تقسیم شدند به‌صورتی که اول و آخر هر فایل گفتار سکوت باشد تا داده‌ای تخریب نشود، همچنین سکوت‌های بیش از ۴ ثانیه حذف شدند تا اطلاعات غنی‌تر شوند. برای یکسان‌سازی فایل‌ها جهت تزریق به شبکه‌ی عصبی، تمامی فایل به فرکانس ۱۶۰۰۰ و فرمت wav. تبدیل شدند. لازم به ذکر است با توجه به زمان‌بر بودن برچسب‌گذاری این مجموعه داده، تنها بخشی از این مجموعه داده برچسب‌گذاری و مورد استفاده قرار گرفت.

ارائه‌ی متن فارسی از گفتار نیز باید به دو صورت برای معماری A و B با توجه به دیکشنری تبدیل حروف فارسی به انگلیسی (فینگلیش) در نظر گرفته شود برای معماری A باید به صورت خوانش واقعی درآید یعنی تمامی حرکت‌گذاری‌ها برای کلمه باید صورت بگیرد. به‌طور مثال کاربر کلمه‌ی "کتاب" را تایپ می‌کند ولی برای این‌که این کلمه به صورت

¹ commonvoice.mozilla.org/en/datasets, 2020

۳-۲- مدل یادگیری

در این پژوهش از دو معماری جهت تولید متن استفاده شده که مکمل یکدیگر بوده و می‌تواند دقت و آیه‌یابی و تولید یک متن متناظر با گفتار را افزایش دهد. به‌صورت کلی ساختار این دو معماری در شکل (۲) نمایش داده شده است. برای ساده‌سازی در ادامه، این دو معماری به‌صورت معماری A و معماری B تعریف می‌شود.



شکل ۲: معماری کلی شبکه عصبی

همان‌طور مشاهده می‌شود، این دو معماری در کلیات شبیه یکدیگر هستند منتها در چینش لایه‌های شبکه‌ی عصبی و نحوه‌ی محاسبه‌ی تابع خطا و به‌روزرسانی با یکدیگر متفاوت هستند. معماری A، بسیار ساده طراحی شده و تنها هدفش تولید واج متناظر با گفتار، بدون داشتن اطلاعات مدل زبانی می‌باشد. ولی معماری B کمی پیچیده‌تر بوده و با لایه‌های متعدد و داشتن مدل زبانی هدفش تولید یک متن صحیح با توجه به دامنه‌ی زبان فارسی می‌باشد. در هر دو معماری ابتدا گفتار اصلی را به قالب‌هایی در سطح واج تقسیم‌بندی می‌شوند و از بعد زمان به بعد فرکانس انتقال داده می‌شود. سپس تعدادی فرکانس مهم که در شنوایی انسان یا ساخت سیگنال صحبت انسان از اهمیت بیشتری برخوردار است انتخاب می‌شوند. مؤلفه‌های فرکانسی به‌دست‌آمده با عبور از لایه‌های کانولوشنی متعدد و استخراج ویژگی‌های گوناگون در انتزاع‌های مختلف، می‌تواند ابعاد پنهان گفتار انسان را مدل‌سازی کند. لایه یکی مانده به آخر تعداد نورون‌ها به‌اندازه‌ی حداکثر تعداد واج در نظر گرفته‌شده و در لایه‌ی آخر با استفاده از مدل زمانی پیوسته (CTC) اختلاف متن به‌دست‌آمده حاصل از پیمایش بر روی شبکه‌ی عصبی، با متن اصلی محاسبه شده و سپس به مقدار خطای تشخیص حاصل شده، وزن‌های شبکه‌ی عصبی به‌هنگام می‌شوند، این روند در پروسه‌ی آموزش و توسعه‌ی سیستم تا هنگامی‌که وزن‌ها دیگر به‌هنگام نشوند و یا دقت موردنظر کسب شود ادامه پیدا خواهد کرد. معماری A به‌گونه‌ای طراحی شده که در انتهای خروجی شبکه‌ی عصبی تنها واج‌های متناظر با گفتار، بدون در نظر گرفتن واج‌های پیشین و یا کلمات پیشین، پیش‌بینی شوند. در این روش از مدل زبانی استفاده نمی‌شود و واج‌ها به‌صورت خام جهت پردازش فرآیند بعدی، یعنی جستجوی کلمه‌ی کلیدی، مورد استفاده قرار می‌گیرند. دلیل استفاده از

این معماری سرعت بالای این معماری بود و با استفاده از معماری A، انتظار می‌رود واج‌هایی که به‌صورت ناقص از کلمات مختلف بیان شود نیز لحاظ شود، بنابراین از این روش در کنار معماری B بهره‌برده شده است. در ادامه معماری A و B و نحوه‌ی ترکیب آن‌ها بیان شده است.

۳-۲-۱- معماری A

برای مدل‌سازی کاراکترهای فارسی در این مدل تمام صورت‌های آوایی در نظر گرفته شده است. به دلیل وجود حروف با صدای یکسان در زبان فارسی، برای بالا بردن دقت مدل، حروف هم‌صدا با یک حرف نشان داده شده و یک حرف نماینده آن دسته از حروف می‌باشد (شکل ۳).

مثال	حروف
آبی==Abi	آ -- A
آسب==asb	ا -- a
احسان==ehsAn	ا -- e
امید==omid	ا -- o
سوله==sule	او -- u
کتبی==ketAbi	ای -- i
یاور==yAvar	ی -- y
بابک==bAbak	ب -- b
پایان==payAn	پ -- p
طناب==tanab	ت ط -- t
صندلی==sandali	ث س ص -- s
چهار==cAhAr	چ -- c
حوله==hosele	ح -- h
خرداد==xordAd	خ -- x
در==dar	د -- d
زبان==zabAn	ذ ز ض ظ -- z
رادیو==rAdiyo	ر -- r
ژاپن==ZApon	ژ -- Z
شیر==Sir	ش -- S
قالب==qAleb	غ ق -- q
فارغ==fAreq	ف -- f
کار==kAr	ک -- k
گور==gur	گ -- g
لاک==lAk	ل -- l
مردم==mardom	م -- m
تاریخ==nArenj	ن -- n
ولی==vali	و -- v

شکل ۳: مدل‌سازی آواهای فارسی در معماری A

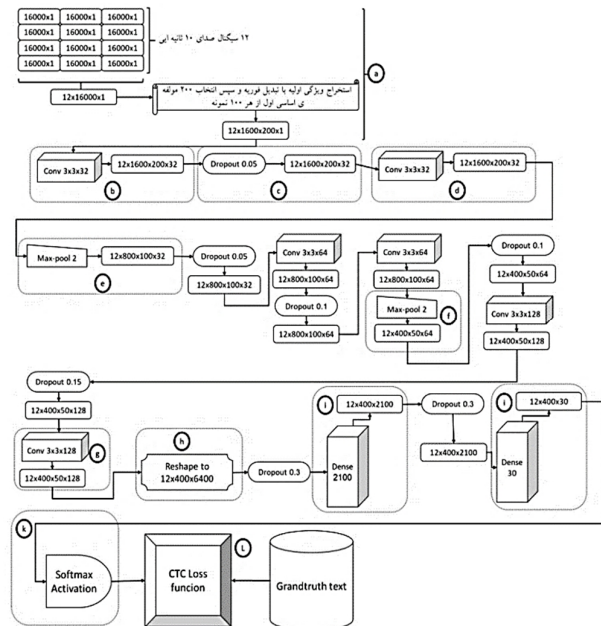
در شکل (۴) دیاگرام اصلی معماری A مشخص شده است. در فاز آموزش فایل‌های آموزشی وارد شبکه‌ی عصبی شده و بعد از محاسبه خطا و اصلاح وزن‌ها، با برچسب فایل‌ها تا رسیدن به دقت مطلوب ادامه می‌یابد.

