



Semnan University

Journal of Modeling in Engineering

<https://modelling.semnan.ac.ir/>

ISSN: 2783-2538



Research Article

A Conditional Generative Chatbot using Deep Learning

Nura Esfandiari ^a, Kourosh Kiani ^{a,*}, Razieh Rastgoo ^a

^a Electrical and Computer Engineering Faculty, Semnan University, Semnan, Iran

PAPER INFO

Paper history:

Received: 2025-01-20

Revised: 2025-03-23

Accepted: 2025-04-26

Keywords:

Generative chatbot;
Deep learning;
Conditional generative
model;
Transformer model;
Dialog system.

ABSTRACT

A Chatbot serves as a communication tool between a human user and a machine to achieve an appropriate answer based on the human input. In more recent approaches, a combination of Natural Language Processing and sequential models are used to build a generative Chatbot. The main challenge of these models is their sequential nature, which leads to less accurate results. To tackle this challenge, in this paper, a novel architecture is proposed using conditional Wasserstein Generative Adversarial Networks and a transformer model for answer generation in Chatbots. While the generator of the proposed model consists of a full transformer model to generate an answer, the discriminator includes only the encoder part of a transformer model followed by a classifier. To the best of our knowledge, this is the first time that a generative Chatbot is proposed using the embedded transformer in both generator and discriminator models. Relying on the parallel computing of the transformer model, the results of the proposed model on the Cornell Movie-Dialog corpus and the Chit-Chat datasets confirm the superiority of the proposed model compared to state-of-the-art alternatives using different evaluation metrics.

DOI: <https://doi.org/10.22075/jme.2025.36654.2797>

© 2025 Published by Semnan University Press.

This is an open access article under the CC-BY 4.0 license. (<https://creativecommons.org/licenses/by/4.0/>)

* Corresponding author.

E-mail address: Kourosh.kiani@semnan.ac.ir

How to cite this article:

N. Esfandiari, K. kiani and R. Rastgoo, "A Conditional Generative Chatbot using Deep Learning," Journal of Modeling in Engineering, 23 82 (2025): 99-113, doi: 10.22075/jme.2025.36654.2797

چت بات مولد شرطی با استفاده از یادگیری عمیق

نورا اسفندیاری^۱، کوروش کیانی^{۱*}، راضیه راستگو^۱

اطلاعات مقاله	چکیده
دریافت مقاله: ۱۴۰۳/۱۱/۰۱	چت بات ها به عنوان ابزاری برای ارتباط بین کاربر انسانی و ماشین عمل می کنند تا براساس ورودی انسان، پاسخ مناسبی ارائه دهند. در رویکردهای اخیر، ترکیبی از پردازش زبان طبیعی و مدل های دنباله ای برای ساخت چت بات های مولد استفاده می شود. چالش اصلی این مدل ها ماهیت دنباله ای آنهاست که منجر به کاهش دقت پاسخ ها می شود. برای مقابله با این چالش، در این مقاله، یک معماری جدید با استفاده از شبکه مولد تخاصمی Wasserstein شرطی (cWGAN) و مدل ترانسفورمر برای تولید پاسخ در چت بات ها پیشنهاد شده است. در حالی که مولد مدل پیشنهادی شامل یک مدل کامل ترانسفورمر برای تولید پاسخ است، تمایزگر تنها شامل قسمت رمزگذار مدل ترانسفورمر به همراه یک طبقه بند است. تا آنجا که ما می دانیم، این اولین باری است که یک چت بات مولد با استفاده از ترانسفورمر تعبیه شده در هر دو مدل مولد و تمایزگر پیشنهاد می شود. با تکیه بر محاسبات موازی مدل ترانسفورمر، نتایج مدل پیشنهادی بر روی مجموعه داده های مکالمات Cornell و Chit-Chat، برتری مدل پیشنهادی را نسبت به روش های قبلی با استفاده از معیارهای ارزیابی مختلف تأیید می کند.
بازنگری مقاله: ۱۴۰۴/۰۱/۰۳	
پذیرش مقاله: ۱۴۰۴/۰۲/۰۶	
واژگان کلیدی:	
چت بات مولد،	
یادگیری عمیق،	
مدل مولد شرطی،	
مدل ترانسفورمر،	
سیستم گفتگو.	

DOI: <https://doi.org/10.22075/jme.2025.36654.2797>

© 2025 Published by Semnan University Press.

This is an open access article under the CC-BY 4.0 license. (<https://creativecommons.org/licenses/by/4.0/>)

۱- مقدمه

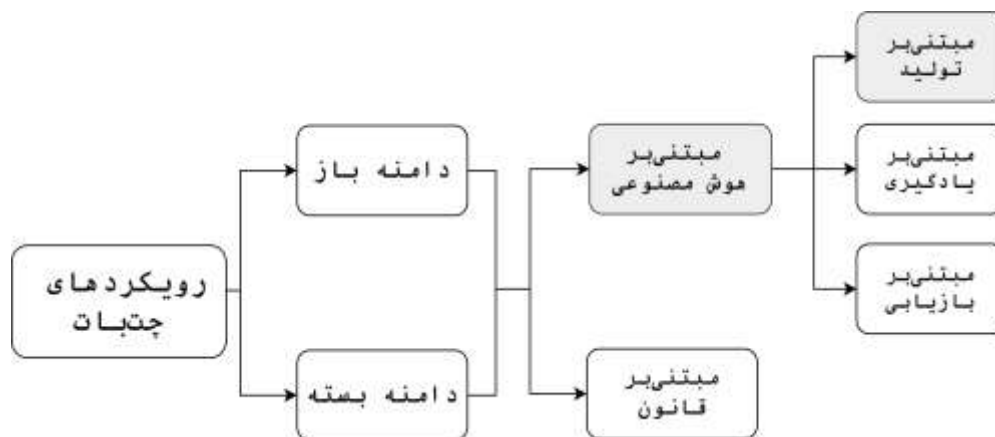
با توجه به افزایش استفاده از شبکه های اجتماعی، فناوری جدیدی به نام چت بات به عنوان ابزاری برای تعامل انسان و کامپیوتر توسعه یافته است. چت بات ها در حوزه های مختلفی مانند تجارت الکترونیک [۱]، آموزش [۲]، بانکداری و بیمه [۳]، جمع آوری و مدیریت داده [۴] و بهداشت [۵] به طور گسترده ای مورد استفاده قرار گرفته اند. چت بات ها به طور خودکار پاسخ های جذاب تری را از طریق ارتباطات آسان و کارآمد به کاربران ارائه می دهند [۶]. عامل کلیدی در طراحی مناسب چت بات ها، ارائه پاسخ های قابل درک به کاربر است [۷]. برای این منظور، رویکردهای مختلفی اخیراً برای توسعه چت بات ها ارائه یافته است. به

طور کلی، رویکردهای توسعه چت بات را می توان به دو دسته تقسیم کرد: دامنه باز و دامنه بسته (شکل ۱). چت بات هایی با توانایی پاسخگویی در بیش از یک حوزه، دامنه باز نامیده می شوند. در مقابل، چت بات های دامنه بسته فقط می توانند به سوالات مربوط به یک دامنه خاص پاسخ دهند. چت بات های دامنه باز و بسته را می توان به دو دسته چت بات های مبتنی بر قانون و مبتنی بر هوش مصنوعی (AI) طبقه بندی کرد. چت بات های مبتنی بر قانون بر ورودی کاربر با یک الگوی پایه تکیه می کنند. این نوع چت بات ها یک پاسخ از پیش تعریف شده را از مجموعه ای از پاسخ ها انتخاب می کنند [۸]. با این حال، آنها انعطاف ناپذیر هستند و به برخی پاسخ های از پیش

* پست الکترونیک نویسنده مسئول: Kourosh.kiani@semnan.ac.ir

۱. دانشکده مهندسی برق و کامپیوتر، دانشگاه سمنان، سمنان، ایران

استناد به این مقاله:



شکل ۱- طبقه‌بندی انواع رویکردهای چتبات

ماهیت دنباله‌ای این مدل‌ها منجر به استفاده از تکنیک‌های یادگیری تقویتی برای تکمیل پاسخ‌های تولیدی می‌شود. این مدل‌ها به دلیل واریانس بالا و سرعت پردازش پایین، دارای مشکل همگرایی کند هستند.

موفقیت اخیر چتبات‌های مبتنی بر مدل‌های زبانی بزرگ (LLM) مانند ChatGPT، پژوهشگران را ترغیب کرده است تا کاربردهای جدید و نوآورانه‌ای برای فناوری چتبات‌ها کاوش کنند. انتشار ChatGPT توسط OpenAI گامی مهم در پیشرفت چتبات‌های مبتنی بر LLM در حوزه هوش مصنوعی بود [۲۳]. مدل‌های زبانی بزرگ، دنیای چتبات‌ها را بهبود بخشیده‌اند و آن‌ها را قادر ساخته‌اند تا به شیوه‌هایی با انسان‌ها تعامل کنند که پیش‌تر قابل تصور نبود. عامل‌های مکالمه‌ای که بیشتر به‌عنوان چتبات‌های هوش مصنوعی شناخته می‌شوند، از حجم عظیمی از داده‌ها برای آموزش مدل‌های زبانی پیشرفته استفاده می‌کنند. این مدل‌های زبانی سپس از آموزش خود برای تولید محتوای جدید در پاسخ به درخواست‌های کاربران بهره می‌برند [۲۴]. اگرچه LLM‌ها روی مجموعه داده‌های عظیمی آموزش دیده‌اند و دانش عمومی گسترده‌ای دارند، اما اغلب برای برتری در وظایف خاص نیاز به تنظیم دقیق (fine-tuning) دارند. در گذشته، مدل‌های کوچک‌تر مانند BERT (با ۴۵۰ میلیون پارامتر) یا RoBERTa (با ۱.۳ میلیارد پارامتر) به‌طور کامل برای وظایف خاص تنظیم می‌شدند. اما با ظهور مدل‌های بسیار بزرگ‌تر مانند GPT-4.0، تنظیم کامل این مدل‌ها از نظر محاسباتی بسیار پرهزینه شده است. علاوه بر این، این فرآیند اغلب به حجم زیادی از داده‌های باکیفیت نیاز دارد [۲۵]. جمع‌آوری داده‌های کافی و مرتبط، می‌تواند دشوار

تعریف شده محدود می‌شوند. ELIZA و ALICE بر اساس این رویکرد هستند [۹].

