

دریافت مقاله: ۱۴۰۲/۱۲/۰۴

پذیرش مقاله: ۱۴۰۳/۰۱/۳۰

نوع مقاله: پژوهشی

ایجاد مدل هوش مصنوعی بر پایه مبدل به منظور پاسخ به سؤالات از داده‌های بدون ساختار

فربرز دلیری*

دانشجوی کارشناسی ارشد، گروه مهندسی کامپیوتر، دانشکده و پژوهشکده رایانه شبکه و ارتباطات، دانشگاه جامع امام حسین(ع)، تهران، ایران
پست الکترونیکی: fariborz.daliri@ihu.ac.ir

محسن نوروزی

پژوهشگر گروه مهندسی کامپیوتر، دانشکده و پژوهشکده رایانه شبکه و ارتباطات، دانشگاه جامع امام حسین(ع)، تهران، ایران
پست الکترونیکی: m.norouzi@ihu.ac.ir

علی ناصری

استاد گروه مهندسی کامپیوتر، دانشکده و پژوهشکده رایانه شبکه و ارتباطات، دانشگاه جامع امام حسین(ع)، تهران، ایران
پست الکترونیکی: a.naseri@ihu.ac.ir

چکیده

متون است. در این پژوهش به منظور پاسخ به سؤالات از روی متون بدون ساختار در زمینه تحلیل امنیت شبکه‌های کامپیوتری سیستمی بر پایه هوش مصنوعی ارایه شده است. سیستم مذکور بر پایه مبدل‌ها پیاده‌سازی شده و با زبان پایتون و در بن‌سازه گوگل کولب پرو کدنویسی و پیاده‌سازی شده است. این سیستم پس از آموزش روی پرسش و پاسخ‌های حوزه امنیت شبکه‌های کامپیوتری قادر به استخراج پاسخ سؤالات از متن بدون ساختار است. به منظور آموزش مدل یک مجموعه داده سفارشی به نام هوشفا توسعه داده شده است که شامل ۱۲,۱۵۷ رکورد است. سیستم مذکور پس از ارزیابی دقت ارایه پاسخ در معیار تطبیق دقیق را بسته به نوع مدل پایه به میزان ۲۰ تا ۴۳ درصد و در معیار اف ۱ به میزان ۸ تا ۱۸ درصد بهبود بخشیده است. همچنین برای تعیین داده‌های مرتبط آموزشی برای آموزش مدل‌ها، سیستم پیشنهادی از یک

به دنبال موفقیت‌های چشم‌گیر مدل‌های زبانی از پیش آموزش دیده و یادگیری انتقالی، اخیراً موج قابل توجهی از تحقیقات هوش مصنوعی در این حوزه انجام شده است. مدل‌های زبانی یکی از مؤلفه‌های اساسی در پردازش زبان طبیعی هستند. این مدل‌ها تنها یک بار روی پیکره‌های بزرگ متنی به صورت خود ناظر آموزش دیده و سپس در حل مسایل مختلف فهم زبانی مانند پاسخ به پرسش، بسط متن و طبقه‌بندی متون به کارگیری می‌شوند. به کارگیری مبدل‌ها برای حل مسایل گوناگون مستلزم فراهم نمودن سه پیش‌نیاز (مدل، مجموعه داده مناسب، استانداردهایی به منظور توصیف نحوه ارزیابی مدل) است. بررسی حجم بالای اسناد بدون ساختار که بیشتر در قالب متن هستند برای متخصصان امری دشوار است. از طرفی پاسخ به بسیاری از سؤالات تحقیق نیازمند بررسی حجم بالایی از

* نویسنده مسئول

مدل طبقه‌بندی داده استفاده می‌کند. نمرات ارزیابی مدل طبقه‌بندی داده نتایج بالای ۸۰٪ در معیارهای اف ۱، پوشش و صحت را نشان داده است. همچنین نمرات ارزیابی تعیین داده‌های غیر مرتبط نتایج بالای ۹۷٪ در معیارهای اف ۱، پوشش و صحت را نشان داده است.

واژه‌های کلیدی: مبدل‌ها، تحلیل داده‌های بدون ساختار، پردازش زبان طبیعی، امنیت شبکه‌های کامپیوتری، مدل‌های زبانی از پیش‌آموزش دیده

۱. مقدمه

اخیراً مدل‌های زبانی^۱ از پیش‌آموزش دیده در حوزه پردازش زبان طبیعی به‌طور قابل‌توجهی افزایش یافته است. سیستم‌های پاسخ به سؤال نوعی از مدل‌های زبانی هستند که در سال‌های اخیر موردتوجه زیادی قرار گرفته است [۳۲]. بیشتر آثار مربوط به پاسخ به سؤال در ابتدا برای زبان انگلیسی پیشنهاد شده است، اما برخی از مطالعات تحقیقاتی اخیراً بر روی زبان‌های غیرانگلیسی انجام شده است [۱۰]. فهم زبان طبیعی زیرشاخه‌ای از پردازش زبان طبیعی و در رده سخت‌ترین مسایل هوش مصنوعی قرار می‌گیرد [۱]. با توجه به کاربردهای عملی این فناوری تمایل زیادی به آن وجود دارد و می‌توان از کاربردهای آن، استدلال خودکار، ترجمه ماشینی، دسته‌بندی اسناد، جمع‌آوری خبر، پاسخ به پرسش‌ها، تشخیص گفتار، پردازش تصویر (مانند نویسه‌خوان نوری)، فهرست بندی و تحلیل محتوا در ابعاد بزرگ اشاره کرد.

از مهم‌ترین کاربرد پردازش زبان طبیعی می‌توان به استخراج دانش نهفته از متون بدون ساختار (نوشته متنی ساده) اشاره نمود. متون بدون ساختار به جهت این که مدل داده‌ای از پیش تعیین شده ندارند، در استخراج دانش با چالش‌های زیادی همراه هستند. طبق منابع علمی مختلف داده‌های بدون ساختار بیش از ۸۰ درصد کل داده‌ها را تشکیل می‌دهند [۱-۴] و با این که سایر پژوهشگران در این زمینه تحقیقاتی انجام داده‌اند اما هنوز نیاز به کار

بیشتر در این حوزه (به خصوص در زبان فارسی) وجود دارد. مبدل‌ها^۲ از سازوکار خودتوجه استفاده می‌کنند و این امکان را به ماشین می‌دهد تا در محدوده‌های محلی و عمومی متمرکز شود و عمل یادگیری را علاوه بر توجه به مکان فعلی کلمه، با در نظر داشتن کل جمله انجام دهد. تحقیقات اخیر در حوزه مدل‌های زبانی از پیش‌آموزش داده شده نشان داده‌اند این مدل‌ها در زمان پیش‌آموزش به‌طور ضمنی و بدون نظارت، دانش‌های زبانی را در بازنمایی‌های خود کد می‌کنند. به‌عنوان مثال برت^۳ این دانش‌ها را به‌طور سلسله‌مراتبی فراگرفته است؛ به‌طوری‌که لایه‌های اولیه در مدل، دانش‌های زبانی سطح پایین مانند اطلاعات مربوط به موقعیت مکانی کلمات را کد می‌کنند و لایه‌های میانی در اطلاعات نحوی مانند بازسازی درخت تجزیه جمله غنی‌تر هستند.

۱-۱. بیان مسئله پژوهش

در این پژوهش با به‌کارگیری تکنیک‌های یادگیری عمیق و پردازش زبان طبیعی، مدلی آموزش داده شده است تا با استفاده از دریافت ورودی در قالب داده‌های بدون ساختار، قادر به بررسی و ارایه پاسخ به موضوعات مرتبط با امنیت شبکه‌های کامپیوتری باشد. مدل مذکور از مقادیر زیادی داده‌های متنی جمع‌آوری شده برای آموزش استفاده نموده و با ایجاد درک از نقش کلمه در جملات (مشابه مغز انسان)، قادر به یادگیری، بسط و پاسخ به سؤالات است. محصول تولید شده قابلیت تنظیم دقیق و به‌کارگیری در همه حوزه‌ها را دارد؛ اما در این تحقیق بر روی مفاهیم امنیت شبکه‌های کامپیوتری به‌خصوص حملات سایبری که منجر به نقض امنیت و خارج از سرویس شدن شبکه‌های کامپیوتری می‌شوند، تمرکز شده است. همچنین برای به‌روز بودن مدل، سابقه برخی موارد مرتبط با امنیت شبکه‌های سایبری مانند گزارش‌های تحلیلی در وب‌گاه‌ها از جمله مراکز آ‌پا نیز به مدل آموزش داده می‌شود. انتظار می‌رود مدل مذکور پس از آموزش،

2-Transformers

3-Bert

1-Language Model

پرداخته شده است. همچنین با تولید یک مجموعه داده به نام هوشفا میزان دقت مدل‌ها در ارایه پاسخ را بسته به مدل‌های مختلف؛ در معیار تطبیق دقیق به میزان ۲۰ تا ۴۳ درصد و در معیار اف ۱ به میزان ۸ تا ۱۸ درصد نسبت به پژوهش‌های قبلی بهبود یافته است.

۴-۱. تعاریف و مفاهیم پایه

یک مدل مبدل جهانی کامل و کارآمد برای پاسخ به سؤال وجود ندارد [۵]. بهترین مدل برای یک مسئله، مدلی است که بهترین خروجی را برای یک مجموعه داده با وظیفه خاص تولید کند. از سال ۲۰۱۷ تا ۲۰۲۳ تعداد پارامترها از 65M به 1/7T در مدل جی‌پی‌تی-۴ افزایش یافته است. روند رشد پارامترها در جدول زیر قابل مشاهده است:

جدول (۱). تکامل تعداد پارامترهای مبدل

پارامتر	مقاله	مدل مبدل
65M	Vaswani et al. (2017)	Transformer Base
213M	Vaswani et al. (2017)	Transformer Big
110M	Devlin et al. (2019)	BERT-Base
340M	Devlin et al. (2019)	BERT-Large
117M	Radford et al. (2019)	GPT-2
345M	Radford et al. (2019)	GPT-2
1.5B	Radford et al. (2019)	GPT-2
175B	Brown et al. (2020)	GPT-2
175B	Brown et al. (2020)	GPT-3
175B	OpenAI.(2022)	GPT-3.5
1/7T	OpenAI.(2023)	GPT-4

با توجه به این که مدل‌ها در اندازه‌های مختلفی ایجاد می‌شوند ممکن است برخی از نسخه‌های مدل‌های قبلی نسبت به مدل‌های جدید دارای پارامترهای بیشتری باشند. به عنوان مثال، مدل کوچک جی‌پی‌تی-۳ شامل 125M پارامتر است که از یک مدل جی‌پی‌تی-۲ که شامل 345M پارامتر است کوچک‌تر است.

۴-۱-۱. معرفی بن‌سازه اشتراک مدل‌های مبدل

هاگینگ فیس به یک بن‌سازه مرکزی اطلاق می‌شود

توانایی‌های تحلیلی و استدلالی مناسبی در مورد تحلیل امنیت شبکه‌های کامپیوتری ارایه نماید.

ارایه خروجی به صورت قراردادی از نیازهای اطلاعاتی انجام می‌شود. به عنوان مثال یک روش این است که کاربر اعلان‌هایی را به سامانه ارسال می‌کند و هوش مصنوعی سامانه که در قالب مدل آموزش داده شده، با توجه به مطالبی که از پیش یاد گرفته، در هر اعلان یک پاراگراف بسط ارایه می‌کند. روش دیگر ارایه نتایج به صورت پرسش و پاسخ است و برای این منظور نیاز است مجموعه داده مرتبط ایجاد و به مدل آموزش داده شود.