در تمام لایه‌های میانی این معماری از تابع فعال‌سازی یک‌سوساز خطی^۱ و از الگوریتم بهینه‌سازی ADAM جهت به‌روزرسانی وزن‌ها استفاده شده

1 Rectified linear units- RELU

کانولوشن ۶۴ فیلتره که در نتیجه ماتریسی به ابعاد $12 \times 800 \times 100 \times 64$ حاصل می‌شود، اعمال maxpooling با ضریب ۲ که ماتریس ورودی را به ابعاد $12 \times 400 \times 50 \times 64$ تبدیل می‌کند (شکل (۴) بخش f)، اعمال لایه‌ی Dropout، اعمال لایه‌ی کانولوشن ۱۲۸ فیلتره که در نتیجه ماتریسی به ابعاد $12 \times 400 \times 50 \times 128$ حاصل می‌شود، اعمال لایه‌ی Dropout با ضریب ۰٫۳، اعمال لایه‌ی کانولوشن ۱۲۸ فیلتره که در نتیجه ماتریسی به ابعاد $12 \times 400 \times 50 \times 128$ حاصل می‌شود (شکل (۴) بخش g)، در مرحله بعد ماتریس ویژگی، جهت متناسب‌سازی با لایه‌ی خروجی، تغییر شکل پیدا می‌کند و به ماتریسی با ابعاد $12 \times 400 \times 6400$ تبدیل می‌شود در حقیقت در این گام دو بعد آخر باهم ترکیب می‌شوند (شکل (۴) بخش h)، در مرحله بعد یک لایه Dropout اعمال می‌شود و سپس به یک لایه‌ی تمام متصل 2100 نورونی، فرستاده می‌شود (شکل (۴) بخش i) و به یک ماتریس با ابعاد $12 \times 400 \times 2100$ تبدیل می‌شود. سپس یک لایه Dropout اعمال و به یک لایه‌ی تمام متصل 30 نورونی که همان تعداد واج‌های زبان فارسی (به همراه سکوت) می‌باشد فرستاده می‌شود (شکل (۴) بخش j) در نتیجه در این گام، ماتریس به ابعاد $12 \times 400 \times 30$ تبدیل می‌شود. به‌طور دقیق‌تر برای هر فایل گفتار ورودی ماتریسی به اندازه‌ی 400×30 پیش‌بینی می‌شود که درایه‌های آن برابر احتمال وجود هر کاراکتر می‌باشد، عدد 400 به معنای توالی کاراکترها یا همان اندازه‌ی متن می‌باشد که ثابت است، قاعدتاً فایل گفتار زیر ۱۰ ثانیه ۴۰۰ کاراکتر را در خود ندارند، به همین جهت در پیاده‌سازی محاسبه‌ی خطای تشخیص برنامه یک کاراکتر blank=# به‌طور مثال واژه‌ی ketAb و kk#eee#t###AAA#b و motehayyer#mmmm#ote#eeehayyy#####er# بعدی ماتریس انتهایی جهت پیش‌بینی کاراکترهای تعریف‌شده به لایه‌ی فعال‌سازی softmax فرستاده می‌شود (شکل (۴) بخش k). در این معماری، ماکزیمم کاراکتر موجود در فایل ۱۰ ثانیه‌ای ۱۷۸ کاراکتر در نظر گرفته‌شده است (با توجه به نمونه‌های موجود). لذا ماتریس $12 \times 400 \times 30$ پیش‌بینی‌شده در مرحله قبل، همراه با برچسب (واقعی)، به لایه‌ی آخر برای تعیین خطا ارسال می‌شود (شکل (۴) بخش L). در لایه‌ی آخر با استفاده از مدل زمانی پیوسته یا همان CTC اختلاف متن به دست‌آمده حاصل از پیمایش بر روی شبکه‌ی عصبی با متن اصلی محاسبه شده و سپس به مقدار خطای تشخیص حاصل‌شده، وزن‌های شبکه‌ی عصبی به هنگام می‌شوند، این روند در پروسه‌ی آموزش سیستم تا هنگامی که وزن‌ها دیگر به هنگام نشوند و یا دقت موردنظر کسب شود ادامه پیدا خواهد کرد. خروجی این معماری در مرحله‌ی آزمایش یک متن خام متناظر با گفتار خواهد بود که واج‌های آن بدون در نظر گرفتن احتمال حضور واج‌های

است. میزان یادگیری^۱ اولیه در این معماری با توجه به آزمایش‌های متعدد صورت گرفته برابر ۰.۰۰۱ در نظر گرفته شد.



شکل ۴: شبکه عصبی طراحی شده در معماری A

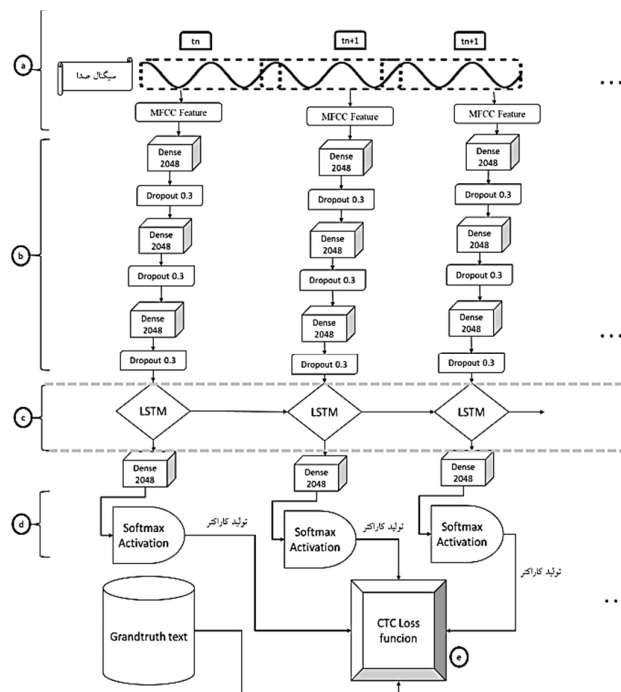
در این معماری دو نوع استخراج ویژگی صورت گرفته در گام اول به ازای هر قاب از گفتار اصلی با تبدیل فوریه زمان کوتاه ۲۰۰ ضریب ابتدایی آن استخراج می‌شود که اگر فرکانس نمونه‌برداری ۱۶ کیلوهرتز باشد، ابعاد ماتریس بردار ویژگی $1600 \times 200 \times 1$ به ازای هر فایل گفتار می‌شود (شکل (۴) بخش a) بعد از به دست آوردن ۱۲ ماتریس (تعداد فایل‌های گفتار که در هر مرحله برای اصلاح وزن‌ها در نظر گرفته می‌شود) بردار ویژگی $1600 \times 200 \times 1$ در مرحله‌ی قبل، در گام دوم این ماتریس ویژگی، جهت کسب ویژگی‌های پنهان از ۳۲ فیلتر کانولوشن دوبعدی 3×3 عبور می‌کند تا ویژگی‌های انتزاعی در سطوح مختلف به دست آید که خروجی آن یک ماتریس $12 \times 1600 \times 200 \times 32$ خواهد بود (شکل (۴) بخش b). در این معماری به ازای هر لایه‌ی کانولوشن یک لایه‌ی Dropout (جهت خاموش‌سازی بخشی نورون‌ها) (شکل (۴) بخش c)؛ و به ازای هر دو لایه‌ی کانولوشن یک لایه‌ی maxpooling (ماکزیمم مقدار هر پنجره) تعریف‌شده است (شکل (۴) بخش e). با حرکت در عمق شبکه، حساسیت به ویژگی‌های جزئی و نویز بیشتر می‌شود، لذا با اعمال خاموش‌سازی بیشتر، از بیش‌برازش شبکه جلوگیری می‌شود. در بخش e یک لایه maxpooling با ضریب ۲ جهت نصف کردن ابعاد ماتریس اعمال می‌شود تا ابعاد ماتریس به صورت $12 \times 800 \times 100 \times 32$ درآید. ادامه فرایندها، چرخه تکراری اعمال لایه کانولوشن، Dropout و maxpooling با ضرایب مختلف است که به ترتیب عبارت‌اند از اعمال لایه‌ی Dropout با ضریب ۰٫۵، اعمال لایه‌ی کانولوشن ۶۴ فیلتره که در نتیجه ماتریسی به ابعاد $12 \times 800 \times 100 \times 64$ حاصل می‌شود، اعمال لایه‌ی Dropout با ضریب ۰٫۱، اعمال لایه‌ی