چتبات‌های مبتنی بر هوش مصنوعی با استفاده از ترکیبی از پردازش زبان طبیعی (NLP) و رویکردهای یادگیری عمیق آموزش داده می‌شوند تا مکالماتی شبیه به انسان داشته باشند. به طور مشخص‌تر، چتبات‌های مبتنی بر هوش مصنوعی به سه دسته تقسیم می‌شوند: روش‌های مبتنی بر بازیابی، مبتنی بر یادگیری و مولد. دسته اول، روش مبتنی بر بازیابی، با استفاده از یک معیار امتیازدهی عملکردی، مشابه‌ترین پاسخ را از یک مجموعه داده انتخاب می‌کند [۱۰]. دسته دوم، روش مبتنی بر یادگیری، معمولاً الگوهایی را از داده‌های آموزشی، حاوی سؤالات و پاسخ‌ها، از طریق مدل‌های یادگیری عمیق یاد می‌گیرد تا پاسخ‌های مرتبط را تولید کند. دسته آخر، روش چتبات‌های مولد، معمولاً بر اساس مدل‌های یادگیری توالی به توالی (Seq2Seq) است [۱۱-۱۴]. با این حال، چالش اصلی این مدل‌ها ماهیت دنباله‌ای آن‌هاست که منجر به کاهش دقت نتایج می‌شود. برای مقابله با این چالش، در سال‌های اخیر، محققان برخی از مدل‌ها را بر اساس مدل‌های ترانسفورمر توسعه داده‌اند [۱۵-۱۹]. با تکیه بر محاسبات موازی و همچنین مکانیزم خودتوجهی مدل ترانسفورمر، عملکرد مدل‌های توسعه‌یافته با استفاده از ترانسفورمرها در مقایسه با مدل‌های Seq2Seq بهبود یافته است [۲۰]. اگرچه مدل‌های مبتنی بر یادگیری قادر به تولید پاسخ‌های متنوع نیستند. برای غلبه بر این کمبود، مدل‌های مولد برای یادگیری توزیع پاسخ‌ها توسعه یافته‌اند. برخی از مدل‌ها مانند شبکه‌های تولیدی متخاصم توالی (SeqGAN) [۲۱] و StepGAN [۲۲] به عنوان دسته سوم، توسعه یافته‌اند.

یک پاسخ از پیش تعریف شده ارائه می شود [۷]. رویکردهای مبتنی بر قانون را می توان به دو دسته تقسیم کرد: روش های تطبیق الگو و سیستم های استاندارد مبتنی بر وظیفه [۲۶]. در روش های تطبیق الگو، چت بات ها، ورودی کاربر را با الگوی قوانین مطابقت می دهند و با استفاده از الگوریتم های تطبیق الگو، یک پاسخ از پیش تعیین شده را از مجموعه پاسخ ها انتخاب می کنند. سیستم های مبتنی بر وظیفه، کاربر را برای تکمیل موارد خاصی راهنمایی می کنند. از دهه ۱۹۹۰، تحقیقات زیادی در زمینه طراحی چت بات ها بر اساس قوانین مشابه برای ارائه خدمات در حوزه های خاص انجام شده است. این چت بات ها به عنوان سیستم های گفتگوی مدولار مبتنی بر وظیفه شناخته می شوند که کاربر را برای انجام برخی وظایف ساختاریافته، مانند رزرو رستوران و فیلم راهنمایی می کنند. از سال ۱۹۶۶، توسعه چت بات Eliza با رویکرد مبتنی بر الگو آغاز شده است. این چت بات، جمله ورودی را بر اساس قوانین تجزیه تعیین شده توسط کلمات کلیدی در جمله، تحلیل می کند [۲۷]. "Pari" متغیرهای تأثیرگذار دیگری مانند "ترس"، "عصبانیت" و "بی اعتمادی" را به قوانین پیچیده تر اضافه می کند. این قوانین، مکالمه را انسانی تر کرده اند [۲۸]. ALICE از زبان نشانه گذاری هوش مصنوعی (AIML) استفاده می کند که دسته ای است که واحد دانش را تشکیل می دهد تا یک الگو و یک فیلد اختیاری را ترکیب کند [۲۹]. علاوه بر این، پلتفرم های متعددی مانند Microsoft LUIS، IBM Watson Assistant و Dialog flow برای کمک به کاربر در ساخت چت بات ها توسعه یافته اند. این نوع چت بات ها علی رغم سادگی، پیاده سازی سریع و مقرون به صرفه بودن، دارای معایبی هستند. انعطاف ناپذیری، عدم یادگیری، ناتوانی در ایجاد پاسخ های جدید و محدود بودن به برخی پاسخ های از پیش تعریف شده، مهم ترین چالش های چت بات های مبتنی بر قانون هستند.

۲-۲- مدل های مبتنی بر بازیابی

چت بات های مبتنی بر بازیابی، با جستجو در یک مخزن مکالمه ای از پیش ساخته شده، بهترین پاسخ مطابق با سؤال کاربر را انتخاب می کنند [۳۰]. Lowe و همکاران [۳۱] یک چت بات مبتنی بر بازیابی را با استفاده از روش فرکانس اصطلاح - فرکانس معکوس سند (TF-IDF) توسعه دادند. بردارهای TF-IDF هر سؤال و پاسخ کاندید با الحاق تمام امتیازهای TF-IDF محاسبه می شوند. پاسخ های کاندید با

باشد و کیفیت داده ها تأثیر مستقیمی بر کارایی مدل تنظیم شده دارد. همچنین، تنظیم دقیق LLM ها می تواند از نظر محاسباتی گران باشد و به قدرت پردازش و فضای ذخیره سازی قابل توجهی نیاز داشته باشد. این موضوع می تواند برای کسب و کارهای کوچک یا کسانی که منابع محدودی دارند، مانعی جدی ایجاد کند. به همین دلیل در این مقاله رویکردی پیشنهاد شده است که مستقل از LLM باشد. به عبارتی دیگر توسعه چت بات مبتنی بر روش پیشنهادی نسبت به چت بات های مبتنی بر LLM نیاز به داده های کمتری داشته و از نظر پردازشی و ذخیره سازی نیز به سخت افزار کمتری نیاز دارند.

برای مقابله با چالش ماهیت دنباله ای و همچنین بهبود دقت مدل های مبتنی بر تولید، یک معماری جدید بر اساس cWGAN با استفاده از مدل ترانسفورمر که در طول فاز آموزش به صورت موازی پردازش می شود، پیشنهاد شده است. این معماری با تولید پاسخ های شبیه به انسان، کارایی را افزایش می دهد.

سهم اصلی این مقاله را می توان به شرح زیر فهرست کرد: (۱) ما یک مدل جدید با استفاده از مدل ترانسفورمر پیشنهاد می کنیم که در هر دو مولد و تمایزگر استفاده می شود. (۲) ما سعی می کنیم چالش دقت پایین در مدل های دنباله ای را با استفاده از شبکه های مولد تخصصی Wasserstein شرطی و یک مدل ترانسفورمر برای تولید پاسخ در چت بات ها را بهبود دهیم. (۳) آزمایش های گسترده ای که بر روی دو مجموعه داده چالش برانگیز انجام شد، نشان می دهد که معماری ما در مقایسه با روش های قبلی در این زمینه عملکرد بهتری دارد.

سازماندهی این مقاله به شرح زیر است: بخش ۲ کارهای مرتبط در چت بات را بررسی می کند. معماری پیشنهادی در بخش ۳ توضیح داده شده است. در ادامه، نتایج تجربی و بحث در بخش های ۴ و ۵ ارائه شده است. در نهایت، نتیجه گیری و کارهای آینده در بخش ۶ ارائه شده است.

۲-۲- مروری بر ادبیات

در این بخش، به طور خلاصه به بررسی کارهای اخیر در چهار دسته زیر می پردازیم: مبتنی بر قانون، مبتنی بر بازیابی، مبتنی بر یادگیری و مولد.