۴-۱-۲. ضرورت و اهمیت پژوهش

در این پژوهش با به کارگیری مبدل‌ها از ترکیب شبکه‌های یادگیری عمیق و تکنیک‌های پردازش زبان طبیعی استفاده شده تا یک مدل هوش مصنوعی برای ارزیابی امنیت شبکه‌های کامپیوتری ایجاد شود. مدل با استفاده از مجموعه داده ایجاد شده در حوزه امنیت شبکه‌های کامپیوتری، آموزش دیده است و از سرعت و دقت بالاتری نسبت به سایر روش‌های مشابه برخوردار است. با توسعه این مدل محققان حوزه سایبری می‌توانند پاسخ سؤالات تحلیلی خود را از انبوهی از متن بدون ساختار استخراج کنند.

۴-۱-۳. انگیزه‌ها و دستاوردهای پژوهش

اصلی‌ترین انگیزه این پژوهش ایجاد یک مدل هوش مصنوعی به منظور ارزیابی متون بدون ساختار در حوزه امنیت شبکه‌های کامپیوتری است. برای این منظور نیاز به یک مجموعه داده آموزشی وجود دارد. با ایجاد مجموعه داده، سایر محققان می‌توانند باهدف تحلیل امنیت شبکه‌های کامپیوتری مدل‌هایی طراحی و پیاده‌سازی نموده و همچنین تصمیم به بهبود و گسترده‌تر کردن مدل و یا مجموعه داده نمایند.

در این پژوهش با ارایه سیستمی به جستجوی پاسخ‌ها در زمینه امنیت شبکه‌های کامپیوتری از متن بدون ساختار

استفاده از این مدل باید توجه کرد که مدل اولیه به راحتی می تواند متنی تولید کند که دارای سوگیری یا محدودیت است و ممکن است حتی حاوی محتوای نژادپرستانه باشد. همچنین ممکن است تنظیم دقیق مدل بر روی داده های جدید نیز نتواند این سوگیری ذاتی را از بین ببرد.

۱-۴-۳. مجموعه داده

در دسترس بودن مجموعه داده های ارزیابی، باکیفیت بالا یک ضرورت برای ارزیابی قابل اعتماد از پیشرفت در وظایف و دامنه های مختلف فهم زبان طبیعی است [۶]. مجموعه داده به مجموعه ای از داده های آماری یا رایانه ای مربوط به یک پایگاه داده گفته می شود، که با هدف یکپارچه نمودن داده ها، محتویات آن در قالب یک جدول پایگاه داده یا یک ماتریس داده ای، تنظیم و مرتب شده و در آن ستون پایگاه داده، نشان دهنده یک متغیر خاص و هر ردیف به یکی از اعضای مجموعه داده موردنظر مرتبط می باشد.

۱-۴-۴. مطالعه موردی

یک آزمایشگاه تحقیقاتی هوش مصنوعی به نام اوپن آی که توسط ایلان ماسک تأسیس شده، آخرین نسخه سیستم پردازش زبان طبیعی مبتنی بر هوش مصنوعی به نام جی پی تی-۳ را راه اندازی کرده که می تواند زبان انسان را تقلید کند. جی پی تی-۳ می تواند برای تفسیر متن پیچیده، ایجاد هشدار، شناسایی متن تکراری، خودکارسازی مواردی که ممکن است برای کاربران انسانی آشکار نباشد، استفاده شود. همچنین در سایر حوزه ها می تواند از تحلیل گزارش های پزشکی تا تولید شعر نیز استفاده کرده و متنی تولید کند که اغلب از متن انسانی قابل تشخیص نیست. به عنوان مثال، یک وب نویس از جی پی تی-۳ برای نوشتن پست های وب نوشتی استفاده کرده و عملکرد فوق العاده ای را در هرکینوز^۹، رقم زده است.

با وجود موارد ذکر شده خروجی ها از عجیب تا عاقلانه

که فرصت کشف، استفاده و کمک به ارتقای پیشرفته ترین مدل ها و مجموعه داده ها را به کاربران می دهد. این بن سازه از طیف گسترده ای از مدل ها میزبانی می کند و بیش از ده هزار مورد آن در دسترس محققان قرار دارد. مدل های موجود در هاگینگ فیس هاب^۵ به مبدل ها یا حتی پردازش زبان طبیعی محدود نیستند. به عنوان نمونه، مدل های مربوط به گفتار، بینایی ماشین و غیره نیز وجود دارند. هرکدام از این مدل ها در قالب یک منبع گیت^۶ میزبانی می شوند. اشتراک گذاری مدل در هاب به معنای این است که عموم کاربران می توانند از آن استفاده کنند. دیگر کاربران مجبور نیستند مدل را آموزش دهند و کار اشتراک و استفاده بسیار آسان شده است. اشتراک گذاری و استفاده از مدل ها در هاب کاملاً رایگان است. البته اگر قصد اشتراک گذاری مدل های خصوصی مطرح باشد، می توان از طرح های پولی استفاده نمود. خط لوله^۷ اصلی ترین شیء در کتابخانه مبدل ها است. خط لوله مدل را به مراحل پیش پردازش متصل کرده و امکان می دهد هر متنی را به طور مستقیم وارد و پاسخی قابل قبول دریافت نمود. خط لوله به صورت پیش فرض یک مدل خاص و آموزش دیده را که برای تحلیل احساسات به زبان انگلیسی تنظیم شده است را انتخاب می کند.

۱-۴-۲. سوگیری و محدودیت ها در مدل ها

مدل ها علی رغم این که ابزارهای قدرتمندی هستند دارای سوگیری^۸ و محدودیت هم هستند. بزرگترین محدودیت این است که برای انجام پیش آموزش روی حجم زیادی از داده ها، اغلب باید تمام محتوای ممکن شامل بهترین و بدترین موارد موجود در اینترنت جستجو و به کار گرفته شوند. اگرچه به ندرت می توان مدلی مثل برت را در میان مدل های مبدل یافت اما حتی چنین مدلی نیز می تواند پیش بینی های دارای سوگیری ارایه دهد. هنگام

5-Hugging Face Hub

6-Git

7-Pipeline

8-Bias

وظیفه امانت‌داری اولیه در قبال بشریت است و شرکت برای برنده شدن در مسابقه هوش مصنوعی، ایمنی را به خطر نمی‌اندازد.

۲. سوابق پژوهش

طی بیش از صدسال گذشته بسیاری از محققان روی انتقال توالی و مدل‌سازی زبانی کارکرده‌اند و ماشین‌ها به‌تدریج یاد گرفتند که چگونه توالی‌های احتمالی کلمات را پیش‌بینی کنند. برای استناد به همه مواردی که این اتفاق را انجام داده است، به‌اندازه یک کتاب کامل زمان لازم است؛ لذا در این بخش مواردی بررسی می‌شوند که زمینه ورود به مباحث مبدل‌ها در پردازش زبان طبیعی فراهم شود. در ادامه، پیشینه پردازش زبان طبیعی که منجر به ظهور مبدل‌ها شد را مورد بررسی قرار داده و مشاهده می‌کنیم که چگونه مدل مبدل اختراع شده توسط تیم تحقیقاتی گوگل چندین دهه تحقیق، توسعه و پیاده‌سازی‌های پردازش زبان طبیعی را متحول کرده است.

در اوایل قرن بیستم، آندری مارکوف^{۱۲} مفهوم مقادیر تصادفی را معرفی و نظریه‌ای از فرآیندهای تصادفی ایجاد کرد. این نظریه در هوش مصنوعی با عناوین فرآیندهای تصمیم‌گیری مارکوف^{۱۳}، زنجیره‌های مارکوف^{۱۴} و فرآیندهای مارکوف^{۱۵} شناخته می‌شوند. در سال ۱۹۰۲، مارکوف نشان داد که می‌توان عنصر بعدی یک زنجیره را تنها با استفاده از آخرین عنصر پیش‌بینی نمود. در سال ۱۹۱۳، مارکوف عمل پیش‌بینی برای حروف آینده را با استفاده از توالی‌های گذشته زنجیره در یک مجموعه داده ۲۰,۰۰۰ حرفی به کار برد. با این که مارکوف رایانه‌ای نداشت اما موفق به اثبات نظریه‌ای شد، که هنوز هم در هوش مصنوعی مورد استفاده قرار می‌گیرد.

کلود شانون^{۱۶} در [۷]، نتایج تحقیقات خود را تحت عنوان «تئوری ریاضی ارتباطات شانون» در سال ۱۹۴۸ منتشر

متفاوت هستند و وجود تناقض در برخی موارد دور از انتظار نیست و نیاز به کارهای بیشتر وجود دارد. این مدل در هسته خود از مبدل استفاده می‌کند و برای توابع تولید متن، پاسخگویی به سؤال، تحلیل احساسات، شناسایی موجودیت‌های نامدار، ترجمه ماشینی و خلاصه کردن متن طراحی شده است. مدل جی‌پی‌تی-۳ شامل یادگیری عمیق با ۱۷۵ میلیارد پارامتر است و بر روی مجموعه داده‌های متنی بزرگ شامل صدها میلیارد کلمه آموزش داده شده است. برای آموزش محتوای متنی صفحات وب از سال ۲۰۱۶ تا ۲۰۱۹ مصرف شده که دارای ۴۵ ترابایت داده متنی فشرده است. مدل‌های قبلی جی‌پی‌تی-۱ از پایگاه داده ۵ گیگابایتی متشکل از محتوای متنی بیش از ۷۰۰۰ کتاب توسط دانشگاه تورنتو و ام‌آی‌تی آموزش دید و تعداد کلمات حدود یک میلیون کلمه بود. همچنین مدل جی‌پی‌تی-۲ برای آموزش محتوای متنی یک مجموعه داده بومی، متشکل از هشت میلیون صفحه وب با ۴۰ گیگابایت داده را دریافت کرد. همان‌طور که ذکر شد جی‌پی‌تی-۳ مهارت فوق‌العاده‌ای دارد و می‌تواند با دریافت یک جمله یا حتی چند کلمه، پنج پاراگراف کامل را نگارش کند. در حال حاضر این مدل انحصاراً در اختیار مالکان (اوپن‌ای‌آی^{۱۰} و مایکروسافت^{۱۱}) است و به سایرین استفاده محدود داده می‌شود تا از پی‌آمدهای ناخواسته مانند جعل عمیق جلوگیری شود. مدل جی‌پی‌تی-۴ در سال ۲۰۲۳ با تعداد پارامترهای بیشتر منتشر شده است. توجه به سرعت به‌عنوان عنصری حیاتی در مدل‌های پردازش زبان طبیعی موجب شد گوگل با استفاده از مبدل مدل برت را ایجاد کند که یکی دیگر از مدل‌های زبانی بسیار موفق است.

این مدل‌ها در همه حوزه‌ها خلاقیت زیادی را ایجاد می‌کنند و می‌توان از آن‌ها در همه زمینه‌ها استفاده نمود. ممکن است برای استفاده از مدل‌ها در حوزه صنعتی نیاز به امضای توافق‌نامه عدم افشای اطلاعات باشد. در منشورهای برخی شرکت‌های هوش مصنوعی بیان شده

12-Andrey Markov

13-Markov Decision Processes (MDPs)

14-Markov Chains

15-Markov Processes

16-Claude Shannon

10-Open AI

11-Microsoft

نمود. او هنگام ایجاد رویکرد خود در مدل‌سازی توالی، چندین بار به نظریه مارکوف اشاره می‌کند. شانون زمینه‌ای را برای یک مدل ارتباطی بر اساس کدگذاری منبع، ارسال و کدگذاری یا کدگذاری معنایی در دریافت انجام داد.

