



دانشگاه صنعتی اصفهان

دانشکده علوم ریاضی

پروژه کارشناسی

انقلاب هوش مصنوعی و ظهور مکانیزم توجه

واکاوی مقاله

Attention is All You Need

نگارش:

محمدجواد عبدالهی و شنوه‌ئی

استاد راهنما:

دکتر بهناز عمومی

پاییز ۱۴۰۴

چکیده

در این گزارش، به تحلیل مقاله "توجه تمام آن چیزی است که نیاز دارید" اثر تحقیقات تیمی از پژوهشگران گوگل و برخی دانشگاه‌ها بوده که یکی از نقاط عطف مهم در زمینه یادگیری عمیق و شبکه‌های عصبی به شمار می‌آید. این مقاله با معرفی معماری ترنسفورمر، وابستگی به شبکه‌های عصبی بازگشتی را به طور کامل حذف کرده و مدلی مبتنی بر مکانیزم توجه ارائه می‌دهد که نه تنها به نتایج پیشرفته‌تری در وظایف ترجمه ماشینی دست یافت، بلکه به دلیل قابلیت موازی‌سازی بالا، فرآیند آموزش ماشین را نیز به شدت تسریع بخشیده است. در این پروژه، ابتدا به بررسی محدودیت‌های مدل‌های پیشین پرداخته، سپس اجزای اصلی معماری ترنسفورمر شامل ضرب داخلی مقیاس‌شده، توجه چند سر و رمزگذاری موقعیتی را تشریح کرده و در نهایت به بررسی نتایج و تأثیرات گسترده این مقاله بر حوزه پردازش زبان طبیعی می‌پردازیم.

فهرست مطالب

۳	مقدمه
۴	۱ عصر بازگشتی و چالش‌های زبانی (دنیای پیش از ۲۰۱۷)
۴	۱.۱ الگو پردازش متوالی: صف‌بندی کلمات
۵	۲.۱ محدودیت‌های ذاتی مدل‌های بازگشتی: تراژدی فراموشی
۵	۱.۲.۱ مشکل اول: پدیده محو شدن گرادیان
۵	۲.۲.۱ مشکل دوم: گلوگاه پردازش سریالی
۶	۳.۱ تلاش‌های اصلاحی: ظهور حافظه‌های پیشرفته
۶	۴.۱ بن‌بست اطلاعاتی در معماری رمزگذار-رمزگشا
۷	۲ انقلاب توجه و معماری ترنسفورمر
۷	۱.۲ مفهوم توجه در ادراک انسانی
۷	۱.۱.۲ تمثيل اثر مهمانی کوکتل
۸	۲.۲ تولد ترنسفورمر: شکستن زنجیر زمان
۸	۱.۲.۲ حذف کامل پردازش سریالی
۸	۲.۲.۲ جایگزینی حافظه با مکانیزم توجه
۱۰	۳ کالبدشکافی معماری ترنسفورمر
۱۰	گذر از انتزاع به ریاضیات
۱۰	۱.۳ معماری مبتنی بر پشته‌های رمزگذار و رمزگشا
۱۱	۲.۳ توکن‌گذاری و تعبیه‌ی کلمات: تبدیل واژه به مختصات معنایی
۱۱	۱.۲.۳ مفهوم توکن: واحد بنیادین پردازش
۱۲	۲.۲.۳ لایه‌ی تعبیه‌ی کلمات: نگاشت به فضای برداری پیوسته
۱۲	۱.۲.۲.۳ ویژگی‌های ریاضی فضای تعبیه
۱۲	۳.۳ کدگذاری موقعیتی: حل معمای آشفتگی زمانی
۱۳	۴.۳ مکانیزم توجه خودکار: قلب تپنده مدل
۱۴	۱.۴.۳ تشریح فنی با مدل سه‌گانه پرس‌وجو، کلید و مقدار
۱۴	۱.۱.۴.۳ فرمول ریاضی توجه ضرب‌داخلی مقیاس‌شده

۱۴	۲.۱.۴.۳	فرآیند محاسبه توجه
۱۵	۵.۳	توجه چندسره: نگرش چندبعدی به زبان
۱۵	۱.۵.۳	تحلیل نقش ماتریس‌های وزن و تبدیل خطی
۱۶	۶.۳	شبکه‌ی پیش‌خور و اتصالات میان‌بر: تثبیت یادگیری
۱۶	۱.۶.۳	تحلیل ساختاری و هندسی لایه پیش‌خور
۱۷	۲.۶.۳	اتصالات میان‌بر و نرمال‌سازی: تضمین پایداری ریاضی
۱۷	۷.۳	تولید خروجی نهایی: لایه‌ی خطی و تابع نرم‌بیشینه
۱۹	۸.۳	تحلیل توپولوژیک و جریان داده در معماری ترنسفورمر
۱۹	۱.۸.۳	مسیر جریان در پشته رمزگذار (سمت چپ)
۱۹	۲.۸.۳	مسیر جریان در پشته رمزگشا (سمت راست)
۲۰	۳.۸.۳	مکانیزم توجه متقاطع: اتصال دو نیمکره
۲۲	۴	چرا ترنسفورمر جهان را تغییر داد؟
۲۲	۱.۴	پیروزی موازی‌سازی و هم‌افزایی با قانون مور
۲۳	۲.۴	مدیریت وابستگی‌های بافاصله: حذف محدودیت‌های مکانی
۲۳	۳.۴	قوانین مقیاس‌پذیری: بزرگ‌تر یعنی هوشمندتر
۲۴	۵	عصر پسا ترنسفورمر و ظهور غول‌های زبانی
۲۴	۱.۵	تبارشناسی مدل‌ها: سه شاخه تکاملی
۲۴	۱.۱.۵	معماری‌های فقط رمزگذار ^۱
۲۴	۲.۱.۵	معماری‌های فقط رمزگشا ^۲
۲۵	۳.۱.۵	معماری‌های ترکیبی رمزگذار رمزگشا ^۳
۲۵	۲.۵	جدول زمانی تکامل: رشد نمایی پارامترها
۲۶	۳.۵	فراتر از متن: تعمیم به بینایی و مدل‌های چندوجهی
۲۷		نتیجه‌گیری
۲۹		منابع
۳۲		واژه نامه انگلیسی به فارسی
۳۷		واژه نامه فارسی به انگلیسی
۴۲		فهرست اختصارات

¹Encoder Only

²Decoder Only

³Encoder Decoder

مقدمه: سپیده دم عصر نوین ادراک ماشینی

در پهنه‌ی گسترده‌ی تاریخ تحولات تکنولوژیک، لحظاتی وجود دارند که مرز میان «پیشین» و «پسین» را چنان پررنگ و قاطع ترسیم می‌کنند که بازگشت به دوران ماقبل آن‌ها، ناممکن می‌نماید. همان‌گونه که اختراع ماشین چاپ، کشف الکتریسیته و ابداع ترانزیستور مسیر تمدن بشری را تغییر دادند، در دنیای پیچیده‌ی هوش مصنوعی^۱ و به‌طور خاص در حوزه‌ی تخصصی پردازش زبان طبیعی (NLP)^۲ نیز شاهد چنین نقطه عطفی بوده‌ایم. انتشار مقاله‌ی جریان‌ساز «توجه تمام آن چیزی است که نیاز دارید»^۳ در سال ۲۰۱۷ توسط تیمی از پژوهشگران برجسته‌ی گوگل، حکمی مشابه را در این عرصه داشت. این پژوهش نه تنها یک معماری نوین تحت عنوان ترنسفورمر^۴ را به جهان علم معرفی کرد، بلکه الگو حاکم بر تعامل ماشین با زبان انسان را از ریشه دگرگون ساخت. این تغییر بنیادین، زیربنای ظهور مدل‌های زبانی قدرتمند و فراگیر زیر شد که امروزه تار و پود زندگی دیجیتال بشر را احاطه کرده‌اند و مرزهای میان خلاقیت ماشینی و انسانی را کمرنگ ساخته‌اند.

Generative Pre-trained Transformer (GPT): این خانواده از مدل‌ها که بر پایه مکانیزم تولیدی بنا شده‌اند، در ساختن متن‌های منسجم و شبیه به انسان تبحر دارند. تمرکز اصلی آن‌ها بر پیش‌بینی کلمه بعدی در یک دنباله است و به همین دلیل در وظایفی همچون نگارش خلاق، ترجمه و مکالمه، عملکردی خیره‌کننده از خود به نمایش می‌گذارند.

Bidirectional Encoder Representations from Transformers (BERT): برخلاف مدل‌های پیشین که متن را تنها در یک جهت پردازش می‌کردند، این مدل با رویکردی دوطرفه به زبان نگاه می‌کند. این قابلیت به ماشین اجازه می‌دهد تا بافت و معنای کلمات را با توجه به واژگان قبل و بعد آن‌ها به صورت هم‌زمان درک کند که این امر منجر به انقلابی در وظایف درک مطلب و طبقه‌بندی متون شده است.