2 Fully connected

1 Learning rate

خروجی این معماری در مرحله‌ی اول آزمایش یک متن خام متناظر با گفتار خواهد بود که واج‌های آن (با توجه به اعمال شبکه LSTM) با احتمال حضور واج‌های دیگر تولید شده‌اند. در مرحله‌ی دوم با استفاده از مدل زبانی، احتمال حضور واج‌ها و کلمات در کنار یکدیگر سنجیده شده و در نهایت متن اصلاح و غلطیابی می‌شود. برای ساخت مدل زبانی می‌بایست یک مجموعه داده از متن‌ها و جملات مختلف در حوزه‌های گوناگون جمع‌آوری شده و سپس با یک الگوریتم n -gram، احتمال حضور هر کلمه در حضور کلمات دیگر محاسبه و ذخیره شود (روش n -gram به این صورت است که در یک جمله به ازای هر کلمه، تعداد تکرار n کلمه‌ی قبلی یا بعدی محاسبه شده و در نهایت، به ازای هر کلمه مشخص می‌شود چه تعداد کلمات دیگری و با چه تکراری حضور دارند). در این پژوهش در مجموع، با تهیه‌ی حدود ۲ میلیون جمله‌ی فارسی یکتا که از ۱۷۰ هزار کلمه‌ی یکتا تشکیل می‌شد و با روش 5-gram مدل زبانی توسعه پیدا کرد (چهار نوع مدل زبانی دو گرام، سه گرام، چهار گرام و پنج گرام) لازم به ذکر است در این پژوهش، روش ساخت مدل زبانی بر اساس kenlm طرح‌ریزی و پیاده‌سازی شد.^۴ کد پیاده‌سازی معماری B نیز در لینک زیر آورده شده است.

<https://github.com/mahdaviFarAsr/asr.git>



شکل ۵: شبکه عصبی طراحی شده در معماری B

دیگر تولید شدند قاعدتاً لایه‌ی محاسبه خطا در بخش آزمایش، حذف می‌شود.^۱

۳-۲-۲- معماری B

در شکل (۵) دیاگرام اصلی معماری B نشان داده شده است. این معماری بر اساس روش هانن که شرکت موزیلا آن را با زبان برنامه‌نویسی پایتون و کتابخانه تنسورفلو پیاده‌سازی کرده^۲ طراحی شده است [15]. هسته‌ی اصلی این سیستم شبکه‌ی عصبی RNN می‌باشد که با توجه به ویژگی‌های اسپکتوگرام صدا، می‌تواند یک متن متناظر به سیگنال صحبت انسان بدهد. شبکه‌ی RNN جهت مدل‌سازی فرآیندهای مبتنی بر زمان مناسب بوده و می‌تواند از اطلاعات لحظاتی $t_n, t_{n-1}, \dots$ برای تصمیم‌گیری در لحظه t_n استفاده کند.

در روش پیشنهادی در این معماری از ویژگی‌های MFCC که پیش‌تر در پیشینه بیان شد جهت استخراج ویژگی استفاده می‌شود. از این‌رو با تنظیم گروه داده‌ها برابر ۶۴، (در یک به‌روزرسانی از ۶۴ فایل گفتار، MFCC گرفته شده) پیمایش بر روی شبکه‌ی عصبی با استفاده از این ویژگی‌ها شروع می‌شود (شکل (۵) بخش a) این معماری از ۵ لایه‌ی اصلی استفاده می‌کند که ۳ لایه اول غیر بازگشتی، لایه‌ی چهارم بازگشتی و لایه‌ی آخر مجدداً غیر بازگشتی می‌باشد و علاوه بر ۵ لایه‌ی اصلی فوق، از لایه‌ی Dropout نیز با پارامتر 0.3 و الگوریتم بهینه‌سازی ADAM با میزان یادگیری اولیه 0.0001 و مقدار آستانه ۲۰ برای به‌روزرسانی وزن‌ها استفاده شده است.

در گام نخست ویژگی‌های MFCC از ۳ لایه‌ی تمام متصل با تعداد نورون ۲۰۴۸ به همراه لایه‌های Dropout عبور داده می‌شوند (شکل (۵) بخش b). بعد از گذر از لایه‌های تمام متصل، لایه‌ی چهارم این شبکه‌ی عصبی، شبکه‌ی بازگشتی LSTM می‌باشد. این شبکه در واقع حافظه‌ی بلندمدت شبکه‌ی عصبی بوده و اطلاعات حال و گذشته را با وزن دهی مناسب تنظیم می‌کند، اطلاعات این لایه علاوه بر پیمایش جهت تولید کاراکتر در زمان t_n ، به لحظه‌ی t_{n+1} نیز ارسال می‌شود (شکل (۵) بخش c).

لایه‌ی پنجم این شبکه یک لایه‌ی تمام متصل با ۲۰۴۸ نورون می‌باشد؛ که بعد از اعمال تابع فعال‌سازی RELU، جهت تخمین احتمال وجود کاراکترها به لایه‌ای با تابع فعال‌سازی softmax فرستاده می‌شود (شکل (۵) بخش d). لایه‌ی آخر، لایه‌ی تخمین خطای CTC می‌باشد ((شکل (۵) بخش e)، در این لایه، اختلاف متن به دست آمده حاصل از پیمایش بر روی شبکه‌ی عصبی با متن اصلی محاسبه شده و سپس به مقدار خطای تشخیص حاصل شده، وزن‌های شبکه‌ی عصبی به هنگام می‌شوند، این روند در پروسه‌ی آموزش و توسعه تا هنگامی که وزن‌ها دیگر به هنگام نشوند و یا دقت مورد نظر کسب شود ادامه پیدا خواهد کرد.