آلن تورینگ^{۱۷} در [۸] نتایج تحقیقات خود تحت عنوان «دستگاه محاسباتی و هوش» را در سال ۱۹۵۰ منتشر کرد. آلن تورینگ این مقاله را بر اساس هوش دستگاه بر روی دستگاه تورینگ فوق‌العاده موفق که پیام‌های آلمانی را رمزگشایی می‌کرد، پایه‌گذاری کرد. واژه هوش مصنوعی برای اولین بار توسط جان مک‌کارتی^{۱۸} در سال ۱۹۵۶ مورد استفاده قرار گرفت. با این حال، آلن تورینگ در دهه ۱۹۴۰ هوش مصنوعی را برای رمزگشایی پیام‌های رمزگذاری شده به زبان آلمانی پیاده‌سازی کرد.

از نخستین سیستم‌های پرسش و پاسخ سیستمی به نام بیس‌بال^{۱۹} بود که در سال ۱۹۶۱ برای پاسخ به سؤالات کاربران در رابطه با لیگ سالانه در آمریکا طراحی شد [۹]. برنامه رایانه‌ای مذکور قادر بود سؤالات کاربران در زمینه تاریخ و زمان برگزاری مسابقات و همچنین امتیازات کسب شده پاسخ دهد. کاربران سؤالات را با استفاده از کارت منگنه مطرح کرده و سیستم کلمات و اصطلاحات موجود را در یک لغت‌نامه جستجو و با تحلیل ساختاری و دستورزبانی پرسش پاسخ مناسب را تهیه می‌نمود.

جان هاپفیلد^{۲۰} در سال ۱۹۸۲، شبکه‌های عصبی بازگشتی^{۲۱} معروف به شبکه‌های هاپفیلد یا شبکه‌های عصبی «انجمنی» را معرفی کرد. هاپفیلد برای این کار از لیتل^{۲۲} که در سال ۱۹۷۴ «وجود حالات مداوم در مغز» را نوشته است الهام گرفته است. شبکه‌های عصبی بازگشتی تکامل یافتند و LSTMها پدیدار شد. شبکه‌های عصبی بازگشتی حالات مداوم یک دنباله را به‌طور مؤثر به خاطر دارد:

یان‌لی‌کان^{۲۳} در دهه ۱۹۸۰ شبکه عصبی هم‌آمیختی^{۲۴} را طراحی کرد. او این شبکه عصبی را برای توالی متن اعمال کرد به‌طوری‌که سپس این شبکه‌های عصبی به‌طور گسترده‌ای برای انتقال و مدل‌سازی توالی استفاده شدند. آن‌ها همچنین بر اساس «وضعیت‌های ماندگار» اطلاعات را بر اساس لایه به لایه جمع می‌کنند.

در سال ۱۹۹۰ یان‌لی‌کان لی‌نت-۵^{۲۵} را ایجاد کرد و منجر به بسیاری از مدل‌های شبکه‌های عصبی هم‌آمیختی شد که امروزه می‌شناسیم. پس از آن توسعه‌دهندگان برای تجزیه و تحلیل توالی‌های طولانی نیاز به افزایش قدرت رایانه داشتند و با استفاده از رایانه‌های قدرتمندتر، راه‌هایی برای بهینه‌سازی شیب^{۲۶} پیدا می‌کردند. در این زمان به نظر می‌رسید هیچ چیز دیگری نمی‌تواند برای پیشرفت بیشتر انجام شود.

پس از سی سال در دسامبر سال ۲۰۱۷، مبدل‌ها نوآوری حیرت‌آوری را به همراه آوردند و امتیازات چشمگیری را در مجموعه داده‌های استاندارد ایجاد کردند. وسوانی^{۲۷} و همکاران در سال ۲۰۱۷ نتایج کارهای خود را با عنوان «توجه، تمام چیزی که نیاز دارید» در [۱۰] منتشر کردند. آن‌ها معماری شبکه عصبی جدیدی به نام مبدل بر اساس سازوکار توجه ارائه نمودند و عملکرد مدل‌هایی مانند جی‌پی‌تی و برت با تکیه بر این معماری کاملاً بهتر از شبکه‌های قبلی بود. این رویکرد به‌طور گسترده از سایر رویکردها پیشی گرفت و تقریباً تمام مدل‌های اخیر مبتنی بر معماری مبدل هستند. این یافته‌ها در تیم تحقیقاتی گوگل و پروژه گوگل برین^{۲۸} انجام شد و سریع‌تر از معماری‌های قبلی، آموزش‌دیده و با کسب نتایج ارزیابی بالاتر به یک مؤلفه اصلی پردازش زبان طبیعی تبدیل شدند.

دولین^{۲۹} و همکاران در [۱۱] بیان کردند پیشرفت‌های تجربی اخیر ناشی از انتقال یادگیری با مدل‌های زبانی

23-Yann Le Cun

24-Convolutional Neural Network (CNN)

25-LENET-5

26-Gradients

27-Vaswani

28-Google Brain

29-Devlin

17-Alan Turing

18-John McCarthy

19-BASEBALL

20-John Hopfield

21-Recurrent Neural Networks (RNNs)

22-Little

مجموعه داده‌ها را انجام تحقیقات در زمینه درک مطلب فارسی و توسعه سیستم‌های پاسخ به سؤال فارسی بیان نموده‌اند. آزمایش‌های انجام شده امتیاز $74/8\%$ تطابق دقیق^{۳۳} و امتیاز $87/6\%$ در معیار اف^{۳۴} را نشان می‌دهد که می‌تواند به عنوان نتایج پایه برای تحقیقات بیشتر در مورد پرسش و پاسخ فارسی مورد استفاده قرار گیرد.

دوان^{۳۵} و همکاران در [۱۴] به بررسی مدل‌های زبان از پیش آموزش‌دیده پرداخته‌اند. بررسی در تحقیقات شامل معرفی تاریخ توسعه و دستاوردهای مدل‌های زبان از پیش آموزش‌دیده، ارایه چندین مدل زبان از پیش آموزش‌دیده مانند برت و مدل‌های چندحالت همراه با نمودار دانش برای تولید زبان طبیعی معرفی شده است. در نهایت بررسی آینده زبان‌های از پیش آموزش‌دیده انجام شده است. انتظار می‌رود این تحقیق یک راهنمای عملی برای درک، استفاده و توسعه زبان‌های از پیش آموزش‌دیده فراهم کند. سینگ^{۳۶} و همکاران در [۱۵] خلاصه‌ای جامع و مفصل از مدل‌های اصلی زبانی ارایه می‌دهند که منجر به ایده‌های نوآورانه در پردازش زبان طبیعی شده است. از زمان راه‌اندازی سازوکار توجه و معماری مبدل پیشرفت پردازش زبان طبیعی نمایی شده است. آن‌ها یک نقشه ذهنی سطح بالا از طبقه‌بندی مدل ارایه نموده‌اند. این طبقه‌بندی در درجه اول مبتنی بر معماری‌های مشتق شده از مبدل‌ها است که برای کارهای تخصصی مانند درک و تولید زبان، کاهش اندازه مدل، بازیابی اطلاعات، مدل‌سازی توالی طولانی و سایر تکنیک‌های کاهش مدل عمومی ایجاد شده است. مدل‌های زبان اخیر در درجه اول با دستیابی به عملکرد بالاتر به‌سوی منابع محاسباتی عظیم هدایت می‌شوند. بنابراین، مقیاس‌بندی مدل مسیر طبیعی در صنعت است. تلاش‌های قابل توجهی برای مهندسی مدل‌های با اندازه معقول و محاسبه توجه کارآمد برای سرعت بخشیدن به همگرایی مدل منجر به تأخیر کمتر در مدل‌ها انجام شده

نشان داده است که پیش آموزشی غنی و بدون نظارت بخش جدایی‌ناپذیر بسیاری از سیستم‌های درک زبان است. همچنین این نتایج وظایف پایین‌دست را قادر می‌سازند تا از معماری‌های تک‌جهتی عمیق بهره‌مند شوند و سهم اصلی ما تعمیم بیشتر این یافته‌ها به معماری‌های دوسویه عمیق است. مدل پیش‌آموزش‌دیده این امکان را فراهم می‌کند تا مجموعه وسیعی از وظایف پردازش زبان طبیعی انجام شود. آن‌ها در ادامه یک مدل چندزبانه برت معرفی کردند که از زبان فارسی نیز پشتیبانی می‌کند و بر پایه مبدل‌ها است. این مدل برخلاف مدل‌های ارایه شده زبانی اخیر برای آموزش دوسویه عمیق، از متن بدون برچسب با شرطی‌سازی مشترک در بخش چپ و راست در همه لایه‌ها طراحی شده است. مدل برت آموزش‌دیده می‌تواند تنها با یک لایه خروجی اضافی تنظیم شده تا مدل برای طیف وسیعی از وظایف به کار گرفته شود. وظایف موردنظر مانند پاسخ به سؤال و استنتاج زبان است.

دایی^{۳۷} و همکاران در [۱۲] نشان دادند که مدل زبانی از قبل آموزش داده شده برای بهبود بسیاری از کارهای پردازش زبان طبیعی مؤثر است. آن‌ها در این پایان‌نامه دو روش برای استفاده از داده‌های بدون برچسب به منظور بهبود شبکه تکرارشونده ارایه نمودند. اولین روش پیش‌بینی آنچه در دنباله بعدی می‌آید که یک مدل زبانی در پردازش زبان طبیعی است و روش دوم استفاده از کدگذار است که توالی ورودی را به یک توالی ورودی برداری قابل پیش‌بینی نگاشت می‌کند.

درویشی^{۳۸} و همکاران در [۱۳] یک مجموعه داده پاسخ به سؤال فارسی به نام پی‌کیو ای دی^{۳۹} از داده‌های درک مطلب در مقالات ویکی‌پدیا فارسی ارایه داده‌اند. پرسش شامل ۸۰,۰۰۰ سؤال همراه با پاسخ‌های است که ۲۵ درصد سؤالات غیرقابل پاسخ هستند. سؤالات غیرقابل پاسخ موجب می‌شوند مدل زمانی که پاسخ مناسبی وجود نداشته باشد، از پاسخ امتناع کند. آن‌ها هدف از انتشار این

33-Exact Match (EM)

34-F1- score

35-Duan

36-Singh

30- Dai

31- Darvishi

32- PQuAD

از این رو پارس‌برت در زبان فارسی می‌تواند گزینه مناسبی برای کارهای پردازش زبان طبیعی باشد.

طاهر^{۴۲} و همکاران در [۱۸] یک مدل به نام بهشتی‌ان‌ای‌آر^{۴۳} بر پایه برت به منظور شناسایی موجودیت‌های نامدار^{۴۴} ارائه نموده‌اند. این مدل در ارزیابی وظیفه ۷ از نسیوآرال^{۴۵} ۲۰۱۹ که در ارتباط با شناسایی موجودیت‌های نامدار در زبان فارسی است، رتبه دوم را کسب کرده است. نتایج شامل ۸۳/۵ و ۸۸/۴ در امتیاز اف ۱ از مجموعه داده کنل^{۴۶} در سطح عبارت و کلمه است.

بلبل‌زبان یک مدل تولید شعر فارسی است که در [۱۹] قابل دسترس برای عموم قرار دارد. این مدل با توجه با ۲۵۶ نویسه قبلی قادر به تولید شعر با درک قافیه و ریتم است اما شعر تولید شده در عین حال ممکن است دارای انسجام معنایی ضعیف باشد.

آلام^{۴۷} و همکاران در [۲۰] یک کتابخانه منبع باز پایتون به‌منظور استخراج اطلاعات تهدید سایبری از متن بدون ساختار از منابع ناهمگن اینترنت ارائه کردند که برای شناسایی موجودیت‌های نامدار در امنیت سایبری کاربرد دارد. این ابزار که سای‌ان‌ای‌آر^{۴۸} نام دارد، مبتنی بر مبدل‌ها است و برای استخراج موجودیت‌های مرتبط با امنیت سایبری شاخص‌های مختلف فنی و مدل‌های موجود عمومی را ترکیب می‌کند. سای‌ان‌ای‌آر رویدادها را از دامنه وسیعی از پیکره اطلاعات تهدید استخراج می‌کند. کاربران می‌توانند برای رفع نیازهایشان پیش‌بینی‌ها را از چندین رویکرد مختلف ترکیب کنند. این کتابخانه به‌طور عمومی منتشر شده است.