۲-۱- مدل های مبتنی بر قانون

در مدل های مبتنی بر قانون، ابتدا متغیرهای مشخصه عبارت ورودی مشخص می شوند. سپس، بر اساس متغیرها و قوانین،

Zhang و همکاران [۴۰] به عنوان یک مدل از پیش آموزش دیده ارائه شده است. این مدل به صورت عمومی منتشر شده است تا توسعه سیستم‌های گفتگوی هوشمندتر را تسهیل کند. مدل پیشنهادی دیگر، یک مدل چت بات با استفاده از مدل نمایش‌های کدگذاری دوطرفه از ترنسفورمرها (BERT) است که فقط دارای یک کدگذار است [۴۱]. Lin و همکاران [۴۲]، یک عامل مکالمه همدلی انتها به انتها به نام CAiRE را معرفی کردند. مدل پیشنهادی از یک مدل زبانی از پیش آموزش دیده در مقیاس بزرگ بهره می‌برد و آن را با استفاده از اهداف چندانگانه، مانند مدل‌سازی زبان پاسخ، پیش‌بینی پاسخ و تشخیص احساسات گفتگو، تنظیم دقیق می‌کند. همچنین، Nuruzzaman و Hussain [۴۳] یک چت بات خاص دامنه به نام IntelliBot را پیشنهاد می‌کنند که یک چت بات مبتنی بر گفتگوی خلاقانه است که از استراتژی‌های متعدد (مبتنی بر قانون، مبتنی بر بازیابی و مبتنی بر یادگیری) برای تولید پاسخ استفاده می‌کند. در حالی که مدل‌های مبتنی بر یادگیری نتایج امیدوارکننده‌ای به دست آورده‌اند، اما به دلیل نداشتن توانایی یادگیری توزیع پاسخ‌ها، نمی‌توانند پاسخ‌های متنوعی تولید کنند.

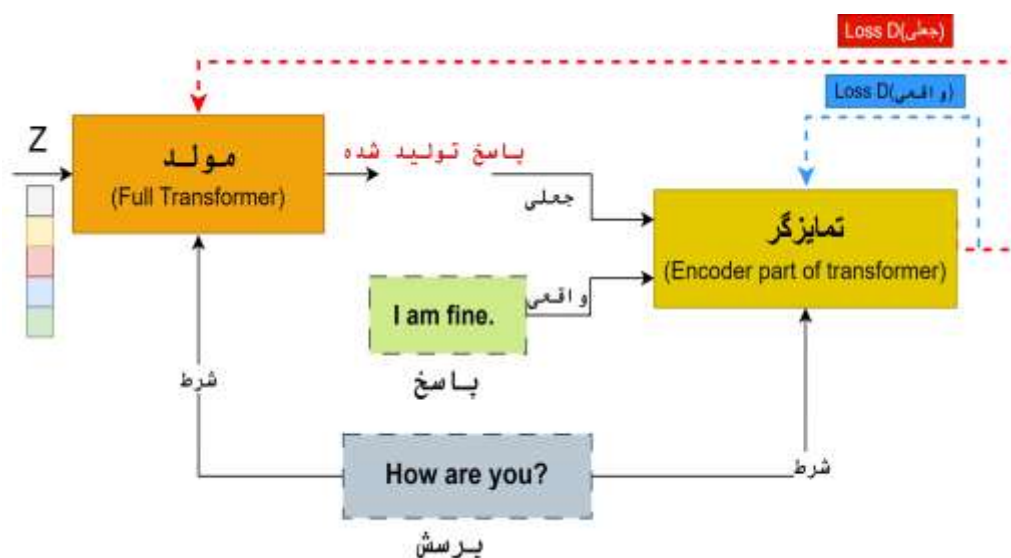
۲-۴- مدل‌های مولد

رویکردهای مولد، توزیع پاسخ‌ها را با استفاده از مدل‌های مولد مانند SeqGAN [۲۱] و StepGAN [۲۲] یاد می‌گیرند. این مدل‌ها از تکنیک‌های یادگیری تقویتی (RL) استفاده می‌کنند. در SeqGAN، سیگنال پاداش RL از تمایزگر می‌آید و با استفاده از جستجوی مونت کارلو به مراحل اقدامات حالت میانی بازخورد داده می‌شود. با این حال، جستجوی درختی مونت کارلو (MCTS) مورد استفاده در SeqGAN دارای پیچیدگی محاسباتی بالایی است. در رویکردهای StepGAN، شبکه GAN به صورت مرحله به مرحله ارزیابی می‌شود. به این ترتیب، تمایزگر برای اختصاص خودکار امتیازها و تعیین مناسب بودن هر زیرتوالی تولید شده، اصلاح می‌شود. Wu و Wang یک تابع زبان جدید را برای دستیابی به متنی تولیدی مشابه با متن‌های واقعی توسعه دادند [۴۴]. همچنین، یک مدل شبکه تمایزگر بر اساس مکانیزم خودتوجهی طراحی شد تا ویژگی‌های معنایی غنی‌تری را به دست آورد. در مقاله دیگری، یک مدل GAN ضمنی نهفته مبتنی بر ترنسفورمر (TILGAN) پیشنهاد شد که یک خودرمن‌نگار ترنسفورمر

بالاترین شباهت کسینوسی با بردار سؤال به عنوان پاسخ‌های نهایی انتخاب می‌شوند. Lu و همکاران [۳۲] معماری‌ای را برای غلبه بر مشکل تطبیق متن کوتاه مدل‌های توسعه‌یافته پیشنهاد کردند. بعداً، شبکه عصبی کانولوشنال (CNN)، شبکه عصبی بازگشتی (RNN) و زیرمجموعه‌های آن مانند حافظه کوتاه‌مدت بلند (LSTM) و واحد بازگشتی دروازه‌دار (GRU) به طور گسترده در این زمینه مورد استفاده قرار گرفتند [۲۶]. به عنوان مثال، Zhou و همکاران رویکردی را با افزودن مکانیزم توجه به شبکه عمیق برای ارائه تطبیق سؤال و پاسخ پیشنهاد کردند [۳۳]. در مقاله دیگری، Shu و همکاران [۳۴] یک مدل مبتنی بر بازیابی را برای جستجوی پاسخ‌های مرتبط با سؤال با ترکیب ماژول‌های استخراج کلمات کلیدی و ترانسفورمر دو مرحله‌ای پیشنهاد کردند. در کل، در حالی که رویکرد مبتنی بر بازیابی به طور گسترده توسط محققان مورد استفاده قرار می‌گیرد، هیچ پاسخ جدیدی تولید نمی‌شود و تنها محتمل‌ترین پاسخ از پایگاه داده بازیابی می‌شود. این امر می‌تواند پاسخ‌های چت بات را محدود کند. علاوه بر این، داشتن یک پایگاه داده در فاز استنتاج ضروری است.

۲-۳- مدل‌های مبتنی بر یادگیری

این رویکرد معمولاً بر اساس مدل یادگیری Seq2Seq است. در مورد جملات و مکالمات طولانی (بیش از ۲۰ کلمه)، تمام اطلاعات ضروری یک جمله منبع باید به یک بردار با طول ثابت فشرده شود که چالش‌برانگیز است [۲۰]. برای مقابله با این چالش، رویکردهای مختلفی توسط محققان پیشنهاد شده است. به عنوان مثال، ایجاد تغییرات ساختاری، مانند افزودن ماتریس تعبیه کلمات [۳۵]، اصلاح کدگذار یا رمزگشا [۳۶، ۳۷] و افزودن مکانیزم توجه [۱۱، ۱۳، ۱۴، ۳۸] برخی از راهکارهای پیشنهادی توسط محققان هستند. علاوه بر این، برخی از محققان از مدل مبتنی بر ترانسفورمر استفاده کرده‌اند که معمولاً لایه‌های بازگشتی مورد استفاده در معماری‌های کدگذار-رمزگشا را با خودتوجهی چندسری جایگزین می‌کند. Rao و همکاران [۳۹]، یک مدل ترکیبی برای تولید متن مبتنی بر یادگیری انتقال عمیق با استفاده از مدل زبانی Elmo برای تعبیه، خودرمن‌نگار تغییری (VAE) و حافظه کوتاه‌مدت بلند دوطرفه (BiLSTM) توسعه داده‌اند. یک مدل تولید پاسخ مبتنی بر ترنسفورمر به نام DIALOGPT نیز توسط



شکل ۲- معماری مدل پیشنهادی

این مقاله، یک معماری جدید بر اساس cWGAN با استفاده از مدل ترانسفورمر پیشنهاد شده است.