هدف غایی این گزارش پژوهشی، ارائه‌ی یک تحلیل جامع، عمیق و در عین حال شفاف برای مخاطبی است که لزوماً پیش‌زمینه‌ی فنی سنگینی در حوزه‌ی یادگیری عمیق^۵ یا هوش مصنوعی ندارد. ما در این نوشتار، سفری اکتشافی را از دوران حکمرانی الگوریتم‌های کلاسیک و محدودیت‌های ساختاری آن‌ها آغاز می‌کنیم؛ دورانی که در آن ماشین‌ها در درک روابط طولانی و پیچیده‌ی کلامی ناتوان بودند. سپس به لحظه‌ی تولد معماری ترنسفورمر می‌رسیم و با زبانی ساده اما دقیق، مکانیزم‌های درونی آن را کالبدشکافی می‌کنیم تا دریابیم چگونه مفهوم «توجه» جایگزین محاسبات خطی و زمان‌بر شد.

این گزارش صرفاً توصیف وقایع تاریخی نیست، بلکه تلاشی تحلیلی است برای درک چرایی و چگونگی وقوع انقلابی که به ماشین‌ها اجازه داد از سطح «خواندن» کلمات عبور کرده و به لایه‌های عمیق‌تر «درک» مفهوم، بافت کلام و ظرافت‌های پنهان زبانی دست یابند. ما بررسی خواهیم کرد که چگونه این جهش تکنولوژیک، راه را برای مدل‌هایی هموار کرد که قادرند شعر بگویند، کد بنویسند و استدلال کنند.

¹ Artificial Intelligence

² Natural Language Processing (NLP)

³ Attention Is All You Need

⁴ Transformer

⁵ Deep Learning

فصل ۱

عصر بازگشتی و چالش‌های زبانی (دنیای پیش از (۲۰۱۷)

برای درک عمیق از ابعاد بنیادین و عظمت تحولی که مقاله «توجه تمام آن چیزی است که نیاز دارید» در پیکره هوش مصنوعی ایجاد کرد، پیش‌نیاز اساسی آن است که وضعیت علم و فناوری را در دوران پیش از آن واکاوی کنیم. باید بدانیم مهندسان و دانشمندان علوم داده با چه ابزارهایی با پدیده زبان برخورد می‌کردند و چرا این ابزارها برای دستیابی به سطوح عالی ادراک ماشینی ناکافی و ناکارآمد بودند.

زبان انسان، برخلاف بسیاری از داده‌های ساختاریافته، پدیده‌ای ذاتاً پیچیده، سیال، ابهام‌آمیز و شدیداً وابسته به بافت است. برخلاف داده‌های تصویری که در آن‌ها پیکسل‌های مجاور معمولاً دارای همبستگی معنایی و ساختاری هستند، در زبان طبیعی، ارتباط میان کلمات می‌تواند در طول جملات طولانی پراکنده باشد و ساختاری غیرخطی را در بستری خطی پنهان کند.

۱.۱) الگو پردازش متوالی: صف‌بندی کلمات

تا پیش از سال ۲۰۱۷، الگو غالب در حوزه پردازش زبان طبیعی بر پایه یک فرض بنیادین و شهودی استوار بود:

«زبان یک توالی زمانی است»

منطق این فرض ساده بود؛ هنگامی که انسان صحبت می‌کند یا متنی را می‌نگارد، کلمات در بستر زمان جاری می‌شوند؛ کلمه اول پیش از کلمه دوم، و کلمه دوم پیش از کلمه سوم قرار می‌گیرد. بر این اساس، دانشمندان به این نتیجه رسیدند که مدل‌های ماشینی نیز باید متن را دقیقاً به همین ترتیب، یعنی به صورت کلمه به کلمه و متوالی پردازش کنند.

این رویکرد منجر به ظهور و توسعه خانواده‌ای از الگوریتم‌ها تحت عنوان شبکه‌های عصبی بازگشتی (RNN)^۱ گردید. این شبکه‌ها با الهام‌گیری از مکانیزم حافظه کوتاه‌مدت در مغز انسان طراحی شده بودند. ایده محوری این بود که شبکه، کلمه نخست را دریافت کرده، پردازش می‌کند و عصاره اطلاعات آن را در یک حافظه داخلی به نام

^۱ Recurrent Neural Network (RNN)

«حالت پنهان»^۱ ذخیره می‌نماید. سپس این اطلاعات به مرحله بعد (گام زمانی بعدی) منتقل می‌شود تا در کنار کلمه دوم پردازش شود. به عبارت دیگر، درک شبکه از کلمه جاری، تابعی از خود کلمه و دانشی بود که از کلمات پیشین کسب کرده بود.

۲.۱) محدودیت‌های ذاتی مدل‌های بازگشتی: تراژدی فراموشی

اگرچه ظهور Recurrent Neural Network (RNN) گامی بلند نسبت به مدل‌های آماری کلاسیک محسوب می‌شد، اما این معماری با چالش‌های ساختاری و مهلکی دست‌وپنج نرم می‌کرد که مانع از پیشرفت جدی در وظایف پیچیده‌ای نظیر ترجمه متون طولانی یا درک اسناد حقوقی و علمی می‌شد.

۱.۲.۱) مشکل اول: پدیده محو شدن گرادیان

بزرگ‌ترین چالش فنی در آموزش RNN‌ها، ناتوانی آن‌ها در حفظ و انتقال اطلاعات در طول توالی‌های طولانی بود. برای تقریب ذهنی، بازی «تلفن‌بازی» را در نظر بگیرید؛ نفر اول پیامی را در گوش نفر دوم زمزمه می‌کند و این زنجیره ادامه می‌یابد. هرچه صف طولانی‌تر باشد، پیام نهایی دچار تحریف بیشتری شده و شباهت خود را به پیام اصلی از دست می‌دهد.

در شبکه‌های عصبی بازگشتی، هنگام استفاده از الگوریتم آموزش پس‌انتشار^۲، اتفاقی مشابه رخ می‌دهد. سیگنال‌های خطا که باید شبکه را اصلاح کنند، در حین حرکت به عقب در زمان (از انتهای جمله به ابتدا)، در اعداد کوچک‌تری ضرب می‌شوند و به تدریج کوچک و کوچک‌تر می‌گردند تا جایی که عملاً صفر یا محو می‌شوند. این پدیده که در ادبیات فنی «محو شدن گرادیان»^۳ نامیده می‌شود، باعث می‌شود شبکه نتواند ارتباط میان کلماتی را که فاصله زیادی از یکدیگر دارند، درک کند. این ناتوانی در مدیریت وابستگی‌های طولانی‌برد^۴ نقطه ضعف مدل‌های بازگشتی بود.

به عنوان مثال، در جمله: «آن گربه‌ای که هفته پیش در حیاط پشتی خانه همسایه قدیمی ما زیر درخت بلوط نشسته بود و با توپی قرمز بازی می‌کرد، گرسنه است.» برای اینکه مدل دریابد چه کسی گرسنه است، باید کلمه «گربه» را که ده‌ها کلمه قبل‌تر آمده، در حافظه داشته باشد. RNN‌های ساده در برقراری این ارتباطات طولانی شکست می‌خورند و ممکن بود به اشتباه تصور کنند که «درخت» یا «توپ» گرسنه است.

۲.۲.۱) مشکل دوم: گلوگاه پردازش سریالی

چالش دوم ماهیتی محاسباتی داشت. در یک شبکه بازگشتی، پردازش کلمه دهم از نظر ریاضیاتی نمی‌تواند آغاز شود مگر آنکه پردازش کلمه نهم به اتمام رسیده باشد. این وابستگی زمانی^۵ دقیقاً مانند یک مسابقه دوی امدادی است؛ دوندۀ دوم باید منتظر بماند تا دوندۀ اول چوب را به او برساند.

^۱Hidden State

^۲Backpropagation

^۳Vanishing Gradient

^۴Long-Range Dependencies

^۵Temporal Dependency

این ویژگی با روند پیشرفت سخت‌افزارهای کامپیوتری در تضاد کامل بود. کارت‌های گرافیک مدرن یا همان Graphics Processing Unit (GPU) ها، برای انجام هزاران محاسبه پردازش موازی^۱ طراحی شده‌اند، نه محاسبات سریالی. ساختار خطی RNN اجازه نمی‌داد که از قدرت عظیم پردازش موازی استفاده شود، زیرا هر مرحله معطل مرحله قبل بود. این گلوگاه باعث می‌شد آموزش مدل‌های بزرگ روی حجم عظیمی از داده‌ها، بسیار کند، زمان‌بر و پرهزینه باشد.