گوا^{۴۹} و همکاران در [۲۱] روشی را برای غلبه بر چالش استخراج اطلاعات حساس از داده‌های حجیم بدون ساختار ارائه کرده‌اند. این روش از استخراج داده‌های حساس

است. در این تحقیقات یک روش مؤثر برای مدل‌های بزرگ به‌منظور دستیابی به راندمان محاسباتی معرفی می‌شود.

سیار^{۳۷} و همکاران در [۱۶] برنامه‌ای به‌منظور شناسایی خطرات امنیت ملی در برنامه‌های وب ارائه نمودند. آن‌ها بیان کردند با توجه به این که برنامه‌های وب شامل طیف وسیعی از اطلاعات شخصی، مالی، دفاعی و سیاسی هستند ممکن است حادثه‌ای مانند ویکی‌لیکس^{۳۸} را رقم بزنند. از این رو ممکن است ثبت چنین اطلاعاتی به یک مشکل حیاتی تبدیل شده و به خطر امنیت ملی تبدیل شود. در تحقیقات آن‌ها یک راه‌حل برای این مشکل با ارائه مدلی جدید انجام شده که قادر به تشخیص درخواست‌های اچ‌تی‌پی‌جی‌جی‌جی عادی و درخواست‌های اچ‌تی‌پی‌جی‌جی‌جی غیرعادی است. این مدل از تکنیک‌های کدگذار پردازش زبان طبیعی دوسویه مبدل برت و روش‌های یادگیری عمیق استفاده می‌کند. نتایج تحقیقات در طبقه‌بندی درخواست‌های عادی و غیرعادی نشان داد روش پیشنهادی به نرخ موفقیت بیش از ۹۹/۹۸٪ و امتیاز اف ۱ بیش از ۹۸/۷۰٪ دست می‌یابد. همچنین زمان تشخیص به‌طور قابل‌توجهی کمتر از سایر روش‌های ارائه شده است.

فراهانی^{۴۰} و همکاران در [۱۷] یک برت یک‌زبانه برای زبان فارسی به نام پارس‌برت^{۴۱} ارائه کردند که عملکرد بهتری در مقایسه با سایر معماری‌ها و مدل‌های چندزبانه دارد. همچنین با توجه به محدود بودن مقادیر داده‌ها برای وظایف پردازش زبان طبیعی در زبان فارسی، یک مجموعه داده بزرگ برای کارهای مختلف و پیش آموزش مدل تدوین شده است. پارس برت در تمام مجموعه داده‌ها نمرات بالاتری کسب کرده و عملکرد آن فراتر از برت چندزبانه و سایر کارهای قبلی در تحلیل احساسات، طبقه‌بندی متن و شناسایی موجودیت‌های نامدار است. با توجه به این که اکثر مدل‌ها معمولاً بر انگلیسی متمرکز هستند، زبان‌های دیگر به مدل‌های چندزبانه با منابع محدود وابسته هستند.

42-Taher

43-Beheshti-NER

44-Named Entity Recognition

45-NSURL

46-CONLL

47-Alam

48-CyNER

49-Guo

37-Seyyar

38-WikiLeaks

39-HTTP

40-Farahani

41-ParsBERT

با توجه به این که نتایج تحقیقات نشان داده‌اند مبدل‌هایی که بر روی داده‌های مهندسی نرم‌افزار آموزش دیده‌اند دارای عملکرد بهتری در حوزه تخصصی هستند، ۱/۹۹۸ رکورد مربوط به داده‌های فناوری اطلاعات و ۶/۸۰۳ رکورد مربوط به داده‌های عمومی از سایر مجموعه داده‌ها برای افزایش درک عمومی مدل به مجموعه داده افزوده شده است. در مجموع رکوردهای مجموعه داده به ۱۲/۱۵۷ رسیده است.

چند مدل مختلف مبدل بر روی این مجموعه داده آموزش داده شده و عملکرد آنها نیز در بخش بعدی مقایسه شده است.

۳-۱. سیستم‌های پاسخ به سوال

یک روش مناسب با یک مدل متوسط اغلب نتایج کارآمدتری نسبت به یک روش ناقص با مدل عالی ایجاد می‌کند. به عنوان مثال، در یک پروژه تولید موشک و فضاپیما، پرسیدن یک سؤال از یک مدل هوش مصنوعی به معنای به دست آوردن یک پاسخ دقیق است. فرض می‌شود کاربر باید در مورد گزارش صد صفحه‌ای در مورد وضعیت نازل و محفظه احتراق مجدد یک موشک، سؤالی بپرسد. این سؤال می‌تواند خاص باشد، مانند این که "آیا وضعیت خنک‌کننده قابل اعتماد است یا خیر؟" یک مدل هوش مصنوعی قابل اعتماد، در پس زمینه به یک پایگاه دانش داده‌ها و قوانین متصل می‌شود تا پاسخ را بررسی کند. مدل هوش مصنوعی یک پاسخ به زبان طبیعی شفاف و قابل اعتماد و احتمالاً با صدای انسانی می‌تواند ایجاد نماید.

۳-۲. سخت‌افزار و محیط کدنویسی

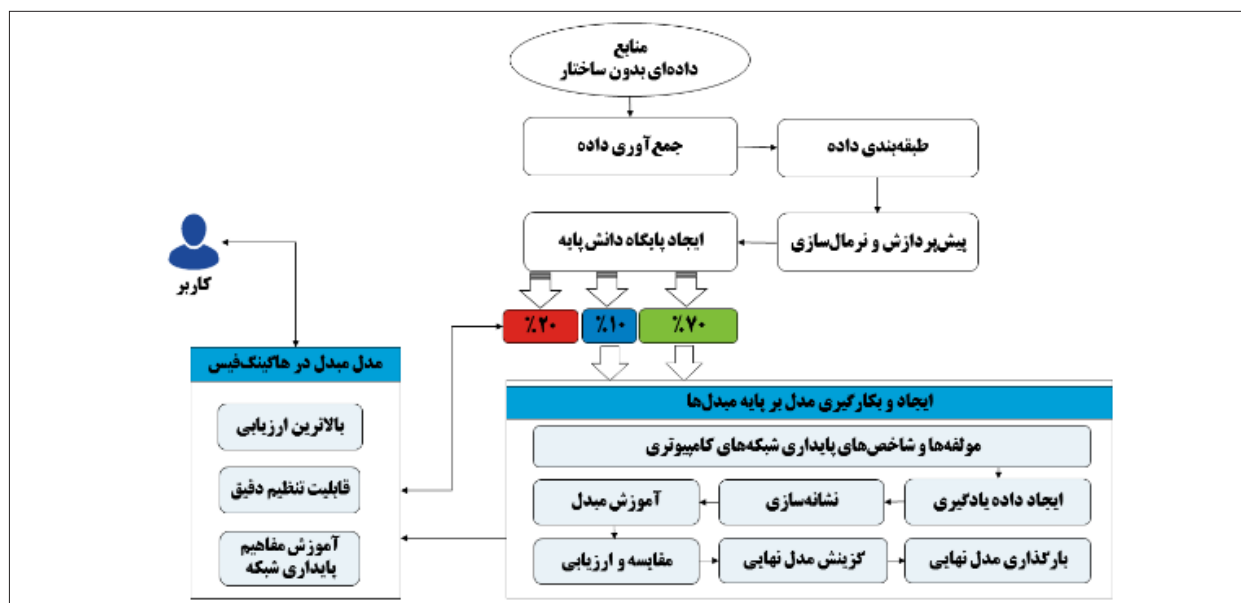
برای دسترسی به قدرت پردازشی مورد نیاز به منظور آموزش مدل از بن‌سازه گوگل کولب^{۵۳} استفاده شده است. برای استفاده از گوگل کولب نیاز به نصب برنامه و یا پیکربندی‌های پیچیده نیست و تنها کافی است یک نوت‌بوک کولب در مرورگر راه‌اندازی و اجرا شود.

مبتنی بر الگو (آی‌پی^{۵۰}، پست الکترونیکی، کلید و ...) و از مدل برت برای استخراج اطلاعات حساس استفاده نموده است. در واقع آخرین نسخه برت برای انجام تعبیه کلمات استفاده شده و اطلاعات حساس را با استفاده از شبکه عصبی BiLSTM و سازوکار توجه استخراج می‌کند. نتایج ارزیابی اف ۱ در این کار نشان داد که روش پیشنهادی با دقت ۹۹/۱۵ درصد بسیار موفق‌تر از روش‌های فردی تحلیل محتوا عمل کرده است. به علاوه حدود یک میلیون متن مورد تحلیل قرار گرفته است و ثابت شده که روش پیشنهادی می‌تواند اطلاعات حساس را از داده‌های بدون ساختار به طور مؤثر استخراج کند.

ژوانگ^{۵۱} و همکاران در [۲۲] مطرح کردند برای ورود مقدار زیادی داده‌ها به کامپیوتر به این امید که بتواند یاد بگیرد. ابتدا باید داده‌ها را آماده کنیم تا برای رایانه راحت باشد الگوها و استدلال‌ها را کشف کند. همچنین نیاز است با اضافه کردن فراداده‌های مربوط به مجموعه داده اصلی به این هدف دست یابیم. در این مقاله استخراج اطلاعات متنی ساخت یافته از متن نامرتب یا غیر ساخت یافته مورد توجه قرار گرفته و به استخراج اطلاعات بر اساس پردازش زبان طبیعی، سیستم بازیابی اطلاعات وب، استخراج روابط متنی و عملکرد پردازش خط لوله‌ای^{۵۲} پرداخته شده است.

۳. سیستم پیشنهادی

سیستم پیشنهادی یک مدل هوش مصنوعی است که از یک مبدل آموزش دیده در زمینه تحلیل داده‌های بدون ساختار در حوزه شبکه‌های کامپیوتری تشکیل شده است. آموزش مبدل بر روی سؤالات احصاء شده در رابطه با امنیت شبکه‌های کامپیوتری انجام شده و سیستم می‌تواند از طریق پایگاه دانش خود و یا با دریافت گزارش به صورت مستقیم، پاسخ سؤالات را از متن بدون ساختار استخراج نماید. سؤالات آموزشی در قالب یک مجموعه داده شامل ۳/۳۵۶ رکورد ایجاد شده است.



شکل (۱). روندنمای سیستم پیشنهادی

این منابع شامل وب، شبکه‌های اجتماعی، مجموعه داده‌های جمع‌آوری شده متنی و سایر موارد هستند. منابع برخط مانند وب توسط روش پیشنهادی جمع‌آوری و ذخیره می‌شوند تا در مراحل بعدی مورد استفاده قرار گیرند.

۳-۳-۲. جمع‌آوری داده‌ها

برخی از مجموعه داده‌ها آماده برای آموزش مدل وجود دارند. اما این مجموعه داده‌ها ایستا هستند و بدون ارایه روشی برای جمع‌آوری داده‌ها امکان به‌روزرسانی اخبار جدید در آن‌ها وجود ندارد. همچنین ممکن است این مجموعه داده‌ها تعداد محدودی از سبک‌های نوشتن و موضوعات را پوشش دهند. بنابراین، برای افزایش دامنه اطلاعات یادگیری روشی برای جمع‌آوری داده‌های جدید ارایه شده است. این روش بدون نیاز به کد نویسی و با استفاده از افزونه وب‌اسکرپر^{۵۶} در مرورگر گوگل کروم^{۵۷} انجام شده است. افزونه وب‌اسکرپر گوگل کروم امکان جمع‌آوری اطلاعات در بستر ابری^{۵۸} را داد. در جمع‌آوری اطلاعات در بستر ابری امکان استفاده از لیست پراکسی

کولب امکان استفاده از سخت‌افزارهای شتاب‌دهنده مانند جی‌پی‌یو^{۵۹} و تی‌پی‌یو^{۶۰} را فراهم می‌کند. استفاده از شتاب‌دهنده برای همیشه رایگان نیست و بعد از آموزش چند مدل شتاب‌دهنده برای مدتی از کاربر گرفته می‌شود و برای پردازش بیشتر کولب‌پرو تهیه شده است. در ایران شرکت‌ها و مؤسساتی وجود دارند که با دور زدن تحریم اکانت کولب‌پرو را روی حساب ایمیل افراد شارژ می‌کنند. برای فضای ذخیره‌سازی از گوگل درایو استفاده شده و مجموعه داده‌ها و نتایج جمع‌آوری اطلاعات بر روی این محیط قرار داده شده است. گوگل درایو ۱۵ گیگابایت فضا روی بستر ابری در اختیار هر کاربر قرار می‌دهد که بر انجام تحقیقات اولیه و آموزش مدل کافی است.