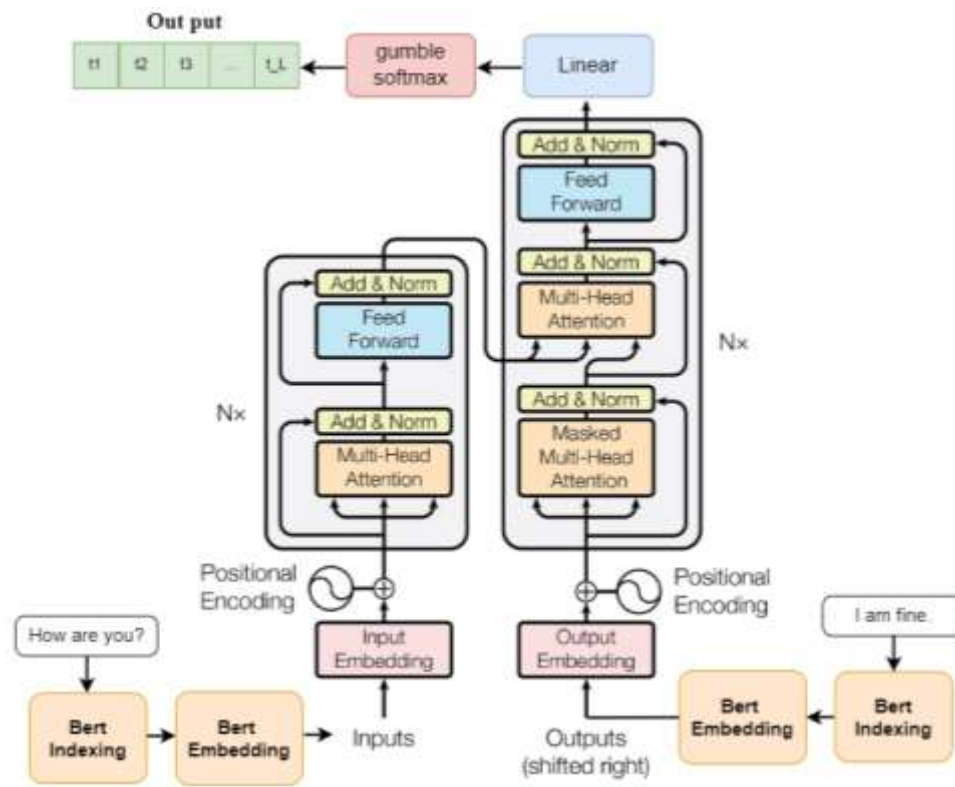
۳- روش پیشنهادی

یک مدل GAN ساده از دو بخش اصلی تشکیل شده است: مولد (G) و تمایزگر (D). در محدوده این پژوهش، ورودی GAN پرسشی به شکل یک دنباله x است. تمایزگر یاد می‌گیرد تا امتیاز $D(x, y^*)$ را بیشینه و امتیاز $D(x, \hat{y})$ را کمینه کند، در حالی که مولد می‌آموزد تا پاسخی به شکل $\hat{y}$ تولید کند تا امتیاز $D(x, \hat{y})$ را بیشینه کند، این رابطه در معادله (۱) بیان شده است:

$$\min \max E[\log D(x, y^*)] + E[\log(1 - D(x, \hat{y}))] \quad (1)$$

در رابطه (۱) (x, y^*) و $(x, \hat{y}) \sim P_R(x, y)$ بیانگر توزیع احتمال مشترک (x, y) و $x \sim P_R(x)$ بوده و نمایانگر توزیع احتمال x از داده‌های آموزشی است. از آنجایی که مدل‌های GAN ممکن است هرگز به پایداری نرسند و دچار مشکل فروپاشی حالت (mode collapses) شوند، محققان انواع مختلفی از مدل GAN را پیشنهاد کرده‌اند [۵۱]. یکی از این انواع، WassersteinGAN (WGAN) است که گسترش یافته‌ای از GAN بوده و به دنبال روشی جایگزین برای آموزش مدل مولد جهت تقریب بهتر توزیع داده‌های مشاهده شده در مجموعه داده آموزشی است. به جای استفاده از یک تمایزگر برای طبقه‌بندی یا پیش‌بینی احتمال واقعی یا ساختگی بودن داده‌های تولید شده، WGAN مدل تمایزگر را با یک منتقد (critic)

و GAN را در فضای نهفته با طراحی جدید و تطبیق توزیع براساس واگرایی Kullback-Leibler (KL) ترکیب می‌کند [۴۵]. برای تولید پاسخ‌های مرتبط‌تر، برخی از محققان از مدل‌های ترکیبی، از جمله ترکیب رویکردهای مولد و بازیابی استفاده کردند. به عنوان مثال، Zhang و همکاران [۴۶] پاسخ‌های بازیابی شده از مدل مبتنی بر بازیابی را به عنوان اطلاعات اضافی برای مدل‌های تمایزگر و مولد تغذیه کردند. Zhu و همکاران [۴۷] از N پاسخ برتر بازیابی شده برای محاسبه پاداش برای مدل مولد استفاده کردند. Zhang و همکاران [۴۸] از پاسخ تولید شده توسط مدل‌های مبتنی بر بازیابی به عنوان اطلاعات اضافی برای آموزش مدل‌های مولد استفاده کردند و یک فیلتر برای انتخاب بهترین پاسخ قرار دادند. علاوه بر این، یادگیری تقویتی عمیق ترکیبی نیز توسط Cuayáhuítl و همکاران [۴۹] برای چت بات‌های مولد توسعه داده شد. به طور خلاصه، موفقیت اخیر ChatGPT نیز محققان را تشویق کرده است تا رویکردهای بیشتری را برای کاربردهای چت بات توسعه دهند [۵۰]. با این حال، علی‌رغم پیشرفت‌های تکنولوژیکی، هنوز چالش‌های متعددی وجود دارد که باید برای ایجاد چت بات‌هایی که به درستی مکالمه انسانی و ویژگی‌های آن را به تصویر می‌کشند، برطرف شوند. در حالی که چت بات‌های مولد قادر به تولید پاسخ‌های مرتبط‌تر هستند، این مدل‌ها با برخی چالش‌ها مانند همگرایی کند به دلیل واریانس بالای آن‌ها و سرعت پردازش پایین، به ویژه در جملات بزرگ، مواجه هستند. برای غلبه بر چالش ماهیت دنباله‌ای و همچنین بهبود دقت در این مدل‌ها، در



شکل ۳- معماری ماژول مولد

آموزش در شکل (۳) نشان داده شده است. معماری ترانسفورمر این شکل از مقاله [۲۰] اقتباس شده است. در مرحله آموزش، لازم است توکن سازی و تبدیل به بردارهای معنایی (embedding) برای ورودی‌های رمزگذار (encoder) و رمزگشا (decoder) انجام شود. بنابراین، برای ورودی رمزگذار، پرسش‌های واقعی توکن سازی شده و توسط مدل از پیش آموزش دیده BERT که دارای ۱۲ لایه، ۷۶۸ ویژگی و ۱۲ هد است، به بردارهای معنایی تبدیل می‌شوند. برای افزایش سرعت آموزش، یک مدل خطی برای کاهش ابعاد ویژگی‌ها به ۱۲۸ مورد استفاده قرار گرفت، این رابطه در معادله (۲) نشان داده شده است:

$$y = f\left(\sum_{i=0}^n w_i x_i + \theta\right) \quad (2)$$

موقعیت کلمات در جمله حفظ شود. کدگذاری موقعیتی یک ماتریس است که متن را بر اساس موقعیت کلمه در جمله ارائه می‌دهد، این رابطه‌ها در معادلات (۳) و (۴) نشان داده شده است:

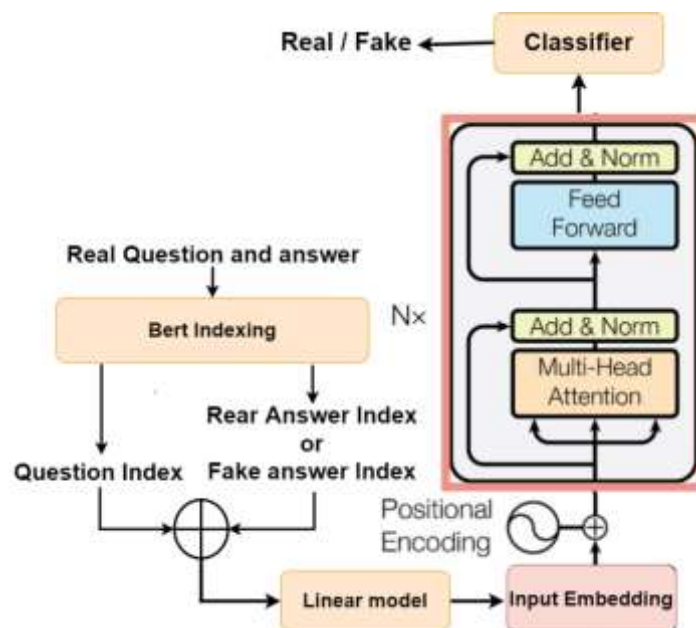
$$PE_{(pos, 2i)} = \sin\left(\frac{pos}{10000^{\frac{2i}{d_{model}}}}\right) \quad (3)$$

$$PE_{(pos, 2i+1)} = \cos\left(\frac{pos}{10000^{\frac{2i}{d_{model}}}}\right) \quad (4)$$

جایگزین می‌کند که امتیاز واقع‌گرایی یا ساختگی بودن داده‌های معین را می‌دهد. هدف، به حداقل رساندن فاصله بین توزیع داده‌ها در مجموعه داده آموزشی و نمونه‌های تولید شده است. این روش می‌تواند به پایداری فرایند آموزش در حین کار با گرادین‌ها کمک کند. با توجه به برتری مدل WGAN نسبت به مدل GAN، ما یک معماری جدید مبتنی بر cWGAN و مدل ترانسفورمر برای تولید پاسخ چت‌بات پیشنهاد می‌کنیم. دلیل انتخاب WGAN در روش پیشنهادی این است که هیچ تابع فعال سازی سیگموئید برای محدود کردن مقادیر به ۰ یا ۱ مطابق با پاسخ های واقعی یا ساختگی وجود ندارد. همچنین، فرایند آموزش زمانی که ترکیبی از WGAN با ترانسفورمر استفاده می‌شود، پایدارتر است. همانطور که در شکل (۲) نشان داده شده است، معماری پیشنهادی شامل دو ماژول است: مولد و تمایزگر که در یک شبکه واحد متصل شده‌اند. دلیل انتخاب WGAN در روش پیشنهادی و ترکیب آن با ترانسفورمر، افزایش پایداری در فرایند آموزش است.