۳.۱) تلاش‌های اصلاحی: ظهور حافظه‌های پیشرفته

پژوهشگران برای غلبه بر مشکل فراموشی، نسخه‌های تکامل‌یافته‌تری از شبکه‌های بازگشتی را با نام‌های Long Short-Term Memory (LSTM) و Gated Recurrent Unit (GRU) توسعه دادند. این معماری‌ها به مکانیزم‌هایی هوشمند مجهز بودند که مانند یک دروازه^۲ عمل می‌کردند؛ آن‌ها قابلیت این را داشتند که به صورت فعال تصمیم بگیرند چه اطلاعاتی را در حافظه نگه دارند و چه اطلاعاتی را دور بریزند. اگر RNN ساده را مانند شخصی در نظر بگیریم که سعی می‌کند همه چیز را طوطی‌وار حفظ کند و در نهایت گیج می‌شود، LSTM مانند شخصی است که یک دفترچه یادداشت ساختاریافته دارد و تنها نکات کلیدی را ثبت می‌کند. این بهبود معماری باعث شد که عملکرد سیستم‌ها در ترجمه ماشینی و تشخیص گفتار در سال‌های ۲۰۱۴ تا ۲۰۱۶ جهشی چشمگیر داشته باشد. با این حال، مشکل دوم یعنی عدم امکان پردازش موازی همچنان پابرجا بود. مدل‌ها دقیق‌تر شده بودند، اما کماکان کند بودند و مقیاس‌پذیری پایینی داشتند.

۴.۱) بن‌بست اطلاعاتی در معماری رمزگذار-رمزگشا

استاندارد طلایی آن دوران برای ترجمه ماشینی، استفاده از معماری رمزگذار-رمزگشا^۳ مبتنی بر شبکه‌های بازگشتی بود. در این ساختار، بخش رمزگذار وظیفه داشت جمله ورودی را کلمه به کلمه بخواند و تمام معنا و مفهوم آن را در یک بسته ریاضی فشرده و با ابعاد ثابت به نام بردار زمینه^۴ فشرده‌سازی کند. سپس بخش رمزگشا می‌بایست تنها با اتکا به همین بسته فشرده، جمله را به زبان مقصد بازسازی نماید.

تصور کنید از شما بخواهند تمام محتوا، احساسات و جزئیات یک فصل از یک رمان را تنها در یک جمله خلاصه کنید و آن را به فرد دیگری بدهید تا او عین متن اصلی را بازنویسی کند! طبیعتاً بسیاری از ظرایف در این فرآیند فشرده‌سازی از بین می‌روند. این محدودیت که به گلوگاه اطلاعاتی^۵ شهرت یافت، نشان می‌داد که فشرده‌سازی تمام اطلاعات یک جمله طولانی در یک بردار ثابت، روشی بهینه برای انتقال معنا نیست. این دقیقاً همان نقطه‌ای بود که نیاز به رویکردی انقلابی همچون مکانیزم توجه احساس می‌شد.

¹Parallel Processing

²Gate

³Encoder-Decoder Architecture

⁴Context Vector

⁵Information Bottleneck

فصل ۲

انقلاب توجه و معماری ترنسفورمر

در فضای علمی سال‌های پیش از ۲۰۱۷ که مملو از چالش‌های محاسباتی و محدودیت‌های حافظه در مدل‌های بازگشتی بود، انتشار مقاله «توجه تمام آن چیزی است که نیاز دارید»^۱ بسان یک صاعقه بر پیکره تحقیقات هوش مصنوعی فرود آمد. عنوان این مقاله، خود به تنهایی یک مانیفست علمی و بیانی‌ای جسورانه علیه الگوهای سنتی حاکم بود.

نویسندگان این مقاله با طرح یک ادعای رادیکال، بیان کردند که برای توسعه مدل‌های زبانی قدرتمند و پیشرفته، دیگر نیازی به استفاده از ساختارهای پیچیده و زمان‌بر RNN یا Convolutional Neural Networks (CNN) نیست. آن‌ها استدلال کردند که می‌توان تمام معماری‌های پیشین را کنار گذاشت و تنها بر یک مفهوم شناختی تمرکز کرد:

«مکانیزم توجه»^۱

این تغییر رویکرد، نقطه عطفی بود که مسیر تکامل یادگیری عمیق را برای همیشه تغییر داد.

۱.۲) مفهوم توجه در ادراک انسانی

پیش از آنکه به تشریح ریاضیاتی و فنی معماری جدید بپردازیم، ضروری است مفهوم انتزاعی توجه را در بستر شناخت انسانی واکاوی کنیم. سیستم عصبی و شناختی انسان دارای ظرفیت محدودی برای پردازش اطلاعات است و نمی‌تواند تمام داده‌های ورودی از محیط (بصری، شنیداری، حسی) را با دقت یکسان و به صورت همزمان پردازش کند. در عوض، مغز ما مجهز به فیلترهای هوشمندی است که به طور ناخودآگاه بر روی بخش‌های مهم تمرکز کرده و سایر اطلاعات را نادیده می‌گیرد.

۱.۱.۲) تمثيل اثر مهمانی کوکتل

روانشناسان شناختی پدیده‌ای را شناسایی کرده‌اند که بهترین مدل برای تبیین مکانیزم توجه است و به «اثر مهمانی کوکتل»^۲ شهرت دارد. برای درک این پدیده، تصور کنید در یک مهمانی شلوغ، پرهیاهو و مملو از صداهای

^۱ Attention Mechanism

^۲ The Cocktail Party Effect

متداخل حضور دارید. ده‌ها نفر به صورت همزمان در حال گفتگو هستند، موسیقی با صدای بلند پخش می‌شود و صدای برخورد ظروف نیز شنیده می‌شود.

در چنین محیط آشفته‌ای، سیستم شنوایی و مغز شما این قابلیت شگفت‌انگیز را دارد که روی صدای دوستان که در مقابل شما ایستاده است قفل کند و تمام صداهای دیگر را به عنوان نویز پس‌زمینه^۱ فیلتر نماید. اما نکته جالب‌تر اینجاست که اگر ناگهان از آن سوی سالن، کسی نام شما را صدا بزند، با وجود اینکه تمرکز شما کاملاً معطوف به دوستان بود، بلافاصله توجهتان به سمت آن صدا جلب می‌شود (تغییر ناگهانی کانون توجه). این قابلیت مغز برای تخصیص وزن^۲ یا وزندهی به ورودی‌های مختلف، اساس کار مدل‌های جدید هوش مصنوعی است. در این فرآیند:

- صدای دوست شما دارای وزن بالا (اهمیت زیاد) است.

- صدای محیط دارای وزن پایین (اهمیت کم) است.

مکانیزم توجه در هوش مصنوعی دقیقاً تلاشی ریاضیاتی برای شبیه‌سازی همین رفتار است: توانایی مدل برای اینکه یاد بگیرد در هر لحظه، کدام بخش از داده‌ها مهم است و کدام بخش‌ها باید نادیده گرفته شوند.

۲.۲) تولد ترنسفورمر: شکستن زنجیر زمان

معماری نوظهوری که در مقاله سال ۲۰۱۷ معرفی شد و ترنسفورمر نام گرفت، دو تغییر ساختاری و بنیادین را در نحوه پردازش داده‌ها ایجاد کرد. این تغییرات باعث شد محدودیت‌های مدل‌های بازگشتی که در فصل پیشین به آن‌ها اشاره شد، کاملاً مرتفع گردند.

۱.۲.۲ حذف کامل پردازش سریالی

مهم‌ترین دستاورد معماری ترنسفورمر، گسستن زنجیر وابستگی زمانی بود. برخلاف RNN که مجبور بود کلمات را به صورت متوالی (یکی پس از دیگری) بخواند، ترنسفورمر تمام جمله یا حتی یک پاراگراف کامل را به صورت یکجا و همزمان پردازش میکند.

این تغییر الگو به معنای حذف گلوگاه پردازشی بود. در این معماری، دیگر نیازی نیست برای شروع پردازش کلمه آخر، منتظر اتمام پردازش کلمات ابتدایی بمانیم. این ویژگی امکان استفاده حداکثری از قابلیت پردازش موازی در سخت‌افزارهای مدرن (GPU) را فراهم کرد. نتیجه این امر، افزایش خیره‌کننده سرعت آموزش مدل‌ها بود که اجازه داد دانشمندان مدل‌هایی با میلیاردها پارامتر را در زمانی معقول آموزش دهند.

۲.۲.۲ جایگزینی حافظه با مکانیزم توجه

دومین تغییر بزرگ، تغییر در نحوه مدیریت اطلاعات بود. در مدل‌های قدیمی مانند LSTM، شبکه تلاش می‌کرد خلاصه اطلاعات گذشته را در یک حافظه محدود و فشرده با نام (Hidden State) با خود حمل کند. اما ترنسفورمر این رویکرد را کنار گذاشت.