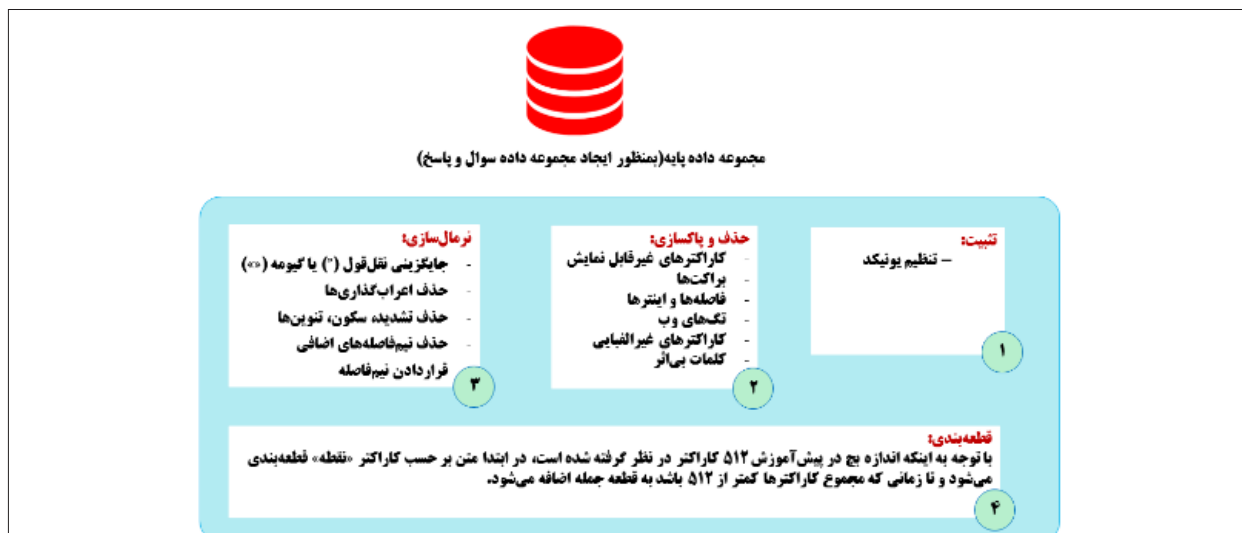
۳-۳-۳. روندنمای سیستم پیشنهادی

روندنمای سیستم پیشنهادی در شکل (۱) ترسیم شده است. همان‌طور که در روندنما مشخص است کاربر با ارایه درخواست به مدل هوش مصنوعی پاسخ سوالات را از پایگاه دانش داده‌های بدون ساختار استخراج می‌کند.

۳-۳-۱. منابع داده‌ای بدون ساختار

56-Web scraper
57-Google chrome
58-Cloud

54-GPU
55-TPU



شکل (۲). پیش پردازش و نرمال سازی داده ها

استفاده از مدل های مختلف بر روی ۹,۶۶۶ عنوان خبر در سه دوره^{۹۹} آموزش دیده است. طبقه بندی به صورت دو رده "مرتبط" و "غیر مرتبط" صورت می گیرد و اخباری که در رده "غیر مرتبط" قرار گیرند از فرآیندهای بعدی سیستم کنار گذاشته می شوند. نتایج ارزیابی این مدل (که روش و نتیجه آن به طور کامل در فصل پنجم تشریح شده است) نشان داده است کارایی این مدل بیش از ۹۵٪ برای اخبار غیرمرتبط و بیش از ۸۰٪ برای اخبار مرتبط است.

۳-۳-۴. پیش پردازش و نرمال سازی داده ها

در ابتدا به هر خبر یک برچسب تخصیص داده می شود. این برچسب بعدها برای یافتن سؤالات مربوطه کاربرد دارد. در مجموعه داده ایجاد شده در حدود ۴۲ نوع برچسب مختلف ایجاد شده است. پس از جمع آوری داده ها و قبل از ورود به مرحله پیش آموزش، یک سلسله مراتب عظیم از مراحل پردازشی از جمله تثبیت یونیکد، حذف و پاک سازی، نرمال سازی و قطعه بندی انجام می شود. مراحل انجام شده در شکل (۲) قابل مشاهده هستند.

۳-۳-۵. ایجاد پایگاه داده پایه

داده های جمع آوری شده پس از پیش پردازش و

وجود دارد. روش استفاده از بستر ابری به این شکل است که ابتدا نقشه وبگاه را در افزونه تنظیم نموده و سپس آن را در رابط وب بستر ابری وارد کرده و جمع آوری اطلاعات در بستر ابری آغاز می شود. برای استفاده بیش از حد تعریف شده نیاز به پرداخت هزینه و تهیه حساب کاربری می باشد.

در مرحله جمع آوری اطلاعات حدود ۲۴,۰۰۰ متن خبر از وبگاه های اعتماد آنلاین، فردانیوز، خبرگزاری فارس و ... جمع آوری شده است. تمام این اخبار در حوزه سایبری نیستند و نیاز به طبقه بندی دارند. همچنین برخی از مجموعه داده های آماده مانند عصر ایران، خبرگزاری تسنیم، ایسنا و ... شناسایی و مورد استفاده قرار گرفت.

۳-۳-۳. طبقه بندی داده ها

در این مرحله یک طبقه بندی برای تفکیک اخبار مرتبط و اخبار غیر مرتبط با موضوع تحقیق کد نویسی شده است. منظور از اخبار غیرمرتبط اخباری هستند که در حوزه سایبری نیستند و آموزش آنها به مدل، خارج از حوزه این تحقیق است. به عنوان مثال از اخبار غیرمرتبط می توان به اخبار مذاکرات هسته ای اشاره نمود. این طبقه بندی که بر مبنای مبدل ها کد نویسی شده است با

جدول (۲). شواهد امنیت در داده‌های بدون ساختار

مؤلفه	سؤال	مؤلفه	سؤال
امنیت زیرساخت و ارتباطات	آیا حادثه ناشی از اقدام سایبری بوده است؟	حکومت	گروه رخنه‌گری به چه نامی است؟
	حادثه منجر به چه پیامدهایی شد؟		رنه‌گر موفق به انجام چه کاری شدند؟
	میزان خسارات چگونه است؟		گروه هکری چه چیزی اعلام کرد؟
	اهمیت زیرساخت مورد حمله چیست؟		فعالیت قبلی این گروه چه بوده است؟
	درباره زیرساخت چه توضیحاتی وجود دارد؟		منتسب به کجا شده است؟
	دشواری بازگردانی زیرساخت در چیست؟		پیچیدگی و گستردگی حادثه چگونه است؟
	چه بخشی حادثه دیده است؟		چه توضیحاتی در مورد گروه هکری وجود دارد؟
	حادثه بطور اتفاق افتاده است؟		از چه بدافزار یا روشی استفاده شد؟
	چه اطلاعاتی به دست رخنه‌گر افتاد؟		عملکرد بدافزار یا روش مورد استفاده چگونه است؟
	چه اطلاعاتی از بین رفته است؟		آیا حادثه از طریق عوامل نفوذی صحت دارد؟
امنیت اطلاعات	رنه‌گرها با اطلاعات سرقت رفته چه کاری انجام دادند؟		چه تعداد افراد و مشتریان تحت تأثیر قرار گرفتند؟
	اطلاعات با چه قیمتی به فروش گذاشته شده است؟		اکنون وضعیت خدمات چگونه است؟
	هدف حمله چه بوده است؟		واکنش متولیان چگونه بوده است؟
	با چه شعاری اقدام به نفوذ شده است؟		واکنش بخش فنی نسبت به حمله چه بوده است؟
حاکمیت	آیا عامل حمله دستگیر شده است؟		قبلاً چه هشدارهای صادر شده بود؟
	واکنش مرکز افتا چگونه بوده است؟		چه اقدامات جبرانی صورت گرفته است؟
	واکنش سازمان پدافند غیر عامل چگونه بوده است؟		چه حادثه‌ای رخ داده است؟
	واکنش مرکز ملی فضای مجازی چگونه بوده است؟		چه زمانی حادثه اتفاق افتاد؟
	واکنش مقامات چه بوده است؟		حادثه مقارن با چه زمانی بوده است؟
	واکنش‌ها در خارج کشور نسبت به حادثه چه بوده؟		در وقایع مرتبط شاهد چه چیزی بودیم؟
کاربران و رسانه	گزارش‌های فضای مجازی و رسانه چه بود؟		
	گزارش کارکنان و مشتریان چه بوده است؟		

شواهد برای آموزش به مدل هوش مصنوعی استفاده می‌شود تا مدل بتواند جواب را برای سایر موارد مشابه ارائه نماید. مؤلفه‌ها شامل پنج بخش به شرح زیر هستند:

- امنیت زیرساخت و ارتباطات
- امنیت اطلاعات
- دفاع و امنیت شبکه
- کاربران و رسانه‌ها
- حاکمیت (قانون‌گذار، سازمان‌ها و مراکز، مسئولین)

در جدول (۲) شواهدی در ارتباط با امنیت شبکه که از داده‌های بدون ساختار قابل احصاء است، فهرست شده است.

۳-۳-۷. ایجاد مجموعه داده (مجموعه داده هوشفا)

با توجه به این که هوش مصنوعی داده‌محور است

نرمال‌سازی در پایگاه داده ذخیره می‌شوند تا در مراحل بعدی به منظور یافتن پاسخ سوالات مورد استفاده قرار گیرند. با توجه به حجم زیاد اطلاعات موجود برای یافتن پاسخ هر سوال قطعاً نیاز است موضوعات مرتبط تشخیص داده شود. برای این منظور هر رکورد از پایگاه داده دارای یک برچسب است که می‌تواند به منظور پالایش اطلاعات مورد نیاز استفاده شود. در مرحله بعد داده‌ها به نسبت ۷۰، ۲۰ و ۱۰ تقسیم شده‌اند. از ۷۰ درصد داده‌ها برای آموزش مدل و ۱۰ درصد برای اعتبارسنجی و ۲۰ درصد باقیمانده برای آزمون مدل استفاده شده است.

۳-۳-۶. شواهد امنیت شبکه در منابع آزاد

برخی از شواهد امنیت شبکه که از منابع آزاد قابل احصاء است در این بخش استخراج شده است. از این

تصحیح می‌شوند تا این که تعداد رکوردها به حد مورد نیاز برسد. در نهایت مجموعه داده مورد استفاده برای آموزش مدل پرسش و پاسخ به ترتیب زیر ایجاد شده است.