۳-۱- ماژول مولد

مولد یک مدل ترانسفورمر کامل است که در مرحله استنتاج، پاسخ‌های ساختگی تولید می‌کند. معماری مولد در مرحله



شکل ۴- معماری ماژول تمایزگر

بعدی استفاده می‌شوند. توزیع Gumbel SoftMax یک توزیع پیوسته است که قادر به تقریب نمونه‌هایی از توزیع گسسته و ارائه خروجی با استفاده از تابع $\arg\max$ است. به طوری که خروجی رمزگشا یک بردار اندیس با طول L خواهد بود که به عنوان ورودی تمایزگر ارائه می‌شود. در مرحله رقابتی، مولد پس از دریافت مدل تمایزگر به‌روزرسانی می‌شود. در مدل GAN، این کار با استفاده از نسبت‌های درست‌نمایی گرادینت تابع هدف که می‌تواند با استفاده از معادله (۵) محاسبه شود، انجام می‌شود.

$$L_G = \nabla_{\theta_G} \frac{1}{m} \sum_{i=1}^m \log[D(G(Q_i))] \quad (5)$$

در رابطه (۵)، L_G تابع زیان مولد را نشان می‌دهد، m تعداد دنباله‌های تولید شده است، در حالی که Q بیانگر پرسش به عنوان داده شرطی است. همانطور که قبلاً بحث شد، ما با هدف غلبه بر چالش‌های مدل GAN و افزایش پایداری از WGAN استفاده می‌کنیم. به این ترتیب، تابع هدف WGAN یعنی L_{WG} می‌تواند با استفاده از معادله (۶) محاسبه شود:

$$L_{WG} = \nabla_{\theta_G} \frac{1}{m} \sum_{i=1}^m f(G(Q_i)) \quad (6)$$

۳-۲- ماژول تمایزگر

تمایزگر شامل تنها بخش رمزگذار ترانسفورمر به همراه یک طبقه‌بندی‌کننده است. این بخش در هر گام زمانی t

در ادامه، خروجی مدل خطی با کدگذاری موقعیتی (position encoding) ترکیب می‌شود تا در رابطه‌های (۳) و (۴) "pos" به موقعیت "کلمه" در دنباله اشاره دارد؛ در حالی که "d" نشان‌دهنده اندازه بردار معنایی کلمه/توکن است. "i" به هر یک از ابعاد بردار معنایی اشاره دارد. فرآیندی مشابه با فرایند رمزگذار با برخی تغییرات برای رمزگشا تکرار می‌شود. در رمزگشا، به جای پرسش، پاسخ واقعی در ورودی رمزگشا قرار دارد. ترانسفورمر در طول مراحل آموزش و استنتاج به شکلی کمی متفاوت عمل می‌کند. در طول استنتاج، تنها پرسش به عنوان دنباله ورودی ارائه می‌شود. هیچ پاسخ واقعی به عنوان دنباله هدف وجود ندارد که بتواند به رمزگشا منتقل شود. از آنجایی که هدف رمزگشا تولید پاسخی به نزدیک‌ترین شکل ممکن به پاسخ واقعی است، خروجی در یک حلقه تولید می‌شود و دنباله خروجی از گام زمانی قبلی در گام زمانی بعدی به رمزگشا تغذیه می‌شود تا زمانی که به توکن پایان برسد. در مولد، یک ترانسفورمر کامل با $N=8$ لایه یکسان و ۱۶ هد استفاده می‌شود. ماژول مولد ابتدا به صورت جداگانه با روش حداکثر درست‌نمایی (MLE) آموزش داده می‌شود تا احتمال همگرایی افزایش یابد. سپس، مدل با استفاده از شبکه رقابتی برای یادگیری توزیع پاسخ‌ها تنظیم دقیق می‌شود. تبدیل خطی و تابع Gumbel SoftMax برای تبدیل خروجی رمزگشا به احتمالات پیش‌بینی شده توکن

نتایج حاصل از معماری پیشنهادی با مدل‌های قبلی موجود مقایسه می‌شود.

۴-۱- جزئیات پیاده‌سازی

پیاده‌سازی بر روی یک پردازنده Core (TM) i5-12600K با ۱۲۸ گیگابایت رم در سیستم عامل ویندوز ۱۱ و با استفاده از نرم‌افزار پایتون و کارت گرافیک NVIDIA GeForce RTX 3090 انجام شدند. کتابخانه PyTorch برای پیاده‌سازی مدل مورد استفاده قرار گرفت. پارامترهای پیاده‌سازی در جدول ۱ فهرست شده‌اند.

جدول ۱- جزئیات پارامترهای تنظیم شده در مدل پیشنهادی

پارامترها	ارزش	پارامترها	ارزش
تعداد لایه‌ها	۸	نرخ یادگیری	۰.۰۰۰۰۵
نرخ جداسازی دیتاست برپا تست	۲۰٪	اندازه دسته‌ها	۶۴
حداکثر طول جمله	۳۰	تعداد مراحل اجرایی	۴۰۰
تعداد Head	۲	نوع پردازش	GPU
اندازه بردار ویژگی BERT	۷۶۸	مقدار Dropout	۰.۵

۴-۲- مجموعه داده‌ها

جهت ارزیابی مدل پیشنهادی، از دو مجموعه داده Cornell Movie Dialogs و Chit-Chat استفاده شده است. مجموعه داده Cornell شامل مجموعه‌ای بزرگ از مکالمات استخراج شده از اسکرپت‌های خام فیلم‌ها است که اطلاعات متا داده غنی نیز دارد. این مجموعه داده حاوی ۲۲۰,۵۷۹ تعامل مکالمه‌ای بین ۱۰,۲۹۲ جفت شخصیت فیلم از ۶۱۷ فیلم و در مجموع ۳۰۴,۷۱۳ عبارت است. مجموعه داده Chit-Chat نیز شامل ۷,۱۶۸ مکالمه از ۲۵۸,۱۴۵ عبارت و ۱,۳۱۵ شرکت‌کننده منحصر به فرد است. این مجموعه داده از چالش Chit-Chat آزمایشگاه کنترل و شناخت دانشگاه BYU استخراج شده است.

۴-۳- معیارهای ارزیابی

ارزیابی مدل پیشنهادی با استفاده از معیارهای BiLingual Evaluation Understudy (BLEU), Recall-Oriented Understudy for Gisting Metric for Evaluation (ROUGE), F-measure و Ordering (Meteor) انجام شده‌است. BLEU در ابتدا برای ارزیابی ترجمه ماشینی توسعه داده شد. در این معیار،

احتمال واقعی یا ساختگی بودن پاسخ‌ها را تعیین می‌کند و به تعداد k دوره (epochs) آموزش می‌بیند. در مدل GAN، تمایزگر با استفاده از نسبت‌های درست‌نمایی گرادیان تابع هدف که در معادله (۷) بیان شده است، به‌روزرسانی می‌شود.

$$L_D = \nabla_{\theta_G} \frac{1}{m} \sum_{i=1}^m -\log[D(A_i, Q_i)] - \log[1 - D(G(Q_i))] \quad (7)$$

در رابطه (۷)، L_D تابع زیان تمایزگر است، A پاسخ واقعی است و Q نشان‌دهنده پرسشی است که به عنوان داده شرطی در نظر گرفته می‌شود.

در مدل WGAN، تمایزگر به عنوان یک مدل منتقد در نظر گرفته می‌شود. L_C را می‌توان با استفاده از معادله (۸) محاسبه کرد.

$$L_C = \nabla_{\theta_G} \frac{1}{m} \sum_{i=1}^m f(q_i) - f(G(q_i)) \quad (8)$$

همانطور که در شکل (۴) نشان داده شده است، در هر دوره از آموزش، تمایزگر یک بار با جفت داده‌های جعلی و یک بار با جفت داده‌های واقعی آموزش داده می‌شود. جفت‌های پرسش و پاسخ واقعی توکن‌سازی شده و توسط مدل از پیش آموزش دیده BERT به اندیس تبدیل می‌شوند. ماتریس اندیس پاسخ جعلی با ماتریس اندیس پرسش به عنوان جفت‌های جعلی ترکیب می‌شود. یک مدل خطی برای کاهش ابعاد ماتریس ورودی و بردار معنایی جمله به کار گرفته می‌شود. سپس، خروجی مدل خطی با ماتریس کدگذاری موقعیتی ترکیب می‌شود تا به عنوان ماتریس ورودی برای لایه اول کدگذار استفاده شود. خروجی لایه آخر کدگذار به یک طبقه‌بندی‌کننده که یک شبکه خطی است، تغذیه می‌شود. این شبکه دو امتیاز خروجی دارد: واقعی یا جعلی بودن یک پاسخ معین. در مرحله استنتاج، ما فقط ماژول مولد را داریم و مدل تمایزگر حذف می‌شود. معماری رمزگشای این شکل (۴) از مقاله [۲۰] اقتباس شده است.