¹ Background Noise

² Weight Allocation

در معماری ترنسفورمر، مفهوم حافظه خطی جای خود را به دسترسی سراسری^۱ داد. مدل در هر لحظه از پردازش، به تمام کلمات جمله دسترسی مستقیم دارد. مکانیزم توجه به مدل اجازه می‌دهد که برای درک هر کلمه، نگاهی به کل جمله بیندازد و تصمیم بگیرد کدام کلمات دیگر برای تفسیر کلمه فعلی ضروری هستند. این یعنی فاصله فیزیکی کلمات در جمله دیگر اهمیتی ندارد؛ کلمه اول می‌تواند به راحتی و با همان دقت کلمه دوم، با کلمه صدم ارتباط برقرار کند.

¹Global Access

فصل ۳

کالبدشکافی معماری ترنسفورمر گذر از انتزاع به ریاضیات

پس از بررسی مبانی نظری و چرایی نیاز به تغییر الگو، اکنون زمان آن فرا رسیده است که ساختار درونی معماری ترنسفورمر را واکاوی کنیم. پرسش بنیادین این است: ماشینی که ماهیت درک آن صرفاً بر پایه عملیات باینری و محاسبات ماتریسی است، چگونه قادر به استخراج ظرایف معنایی از ادبیات کلاسیک یا تحلیل پیچیدگی‌های متون حقوقی می‌باشد؟ پاسخ در معماری مدولار و سلسله‌مراتب ریاضیاتی این مدل نهفته است.

۱.۳) معماری مبتنی بر پشته‌های رمزگذار و رمزگشا

ساختار کلان ترنسفورمر بر پایه الگوی کلاسیک Sequence to Sequence (Seq2Seq) بنا شده است، اما با پیاده‌سازی متفاوت. این معماری از دو مؤلفه یا پشته^۱ اصلی تشکیل شده است:

۱. پشته رمزگذار^۲: این بخش به عنوان تحلیل‌گر عمل می‌کند. وظیفه آن دریافت رشته ورودی (متن خام) و نگاشت آن به یک فضای نمایش انتزاعی پیوسته است. رمزگذار با خواندن متن و تحلیل روابط میان کلمات، یک درک عمیق ریاضی از محتوا ایجاد می‌کند.

۲. پشته رمزگشا^۳: این بخش نقش تولیدگر را ایفا می‌کند. رمزگشا با دریافت خروجی رمزگذار (که حاوی عصاره معنایی متن است)، اقدام به تولید خروجی نهایی (مانند ترجمه متن یا پاسخ به سوال) به صورت گام‌به‌گام می‌نماید.

^۱ Stack

^۲ Encoder

^۳ Decoder

۲.۳) توکن‌گذاری و تعبیه‌ی کلمات: تبدیل واژه به مختصات معنایی

ماشین‌ها فاقد درک ذاتی از زبان هستند و تنها اعداد را پردازش می‌کنند. گام نخست در فرآیند پردازش، نگاشت کلمات گسسته به بردارهای پیوسته در یک فضای ریاضی چندبعدی است که به آن تعبیه کلمات^۱ اطلاق می‌شود. در این فرآیند، هر کلمه به یک بردار^۲ (لیستی از اعداد اعشاری) تبدیل می‌شود. ویژگی کلیدی این فضا آن است که فاصله هندسی بردارها با شباهت معنایی کلمات همبستگی دارد. به عنوان مثال، در این فضای برداری، مختصات کلمه “پادشاه” به “ملکه” و “سیب” به “پرتقال” بسیار نزدیک است و یا اختلاف مختصات کلمات “مادر” و “دختر” با اختلاف کلمات “پدر” و “پسر” یکسان است. این تبدیل، زیربنای درک ماشین از روابط معنایی واژگان است. پیش از آنکه مدل بتواند کلمات را به بردارهای ریاضی تبدیل کند، متن خام باید به واحدهای معنادار کوچکتری به نام توکن^۳ تقسیم شود. این فرآیند که توکن‌گذاری^۴ نامیده می‌شود، پل ارتباطی میان زبان انسانی و دنیای اعداد است.

۱.۲.۳) مفهوم توکن: واحد بنیادین پردازش

توکن‌ها لزوماً کلمات کامل نیستند؛ بسته به الگوریتم مورد استفاده، یک توکن می‌تواند یک کلمه، یک پیشوند، یک پسوند یا حتی تک‌حروف باشد.

● چرا از کلمه کامل استفاده نمی‌شود؟ اگر هر کلمه یک توکن واحد باشد، مدل در مواجهه با کلمات جدید یا مشتقات پیچیده مانند “ورشکستگی‌ها” دچار مشکل می‌شود. این پدیده باعث بزرگ شدن بی‌رویه فضای دایره لغات^۵ و افزایش پارامترهای مدل می‌گردد.

● رویکرد زیر-واحد^۶: مدل‌های مدرن متن را به قطعات کوچکتری می‌شکنند. به عنوان مثال، کلمه “بی‌برنامه” ممکن است به دو توکن “بی” و “برنامه” تقسیم شود. این کار به مدل اجازه می‌دهد با تعداد محدودی توکن در دایره لغات ثابت، بی‌نهایت کلمه و ترکیب جدید را درک و معناییابی کند.

پس از آنکه متن به توکن‌ها تجزیه شد، هر توکن به یک عدد صحیح منحصر به فرد (شناسه^۷) اختصاص می‌یابد. اما از آنجا که این اعداد صرفاً شناسنامه هستند و هیچ رابطه ریاضی با هم ندارند (مثلاً عدد اختصاص یافته به کلمه “گربه” لزوماً نزدیک به عدد “سگ” نیست)، گام بعدی نگاشت این اعداد به یک فضای برداری پیوسته^۸ یا همان لایه تعبیه^۹ است.

^۱ Word Embeddings

^۲ Vector

^۳ Token

^۴ Tokenization

^۵ Vocabulary

^۶ Sub-word

^۷ ID

^۸ Continuous Vector Space

^۹ Embedding

۲.۲.۳) لایه‌ی تعبیه‌ی کلمات: نگاشت به فضای برداری پیوسته

پس از مرحله توکن‌گذاری، ما مجموعه‌ای از اعداد صحیح (شناسه‌ها) داریم. اما این اعداد فاقد ارزش ریاضی هستند؛ برای مثال عدد اختصاص یافته به کلمه «مرد» و «زن» ممکن است بسیار دور از هم باشد، در حالی که از نظر معنایی به هم نزدیک‌اند. لایه‌ی تعبیه کلمات^۱ وظیفه دارد هر توکن را به یک بردار با ابعاد ثابت ($d_{model} = 512$) در یک فضای پیوسته تصویر کند.

۱.۲.۲.۳) ویژگی‌های ریاضی فضای تعبیه

این لایه صرفاً یک جدول جستجو^۲ بزرگ با ابعاد $V \times d_{model}$ است که در آن V نشان‌دهنده‌ی تعداد کل توکن‌های لغت‌نامه است. ویژگی‌های کلیدی این لایه عبارتند از:

۱. **تشابه معنایی و فاصله اقلیدسی^۳:** در این فضای ۵۱۲ بعدی، بردارها به گونه‌ای آموزش می‌بینند که توکن‌های با معنای مشابه، در فاصله نزدیک‌تری از یکدیگر قرار بگیرند. برای مثال، تفاضل برداری میان «ملکه» و «زن» باید تقریباً معادل تفاضل برداری میان «پادشاه» و «مرد» باشد.

۲. **توزیع اطلاعات^۴:** هر بعد از این بردار ۵۱۲ بعدی، جنبه‌ای از مفهوم توکن را رمزگذاری می‌کند؛ برای نمونه یک بعد ممکن است مربوط به زمان، بعدی دیگر مربوط به جنسیت و موارد مشابه باشد، هرچند که این ابعاد برای انسان مستقیماً قابل تفسیر نیستند.

۳. **هم‌مقیاس شدن^۵ متوازن:** نویسندگان مقاله بردار کلمات را در $\sqrt{d_{model}}$ ضرب می‌کنند. این کار باعث می‌شود مقادیر بردار کلمات با مقادیر «کدگذاری موقعیتی» هم‌مقیاس شوند تا هنگام جمع شدن، اطلاعات مکانی باعث محو شدن اطلاعات معنایی نشود.

در حقیقت، لایه‌ی تعبیه، کلمات را از یک بازنمایی گسسته و بی‌معنا به یک نمایش هندسی^۶ غنی تبدیل می‌کند که در آن مفاهیم زبانی با استفاده از روابط برداری قابل محاسبه هستند.