³ Learning rate

⁴ <https://kheafeld.com/code/kenlm/>

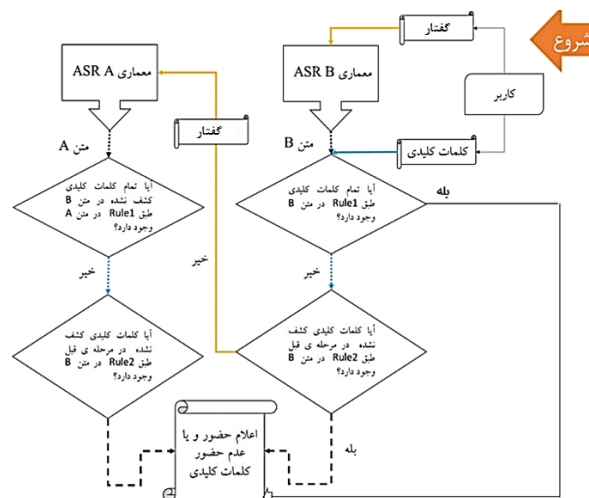
^۱ لینک پیاده‌سازی معماری A: <https://github.com/mahdaviFarAsr/asr.git>

² <https://github.com/mozilla/DeepSpeech/>

۳-۳- واژه‌یابی

بعد از آموزش شبکه‌ی عصبی در دو معماری A و B و به دست آوردن وزن‌های این دو شبکه، در مرحله‌ی آزمایش دو متن به ازای هر فایل گفتار خواهیم داشت. در متن حاصل از معماری A، حضور یک واج به حضور واج‌های دیگر وابسته نیست و تنها از مدل آکوستیکی استفاده شده است؛ بنابراین در این معماری، انتظار می‌رود، به ازای هر کلمه‌ای که بیان می‌شود، تعدادی واج متوالی صحیح در آن کلمه پیش‌بینی شود.

در متن حاصل از معماری B، حضور واج‌ها و کلمات به کلمات پسین و پیشین وابسته می‌باشد، لذا در این معماری کلمات با معنی که در مدل زبانی قبلاً وجود داشتند ارائه خواهد شد. به صورت کلی همان‌طور که در قسمت بعد بررسی خواهد شد، دقت معماری B با توجه به فرآیندهای بیان‌شده، بیشتر از معماری A بوده، منتهی در صورت بیان کلماتی خارج از دامنه‌ی زبانی آموزش داده شده در معماری B، روش معماری A می‌تواند تا حد مطلوبی تعدادی از واج‌های این کلمات را تشخیص داده و در مجموع مکمل خوبی برای معماری B می‌باشد. در شکل (۶) نحوه‌ی استفاده از این دو معماری جهت واژه‌یابی نشان داده شده است.



شکل ۶: فرآیند واژه‌یابی

همان‌طور که در شکل (۶) مشخص می‌باشد، جهت واژه‌یابی در گفتار، ابتدا کاربر یک سیگنال گفتار به همراه یک مجموعه کلمات کلیدی در سیستم بارگذاری می‌کند. سپس ابتدا کلمات کاربر در متن معماری B جستجو شده و در صورت پیدا نکردن کلمات کاربر، این واژه‌ها در متن معماری A جستجو خواهند شد. برای تشخیص واژه در گفتار دو قانون به نام Rule1 و Rule2 طبق شکل (۶) طراحی گردیده است که Rule1 به دنبال شکل صحیح و کامل کلمات در گفتار می‌باشد. در بسیاری از موارد شکل کامل کلمات ممکن است بیان نشود، به‌طور مثال ممکن در متنی به جای کلمه "هستیم"، کلمه "هستم" پیش‌بینی شود و یا تعدادی بیشتری از واج یک کلمه به نادرستی بیان شوند. جهت حل این مشکل قانون Rule2 طراحی شد. در این قانون به ازای یک بازه‌ی اطمینان از شباهت دو کلمه، وجود و یا عدم وجود یک کلیدواژه در گفتار اعلام می‌شود.

شود. الگوریتم لون اشتیان متداول‌ترین روش جهت تشخیص فاصله یا همان شباهت دو کلمه را می‌باشد. با استفاده از این الگوریتم می‌توان کمترین تعداد عملیات لازم جهت تبدیل یک کلمه به کلمه‌ی دیگر را محاسبه کرد. عملیات تعریف‌شده در فاصله‌ی لون اشتیان در این بخش، درج، جایگذاری و حذف می‌باشد. به‌عنوان مثال فاصله لون اشتیان بین "kitten" و "sitting" برابر ۳ است چراکه حداقل سه ویرایش برای تبدیل یک‌رشته به رشته‌ی دیگر وجود دارد. در این پژوهش از بازه‌ی اطمینان ۱ استفاده شده است یعنی کلمات کشف‌شده الزاماً باید در یک واج تفاوت داشته باشند.

۴- پیاده‌سازی و ارزیابی نتایج

پیاده‌سازی روش پیشنهادی تحت زبان برنامه‌نویسی پایتون و با استفاده از کتابخانه‌های Tensorflow و Keras برای پیاده‌سازی معماری‌های A و B برای کار روی فایل‌های گفتار از کتابخانه‌های wave, scipy, ffmpeg, Pydub, PyAudio استفاده شده است.

جهت سنجش کارایی روش واژه‌یابی می‌بایست دقت پیدا کردن واژه موجود در گفتار و دقت در پیدا نکردن واژه ناموجود در گفتار، به صورت مستقل بررسی شود تا معیاری صحیح و همه‌جانبه از کارایی واژه‌یاب طراحی شده به دست بیاید. همچنین در ادامه، به تأثیر پارامترهای شبکه عصبی و مدل زبانی در کارایی روش پیشنهادی با ارائه نمودارها و جداول مختلف پرداخته شده است.

۴-۱- معیارهای سنجش کارایی

هسته اصلی معیار ارزیابی تعریف شده در تبدیل گفتار به متن همان‌طور که پیش‌تر بیان شد بر اساس الگوریتم لون اشتیان می‌باشد. درواقع با استفاده از این الگوریتم می‌توان کمترین تعداد عملیات لازم (درج، جایگذاری و حذف) جهت تبدیل یک کلمه به کلمه‌ی دیگر را محاسبه کرد. جهت بیان دقیق‌تر و قابل تفسیر کارایی، فاصله لون اشتیان اگر کمتر از طول رشته حقیقت یا متن اصلی باشد تقسیم بر این طول شده و در عدد ۱۰۰ ضرب می‌شود تا عددی بر مبنای درصد بیان شود. اگر میزان فاصله لون اشتیان بیشتر از طول رشته حقیقت باشد، آنگاه، میزان خطا ۱۰۰ درصد تعریف می‌شود، بنابراین هرچه عدد به دست آمده کمتر باشد، کارایی بهتر خواهد بود.

جهت آزمایش کارایی دقت واژه‌یابی از ماتریس درهم‌ریختگی استفاده شد. این ماتریس جهت ارزیابی عملکرد الگوریتم‌های هوش مصنوعی در دسته‌بندی داده‌ها استفاده فراوانی داشته و طبق مقادیر آن می‌توان به نتایج قابل تفسیر رسید. نتایج مثبت، یعنی الگوریتم پیشنهادی اعلام می‌کند واژه در گفتار وجود دارد و منفی یعنی واژه در گفتار وجود ندارد؛ بنابراین می‌توان ماتریس درهم‌ریختگی را به صورت زیر بیان کرد [18].

این مجموعه داده، از 143.25 ساعت داده گفتار فارسی جهت آموزش، 2.77 ساعت داده جهت اعتبارسنجی و 12.38 ساعت جهت آزمایش استفاده شد.

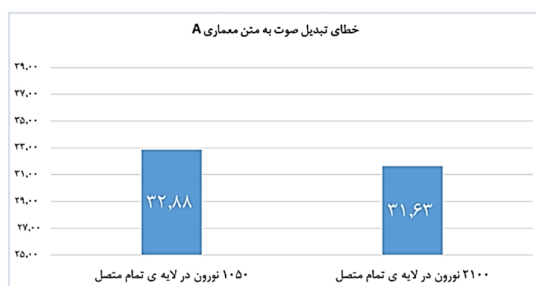
جدول ۳: جزئیات مجموعه داده Speech168H

مرحله	تعداد فایل	مجموع ساعت
آموزش	۵۵۲۲۸	۱۴۳،۲۵
اعتبارسنجی	۱۰۰۰	۲،۷۷
آزمایش	۴۴۶۴	۱۲،۳۸

علاوه بر مجموعه داده ذکر شده در جدول (۳)، یک آزمایش دیگر بر روی مجموعه ی داده Mozilla7H جهت سنجش دقیق تر کارایی سیستم نهایی صورت گرفته است. در این مجموعه داده سعی شده حتی الامکان گویندگان متفاوت باشند و هیچ متن مشابه ای در فاز آموزش و آزمایش نباشد. این مجموعه شامل اصوات محاوره ای بوده و درجه سختی آن بسیار بالاتر از مجموعه داده Speech168H می باشد. مجموعه داده موزیلای مورد استفاده در این مقاله شامل ۸۴۵۱ فایل گفتار به طول ۷،۶۷ ساعت جهت آزمایش در روش پیشنهادی می باشد که از مجموعه داده موزیلای فارسی گردآوری شده است.