۴-۴- نتایج تجربی

در این بخش، به منظور نشان دادن کارایی مدل پیشنهادی، جزئیات مجموعه داده و نتایج بدست آمده، به همراه مقایسه با سایر روش‌ها، ارائه می‌گردد. ابتدا، جزئیات پیاده‌سازی مدل به طور مفصل شرح داده می‌شود. سپس، دو مجموعه داده مورد استفاده در ارزیابی، به طور مختصر معرفی شده و معیارهای ارزیابی به کار رفته نیز بیان می‌گردد. در نهایت،

جدول ۳- مقایسه نتایج مدل پیشنهادی روی مجموعه داده

Cornell		
معیار ارزیابی ROUGE	رویکرد	الگوریتم
0.337	Seq2Seq	S2S+Attention
0.399	BERT +LSTM	BERT
0.427	Retrieval model+ BERT	RAG
0.452	Retrieval model+ Transformer	OPRG
0.74	Seq2Seq	BILSTM
0.815	Transformer + Seq2Seq	DTGEN
0.818	GAN+ Transformer	روش پیشنهادی

جدول ۴ نتایج مقایسه مدل پیشنهادی با MLE [۲۲]، SeqGAN [۲۱]، StepGAN [۲۲] و CaiRE [۴۲] در مجموعه داده Chit-Chat با استفاده از معیار BLEU را نشان می‌دهد. رویکردهای مورد استفاده برای مقایسه از مدل GAN برای تولید داده‌های دنباله‌ای مانند متن استفاده می‌کنند. با توجه به جدول ۴، مدل پیشنهادی با استفاده از معیار BLEU در مجموعه داده Chit-Chat نسبت به تمام رویکردها عملکرد بهتری دارد که به دلیل استفاده از ترکیبی گسترده از مدل ترانسفورمر و روش‌های cWGAN است. در مقایسه با روش‌های دیگر، مدل پیشنهادی از مزیت یادگیری توزیع داده در تولید پاسخ برای اصلاح نتایج بهره می‌برد.

جدول ۴- مقایسه نتایج مدل پیشنهادی روی مجموعه داده

Chit-Chat		
معیار ارزیابی BLEU	رویکرد	الگوریتم
0.25	Seq2Seq	MLE
0.26	GAN +MC	SeqGAN
0.23	GAN +QN	StepGAN
0.70	Transformer	CAiRE
0.96	GAN+ Transformer	روش پیشنهادی

میزان همپوشانی بین جمله تولید شده و جمله حقیقی بر اساس n گرام محاسبه می‌شود. برخلاف BLEU که بر دقت تمرکز دارد، ROUGE به دلیل نیاز به محاسبه شباهت جمله تولید شده و جمله حقیقی، بر فراخوانی تمرکز دارد. F-measure یک معیار ارزشمند برای ارزیابی عملکرد الگوریتم‌های طبقه‌بندی است که می‌تواند به عنوان سازشی بین فراخوانی و دقت تعریف شود. METEOR یک معیار ارزیابی رایج است که به طور گسترده برای اندازه‌گیری شباهت بین جملات استفاده می‌شود. این معیار نتایج نزدیک‌تری به قضاوت‌های انسانی ارائه می‌دهد.

۴-۴- نتایج

در این بخش، اثربخشی معماری پیشنهادی بر روی دو مجموعه داده استفاده شده با استفاده از معیارهای BLEU، ROUGE، F-measure و Meteor ارزیابی شده است (جدول ۲ را مشاهده کنید). همانطور که نتایج این جدول نشان می‌دهد، مدل پیشنهادی عملکرد بهتری را بر روی مجموعه داده Chit-Chat دارد. این امر به این دلیل است که این مجموعه داده بر اساس مکالمه واقعی است در حالی که مجموعه داده Cornell بر اساس دیالوگ‌های فیلم است. برای ارزیابی و مقایسه مدل پیشنهادی با سایر روش‌ها، از چندین رویکرد از جمله S2S+Attention [۳۴]، BERT [۳۴]، تولید پاسخ با بازبازی (RAG) [۵۲]، تولید پاسخ حوزه باز (OPRG) [۳۴] و BILSTM و تولید متن مبتنی بر یادگیری انتقال عمیق (DTGEN) [۳۹] به همراه معیار ROUGE بر روی مجموعه داده Cornell استفاده شده است (جدول ۳).

جدول ۲- نتایج مدل پیشنهادی روی دو مجموعه داده

دیتاست	BLEU	ROUGE	F-measure	Meteor
Cornell	0.71	0.818	0.622	0.52
Chit-Chat	0.96	0.965	0.989	0.87

طبق جدول ۳، مدل پیشنهادی به دلیل استخراج ویژگی‌های غنی‌تر توسط ترانسفورمر و همچنین یادگیری توزیع داده توسط GAN، عملکرد بهتری نسبت به تمام مدل‌ها دارد.

what kind of thing can you respond to?	I am here to help answer your questions	I am here to help answer your questions
how many sisters do you have?	I do not have a family the same way humans do	I do not have a family the same way humans do
do you have a gender identity?	since I am digital, I do not actually have a gender	since I am digital, I do not actually have a gender

۵- بحث

چالش اصلی بسیاری از رویکردهای اخیر چت بات، ماهیت دنباله ای آنها است که منجر به کاهش دقت نتایج آنها می شود. مدل پیشنهادی با استفاده از ترکیبی از cWGAN و مدل ترانسفورمر، این چالش را برطرف می کند. طبق نتایج تجربی، مدل پیشنهادی می تواند پاسخ های دقیق تر و از نظر معنایی مرتبط تری را برای چت بات تولید کند. در حالی که مدل پیشنهادی از یادگیری توزیع پاسخ ها بهره مند می شود، پاسخ های تولید شده تنوع کمی دارند. این امر به دلیل ماهیت مجموعه داده های مورد استفاده در مرحله آموزش است. هر دو مجموعه داده فاقد پاسخ های متعدد برای هر پرسش هستند. در این بخش، ما مدل پیشنهادی را از دو دیدگاه زیر مورد بحث قرار می دهیم.

- **تحلیل پیش آموزش مولد:** قبل از آموزش رقابتی مدل پیشنهادی، مدل مولد به صورت جداگانه توسط یک مدل ترانسفورمر آموزش داده می شود. نکته این است که پیش آموزش مدل مولد تأثیر مستقیمی بر عملکرد کل مدل پیشنهادی دارد. در بهترین تلاش، از ۲۰۰ دوره برای آموزش ترانسفورمر به عنوان یک مدل از پیش آموزش دیده برای مولد استفاده شد. پس از آن، مدل پیشنهادی با استفاده از یادگیری رقابتی در ۴۰۰ دوره آموزش داده می شود. طبق جدول ۷، مدلی که فقط با ترانسفورمر آموزش داده شده است و یادگیری رقابتی ندارد، کمترین عملکرد را دارد. ما این مدل را به عنوان یک مدل از پیش آموزش دیده برای شروع آموزش در مرحله رقابتی در نظر گرفتیم. نتایج نشان می دهند که این امر عملکرد چت بات را در تمام معیارهای مجموعه داده Cornell بهبود می بخشد.

برای بررسی پاسخ های تولید شده توسط چت بات پیشنهادی، پرسش های مختلفی از دامنه های متفاوت به چت بات داده شد. جدول ۵ نتایج چت بات آموزش داده شده با استفاده از مجموعه داده Cornell را نشان می دهد.

جدول ۵- نمونه هایی از پاسخ های تولید شده توسط مدل پیشنهادی آموزش داده شده با مجموعه داده Cornell

مکالمات احوال پرسی		
پاسخ چت بات	پاسخ واقعی	ورودی کاربر
hello Ju thank you	hello Jo thanks you	hello Mr. parker how are you?
well, it was good seeing you	well, it was nice seeing you	I heard that was you good?
I do not know truly	I do not know exactly	how many girlfriends did you have?
مکالمات عمومی		
پاسخ چت بات	پاسخ واقعی	ورودی کاربر
respect yes	unfortunately, yes	is this for like real?
he told me, maybe they have in jail	she told me, he was in jail	how do you know that?
there is another answer to this place	but there is another entrance to this place	there is no other way?

علاوه بر این، جدول ۶ نتایج چت بات پیشنهادی آموزش داده شده بر روی مجموعه داده Chit-Chat را در دو دامنه نشان می دهد.

جدول ۶- نمونه هایی از پاسخ های تولید شده توسط مدل پیشنهادی آموزش داده شده با مجموعه داده Chit-Chat

مکالمات احوال پرسی		
پاسخ چت بات	پاسخ واقعی	ورودی کاربر
Hello there	Hello there	can you hear me?
I am great thanks for asking	I am great thanks for asking me	how are you today?
sure, ask me a question	sure, ask me a question	can you talk with me?
مکالمات احوال پرسی		
پاسخ چت بات	پاسخ واقعی	ورودی کاربر

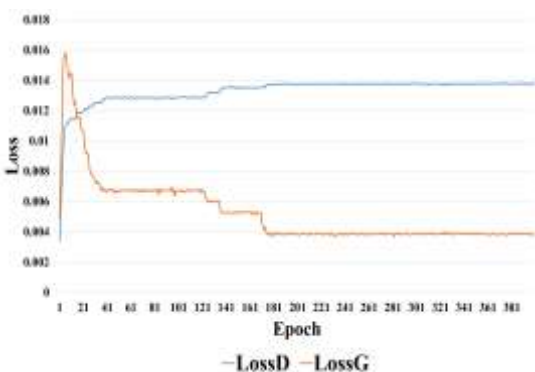
پیشنهادی به ترتیب در معادلات (۱۱) و (۱۲) خلاصه شده است:

$$GLoss = -[avg\ critic\ score\ on\ fake\ answer] \quad (10)$$

$$CLoss = [avg\ critic\ score\ on\ real\ answer] - [avg\ critic\ score\ on\ fake\ answer] \quad (11)$$

با توجه به شکل (۵)، در مجموعه داده Cornell، تابع زیان مولد در ابتدا به دلیل داشتن یک مدل مولد از پیش آموزش دیده، پایین است. پس از چند دوره، تابع زیان افزایش یافته و سپس با شیب ملایمی به مقداری کمتر از مقدار اولیه کاهش می‌یابد. این امر نشان‌دهنده یادگیری توزیع داده‌ها در فرآیند یادگیری رقابتی است. علاوه بر این، تمایزگر در ابتدا دارای تابع زیان پایینی است. با یادگیری توزیع داده‌ها توسط مولد و تولید داده‌های ساختگی بهتر، مقدار تابع زیان تمایزگر به آرامی افزایش می‌یابد. در نهایت، توابع زیان مولد و تمایزگر به همگرایی می‌رسند.