۳.۳) کدگذاری موقعیتی: حل معمای آشفستگی زمانی

یکی از چالش‌های ناشی از حذف پردازش سریالی (که در فصل قبل اشاره شد)، از دست رفتن مفهوم بود. از آنجا که مدل ترنسفورمر تمام کلمات را به صورت همزمان پردازش می‌کند، ذاتاً نمی‌تواند تفاوت میان دو جمله «علی زد حسن را» و «حسن زد علی را» را تشخیص دهد، زیرا مجموعه کلمات در هر دو یکسان است. برای غلبه بر این چالش، نویسندگان مقاله راهکاری هوشمندانه تحت عنوان رمزگذاری موقعیتی^۷ ارائه دادند. در این روش، اطلاعات مربوط به جایگاه هر کلمه در جمله، از طریق جمع کردن بردار کلمه با یک بردار موقعیت

¹ Word Embedding

² Lookup Table

³ Euclidean Distance

⁴ Information Distribution

⁵ Scaling

⁶ Geometric Representation

⁷ Positional Encoding

خاص (PE)، به مدل تزریق می‌شود. نویسندگان به جای استفاده از اعداد صحیح ساده که در جملات طولانی باعث ناپایداری محاسباتی^۱ و رشد نامحدود مقادیر می‌شوند، از توابع مثلثاتی با فرکانس‌های متفاوت استفاده کردند.

فرمول‌های ریاضی این کدگذاری برای هر جایگاه pos و هر بعد^۲ i به صورت زیر تعریف می‌شوند:

$$PE_{(pos,2i)} = \sin\left(\frac{pos}{10000^{2i/d_{model}}}\right) \quad (۱.۳)$$

$$PE_{(pos,2i+1)} = \cos\left(\frac{pos}{10000^{2i/d_{model}}}\right) \quad (۲.۳)$$

در این روابط، pos نشان‌دهنده موقعیت توکن در دنباله و i نشان‌دهنده شاخص بعد در فضای نمایش است. انتخاب این ساختار سینوسی و کسینوسی بر پایه یک استدلال منطقی بنا شده است:

این توابع سینوسی و کسینوسی به مدل اجازه می‌دهند تا با درک فرکانس سیگنال هر کلمه، جایگاه نسبی و مطلق آن را در جمله شناسایی کند.

از منظر تحلیل ریاضی، مهم‌ترین ویژگی این فرمول‌بندی آن است که برای هر فاصله ثابت k ، بردار موقعیت PE_{pos+k} می‌تواند به صورت یک تبدیل خطی^۳ از بردار PE_{pos} بیان شود. این ویژگی برای یک مدل ریاضیاتی حیاتی است، زیرا به شبکه عصبی اجازه می‌دهد به سادگی یاد بگیرد که چگونه با توجه به فواصل نسبی^۴، به اطلاعات مکانی^۵ میان کلمات دسترسی پیدا کند. در واقع، با این روش، مدل نه تنها می‌داند کلمه در کجا قرار دارد، بلکه نسبت فاصله آن با سایر کلمات را نیز با دقت ریاضی بالایی درک می‌کند.

۴.۳ مکانیزم توجه خودکار: قلب تپنده مدل

محوری‌ترین نوآوری این مقاله، مکانیزم توجه خودکار^۶ است. این مکانیزم به کلمات اجازه می‌دهد تا با یکدیگر تعامل داشته باشند و میزان وابستگی معنایی خود را در یک بافتار مشخص کنند.

برای تبیین این مفهوم، جمله زیر را در نظر بگیرید: “بانک پول نداشت چون آن ورشکسته شده بود.” ما به عنوان انسان می‌دانیم که ضمیر “آن” به “بانک” اشاره دارد. اما برای ماشین، کلمه “آن” می‌تواند از نظر ساختاری به “پول” نیز اشاره داشته باشد. مکانیزم توجه خودکار، هنگام پردازش کلمه “آن”، تمام کلمات دیگر جمله را اسکن می‌کند تا مرتبط‌ترین واژه را بیابد. در اینجا، وجود کلمه “ورشکسته” که صفتی برای نهادهای مالی است، باعث می‌شود مدل وزن توجه بالایی را میان “آن” و “بانک” برقرار کرده و ابهام ضمیر را برطرف سازد.

^۱Computational Instability

^۲Dimension

^۳Linear Transformation

^۴Relative Distances

^۵Spatial Information

^۶Self-Attention

۱.۴.۳) تشریح فنی با مدل سه گانه پرس وجو، کلید و مقدار

نویسندگان برای مدل سازی ریاضی این تعامل، از مفاهیم بازیابی اطلاعات در پایگاه های داده الهام گرفتند. برای هر کلمه، سه بردار مجزا از طریق ضرب بردار ورودی در سه ماتریس وزن آموزش پذیر (W^Q, W^K, W^V) تولید می شود:

پرس وجو^۱ (Q): نشان دهنده سیگنالی است که کلمه برای جستجوی اطلاعات ارسال می کند.

کلید^۲ (K): نقش شناسه یا برجسب کلمه را برای سایر کلمات ایفا می کند.

مقدار^۳ (V): حاوی محتوای اصلی و عصاره معنایی کلمه است.

۱.۱.۴.۳) فرمول ریاضی توجه ضرب داخلی مقیاس شده

فرآیند محاسبه توجه در مقاله تحت عنوان توجه ضرب داخلی مقیاس شده^۴ معرفی شده است. این عملیات به صورت ماتریسی با فرمول زیر بیان می گردد:

$$\text{Attention}(Q, K, V) = \text{softmax} \left(\frac{QK^T}{\sqrt{d_k}} \right) V \quad (۳.۳)$$

در این معادله، ضرب داخلی QK^T میزان شباهت میان پرس وجوها و کلیدها را محاسبه می کند. اما نکته حائز اهمیت برای تحلیل ریاضی، حضور عبارت $\frac{1}{\sqrt{d_k}}$ است. نویسندگان استدلال کردند که در ابعاد بالا (d_k بزرگ)، حاصل ضرب داخلی می تواند به اعداد بسیار بزرگی ختم شود. این بزرگی مقادیر باعث می شود تابع سافت مکس^۵ وارد نواحی تخت با گرادیان^۶ بسیار کوچک شود که اصطلاحاً به آن اشباع شدن^۷ می گویند. تقسیم بر $\sqrt{d_k}$ باعث نرمال سازی و پایداری واریانس این مقادیر شده و فرآیند یادگیری را بهبود می بخشد. در نهایت، با اعمال تابع سافت مکس، وزن های توجه به دست می آیند که به صورت یک ترکیب خطی^۸ روی بردارهای مقدار اعمال می شوند تا نمایش نهایی کلمه در آن بافت خاص شکل بگیرد.

۲.۱.۴.۳) فرآیند محاسبه توجه

فرآیند ریاضیاتی توجه را می توان به یک مکالمه تشبیه کرد:

۱. کلمه "آن" (به عنوان پرس وجو) سیگنالی ارسال می کند: «من به دنبال مرجعی هستم که قابلیت ورشکستگی داشته باشد.»

۲. کلمه "بانک" (به عنوان کلید) پاسخ می دهد: «من یک نهاد مالی هستم و تطابق بالایی با درخواست تو دارم.»

۳. کلمه "پول" (به عنوان کلید) پاسخ می دهد: «من شیء فیزیکی هستم و تطابق کمی دارم.»

^۴Scaled Dot-Product Attention

^۵Softmax

^۶Gradient

^۷Saturation

^۸Linear Combination

مدل با محاسبه شباهت برداری (ضرب داخلی) میان Q و K ها، وزن‌های توجه را محاسبه می‌کند. در نهایت، ترکیبی وزن‌دار از بردارهای مقدار ایجاد می‌شود که معنای جدید کلمه “آن” را در بافت این جمله خاص شکل می‌دهد. این عملیات در فرمول‌بندی مقاله «توجه ضرب داخلی مقیاس‌شده» نامیده شده است.