جدول (۳). رکوردهای موضوعی مجموعه داده هوشفا

امنیت شبکه	فناوری اطلاعات	عمومی	مجموع رکوردها
۳/۳۵۶	۱/۹۹۸	۶/۸۰۳	۱۲/۱۵۷

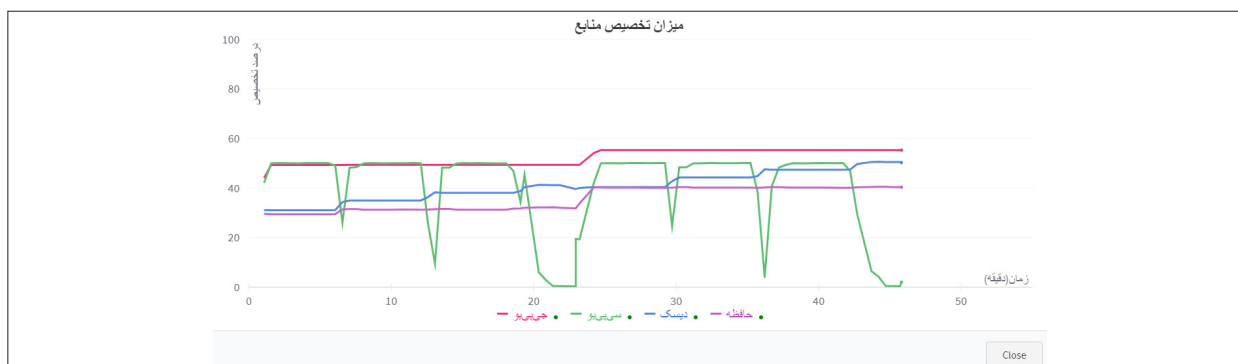
رکوردهای پایه شبکه مربوط به پرسش و پاسخ در زمینه امنیت شبکه‌های کامپیوتری است که با استفاده از روش پیشنهادی تولید شده است. رکوردهای فناوری اطلاعات مربوط به سایر مجموعه داده‌ها است که در حوزه فناوری و اطلاعات به صورت پرسش و پاسخ از قبل توسط سایر محققان تولید شده‌اند و در ترکیب مجموعه داده قرار گرفته است. رکوردهای عمومی مربوط به سایر مجموعه داده‌ها است که در حوزه عمومی به صورت پرسش و پاسخ از قبل توسط سایر محققان تولید شده‌اند و در ترکیب مجموعه داده قرار گرفته‌اند.

۳-۸. آموزش مدل با مجموعه داده هوشفا

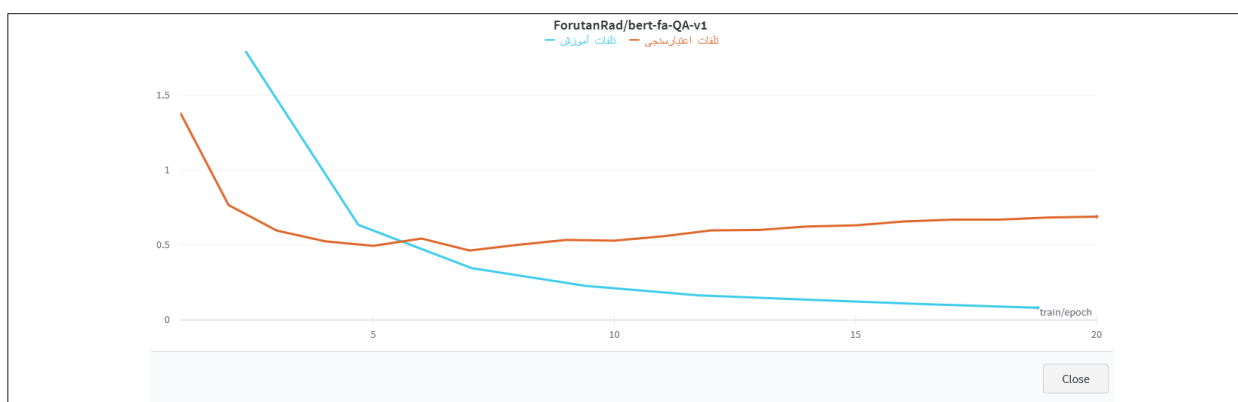
پس از ایجاد مجموعه داده مناسب یادگیری با زبان برنامه‌نویسی پایتون^{۶۱} اقدام به آموزش مدل مبدل شده است. آموزش مدل بر روی ۷۰ درصد از داده‌های موجود انجام شده است و از ۱۰ درصد داده‌ها برای اعتبارسنجی استفاده شده است. همچنین ۲۰ درصد داده‌ها نیز برای آزمون در نظر گرفته شده است. در این بخش مدل‌های هوش مصنوعی با استفاده از مجموعه داده‌ای که توسط مؤلفه‌ها و شاخص‌های امنیت شبکه‌های کامپیوتری در بخش‌های قبل ایجاد شده است آموزش داده می‌شوند. پس از آموزش مدل‌های مختلف (سه مدل در این تحقیق) ارزیابی عملکرد مدل‌ها با یکدیگر صورت می‌گیرد و بهترین مدل مشخص می‌گردد. در ابتدای آموزش و در دوره‌های

معمولاً محققان بخشی از وقت خود را صرف تهیه مجموعه داده مناسب برای مدل‌های مبدل می‌کنند. یک مدل آموزش‌دیده مانند شخصی که زبان آموخته رفتار خواهد کرد. داده باید با طی مراحل، اطلاعات جدیدی به مجموعه داده آموزش اضافه کنند و در نهایت مجموعه داده‌های آموزش می‌تواند به یک پایگاه دانش تبدیل شوند. رویکرد درست آموزش مبدل، پیاده‌سازی آن و ارزیابی نتایج قابل قبول در گرو دسترسی به مجموعه داده مناسب است. در این بخش یک مجموعه داده با استفاده از مؤلفه‌ها و شاخص‌های احصاء شده ایجاد شده است که از این مجموعه داده برای آموزش و ارزیابی مدل هوش مصنوعی استفاده می‌شود. با ایجاد این مجموعه داده مدل قادر است پاسخ سوالاتی که در حوزه امنیت شبکه‌های کامپیوتری احصاء شده است را از متن بدون ساختار، احصاء نموده و به کاربر نمایش دهد.

با توجه به این که انجمن‌های جمع‌سپاری^{۶۲} برای تولید مجموعه داده فارسی وجود ندارد و از طرفی ایجاد مجموعه داده به صورت دستی بسیار زمان‌بر است در این بخش روشی پیشنهاد می‌شود که بتوان سرعت ایجاد مجموعه داده را تا چند برابر افزایش داد. روش پیشنهادی به این صورت است که ابتدا یک مجموعه ۱۵۰ رکوردی از پرسش و پاسخ تولید می‌شود و سپس یک مدل با این مجموعه داده آموزش داده می‌شود. در ادامه مدل با استفاده از مجموعه داده پایه پرسش و پاسخ جدید تولید می‌کند. ملاک این که یک رکورد برای ذخیره‌سازی انتخاب شود، بالا بودن نمره است. همچنین رکوردهایی که نمره بسیار پایینی دارند به عنوان سوالات بدون پاسخ در مجموعه داده ذخیره می‌شوند تا در مراحل بعدی آموزش، مدل تشخیص دهد چه سوالاتی را نباید پاسخ بدهد. در ادامه خروجی تولید شده به صورت دستی تصحیح شده و مواردی که دارای خطا هستند، حذف می‌شوند. این فرآیند به طور نمایی ادامه می‌یابد و در گام دوم ۳۰۰ رکورد و در گام سوم ۶۰۰ رکورد تولید و



شکل (۳). منابع (مدل ForutanRad/bert-fa-QA-v1)



شکل (۴). دوره‌ها مدل (ForutanRad/bert-fa-QA-v1)

همان‌طور که در شکل (۴) نمایش داده شده است، تعداد دوره‌های آموزش بر روی ۲۰ دوره تنظیم شده است. با توجه به تقاطع نمودار تلفات اعتبارسنجی و تلفات آموزش در دوره ششم نتیجه‌گیری می‌شود که مدل در آموزش شش دوره دارای بهترین نتیجه می‌باشد. لذا مدل نهایی بر روی آموزش بر مبنای شش دوره انجام می‌شود.

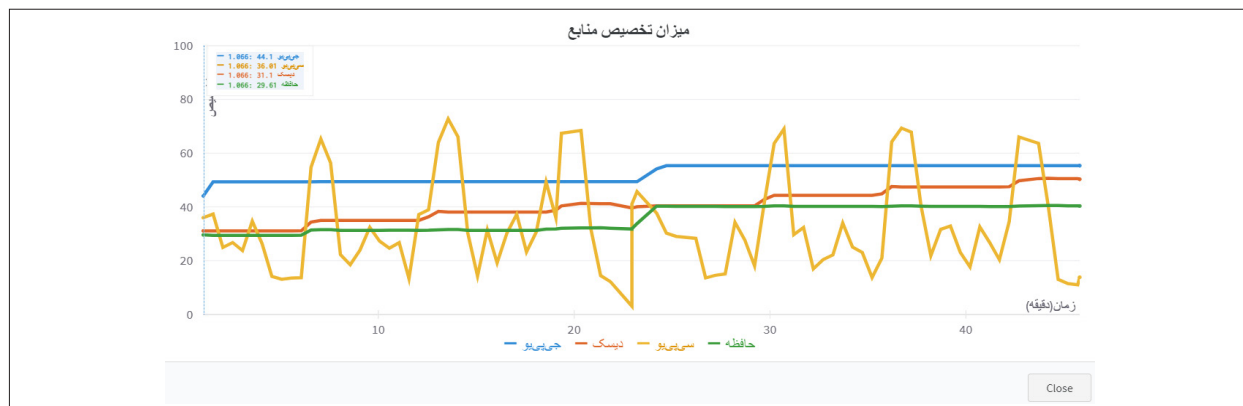
● آموزش مدل SajjadAyoubi_bert-base-fa-qa
یک مدل فارسی با 162M پارامتر است. در آموزش این مدل از سخت‌افزار شتاب‌دهنده گرافیکی با مدل NVIDIA A100 (40 GB) دارای قدرت پردازشی ۱۹٫۵ TeraFlops استفاده شده است. با قدرت پردازشی مذکور آموزش مدل در زمان ۴۵ دقیقه انجام شده است. میزان مشغولیت حافظه، دیسک، سی‌پی‌یو و جی‌پی‌یو در شکل (۵) ترسیم شده است.

پایین، نرخ تلفات و اعتبارسنجی فاصله زیادی داشتند که با انجام اصلاحاتی از قبیل تنظیم نرخ یادگیری، استفاده از سایر توابع بهینه‌ساز، رفع نوفه^{۶۲} داده و سایر روش‌ها فاصله نرخ تلفات و اعتبارسنجی به حد قابل مناسب رسیده است.

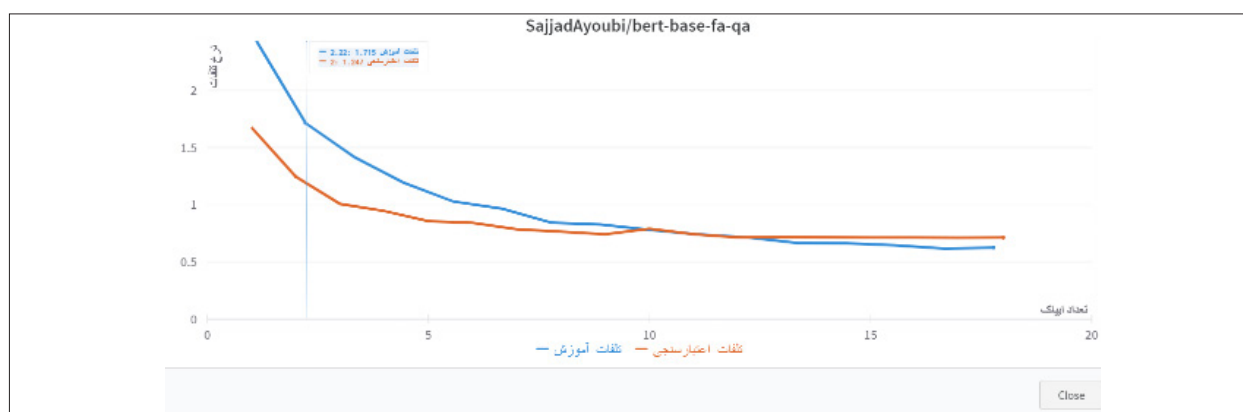
● آموزش مدل ForutanRad/bert-fa-QA-v1

این مدل یک نمونه فارسی از مدل برت است. در آموزش این مدل از سخت‌افزار شتاب‌دهنده گرافیکی با مدل NVIDIA A100 (40 GB) دارای قدرت پردازشی 19.5 TeraFlops استفاده شده است. با قدرت پردازشی مذکور آموزش مدل در زمان ۴۵ دقیقه انجام شده است. میزان مشغولیت حافظه، دیسک، سی‌پی‌یو و جی‌پی‌یو در شکل (۳) ترسیم شده است.