۴-۳- نتایج تبدیل گفتار به متن

با توجه به آزمایش های صورت گرفته در این تحقیق، سعی شده تأثیر پارامترهای مختلف در عملکرد روش های پیشنهادی بررسی شود، لذا در ادامه تأثیر به کارگیری تعداد نورون های مختلف در لایه ها و تأثیر به کارگیری مدل زبانی در معماری B و ... نشان داده شود. یکی از پارامترهای تأثیرگذار در دقت شبکه عصبی، تعداد نورون لایه های تمام متصل هستند، به همین دلیل آزمایشی طراحی شد تا با تغییر نورون های لایه های تمام متصل، کارایی معماری A و B مشخص شود. نتایج آزمایش های صورت گرفته بر روی معماری A، در نمودار شکل (۷) نشان داده شده است.



شکل ۷: کارایی معماری A با تغییر تعداد نورون ها بر روی Speech168H همان طور که در نمودار شکل (۷) مشخص می باشد، خطای تشخیص معماری A با در نظر گرفتن ۱۰۵۰ و ۲۱۰۰ نورون در لایه آخر نشان داده

معیار اول ^۱ (TP) می باشد که این معیار، صحت وجود واژه ها در متن پیش بینی شده با توجه به تک تک واژه های موجود در متن برچسب خورده و اصلی، سنجیده می شود.

معیار دوم، (FP) می باشد و نشان می دهد چه تعداد واژه که در گفتار وجود نداشتند، به اشتباه در دسته مثبت قرار گرفتند. جهت آزمایش این معیار برای هر فایل گفتار ۲۵ واژه به صورت تصادفی استخراج شد تا سنجیده شود آیا این واژه هایی که در اصل در گفتار موجود نیستند، سیستم واژه یاب آن ها را پیدا خواهد کرد یا نه. این ۲۵ واژه به هیچ عنوان با واژه های اصلی فایل گفتار برابر نیستند.

معیار سوم ^۲ (FN) می باشد و نشان می دهد چه تعداد واژه که در گفتار وجود داشتند، به اشتباه در دسته منفی قرار گرفتند.

معیار چهارم ^۴ (TN) می باشد و نشان می دهد چه تعداد واژه که در گفتار وجود نداشتند، به درستی در دسته منفی قرار گرفتند. جهت آزمایش این معیار از همان آزمایش معیار دوم که نرخ تشخیص واژه های تصادفی ناموجود در متن اصلی می باشد، استفاده شده است.

به جای استفاده مستقیم از این ۴ معیار، به طور متداول از معیارهای حساسیت ^۵ (TPR، ^۶ TPR، ^۷، تخصیص ^۸، نرخ تشخیص نادرست دسته مثبت (FBR) ^۹ و دقت کل ^{۱۰} جهت ارزیابی دقیق تر استفاده می شود که طریقه محاسبه هر یک در جدول (۲) آورده شده است.

جدول ۲: معیارهای ارزیابی

$\text{حساسیت (R)} = \frac{TP}{FN+TP}$	این معیار مشخص می کند چه تعداد از واژه های موجود در گفتار به درستی تشخیص داده شدند.
$\text{تخصیص} = \frac{TN}{TN+FP}$	این معیار در مشخص می کند چه تعداد از واژه های ناموجود در گفتار به درستی در دسته منفی قرار گرفتند.
$FBR = \frac{FP}{TN+FP}$	این معیار در مشخص می کند چه تعداد از واژه های ناموجود در گفتار به نادرستی در کلاس مثبت قرار گرفتند.
$\text{دقت کل} = \frac{TP+TN}{TP+FP+TN+FN}$	

طبق تعاریف انجام شده در جدول (۲)، کارایی سیستم وقتی مناسب است که مقادیر حساسیت، تخصیص و دقت کل بالا بوده و درعین حال معیار FBR کم باشد.

۴-۲- دادگان مورد استفاده

جزئیات دادگان مورد سنجش که Speech168H نام گذاری شده در جدول (۳) آورده شده است. به صورت کلی در آزمایش های صورت گرفته بر روی

6 Precision

7 True positive rate

8 Specificity

9 False positive rate

10 Total accuracy

1 True Positive

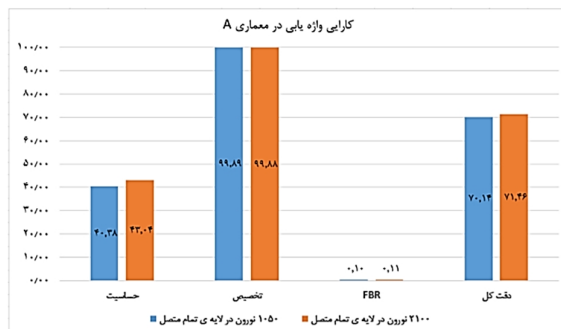
2 False Positive

3 False Negative

4 True Negative

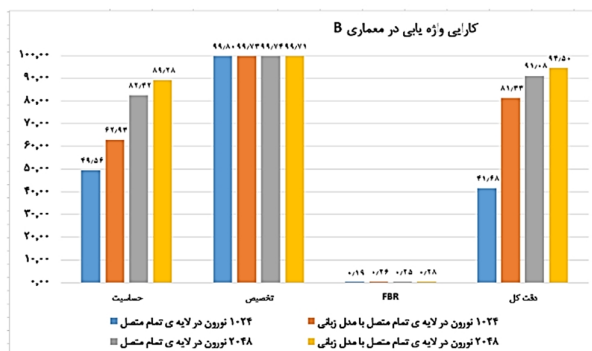
5 Sensitive

وجود کلمه در گفتار، با دقت ۹۹ درصد، اعلام می‌کند که کلمه در متن پیدا نشده است و نرخ FBR این معماری در حدود ۱ درصد می‌باشد. در مجموع معماری A می‌تواند با دقت ۷۱,۴۶ درصد واژه‌های کاربر در گفتار را شناسایی کند و یا در صورت عدم وجود واژه در گفتار آن را اعلام کند.



شکل ۹: کارایی معماری A در واژه‌یابی بر روی Speech168H

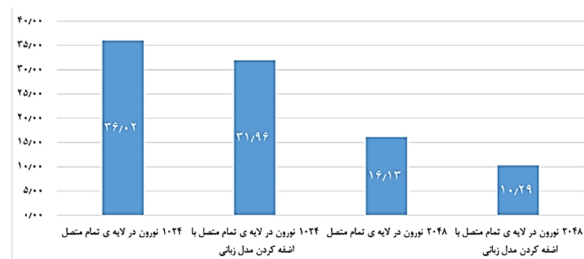
در نمودار شکل (۱۰)، معیارهای حساسیت، تخصیص، FBR و دقت کل در معماری B مشخص شده است.