۵.۳) توجه چندسر: نگرش چندبعدی به زبان

آیا یک کلمه در جمله تنها یک نقش ایفا می‌کند؟ قطعاً خیر. کلمات همزمان دارای نقش‌های نحوی (مانند فاعل یا مفعول)، معنایی و ارجاعی هستند. یک مکانیزم توجه واحد ممکن است تنها بر روی یکی از این جنبه‌ها تمرکز کند. برای حل این محدودیت، ترنسفورمر از تکنیک توجه چند سر^۱ بهره می‌برد. در این روش، مدل به جای یک بار، چندین بار (مثلاً $h = 8$ بار) و به صورت موازی فرآیند توجه را انجام می‌دهد. فرمول‌بندی ریاضی این لایه به صورت زیر است:

$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_h)W^O \quad (۴.۳)$$

که در آن هر سر به صورت مستقل در یک زیرفضای متفاوت محاسبه می‌گردد:

$$\text{head}_i = \text{Attention}(QW_i^Q, KW_i^K, VW_i^V) \quad (۵.۳)$$

۱.۵.۳) تحلیل نقش ماتریس‌های وزن و تبدیل خطی

در این ساختار، دو گام اساسی از منظر جبر خطی رخ می‌دهد که درک آن‌ها برای تبیین مدل ضروری است:

۱. چرا Q, K, V در ماتریس‌های وزن ضرب می‌شوند؟ بردارهای ورودی در ابتدا ثابت هستند. ضرب کردن آن‌ها در ماتریس‌های W_i^Q, W_i^K, W_i^V در واقع یک نگاشت خطی^۲ است. این کار باعث می‌شود که هر سر توجه^۳، بردارها را به یک زیرفضای برداری^۴ جدید با ابعاد کمتر ($d_k = d_{\text{model}}/h$) ببرد. در این زیرفضاهای جدید، مدل آزاد است تا روابط متفاوتی را یاد بگیرد؛ برای مثال، یک سر یاد می‌گیرد که در فضای خود، روی فواصل کوتاه و روابط دستوری تمرکز کند و سر دیگر روی روابط معنایی دوربرد. بدون این ضرب ماتریسی، تمام سرها عملاً یک چیز واحد را تماشا می‌کردند.

۲. نقش ماتریس نهایی W^O چیست؟ پس از آنکه هر سر خروجی خود را تولید کرد، ما با مجموعه‌ای از بردارها روبرو هستیم که هر کدام از دیدگاه خاصی به متن نگریسته‌اند. با عمل متصل کردن^۵، این اطلاعات در کنار هم چیده می‌شوند. اما برای آنکه این اطلاعات پراکنده دوباره با هم تعامل کرده و به ابعاد اصلی مدل بازگردند، از ماتریس W^O استفاده می‌شود. این ماتریس وظیفه تلفیق^۶ دانش استخراج شده توسط

^۱Multi-Head Attention

^۲Linear Projection

^۳Attention Head

^۴Vector Subspace

^۵Concatenate

^۶Aggregation

تمام سرها را بر عهده دارد و وزنهای هر سر را به گونه‌ای ترکیب می‌کند که خروجی نهایی، عصاره‌ای منسجم از تمام نگاه‌های تخصصی باشد.

در واقع، هر سر مسئول استخراج نوع خاصی از روابط است:

- **سر اول:** روی روابط دستوری و نحوی تمرکز می‌کند.

- **سر دوم:** روابط زمانی افعال را بررسی می‌کند.

- **سر سوم:** مترادف‌ها و قرابت‌های معنایی را می‌یابد.

گویی چندین متخصص همزمان در حال تحلیل جمله از زوایای مختلف هستند و در نهایت ماتریس W^O مانند یک سردبیر، نظرات تمام آن‌ها را برای رسیدن به یک نتیجه واحد و جامع ترکیب می‌کند.

۶.۳) شبکه‌ی پیش‌خور و اتصالات میان‌بر: تثبیت یادگیری

پس از آنکه مکانیزم توجه، روابط پیچیده و بافتی میان کلمات را استخراج کرد، بردار خروجی به یک شبکه‌ی پیش‌خور نقطه-به-نقطه^۱ ارسال می‌شود. این شبکه به صورت مستقل و کاملاً یکسان روی هر جایگاه (توکن) اعمال می‌گردد؛ بدین معنا که پارامترها برای تمام توکن‌های یک لایه مشترک است، اما از لایه‌ای به لایه دیگر تغییر می‌کند.

۱.۶.۳) تحلیل ساختاری و هندسی لایه پیش‌خور

از منظر ساختاری، این بخش شامل دو تبدیل خطی (نگاشت آفین) است که یک تابع فعال‌ساز غیرخطی تابع فعال‌ساز خطی اصلاح‌شده^۲ میان آن‌ها قرار دارد. فرمول ریاضی این فرآیند به صورت زیر تدوین می‌گردد:

$$\text{FFN}(x) = \max(0, xW_1 + b_1)W_2 + b_2 \quad (۶.۳)$$

در این معادله، مؤلفه‌های ریاضی نقش‌های زیر را ایفا می‌کنند:

- **انبساط ابعادی^۳:** نگاشت اول $(xW_1 + b_1)$ بردار ورودی را از فضای مدل $(d_{\text{model}} = 512)$ به یک فضای میانی بسیار بزرگ‌تر $(d_{\text{ff}} = 2048)$ منتقل می‌کند. این افزایش ابعاد به مدل اجازه می‌دهد تا ویژگی‌های استخراج شده توسط مکانیزم توجه را در یک فضای با ظرفیت بالاتر، بازنمایی و تفکیک کند.

- **غیرخطی‌سازی:** تابع $\max(0, xW_1 + b_1)$ با حذف مقادیر منفی، ویژگی‌های غیرخطی را به مدل تزریق می‌کند که برای یادگیری الگوهای پیچیده زبانی ضروری است.

^۱Position-wise Feed-Forward Network

^۲ReLU

^۳Dimensional Expansion

● **انقباض و بازگشت:** نگاشت دوم (W_2) وظیفه دارد داده‌ها را مجدداً به ابعاد اصلی مدل بازگرداند تا برای ورود به لایه‌های بعدی آماده شوند.

در حالی که لایه‌ی توجه وظیفه‌ی برقراری ارتباطات «افقی» میان کلمات مختلف را بر عهده داشت، لایه Feed-Forward Network (FFN) وظیفه دارد ویژگی‌های استخراج شده برای هر کلمه را به صورت «عمودی» و مجزا پردازش، پالایش و تقویت کند.

۲.۶.۳) اتصالات میان‌بر و نرمال‌سازی: تضمین پایداری ریاضی

یک چالش بنیادین در شبکه‌های عصبی عمیق، پدیده تأثیر محوشدگی گرادیان^۱ است. در فرآیند پس‌انتشار، مشتقات متوالی ممکن است به سمت صفر میل کنند و مانع از بروزرسانی وزن‌های لایه‌های اولیه شوند. ترنسفورمر برای حل این بن‌بست ریاضی، از ساختار «جمع و نرمال‌سازی» بهره می‌برد:

$$\text{Output} = \text{LayerNorm}(x + \text{Sublayer}(x)) \quad (۷.۳)$$

این ساختار از دو بخش مکمل تشکیل شده است:

۱. **اتصال میان‌بر^۲:** با جمع کردن ورودی (x) با خروجی زیرلایه، یک بزرگراه اطلاعاتی ایجاد می‌شود. از نظر محاسباتی، این کار باعث می‌شود که در هنگام مشتق‌گیری، عبارت $\frac{\partial(x+f(x))}{\partial x} = 1 + f'(x)$ حاصل شود. وجود عدد یک تضمین می‌کند که حتی اگر مشتق لایه‌ها (f') بسیار کوچک باشد، گرادیان همچنان جریان یافته و به لایه‌های ابتدایی برسد.

۲. **نرمال‌سازی لایه^۳:** این عملیات با ثابت نگاه داشتن میانگین و واریانس فعالیت‌های عصبی در هر لایه، از نوسانات شدید داده‌ها جلوگیری می‌کند. این پایداری عددی، سرعت همگرایی مدل را به شدت افزایش داده و اجازه می‌دهد نرخ یادگیری^۴ بالاتری در آموزش استفاده شود.

این ترکیب هوشمندانه، امکان آموزش مدل‌هایی با ده‌ها و صدها لایه را فراهم می‌سازد، بدون آنکه پایداری ریاضی مدل بر اثر افزایش عمق از دست برود.

۷.۳) تولید خروجی نهایی: لایه‌ی خطی و تابع نرم‌بیشینه

پس از آنکه داده‌ها از آخرین بلوک پشته‌ی رمزگشا عبور کردند، خروجی همچنان به صورت مجموعه‌ای از بردارهای پیوسته در فضای نمایش انتزاعی قرار دارد. برای آنکه این بردارها به کلمات قابل فهم و توزیع احتمال تبدیل شوند، مدل دو گام ریاضی زیر را طی می‌کند:

^۱ Vanishing Gradient Effect

^۲ Residual Connection

^۳ Layer Normalization

^۴ Learning Rate

۱. **لایه خطی**^۱: این لایه یک شبکه‌ی پیش‌خور ساده است که وظیفه‌ی نگاشت^۲ بردار خروجی رمزگشا را به فضای دایره لغات بر عهده دارد. اگر تعداد کلمات موجود در دایره لغات مدل را V در نظر بگیریم، این لایه بردار خروجی را به برداری با طول V تبدیل می‌کند. هر مؤلفه در این بردار، نشان‌دهنده یک امتیاز خام^۳ برای هر کلمه در فرهنگ لغت است.