همان‌طور که در شکل (۳) قابل مشاهده می‌باشد در این آموزش از جی‌پی‌یو استفاده شده است و میزان مشغولیت جی‌پی‌یو بیشتر از سایر منابع می‌باشد.



شکل (۵). منابع (مدل SajjadAyoubi_bert-base-fa-qa)



شکل (۶). دوره‌ها مدل (SajjadAyoubi_bert-base-fa-qa)

میزان مشغولیت جی‌پی‌یو نزدیک به ۱۰۰ درصد بوده است و در شکل (۷) ترسیم شده است.

همان‌طور که در شکل (۸) نمایش داده شده است، تعداد دوره‌های آموزش بر روی ۲۰ دوره تنظیم شده است.

با توجه به تقاطع نمودار تلفات اعتبارسنجی و تلفات آموزش در دوره ۱۲ نتیجه‌گیری می‌شود که مدل در آموزش ۱۲ دوره دارای بهترین نتیجه می‌باشد. لذا مدل نهایی بر روی آموزش بر مبنای ۱۲ دوره انجام می‌شود.

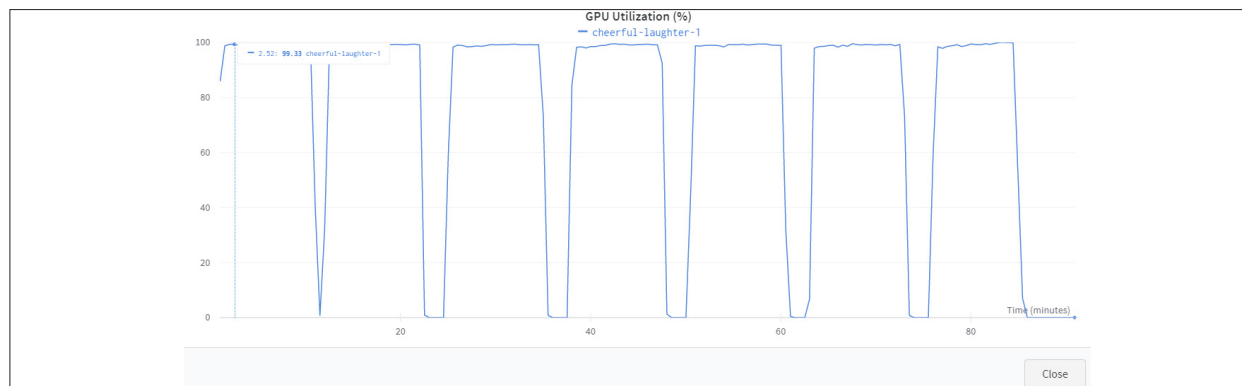
۳-۳-۹- کاربری سیستم پیشنهادی

بسته به این که کاربر، یک محقق، تحلیلگر و یا دانشمند داده باشد، در مواقعی نیاز دارد برای یافتن اطلاعات موردنظر اقیانوسی از اسناد را بررسی کند. پاسخ به سؤال شامل سؤالاتی است که پاسخ آن‌ها می‌تواند به عنوان یک متن در سند وجود داشته باشد. مدل آموزش داده شده، به منظور پاسخ به سوال بر

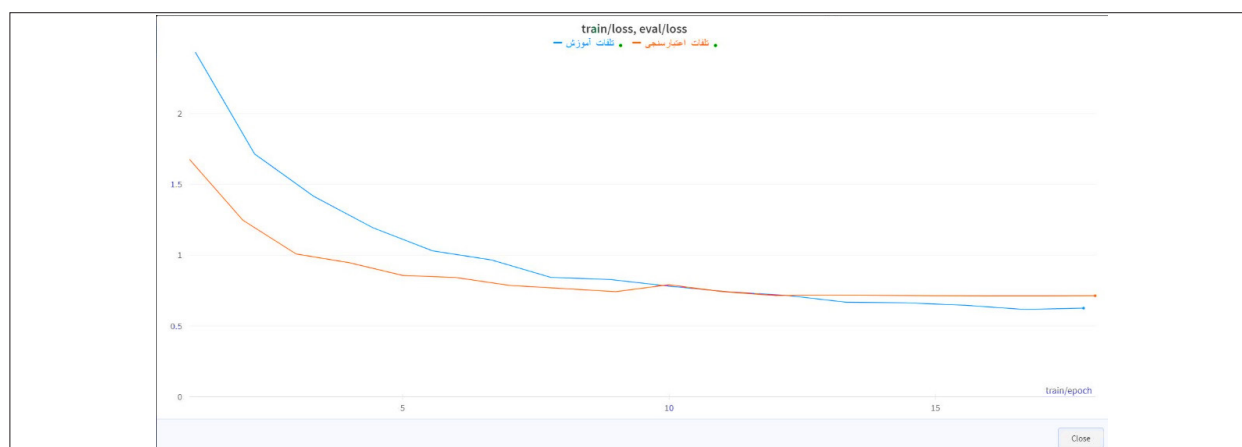
همان‌طور که در شکل (۵) قابل مشاهده می‌باشد در این آموزش از جی‌پی‌یو استفاده شده است و پیوستگی کارکرد و میزان مشغولیت جی‌پی‌یو بیشتر از سایر منابع می‌باشد. همان‌طور که در شکل (۶) نمایش داده شده است، تعداد دوره‌های آموزش بر روی ۱۸ دوره تنظیم شده است. با توجه به تقاطع نمودار تلفات اعتبارسنجی و تلفات آموزش در دوره دهم نتیجه‌گیری می‌شود که مدل در آموزش ده دوره دارای بهترین نتیجه می‌باشد. لذا مدل نهایی بر روی آموزش بر مبنای ده دوره انجام می‌شود.

● آموزش مدل SajjadAyoubi_xlm-roberta-large-fa-qa

یک مدل فارسی بزرگ با 558M پارامتر است. در آموزش این مدل از سخت‌افزار شتاب‌دهنده گرافیکی با مدل NVIDIA A100 (40 GB) دارای قدرت پردازشی 19.5 TeraFlops استفاده شده است. با قدرت پردازشی مذکور آموزش مدل در زمان ۹۰ دقیقه انجام شده است.



شکل (۷). منابع (مدل SajjadAyoubi_xlm-roberta-large-fa-qa)



شکل (۸). دوره‌ها (مدل SajjadAyoubi_xlm-roberta-large-fa-qa)

۴. ارزیابی سیستم پیشنهادی

ارزیابی‌های استاندارد می‌توانند روشی اطمینان‌بخش برای ارزیابی و مقایسه کارایی مدل‌های مبدل ارائه دهند. همان‌طور که در بخش سوم اشاره شده است ۲۰ درصد داده‌ها به عنوان داده‌های آزمون در نظر گرفته شده است.

۴-۱. معیارهای ارزیابی

مقایسه یک مدل مبدل با سایر مدل‌ها بدون معیارهای اندازه‌گیری استاندارد، غیرممکن است. از این رو در این پژوهش از میان انبوه معیارهایی که برای ارزیابی مدل‌ها رایج شده، از یک یا چند معیار دقت^{۶۳}، صحت^{۶۴}، پوشش^{۶۵} اف^{۶۶}، تطبیق دامنه دقیق^{۶۷} استفاده شده است.

63-Accuracy
64-Precision
65-Recall
66-F1
67-Exact_match

روی هاگینگ‌فیس بارگذاری شده است و امکان استفاده از آن توسط خط لوله وجود دارد. در ادامه مثالی از کاربری مدل ارائه شده است:

توسط پرسمان زیر می‌توان پاسخ را از متون بدون ساختار به‌دست آورد:

```
nlp_qa(context=sequence, question=
'اکنون وضعیت خدمات چگونه است؟')
```

خروجی به‌صورت زیر نمایش داده می‌شود:

```
{ 'answer': '۳۱۰۰ جایگاه سوخت‌رسانی به سامانه
هوشمند وصل شده‌اند اما ۹۰۰ جایگاه همچنان امکان
فروش بنزین با نرخ آزاد دارند،
'end':66,'score':0.918201259751591,'start':55}
```

دقت و فراخوانی به ترتیب زیر استفاده می‌کند:

$$F1 - score = 2 * (precision * recall) / (precision + recall)$$

$$P = \frac{T_p}{T_p + F_p}$$

$$R = \frac{T_p}{T_p + F_n}$$

از این‌رو می‌توان نمره اف ۱ را به‌عنوان میانگین هارمونیک^{۶۸} (میانگین حسابی) دقت (P) و فراخوان (R) مشاهده کرد:

$$F1Score = \frac{PxR}{P + R}$$

۴-۱-۵. معیار دامنه تطبیق دقیق

روش مورد استفاده این معیار در ارزیابی بسیار ساده است. در مقایسه بین رشته‌ها اگر دقیقاً همان رشته پیش‌بینی شده باشد نمره یک و در غیر این صورت نمره صفر لحاظ می‌گردد.

۴-۲. استعدادیابی مدل‌ها

با توجه به تعدد مدل‌های موجود، دور از انتظار نیست که برخی از آنها بهتر از سایرین در آموزش باشند. از این رو نیاز است توسط آزمون بهترین مدل‌ها انتخاب شوند. برای انتخاب مناسب‌ترین مدل‌ها با استفاده از مجموعه داده آزمون هوشفا عملکرد مدل مورد آزمون قرار گرفته تا استعداد مدل مورد سنجش قرار گیرد.

از میان مدل‌های بررسی شده مشخص است که سه مدل از استعداد خوبی به منظور آموزش در مجموعه داده هوشفا دارند.