همچنین با توجه به شکل (۶)، در مجموعه داده Chit-Chat، تابع زیان مولد با مقداری مثبت نزدیک به صفر آغاز می‌شود، زیرا از مدل مولد از پیش آموزش دیده استفاده شده است. پس از چند دوره، تابع زیان افزایش می‌یابد و سپس به آرامی به مقادیر نزدیک به صفر کاهش می‌یابد. این مقدار کمتر از مقدار اولیه است که نشان‌دهنده بهبود کارایی چت بات با افزودن مدل رقابتی است. علاوه بر این، تمایزگر در ابتدا دارای تابع زیان نزدیک به صفر است که به تدریج با یادگیری توزیع داده‌ها توسط مولد افزایش می‌یابد. این نشان‌دهنده تولید پاسخ‌های ساختگی بهتر توسط مولد است.



شکل ۶- تابع زیان مولد و تمایزگر در Chit-Chat مجموعه داده

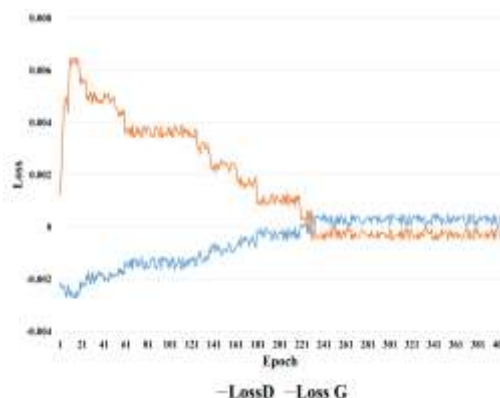
جدول ۷- عملکرد مدل پیشنهادی با پیکربندی‌های مختلف بر روی مجموعه داده Cornell

نوع مدل	BLEU	ROUGE	F-measure
Only pretraining	0.677	0.789	0.531
Only Adversarial	0.689	0.801	0.573
pretraining + Adversarial	0.71	0.818	0.62

• **تحلیل تابع زیان:** تابع زیان نقش کلیدی در عملکرد مدل‌های عمیق ایفا می‌کند. در این بخش، توابع زیان مولد و تمایزگر در مرحله رقابتی در نظر گرفته شده و مورد تجزیه و تحلیل قرار می‌گیرند. در مرحله آموزش رقابتی، از WGAN به جای مدل GAN استفاده می‌شود. WGAN از یک تابع هزینه جدید با استفاده از فاصله Wasserstein استفاده می‌کند که دارای گرادینان نرم‌تری است و صرف نظر از پیاده‌سازی مولد، آموزش داده می‌شود. فاصله Wasserstein با استفاده از معادله (۹) محاسبه می‌شود.

$$W(P_{real}, P_{fake}) = \sup_{f: \mathcal{X} \rightarrow \mathbb{R}} [E_{x \sim P_{real}} [f(x)] - E_{x \sim P_{fake}} [f(x)]] \quad (9)$$

در رابطه (۹)، $\sup$ بیانگر کران فوقانی، f نشان‌دهنده تابع Lipschitz و x نشان‌دهنده پاسخ واقعی یا ساختگی است. این شبکه مشابه تمایزگر است، اما فاقد تابع سیگموئید بوده و امتیاز اسکالر را به جای احتمال خروجی می‌دهد. این امتیاز می‌تواند به عنوان واقع‌گرایی داده ورودی تفسیر شده و به عنوان منتقد در نظر گرفته شود. تابع زیان Wasserstein برای مولد و منتقد (تمایزگر) در مدل



شکل ۵- تابع زیان مولد و تمایزگر در Cornell مجموعه داده

۶- نتیجه‌گیری و کارهای آتی

در این مقاله، یک مدل جدید مبتنی بر ترکیب cWGAN و مدل ترانسفورمر ارائه شده است. مدل پیشنهادی شامل دو شبکه است: مولد و تمایزگر. مولد یک مدل ترانسفورمر کامل است و تمایزگر تنها شامل بخش کدگذار یک مدل ترانسفورمر به همراه یک طبقه‌بند است. مدل پیشنهادی بر روی دو مجموعه داده با استفاده از معیارهای ارزیابی مختلف ارزیابی شده است. نتایج نشان داد که مدل پیشنهادی نسبت به رویکردهای قبلی موجود بر اساس معیارهای BLEU، ROUGE، F-measure و Meteor عملکرد بهتری دارد. با تکیه بر قابلیت‌های cWGAN و مدل ترانسفورمر، مدل پیشنهادی قادر به تولید پاسخ‌های دقیق، مرتبط از نظر معنایی و شبیه به انسان است. در کارهای آینده می‌توان از مدل‌های زبانی بزرگ برای افزایش دامنه پاسخ‌های تولیدی توسط چت بات با حفظ روابط معنایی بین پرسش و پاسخ استفاده کرد.

تأییدیه اخلاقی

نویسندگان متعهد می‌شوند که مطالب این مقاله را در

هیچ مجله دیگری به چاپ نرسانده اند.

مشارکت‌های نویسندگان

خانم نورا اسفندیاری: ایده، پیاده‌سازی، نوشتن نسخه ابتدایی مقاله، ویرایش، ارائه نتایج
دکتر کوروش کیانی: ایده، نظارت، ویرایش، بررسی نتایج، اصلاحات
دکتر راضیه راستگو: نظارت، ویرایش، بررسی نتایج، اصلاحات

منابع مالی

این پژوهش از هیچ منبع تامین مالی دولتی، تجاری یا غیرانتفاعی، کمک مالی خاصی دریافت نکرده است.

تعارض منافع

نویسندگان اعلام می‌دارند که هیچ گونه تضاد منافع ندارند. دکتر راضیه راستگو، یکی از نویسندگان در این مقاله، کمک سردبیر و مدیر داخلی فعلی مجله مدل‌سازی در مهندسی است، اما هیچ دخالتی در فرآیند ارزیابی این مقاله نداشته است. فرآیند داوری این مقاله توسط دکتر محمد دانایی، یکی از دبیران در مجله مذکور مدیریت شد.

مراجع

- [1] Tran, Anh D., Jason I. Pallant, and Lester W. Johnson. "Exploring the impact of chatbots on consumer sentiment and expectations in retail." *Journal of Retailing and Consumer Services* 63 (2021): 102718.
- [2] Miklosik, Andrej, Nina Evans, and Athar Mahmood Ahmed Qureshi. "The use of chatbots in digital business transformation: A systematic literature review." *Ieee Access* 9 (2021): 106530-106539.
- [3] Okonkwo, Chinedu Wilfred, and Abejide Ade-Ibijola. "Chatbots applications in education: A systematic review." *Computers and Education: Artificial Intelligence* 2 (2021): 100033.
- [4] Mogaji, Emmanuel, Janarthanan Balakrishnan, Arinze Christian Nwoba, and Nguyen Phong Nguyen. "Emerging-market consumers' interactions with banking chatbots." *Telematics and Informatics* 65 (2021): 101711.
- [5] Tsai, Meng-Han, Cheng-Hsuan Yang, James Yichu Chen, and Shih-Chung Kang. "Four-stage framework for implementing a chatbot system in disaster emergency operation data management: A flood disaster management case study." *KSCE Journal of Civil Engineering* 25, no. 2 (2021): 503-515.
- [6] Ayanouz, Soufyane, Boudhir Anouar Abdelhakim, and Mohammed Benhmed. "A Smart Chatbot Architecture Based NLP and Machine Learning for Health Care Assistance." *Niss* (2020).
- [7] Esfandiari, Nura, Kourosh Kiani, and Razieh Rastgoo. "Development of a Persian Mobile Sales Chatbot based on LLMs and Transformer." *Journal of AI and Data Mining*, no. 12 (2025): 465-472.
- [8] Esfandiari, Nura, Kourosh Kiani, and Razieh Rastgoo. "Transformer-based Generative Chatbot Using Reinforcement Learning." *Journal of AI and Data Mining*, no. 12 (2025): 349-358.
- [9] Esfandiari, Nura, Kourosh Kiani, and Razieh Rastgoo. "A conditional generative chatbot using transformer model." *arXiv:2306.02074*, (2023).
- [10] Zhu, Yutao, Jian-Yun Nie, Kun Zhou, Pan Du, Hao Jiang, and Zhicheng Dou. "Proactive retrieval-based chatbots based on relevant knowledge and goals." In *Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pp. 2000-2004. 2021.