۲. **تابع سافت‌مکس**: امتیازات حاصل از لایه‌ی خطی، که در ادبیات یادگیری ماشین تحت عنوان امتیازات خام^۴ شناخته می‌شوند، اعدادی حقیقی در بازه‌ی $(-\infty, +\infty)$ هستند. این مقادیر به تنهایی فاقد تفسیر احتمالی مستقیم هستند، چرا که نه در بازه‌ی واحد قرار دارند و نه مجموع آن‌ها لزوماً برابر با یک است. برای تبدیل این بردار از امتیازات خام به یک توزیع احتمالاتی مجاز، از تابع نگاشت سافت‌مکس استفاده می‌شود.

این تابع، بردار $Z = (z_1, z_2, \dots, z_V)$ را به بردار احتمالات $P = (p_1, p_2, \dots, p_V)$ نگاشت می‌کند، به طوری که برای هر مؤلفه z_i داریم:

$$\sigma(z)_i = \frac{e^{z_i}}{\sum_{j=1}^V e^{z_j}} \quad (۸.۳)$$

بسط تحلیلی این فرآیند از چند منظر حائز اهمیت است:

- **تضمین اصل موضوع‌پذیری احتمال**: تابع سافت‌مکس به دلیل ماهیت تابع نمایی، تمام خروجی‌ها را مثبت گردانده و با تقسیم بر مجموع (برای نرمال‌سازی)، تضمین می‌کند که $\sum p_i = 1$ برقرار باشد. این امر خروجی مدل را به یک توزیع احتمال گسسته^۵ روی دایره لغات تبدیل می‌کند.
- **تقویت تفاوت‌ها**: از منظر ریاضی، سافت‌مکس یک نسخه‌ی نرم و مشتق‌پذیر از تابع بیشینه است. این تابع به گونه‌ای عمل می‌کند که بزرگترین امتیاز در ورودی (z_{max}) ، بیشترین سهم از احتمال را به خود اختصاص می‌دهد و سایر امتیازات به صورت نمایی تضعیف می‌شوند. این ویژگی برای فرآیند آموزش و پسانتشار حیاتی است، زیرا اجازه می‌دهد سیگنال خطا از میان این تابع عبور کند.
- **تفسیر هندسی**: این تابع در واقع بردار خروجی لایه‌ی خطی را بر روی یک سیمپلکس احتمالی^۶ در فضای V - بعدی تصویر می‌کند. در مرحله‌ی تولید متن، مدل با نمونه‌برداری از این توزیع یا انتخاب کلمه‌ای با بیشترین احتمال از طریق رمزگشایی حریصانه^۷، توکن بعدی را مشخص می‌سازد.

این عملیات، گام نهایی در تبدیل محاسبات انتزاعی پشته‌ی رمزگشا به تصمیمات زبانی ملموس است که دقت آن مستقیماً بر کیفیت متن تولید شده تأثیر می‌گذارد.

^۱Linear Layer

^۲Projection

^۳Logit

^۴Logits

^۵Discrete Probability Distribution

^۶Probability Simplex

^۷Greedy Decoding

احتمالات خروجی^۱: در نهایت، کلمه‌ای که بالاترین مقدار احتمال را به خود اختصاص داده است، به عنوان پیش‌بینی مدل برای گام زمانی فعلی انتخاب می‌شود. این کلمه برای تولید توکن‌های بعدی، مجدداً به بخش ورودی‌های رمزگشا (خروجی‌های شیف‌ت یافته به راست) بازگردانده می‌شود تا فرآیند به صورت خودرگرسیو^۲ ادامه یابد.

۸.۳) تحلیل توپولوژیک و جریان داده در معماری ترنسفورمر

اکنون که مؤلفه‌های بنیادین ریاضیاتی نظیر مکانیزم توجه و شبکه‌های پیش‌خور را واکاوی کردیم، زمان آن است که نگاهی جامع به نحوه اتصال این اجزا و تشکیل گراف محاسباتی ترنسفورمر بیاندازیم. نمودار زیر (شکل ۱.۳) ساختار کلی مدل را نمایش می‌دهد که می‌توان آن را به سه بخش اصلی تقسیم کرد: ستون رمزگذار (چپ)، ستون رمزگشا (راست) و پل ارتباطی میان آن‌ها.

۱.۸.۳) مسیر جریان در پشته رمزگذار (سمت چپ)

فرآیند پردازش در ستون سمت چپ که معادل رمزگذار است، طی مراحل زیر انجام می‌پذیرد:

۱. **فضای برداری اولیه:** توکن‌های ورودی ابتدا از طریق لایه تعبیه ورودی^۲ به بردارهای متراکم تبدیل شده و بلافاصله با سیگنال‌های رمزگذاری موقعیتی جمع می‌شوند تا اطلاعات مکانی به آن‌ها افزوده گردد.

۲. **بلوک‌های تکرارشونده:** داده‌ها وارد بلوکی می‌شوند که N بار تکرار می‌شود (در نسخه اصلی مقاله $N = 6$). در هر بلوک، جریان داده از دو زیرلایه اصلی عبور می‌کند:

- ابتدا لایه توجه چند سر روابط درونی متن را استخراج می‌کند.
- سپس لایه شبکه پیش‌خور^۳ پردازش غیرخطی را انجام می‌دهد.

۳. **اتصالات میانبر و نرمال‌سازی:** نکته حیاتی در این معماری که در نمودار با فلش‌های کناری مشخص شده است، وجود اتصالات میانبر^۴ است. خروجی هر زیرلایه با ورودی همان زیرلایه جمع شده و سپس وارد لایه نرمال‌سازی لایه (بلوک‌های زرد رنگ) می‌شود. این مکانیزم از نظر ریاضیاتی گرادینان را در طول شبکه حفظ کرده و همگرایی را تضمین می‌کند.

۲.۸.۳) مسیر جریان در پشته رمزگشا (سمت راست)

ستون سمت راست یا رمزگشا، ساختاری مشابه دارد اما با دو تفاوت ساختاری مهم که ناشی از ماهیت خودرگرسیو^۵ (تولید کلمه به کلمه) است:

¹ Output Probabilities

² Auto-regressive

² Input Embedding

³ Feed-Forward

⁴ Residual Connections

⁵ Auto-Regressive

۱. پوشش اطلاعات آینده: در اولین لایه توجه (بلوک پایین سمت راست)، از نوعی خاص از توجه به نام توجه چندسر پوشیده^۱ استفاده شده است. این لایه تضمین می‌کند که در هنگام پیش‌بینی کلمه t ، مدل هیچ دسترسی به کلمات $t + 1$ به بعد نداشته باشد.

۲. تولید توزیع احتمال: در انتهای پشته رمزگشا، خروجی نهایی از یک لایه خطی عبور کرده تا به ابعاد دایره لغات پروژه شود و سپس تابع سافت‌مکس آن را به یک توزیع احتمالاتی معتبر تبدیل می‌کند که کلمه بعدی بر اساس آن انتخاب می‌شود.