۴-۳. نتایج ارزیابی مدل‌ها پس از آموزش با هوشفا

ارزیابی مجموعه داده با استفاده از کدها و کتابخانه‌های زبان برنامه‌نویسی پایتون انجام شده است. داده‌هایی که به

جدول (۴). استعدادیابی مدل‌ها

داده	نام مدل	تطبیق دقیق	اف ۱
مجموعه داده هوشفا	HooshvareLab_bert-fa-base-uncased	۱/۸۶	۲۰/۸۱
	SajjadAyoubi_bert-base-fa-qa	۴۷/۸۸	۲۱/۶۷
	SajjadAyoubi_xlm-roberta-large-fa-qa	۲۸/۱۵	۵۴/۵۷
	ForutanRad/bert-fa-QA-v1	۱۸/۷۱	۴۵/۹۸
	Koala_xlm-roberta-large-fa	۰/۰	۲۱/۲۶
	Koala_bert-large-uncased-fa	۱/۸۶	۲۶/۲۹
	xlm-roberta-base	۱۷/۳۹	۷/۲۳
	xlm-roberta-large	۰/۰۵	۲۲/۰۷
	bert-base-multilingual-cased	۰/۶۵	۲۴/۱۵

۴-۱-۱. معیار دقت

این معیار نمایانگر میزان پیش‌بینی صحیح خروجی توسط مدل است و می‌توان توسط آن حد آموزش صحیح مدل را مشاهده نمود. این معیار میزان نتایج صحیح را نسبت به کل موارد مورد سنجش قرار می‌دهد. در این معیار اگر نتیجه برای زیرمجموعه مثبت باشد مقدار یک و اگر نادرست باشد مقدار صفر بازگردانده می‌شود.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN}$$

۴-۲-۲. معیار صحت

معیار صحت معیاری است که مشخص می‌کند الگوریتم چند درصد مثبت‌ها را درست پیش‌بینی کرده است. فرمول محاسبه این معیار به صورت زیر است:

$$Precision = \frac{T_p}{T_p + F_p}$$

۴-۳-۳. معیار پوشش

معیار پوشش برای محاسبه میزان پوشش روی کل داده‌ها استفاده می‌شود. فرمول محاسبه این معیار به صورت زیر است:

$$Recall = \frac{T_p}{T_p + F_N}$$

۴-۱-۴. معیار اف ۱

معیار اف ۱ یک رویکرد انعطاف‌پذیرتر را معرفی می‌کند که می‌تواند در مواجهه با مجموعه داده‌های حاوی توزیع رده ناهموار مفید واقع شود. معیار اف ۱ از میانگین وزنی

۵. بحث و نتیجه‌گیری

با توجه به چالش رشد سریع داده‌های بدون ساختار، در این پژوهش سیستمی بر پایه هوش مصنوعی به منظور تحلیل داده‌های بدون ساختار با تمرکز بر حوزه امنیت شبکه‌های کامپیوتری ارائه شده است. سیستم مذکور قادر به یادگیری و پاسخ به سوالات در حوزه‌های مختلف می‌باشد و در این پژوهش بر روی امنیت شبکه‌های کامپیوتری تمرکز شده است. سیستم مذکور از مقادیر زیادی داده‌های متنی برای آموزش استفاده نموده و با ایجاد درک از نقش کلمه در جمله (مشابه مغز انسان)، قادر به یادگیری، و پاسخ به سوالات است.

برای طراحی و پیاده‌سازی سیستم مذکور در ابتدا جمع‌آوری ماشینی داده‌های بدون ساختار از منابع وب انجام شده است و یک مدل طبقه‌بندی برای شناسایی داده‌های مرتبط با حوزه امنیت شبکه‌های کامپیوتری آموزش داده شده است. پس از انتخاب برخی موارد مرتبط به امنیت شبکه‌های کامپیوتری، تشکیل مجموعه داده سفارشی به نام هوشفا توسط پژوهشگر انجام شده است. این مجموعه داده سفارشی، شامل ۳/۳۵۶ رکورد جدید و ۸۸۰۱ رکورد از مجموعه داده‌های توسعه داده شده توسط سایر پژوهشگران است. در نهایت آموزش مدل‌های مختلف با استفاده از مجموعه داده سفارشی انجام شده است. نتایج آموزش در سیستم پیشنهادی توسط مجموعه داده سفارشی، بسته به مدل‌های انتخابی، میزان درک ارائه پاسخ را در معیار تطبیق دقیق به میزان ۲۰ تا ۴۳ درصد و در معیار اف ۱ به میزان ۸ تا ۱۸ درصد بهبود بخشیده است.

۶. کارهای آتی

کارهایی را که در آینده می‌توان تحقیقات را ادامه و بهبود بخشید می‌توان به شرح زیر برشماری نمود:

۱. افزایش کمی و کیفی شاخص‌ها
- شاخص‌های احصاء شده در زمینه امنیت شبکه‌های کامپیوتری، مبنای ایجاد مجموعه داده یادگیری هستند و

صورت مستقل از سایر مجموعه‌داده‌ها در مجموعه داده هوشفا تولید شده‌اند تعداد ۳,۳۵۶ رکورد است. در جدول زیر نتایج ارزیابی مدل‌های مختلف بر روی این بخش قابل مشاهده است. "نمره قبل از آموزش"، مربوط به زمانی است که مدل با استفاده از داده‌های هوشفا آموزش داده نشده و "نمره بعد از آموزش" مربوط به زمان است که مدل با استفاده از داده‌های هوشفا آموزش داده شده است. "میزان تغییرات" مربوط به نتایج بهبود یافته مدل بعد از آموزش می‌باشد.

جدول (۵). نتایج ارزیابی با آموزش مجموعه داده هوشفا

نام مدل	قبل از آموزش		بعد از آموزش		میزان بهبود	
	تطبیق دقیق	اف ۱	تطبیق دقیق	اف ۱	تطبیق دقیق	اف ۱
SajjadAyoubi/xlm-roberta-large-fa-qa	۱۸/۱۲	۵۰/۷۵	۶۱/۳۵	۶۹/۲۲	۴۳/۲۳	۱۸/۴۷
SajjadAyoubi/bert-base-fa-qa	۵/۹۷	۴۳/۵۴	۲۶/۶۹	۵۱/۰۶	۲۰/۷۲	۷/۵۲
ForutanRad/bert-fa-QA-v1	۴/۵۸	۴۲/۶۹	۲۶/۴۹	۵۱/۵۹	۲۱/۹۱	۸/۹

۴-۴. نتایج ارزیابی مدل طبقه‌بندی اخبار

چنان که بیان شد از مدل طبقه‌بندی اخبار برای تشخیص اخبار مرتبط و غیرمرتبط استفاده شده است. نتایج ارزیابی این مدل روی عناوین اخبار کارایی این مدل بیش از ۹۵٪ برای اخبار غیرمرتبط و بیش از ۸۰٪ برای اخبار مرتبط را نشان داده است. برچسب صفر مربوط به اخبار غیر مرتبط و برچسب یک مربوط به اخبار مرتبط است.

جدول (۶). نتایج ارزیابی مدل طبقه‌بندی اخبار

نام مدل	دقت	پارامتر	اف ۱	پوشش	صحت
HooshvareLab_bert-fa-zwnj-base	۰/۹۶	۰	۰/۹۸	۰/۹۸	۰/۹۷
		۱	۰/۸۳	۰/۸۰	۰/۸۷
HooshvareLab_gpt2-fa	۰/۹۵	۰	۰/۹۷	۰/۹۹	۰/۹۶
		۱	۰/۸۰	۰/۷۱	۰/۹۱
xlm-roberta-base	۰/۹۲	۰	۰/۹۶	۱/۰	۰/۹۲
		۱	۰/۵۷	۰/۴۰	۰/۹۸

- processing with transformers. "O'Reilly Media Inc.", 2022.
- [6] ParsiNLU, "ParsiNLU." [Online]. Available: <https://github.com/persiannlp/parsinlu>.
- [7] C. E. Shannon, "A mathematical theory of communication," The Bell system technical journal, vol. 27, no. 3, pp. 379-423, 1948.
- [8] A. M. Turing and J. Haugeland, "Computing machinery and intelligence," The Turing Test: Verbal Behavior as the Hallmark of Intelligence, pp. 29-56, 1950.
- [9] O. Kolomiyets and M.-F. Moens, "A survey on question answering technology from an information retrieval perspective," Information Sciences, vol. 181, no. 24, pp. 5412-5434, 2011.
- [10] A. Vaswani et al., "Attention is all you need," Advances in neural information processing systems, vol. 30, 2017.
- [11] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "Bert: Pre-training of deep bidirectional transformers for language understanding," arXiv preprint arXiv:1810.04805, 2018.
- [12] A. M. Dai and Q. V. Le, "Semi-supervised sequence learning," Advances in neural information processing systems, vol. 28, 2015.
- [13] K. Darvishi, N. Shahbodagh, Z. Abbasiantaeb, and S. Momtazi, "PQuAD: A Persian Question Answering Dataset," arXiv preprint arXiv:2202.06219, 2022.
- [14] J. Duan, H. Zhao, Q. Zhou, M. Qiu, and M. Liu, "A study of pre-trained language models in natural language processing," in 2020 IEEE International Conference on Smart Cloud (SmartCloud), 2020: IEEE, pp. 116-121.
- [15] S. Singh and A. Mahmood, "The NLP cookbook: modern recipes for transformer based deep learning architectures," IEEE Access, vol. 9, pp. 68675-68702, 2021.
- [16] Y. E. Seyyar, A. G. Yavuz, and H. M. Ünver, "An attack detection framework based on BERT and deep learning," IEEE Access, vol. 10, pp. 68633-68644, 2022.
- [17] M. Farahani, M. Gharachorloo, M. Farahani, and M. Manthouri, "Parsbert: Transformer-based model for persian language understanding," Neural Processing Letters, vol. 53, no. 6, pp. 3831-3847, 2021.
- [18] E. Taher, S. A. Hoseini, and M. Shamsfard, "Beheshtiner: Persian named entity recognition Using BERT," arXiv preprint arXiv:2003.08875, 2020.
- [19] A. Khashei, "Bolbol-Zaban." [Online]. Available: <http://www.bolbolzaban.com/>.
- [20] M. Tanvirul Alam, D. Bhusal, Y. Park, and N. Rastogi, "CyNER: A Python Library for Cybersecurity Named Entity Recognition," arXiv e-prints, p. arXiv: 2204.05754, 2022.
- [21] Y. Guo, J. Liu, W. Tang, and C. Huang, "Exsense: Extract sensitive information from unstructured data," Computers & Security, vol. 102, p. 102156, 2021.
- [22] W. Zhuang, "Architecture of Knowledge Extraction System based on NLP," in 2021 International Conference on Aviation Safety and Information Technology, 2021, pp. 294-297.

افزایش کمی و کیفی شاخص‌های منجر به ارایه نتایج بهتر می‌شود.

۲. ارتقای کمی و کیفی مجموعه داده هوشفا

مجموعه داده یادگیری هوشفا به طور مستقیم در ایجاد درک مدل از مسایل مختلف تاثیر دارد و با افزایش کمی و کیفی رکوردها درک مدل از مسایل مختلف افزایش می‌یابد.

۳. آموزش مدل‌های مختلف جدید

مدل‌های مبدل مختلف دارای پارامترهای متفاوتی هستند و نتایج متفاوتی را ارایه می‌کنند. آموزش مدل‌های مختلف با مجموعه داده هوشفا و بررسی و مقایسه کارایی می‌تواند به بهبود محصول نهایی منجر شود.

۴. بهبود پارامترهای یادگیری

به‌کارگیری توابع بهینه‌ساز مختلف و دوره‌های متفاوت و ارزیابی نتایج حاصله می‌تواند منجر به بهبود درک مدل شود.

۵. آموزش مدلی به منظور طبقه‌بندی اخبار و منابع

مرتبط با امنیت شبکه‌های کامپیوتری

استفاده مدل هوش مصنوعی از مطالب مرتبط با موضوع مورد آموزش می‌تواند منجر به ارایه نتایج بهتر شود. این موضوع به‌خصوص در آموزش مدل‌های بر پایه متن مانند برت محسوس است.

۷. منابع

- [1] M. A. Hassan, "Relational and NoSQL Databases: The Appropriate Database Model Choice," in 2021 22nd International Arab Conference on Information Technology (ACIT), 2021: IEEE, pp. 1-6.
- [2] A. Agarwal, C. Baechle, R. Behara, and X. Zhu, "A natural language processing framework for assessing hospital readmissions for patients with COPD," IEEE journal of biomedical and health informatics, vol. 22, no. 2, pp. 588-596, 2017.
- [3] N. N. Vo, S. Liu, X. Li, and G. Xu, "Leveraging unstructured call log data for customer churn prediction," Knowledge-Based Systems, vol. 212, p. 106586, 2021.
- [4] S. Crockett and C. Lee, "Does it Matter What They Said? A Text Mining Analysis of the State of the Union Addresses of USA Presidents," in 2012 13th ACIS International Conference on Software Engineering, Artificial Intelligence, Networking and Parallel/Distributed Computing, 2012: IEEE, pp. 77-82.
- [5] L. Tunstall, L. von Werra, and T. Wolf, "Natural language