شکل ۱۰: کارایی معماری B در واژه‌یابی بر روی Speech168H

همان‌طور که در نمودار شکل (۱۰) مشخص است، معماری B با ۲۰۴۸ نورون نسبت به ۱۰۲۴ نورون و استفاده از مدل زبانی در مجموع توانسته دقت جستجو را بالا ببرد، در این حالت با حساسیت ۸۹,۲۸ درصد توانسته واژه‌های موجود در گفتار را پیدا کند و همچنین در صورت عدم وجود واژه در گفتار، با دقت ۹۹,۷۱ درصد، پیغام عدم وجود را به کاربر اعلام کند. FBR معماری B مانند معماری A بسیار ناچیز بوده و این بدین معنی است که کلماتی که واقعاً در گفتار نیستند با درصد کمی به کاربر اعلام وجود می‌شوند. معماری B با دقت کل ۹۴,۵۰ درصد می‌تواند واژه‌های موجود در گفتار را تشخیص داده و یا واژه‌های ناموجود را تعیین کند. معماری B نسبت به معماری A حدوداً ۲۳ درصد دقت بهتری دارد اما ترکیب این دو معماری همان‌طور که شکل (۱۱) نشان داده شده، حدود ۹۸,۵۱ درصد دقت در TP داشته است در حالی که این مقدار برای معماری B ۸۹,۲۸ می‌باشد. منتهی مدل ترکیبی معماری A و B در معیار FN دارای ۱۷,۲۸ درصد خطا می‌باشد، در حالی که این معیار برای معماری B ۱۰,۷۱ می‌باشد. استفاده از مدل ترکیبی جایی اهمیت پیدا می‌کند که کاربر

شده است. افزایش تعداد نورون‌های لایه آخر توانسته حدود یک درصد خطای تشخیص کاهش داده تا میزان آن به ۳۱,۶۳ برسد. افزایش نورون‌ها تا حدی در کاهش خطا تأثیر داشته، منتهی در صورت افزایش این نورون‌ها پیچیدگی شبکه عصبی نسبت به داده‌های موجود نیز افزایش داشته و سبب بیش برآزش شبکه عصبی می‌شود. به همین دلیل افزایش پارامترها بیش از 2100 نورون در لایه انتهایی پیشنهاد نمی‌شود. نتایج آزمایش‌های صورت گرفته بر روی معماری B، نیز در نمودار شکل (۸) نشان داده شده است. در این نمودار، خطای تشخیص معماری B با در نظر گرفتن ۱۰۲۴ و ۲۰۴۸ نورون در لایه‌های تمام متصل و همچنین تأثیر اضافه شدن مدل زبانی بر هر یک از این شبکه‌ها نشان داده شده است.



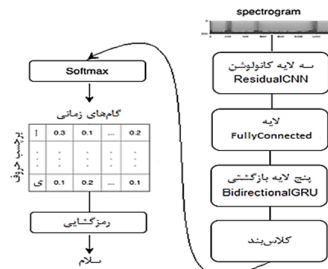
شکل ۸: کارایی معماری B با تغییر تعداد نورون‌ها و مدل زبانی

همان‌طور که در نمودار شکل (۸) مشخص است، افزایش تعداد نورون‌های لایه‌های تمام متصل در معماری B از ۱۰۲۴ نورون به ۲۰۴۸ نورون، موجب کاهش ۲۰ درصدی خطا شده و همچنین اضافه شدن مدل زبانی باعث کاهش خطای تشخیص ۵ درصدی می‌شود. وجود مدل زبانی، در هر دو مدل ۱۰۲۴ و ۲۰۴۸ نورونی، باعث تصحیح املائی کلمات شده و در صورت لزوم می‌تواند کلماتی معنی‌دار را با توجه به کلمات پیشین و پسین در متن پیش‌بینی شده، بگنجاند. بهترین عملکرد در معماری B، شبکه عصبی با ۲۰۴۸ نورون در لایه‌های تمام متصل به همراه مدل زبانی می‌باشد که توانسته با خطای تشخیص ۱۰,۲۹ درصد گفتار مورد آزمایش را به متن تبدیل کند.

ساختار معماری B با وجود شبکه عصبی بازگشتی LSTM، وابستگی بین ورودی‌ها را در نظر می‌گیرد، بنابراین، هر چه تعداد نورون‌های لایه‌های تمام متصل، بیشتر باشد، LSTM در نهایت بهتر تغذیه شده و می‌تواند ویژگی‌های مناسب‌تری جهت تشخیص واج ارائه دهد. این معماری، مانند معماری A، در صورت افزایش بیش از حد نورون‌های تمام متصل به سمت بیش برآزش حرکت می‌کند.

۴-۴-۴ نتایج واژه‌یابی

در این بخش، به بررسی نتایج کارایی واژه‌یابی پرداخته خواهد شد. در نمودار شکل (۹)، معیارهای حساسیت، تخصیص، FBR و دقت کل در معماری A مشخص شده است. همان‌طور که در شکل (۹) مشخص است، معماری A با ۲۱۰۰ نورون در لایه تمام متصل، با دقت ۴۳,۰۴ درصد می‌تواند واژه‌های موجود در گفتار را پیدا کند، همچنین در صورت عدم



شکل ۱۲: معماری سامانه شناسه گفتار فارسی

در مقاله "تشخیص عبارات‌های گفتاری برای اخبار فارسی صداوسیما جمهوری اسلامی ایران" با مجموعه داده آموزشی فارس‌دات کوچک و بزرگ و دادگان اخبار فارسی صداوسیما جمهوری اسلامی به درصد خطای تشخیص کلمه ۷،۹۳ و دقت واژه‌یابی ۸۳ درصد با مدل آموزش دیده با دادگان فارس‌دات کوچک بر روی همین دادگان و درصد خطای تشخیص کلمه ۲،۷۱ دقت واژه‌یابی ۹۲ درصد با مدل آموزش دیده با فارس‌دات بزرگ بر روی همین دادگان و درصد خطای تشخیص کلمه ۲۸،۲۳ درصد دقت واژه‌یابی ۷۹ درصد با مدل آموزش دیده با فارس‌دات بزرگ بر روی دادگان اخبار فارسی با استفاده از مدل SGMM رسید. [20]

در جدول (۵) نتایج مقایسه روش پیشنهادی با الگوریتم‌های موجود در زبان فارسی نشان داده شده است.

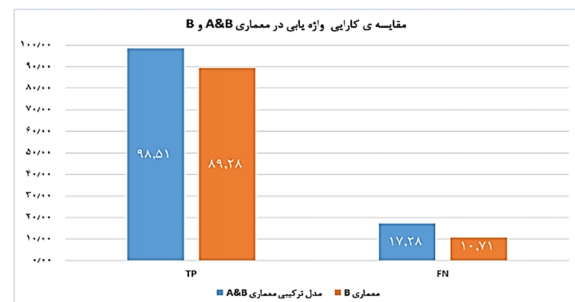
جدول ۵: مقایسه کارایی سامانه با محصول مشابه

سیستم	دادگان	آموزش مقدار داده	گفتار به متن، خطای تبدیل واژه‌یابی	دقت کل
روش پیشنهادی	Mozilla, Persian corpus و مختص همین پژوهش	۱۵۲	۱۴،۳۷	۸۸
[19]	Mozilla	۱۰۰	۴۶،۶	ندارد
[20]	فارس‌دات کوچک و بزرگ و دادگان اخبار فارسی	۷۷،۵	۲۸	۷۹

در رابطه با اهمیت پایگاه داده و موارد مشابه در قسمت قبل همان‌طور که در جدول (۵) مشاهده می‌شود روش پیشنهادی توانسته دقت قابل قبولی در مقایسه با کارهای مشابه داشته باشد. لذا علاوه بر مقایساتی که برای اجزای مختلف روش پیشنهادی (تعداد نوروها، نحوه به‌کارگیری معماری ها و ...) در قسمت‌های قبل در این قسمت نیز روش پیشنهادی با روش‌های مشابه مورد مقایسه قرار گرفت و نتایج قابل قبولی از خود نشان داد.