- [11] Dhyani, Manyu, and Rajiv Kumar. "An intelligent Chatbot using deep learning with Bidirectional RNN and attention model." *Materials today: proceedings* 34 (2021): 817-824.
- [12] Wang, Yanmeng, Wenge Rong, Yuanxin Ouyang, and Zhang Xiong. "Augmenting dialogue response generation with unstructured textual knowledge." *IEEE Access* 7 (2019): 34954-34963.
- [13] Peng, Yehong, Yizhen Fang, Zhiwen Xie, and Guangyou Zhou. "Topic-enhanced emotional conversation generation with attention mechanism." *Knowledge-Based Systems* 163 (2019): 429-437.
- [14] Yang, Min, Wenting Tu, Qiang Qu, Zhou Zhao, Xiaojun Chen, and Jia Zhu. "Personalized response generation by dual-learning based domain adaptation." *Neural Networks* 103 (2018): 72-82.
- [15] Lin, Tzu-Hsuan, Yu-Hua Huang, and Alan Putranto. "Intelligent question and answer system for building information modeling and artificial intelligence of things based on the bidirectional encoder representations from transformers model." *Automation in Construction* 142 (2022): 104483.
- [16] Masum, Abu Kaisar Mohammad, Sheikh Abujar, Sharmin Akter, Nushrat Jahan Ria, and Syed Akhter Hossain. "Transformer based bengali chatbot using general knowledge dataset." In *2021 20th IEEE International Conference on Machine Learning and Applications (ICMLA)*, pp. 1235-1238. IEEE, 2021.
- [17] Peng, Baolin, Michel Galley, Pengcheng He, Chris Brockett, Lars Liden, Elnaz Nouri, Zhou Yu, Bill Dolan, and Jianfeng Gao. "Godel: Large-scale pre-training for goal-directed dialog." *arXiv preprint arXiv:2206.11309* (2022).
- [18] Shao, Taihua, Yupu Guo, Honghui Chen, and Zepeng Hao. "Transformer-based neural network for answer selection in question answering." *IEEE Access* 7 (2019): 26146-26156.
- [19] Shang, Shengjie, Jin Liu, and Yihe Yang. "Multi-layer transformer aggregation encoder for answer generation." *IEEE Access* 8 (2020): 90410-90419.
- [20] Vaswani, Ashish, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Łukasz Kaiser, and Illia Polosukhin. "Attention is All You Need." *Nips'17*, (2017): 6000–6010.
- [21] Yu, Lantao, Weinan Zhang, Jun Wang, and Yong Yu. "Seqgan: Sequence generative adversarial nets with policy gradient." In *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 31, no. 1. 2017.
- [22] Tuan, Yi-Lin, and Hung-Yi Lee. "Improving Conditional Sequence Generative Adversarial Networks by Stepwise Evaluation." *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, (2019).
- [23] Lin, Chien-Chang, Anna YQ Huang, and Stephen JH Yang. "A review of ai-driven conversational chatbots implementation methodologies and challenges (1999–2022)." *Sustainability* 15, no. 5 (2023): 4012.
- [24] Jeong, Cheonsu. "Fine-tuning and utilization methods of domain-specific llms." *arXiv preprint arXiv:2401.02981* (2024).
- [25] Adamopoulou, Eleni, and Lefteris Moussiades. "Chatbots: History, technology, and applications." *Machine Learning with Applications* 2 (2020): 100006.
- [26] Peng, Zhenhui, and Xiaojuan Ma. "A survey on construction and enhancement methods in service chatbots design." *CCF Transactions on Pervasive Computing and Interaction* 1, no. 3 (2019): 204-223.
- [27] Weizenbaum, Joseph. "ELIZA—a computer program for the study of natural language communication between man and machine." *Communications of the ACM* 9, no. 1 (1966): 36-45.
- [28] Colby, Kenneth Mark. "Artificial Paranoia: A Computer Simulation of Paranoid Processes." Elsevier Science Inc: New York, (1975).
- [29] Wallace, R.S. "The Anatomy of A.L.I.C.E, in *Parsing the Turing Test: Philosophical and Methodological Issues in the Quest for the Thinking Computer*." Springer Netherlands: Dordrecht. (2009): p. 181-210.
- [30] Yan, Rui, Yiping Song, and Hua Wu. "Learning to respond with deep neural networks for retrieval-based human-computer conversation system." In *Proceedings of the 39th International ACM SIGIR conference on Research and Development in Information Retrieval*, pp. 55-64. 2016.
- [31] Lowe, Ryan, Nissan Pow, Iulian Serban, and Joelle Pineau. "The ubuntu dialogue corpus: A large dataset for research in unstructured multi-turn dialogue systems." *arXiv preprint arXiv:1506.08909* (2015).
- [32] Lu, Zhengdong, and Hang Li. "A deep architecture for matching short texts." *Advances in Neural Information Processing Systems* 26 (2013).

- [33] Zhou, Xiangyang, Lu Li, Daxiang Dong, Yi Liu, Ying Chen, Wayne Xin Zhao, Dianhai Yu, and Hua Wu. "Multi-turn response selection for chatbots with deep attention matching network." In Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), pp. 1118-1127. 2018.
- [34] Shu, Chang, Zijian Zhang, Youxin Chen, Jing Xiao, Jey Han Lau, Qian Zhang, and Zheng Lu. "Open domain response generation guided by retrieved conversations." IEEE Access 11 (2022): 99365-99375.
- [35] Serban, Iulian, Alessandro Sordoni, Ryan Lowe, Laurent Charlin, Joelle Pineau, Aaron Courville, and Yoshua Bengio. "A Hierarchical Latent Variable Encoder-Decoder Model for Generating Dialogues." Aaai'17, (2017): 3295-3301.
- [36] Zhang, Wei-Nan, Qingfu Zhu, Yifa Wang, Yanyan Zhao, and Ting Liu. "Neural personalized response generation as domain adaptation." World Wide Web 22, no. 4 (2019): 1427-1446.
- [37] u, Jiatao, Zhengdong Lu, Hang Li, and Victor OK Li. "Incorporating Copying Mechanism in Sequence-to-Sequence Learning." in Proc. 54th Annu. Meeting Assoc. Comput. Linguistics. (2016).
- [38] Palasundram, Kulothunkan, Nurfadhlin Mohd Sharef, Khairul Azhar Kasmiran, and Azreen Azman. "SEQ2SEQ++: A multitasking-based Seq2Seq model to generate meaningful and relevant answers." IEEE Access 9 (2021): 164949-164975.
- [39] Rao, K. Yogeswara, and K. Srinivasa Rao. "Modeling text generation with contextual feature representation and di-mensionalitreduction using deep transfer learning and bi-lstm." Journal of Theoretical and Applied Information Technology 100, no. 9 (2022).
- [40] Zhang, Yizhe, et al. "DIALOGPT : Large-Scale Generative Pre-training for Conversational Response Generation." in Annual Meeting of the Association for Computational Linguistics. (2019).
- [41] Yu, Shi, Yuxin Chen, and Hussain Zaidi. "AVA: A financial service chatbot based on deep bidirectional transformers." Frontiers in Applied Mathematics and Statistics 7 (2021): 604842.
- [42] Lin, Zhaojiang, Peng Xu, Genta Indra Winata, Farhad Bin Siddique, Zihan Liu, Jamin Shin, and Pascale Fung. "Caire: An end-to-end empathetic chatbot." In Proceedings of the AAAI conference on artificial intelligence, vol. 34, no. 09, pp. 13622-13623. 2020.
- [43] Nuruzzaman, Mohammad, and Omar Khadeer Hussain. "IntelliBot: A Dialogue-based chatbot for the insurance industry." Knowledge-Based Systems 196 (2020): 105810.
- [44] Wu, Yuxi, and Junli Wang. "Text generation service model based on truth-guided SeqGAN." IEEE Access 8 (2020): 11880-11886.
- [45] Diao, Shizhe, Xinwei Shen, Kashun Shum, Yan Song, and Tong Zhang. "TILGAN: transformer-based implicit latent GAN for diverse and coherent text generation." In Findings of the Association for Computational linguistics: ACL-IJCNLP 2021, pp. 4844-4858. 2021.
- [46] Zhang, Jiayi, Chongyang Tao, Zhenjing Xu, Qiaojing Xie, Wei Chen, and Rui Yan. "Ensemblegan: Adversarial learning for retrieval-generation ensemble model on short-text conversation." In Proceedings of the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 435-444. 2019.
- [47] Zhu, Qingfu, Lei Cui, Weinan Zhang, Furu Wei, Ting Liu. "Retrieval-Enhanced Adversarial Training for Neural Response Generation." In Annual Meeting of the Association for Computational Linguistics. (2018).
- [48] Zhang, Liang, Yan Yang, Jie Zhou, Chengcai Chen, and Liang He. "Retrieval-polished response generation for chatbot." IEEE Access 8 (2020): 123882-123890.
- [49] Cuayáhuítl, Heriberto, Donghyeon Lee, Seonghan Ryu, Yongjin Cho, Sungja Choi, Satish Indurthi, Seunghak Yu, Hyungtak Choi, Inchul Hwang, and Jihie Kim. "Ensemble-based deep reinforcement learning for chatbots." Neurocomputing 366 (2019): 118-130.
- [50] Lin, Chien-Chang, Anna YQ Huang, and Stephen JH Yang. "A review of ai-driven conversational chatbots implementation methodologies and challenges (1999-2022)." Sustainability 15, no. 5 (2023): 4012.
- [51] Brownlee, Jason. Generative adversarial networks with python: deep learning generative models for image synthesis and image translation. Machine Learning Mastery, 2019.
- [52] Lewis, Patrick, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler et al. "Retrieval-augmented generation for knowledge-intensive nlp tasks." Advances in neural information processing systems 33 (2020): 9459-9474.