یکی دیگر از معیارهای ارزیابی الگوریتم‌های واژه‌یاب گفتار، نمودار ROC^۱ می‌باشد. این معیار در قالب یک منحنی دوبعدی نمایش داده می‌شود که محور افقی آن نماینده تعداد اخطارهای اشتباه^۲ بر کلمه کلیدی و محور عمودی آن نماینده تعداد تشخیص‌های درست^۳ است.

می‌خواهد با حساسیت بیشتری نتایج جستجوی کلیدواژه‌ها را مشاهده کند هرچند ممکن است، سیستم واژه‌هایی را نمایش دهد که در گفتار موجود نباشند.



شکل ۱۱: کارایی معماری B و A&B بر روی Speech168H

۴-۵- آزمایش بر روی مجموعه داده Mozilla7H

در جدول ۴ نتایج آزمایش روش پیشنهادی بر روی مجموعه داده Mozilla7H نشان داده شده است.

جدول ۴: نتایج آزمایش نرم‌افزار نهایی بر روی دیتاست Mozilla7H

تبدیل گفتار به متن	خطای تشخیص کلمه	۱۴،۳۷
واژه‌یابی	حساسیت	۷۶،۲۱
	تخصیص	۹۹،۸۰
	FBR	۰،۱۹
	دقت کل	۸۸،۰۱

با توجه به آزمایش‌های صورت گرفته روش پیشنهادی تبدیل گفتار به متن بر روی این مجموعه داده ۱۴،۳۷ درصد خطا دارد که با توجه ساعات مجموعه داده آموزشی شبکه عصبی قابل قبول می‌باشد. در بخش واژه‌یابی روش پیشنهادی با دقت ۷۶،۲۱ درصد می‌تواند واژه‌های موجود در گفتار را پیدا کند و با دقت ۹۹،۸۰ درصد نیز می‌تواند در صورت عدم وجود واژه آن را به کار اعلام کند. خطای FBR بر روی این مجموعه حدود ۰،۱۹ درصد بوده و این بدین معنی است که کلماتی که واقعاً در گفتار نیستند با درصد کمی به کاربر اعلام وجود می‌شوند. در مجموع روش پیشنهادی می‌تواند با دقت کل ۸۸،۰۱ درصد واژه‌های موجود در گفتار را تشخیص داده و یا واژه‌های ناموجود را تعیین کند.

۴-۶- مقایسه با روش‌های مشابه

در مقاله "سامانه شناسایی گفتار فارسی انتها به انتها با استفاده از پس‌پردازش دادگان پروژه موزیلا" با معماری شکل زیر، با ۱۰۰ ساعت داده موزیلا (۳۲۹۹ گوینده) و همچنین حدود ۵ ساعت داده برای توسعه و ۵ ساعت داده برای آزمون، به درصد خطای تشخیص کلمه ۴۶،۶ درصد رسیده است. [19]

¹ Receiver Operating Characteristic

² Specificity

³ sensitivity

روش پیشنهادی با دقت ۸۸ درصد می‌تواند کلیدواژه‌های موجود در گفتار را تشخیص دهد و همچنین با دقت ۹۹،۸۰ می‌تواند عدم وجود واژه‌ها را در گفتار پیغام دهد.

۲- تبدیل گفتار به متن با خطای تشخیص کلمه ۱۴،۳۷ درصد:

یکی از دستاوردهای جانبی این پژوهش ساخت یک سیستم تبدیل گفتار به متن است. این سیستم با داشتن الگوریتم‌های مدل زبانی توانسته متن حاصل از فایل گفتار را اصلاح کرده و خروجی قابل‌قبولی ارائه دهد.

۳- ساخت یک پایگاه داده گفتار ۱۷۵ ساعته:

یکی از محصولات مهم این پژوهش که در آینده بیشتر قابل استفاده خواهد بود ساخت پایگاه داده اختصاصی گفتار بود. در این پایگاه داده صوتی با مدت‌زمان زیر ۱۰ ثانیه از گویندگان مختلف وجود دارند که به ازای هر کدام از آن‌ها یک متن نوشته شده است. در این پژوهش به صورت اختصاصی ۴۲ ساعت گفتار، برچسب‌گذاری و همچنین از مجموعه داده موزیلا نیز ۱۲۹ ساعت و از مجموعه داده corpus Persian نیز ۳ ساعت، پالایش و بهینه‌سازی شد تا در مجموع ۱۷۵ ساعت فایل گفتار به همراه متن به وجود آید.

۵-۲- محدودیت‌ها و راه‌کارها

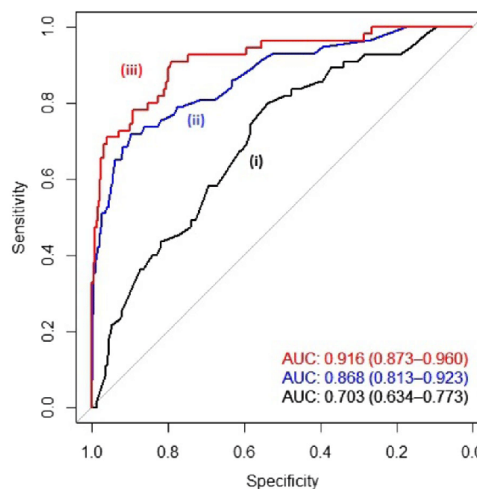
محققین این حوزه با توسعه روزافزون الگوریتم‌ها می‌توانند به پوشش انواع زبان‌ها و لهجه‌ها در زبان مردم ایران زمین بپردازند که در این راستا پیشنهادها زیر عنوان می‌گردد.

الف- مهم‌ترین عامل در کارایی سامانه واژه‌یاب و تبدیل گفتار به متن پیشنهادی، افزایش منابع محاسباتی و داده‌ای می‌باشد. این گسترش در سه بُعد باید صورت پذیرد. اولین بُعد، متناسب‌سازی تعداد گفتار محاوره و رسمی می‌باشد. با توجه به اینکه زبان فارسی محاوره بخش اعظمی از مکالمات روزمره را در برمی‌گیرد و نوع گویش و نگارش آن متفاوت از زبان معیار می‌باشد، می‌بایست تعداد اصوات محاوره و رسمی برابر باشد تا در هنگام آزمایش، کارایی بالا رود. دومین بُعد، متناسب‌سازی بین جنسیت مرد و زن است که باید از هر نوع، به صورتی یکسان، گسترش مجموعه داده صورت پذیرد. سومین بُعد، متناسب‌سازی بین لهجه مختلف زبانی است تا در هنگام آزمایش، تفاوت لهجه‌ها موجب کاهش دقت نشود.

ب- مسئله مهم دیگر گسترش مجموعه داده جهت تولید مدل زبانی می‌باشد، هر چه قدر متن‌های منحصربه‌فرد در ساخت مدل زبانی بیشتر باشد، دامنه زبانی فراگیرتر شده و موجب افزایش دقت می‌شود.

ج- استفاده از شبکه‌های از پیش آموزش دیده، در زبان‌های دیگر. با توجه به فراوانی مجموعه داده در زبان‌های دیگر و بهینه‌سازی وزن‌های شبکه‌های عصبی حاصل از آموزش این مجموعه داده‌ها، می‌توان از این شبکه‌ها جهت آموزش هر چه بهتر شبکه عصبی در زبان فارسی استفاده نمود. استفاده از این ایده موجب همگرا شدن سریع‌تر شبکه عصبی شده و به نحوی کمبود منابع داده‌ای تا حدی جبران می‌شود.

AUC متوسط سطح زیر منحنی ROC برای نرخ اخطار اشتباه مابین ۰ تا ۱ می‌باشد و هر چه به ۱ نزدیک‌تر باشد بیانگر عملکرد بهتر الگوریتم می‌باشد.



شکل ۱۳: نمودار ROC برای ارزیابی دقت چارچوب‌های تبدیل به متن

همان‌طور که در شکل (۱۳) نشان داده شده (i: روش [۱۹]، iii: روش [۲۰]، iii: روش پیشنهادی) روش پیشنهادی در مقایسه با روش‌های دیگر در نرخ اخطار اشتباه یکسان از دقت تشخیص بالاتری برخوردار بوده و متوسط سطح زیر منحنی ROC در روش پیشنهادی نسبت به روش‌های مورد مقایسه بیشتر می‌باشد.

۵- نتیجه‌گیری و پیشنهادها

به صورت خلاصه می‌توان از آزمایش‌های صورت گرفته، نتایج زیر را کسب کرد:

- دقت معماری B، از معماری A بیشتر بوده که دلیل اصلی برتری معماری B نسبت به معماری A استفاده از شبکه عصبی بازگشتی و اضافه کردن مدل زبانی به عنوان احتمال پیشین به مدل آکوستیکی می‌باشد.
- استفاده از مدل ترکیبی معماری A و B موجب نمایش تعداد بیشتری از نتایج اعلام وجود واژه شده و هنگامی کاربر حساسیت زیادی بر روی پیدا کردن واژه موردنظر خود دارد انتخاب آن مفید می‌باشد.
- پارامترهای شبکه عصبی در صورت افزایش، می‌تواند به صورت متناسب افزایش پیدا کند و در نهایت، دقت شبکه عصبی را افزایش دهد، منتها افزایش غیرمتناسب با نمونه‌های آموزشی، موجب بیش برآزش می‌شود.
- مدل زبانی می‌تواند متن‌های پیش‌بینی‌شده را غلطیابی کرده و متناسب با احتمال وجود جمعی کلمات، به یک متن قابل درک برای کاربر برسد؛ بنابراین اعمال این مدل، موجب افزایش دقت واژه‌یابی و تبدیل گفتار به متن می‌شود.

۵-۱- دستاوردها و نوآوری‌ها:

دستاوردها و نوآوری‌های حاصل از این پژوهش عبارت‌اند از:

- ۱- واژه‌یابی با دقت ۸۸،۰۱ درصد:

مراجع

- [11] Weninger F, Erdogan H, Watanabe S, Vincent E, Le Roux J, Hershey JR, Schuller B. Speech enhancement with LSTM recurrent neural networks and its application to noise-robust ASR. In *Latent Variable Analysis and Signal Separation: 12th International Conference, LVA/ICA 2015, Liberec, Czech Republic, August 25-28, 2015, Proceedings 12* 2015 (pp. 91-99). Springer International Publishing.
- [12] Kuleshov V, Enam SZ, Ermon S. Audio super resolution using neural networks. arXiv preprint arXiv:1708.00853. 2017 Aug 2.
- [13] Alumäe T, Karakos D, Hartmann W, Hsiao R, Zhang L, Nguyen L, Tsakalidis S, Schwartz R. The 2016 BBN Georgian telephone speech keyword spotting system. In *2017 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* 2017 Mar 5 (pp. 5755-5759). IEEE.
- [14] Kaur A, Sachdeva R, Singh A. Latest trends in deep learning for automatic speech recognition system. In *Artificial Intelligence and Speech Technology: Third International Conference, AIST 2021, Delhi, India, November 12-13, 2021, Revised Selected Papers* 2022 Jan 29 (pp. 62-72). Cham: Springer International Publishing.
- [15] Tabibian S, Akbari A, Nasersharif B. Discriminative keyword spotting using triphones information and N-best search. *Information Sciences*. 2018 Jan 1;423:157-71.
- [16] Arik S, Kliegl M, Child R, Hestness J, Gibiansky A, Fougner C, Prenger R, Coates A, inventors; Baidu USA LLC, assignee. Convolutional recurrent neural networks for small-footprint keyword spotting. United States patent US 10,540,961. 2020 Jan 21.
- [17] Hannun A, Case C, Casper J, Catanzaro B, Diamos G, Elsen E, Prenger R, Sathesh S, Sengupta S, Coates A, Ng AY. Deep speech: Scaling up end-to-end speech recognition. arXiv preprint arXiv:1412.5567. 2014 Dec 17.
- [18] V. Yathish, "Loss Functions and Their Use In Neural Networks," 2022. [Online].
- [19] س. زارعی، ی. شکفته، "سامانه شناسایی گفتار فارسی انتها به انتها با استفاده از پس پردازش دادگان پروژه موزیلا"، پنجمین کنفرانس ملی فناوری در مهندسی برق و کامپیوتر، دانشگاه گلستان، گلستان، ایران، ۱۴۰۰.
- [1] Young S, Evermann G, Gales M, Hain T, Kershaw D, Liu X, Moore G, Odell J, Ollason D, Povey D, Valtchev V. The HTK book. Cambridge university engineering department. 2002 Dec; 3(175):12.
- [2] Han W, Zhang Z, Zhang Y, Yu J, Chiu CC, Qin J, Gulati A, Pang R, Wu Y. Contextnet: Improving convolutional neural networks for automatic speech recognition with global context. arXiv preprint arXiv:2005.03191. 2020 May 7.
- [3] Zhang Y, Qin J, Park DS, Han W, Chiu CC, Pang R, Le QV, Wu Y. Pushing the limits of semi-supervised learning for automatic speech recognition. arXiv preprint arXiv:2010.10504. 2020 Oct 20.
- [4] Sameti H, Veisi H, Bahrani M, Babaali B, Hosseinzadeh K. A large vocabulary continuous speech recognition system for Persian language. *EURASIP Journal on Audio, Speech, and Music Processing*. 2011 Dec;2011:1-2.
- [5] Veisi H, Haji Mani A. Persian speech recognition using deep learning. *International Journal of Speech Technology*. 2020 Dec;23:893-905.
- [6] م. مرادنژاد و ش. طبیبیان، "یک رویکرد پایانه به پایانه برای تشخیص موثر فرامین صوتی فارسی در خودروی هوشمند"، در بیست و هشتمین کنفرانس بین المللی کامپیوتر، انجمن کامپیوتر ایران، تهران، ۱۴۰۱.
- [7] Fendji, J. L. K. E., Tala, D. C., Yenke, B. O., & Atemkeng, M. (2022). Automatic speech recognition using limited vocabulary: A survey. *Applied Artificial Intelligence*, 36(1), 2095039.
- [8] Mohanty, P., & Nayak, A. K. (2022). CNN based keyword spotting: an application for context based voiced Odia words. *International Journal of Information Technology*, 14(7), 3647-3658.
- [9] s. mehra and s. susan, "deep fusion framework for speech command recognition using acoustic and linguistic features," *multimedia tools and applications*, vol. 23, pp. 1-25, 2023 Mar.
- [10] M. BİLĞİLİ, "Application of long short-term memory (LSTM) neural network," *Turkish Journal of Electrical Engineering and Computer Sciences*, 2022.
- [20] Veisi H, Ghoreishi SA, Bastanfard A. "Spoken term detection for Persian news of Islamic Republic of Iran broadcasting". *Signal and Data Processing*. 2021 Feb 10;17(4):67-88.