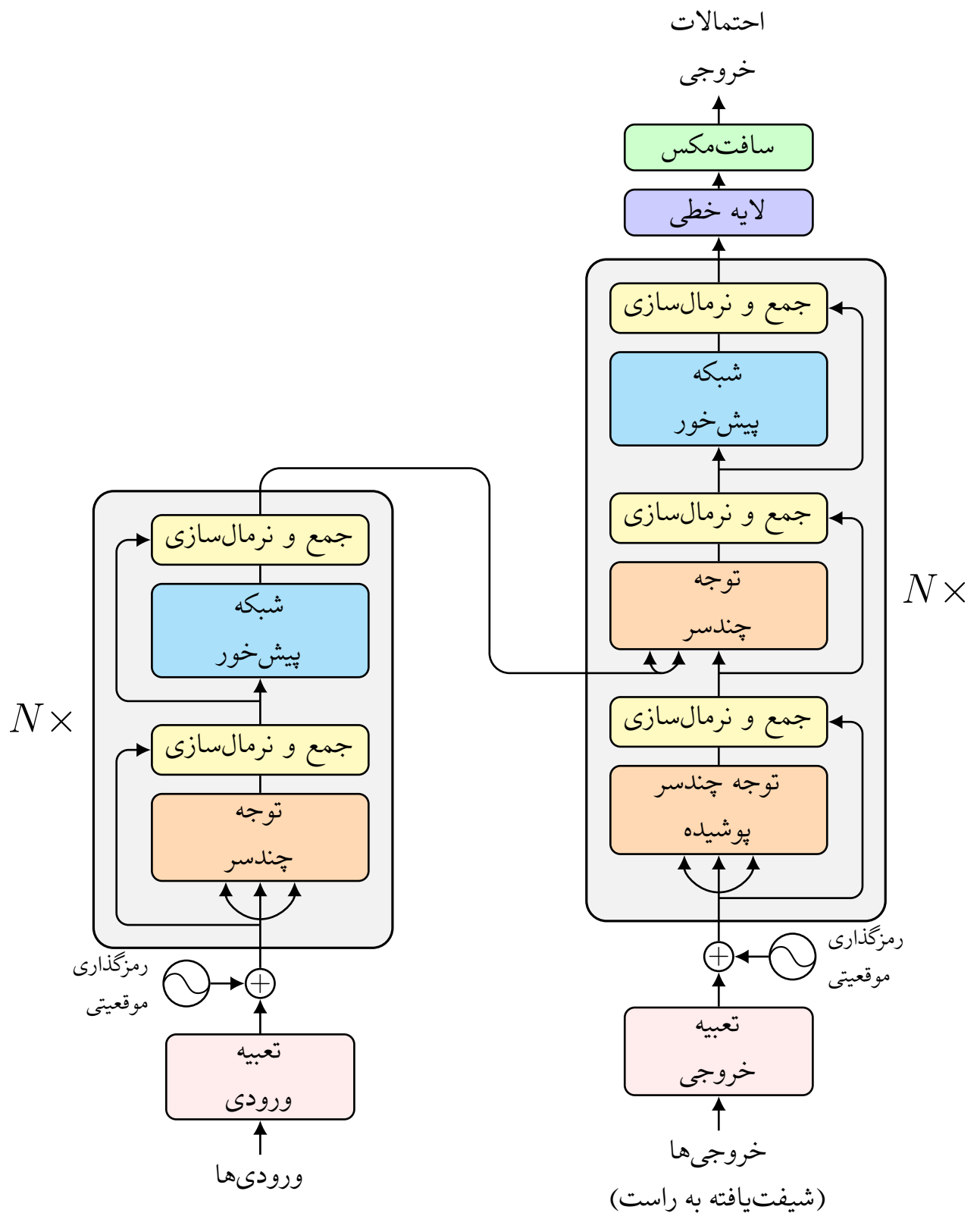
۳.۸.۳ مکانیزم توجه متقاطع: اتصال دو نیمکره

مهم‌ترین بخش نمودار، جایی است که دو ستون به هم متصل می‌شوند. اگر به بلوک میانی در سمت راست دقت کنید، متوجه می‌شوید که این لایه توجه (توجه چند سر) ورودی‌های خود را از دو منبع متفاوت دریافت می‌کند:

- ورودی کوئری (Q): از لایه پایین دستی در خود رمزگشا می‌آید (آنچه مدل قصد دارد بگوید).
- ورودی کلید و مقدار (K, V): از خروجی نهایی پشته رمزگذار (سمت چپ) می‌آید (آنچه متن اصلی بیان کرده است).

این اتصال همان نقطه‌ای است که مدل «زبان مقصد» را با «مفاهیم زبان مبدأ» هم‌تراز می‌کند.

¹Masked Multi-Head Attention



شکل ۱.۳: نمای شماتیک معماری ترنسفورمر. ستون سمت چپ (رمزگذار) وظیفه استخراج ویژگی‌ها و ستون سمت راست (رمزگشا) وظیفه تولید دنباله خروجی را بر عهده دارد. بردارهای خروجی رمزگذار به عنوان ورودی‌های K و V به لایه‌های میانی رمزگشا تغذیه می‌شوند.

فصل ۴

چرا ترنسفورمر جهان را تغییر داد؟

اهمیت و جایگاه مقاله «توجه تمام آن چیزی است که نیاز دارید» در تاریخ علم داده، فراتر از معرفی صرف یک معماری جدید است؛ این پژوهش، دروازه‌ای را به سوی امکاناتی گشود که تا پیش از آن در زمره رویاهای محاسباتی دست‌نیافتنی طبقه‌بندی می‌شدند. این تغییر الگو، نتیجه همگرایی چندین عامل فنی و سخت‌افزاری بود که در ادامه به تشریح آن‌ها می‌پردازیم.

۱.۴) پیروزی موازی‌سازی و هم‌افزایی با قانون مور

همان‌گونه که در فصول پیشین تبیین شد، نقطه‌ضعف مدل‌های بازگشتی (RNN)، ماهیت سریالی آن‌ها بود. این وابستگی، گلوگاهی ساختاری ایجاد می‌کرد که با پیشرفت‌های سخت‌افزاری ناسازگار بود. از منظر پیچیدگی محاسباتی^۱، شبکه‌های بازگشتی دارای «پیچیدگی زمانی^۲» خطی یا $O(N)$ هستند؛ بدین معنا که زمان لازم برای پردازش، رابطه‌ای مستقیم و غیرقابل شکست با طول جمله دارد. اما معماری ترنسفورمر با حذف این قید زمانی، اجازه داد تا محاسبات از حالت توالی خارج شده و به فرم موازی‌سازی ماتریسی^۳ درآیند. این ویژگی انقلابی باعث شد پژوهشگران بتوانند از تمام ظرفیت کلاستر^۴ های عظیم GPU و نسل جدید سخت‌افزارها مانند «Tensor Processing Units (TPU)» بهره‌برداری کنند. این هم‌افزایی با قانون مور^۵ (که پیش‌بینی‌کننده پیشرفت سخت‌افزار است) باعث شد ناگهان آموزش مدل‌هایی با میلیاردها پارامتر بر روی کل داده‌های موجود در ویکی‌پدیا^۶ و اینترنت امکان‌پذیر شود.

اگر RNN‌ها را به مانند پر کردن یک استخر با سطل آب در نظر بگیریم، ترنسفورمرها مانند باز کردن همزمان هزاران شیر فشار قوی عمل می‌کنند.

^۱Computational Complexity

^۲Time Complexity

^۳Matrix Parallelization

^۴Cluster

^۵Moore's Law

^۶Wikipedia

۲.۴) مدیریت وابستگی‌های بافاصله: حذف محدودیت‌های مکانی

در الگوهای کلاسیک پردازش زبان، فاصله فیزیکی میان کلمات، بزرگترین دشمن درک معنایی بود. هرچه فاصله بین فاعل و فعل یا ضمیر و مرجع آن بیشتر می‌شد، احتمال خطا و فراموشی مدل به صورت نمایی افزایش می‌یافت. اما در معماری ترنسفورمر، مفهوم فاصله بازتعریف شده است. به مدد مکانیزم توجه، طول مسیر بیشینه^۱ یا مسافتی که سیگنال باید طی کند تا از کلمه اول به کلمه آخر برسد، به $O(1)$ تقلیل یافته است. این بدان معناست که دسترسی کلمه اول به کلمه صدم، همان‌قدر مستقیم، سریع و شفاف است که دسترسی به کلمه دوم. این ویژگی باعث شده است که ترنسفورمرها در پردازش:

- مقالات علمی پیچیده با ارجاعات متعدد
 - داستان‌های بلند با شخصیت‌های گوناگون
 - و به‌ویژه کدهای برنامه‌نویسی که در آن‌ها تعریف یک متغیر در ابتدای فایل باید در صدها خط بعد شناسایی شود
- عملکردی خیره‌کننده و بی‌رقیب از خود نشان دهند.

۳.۴) قوانین مقیاس‌پذیری: بزرگ‌تر یعنی هوشمندتر

یکی از شگفت‌انگیزترین یافته‌های تجربی پس از ظهور ترنسفورمرها، کشف و فرمول‌بندی «قوانین مقیاس‌پذیری»^۲ بود. در مدل‌های عصبی قدیمی (مانند LSTM)، نمودار عملکرد مدل پس از رسیدن به حجم مشخصی از داده و اندازه شبکه، دچار اشباع عملکرد^۳ می‌شد؛ یعنی افزودن داده بیشتر یا بزرگ کردن مدل، تأثیر چندانی در بهبود نتایج نداشت. اما محققان دریافتند که در معماری ترنسفورمر، این سقف شیشه‌ای شکسته شده است. عملکرد این مدل‌ها با افزایش:

- حجم داده‌های آموزشی
 - تعداد پارامترهای مدل
 - و توان محاسباتی
- به صورت لگاریتمی و پیش‌بینی‌پذیری بهبود می‌یابد. این کشف مهم، آتش رقابت جهانی برای ساخت مدل‌های زبانی بزرگ^۴ یا همان Large Language Models (LLM) را شعله‌ور کرد و منجر به ظهور مدل‌های غول‌پیکری شد که امروزه مرزهای هوش ماشینی را جابجا کرده‌اند.

¹Maximum Path Length

²Scaling Laws

³Performance Saturation

⁴Large Language Models

فصل ۵

عصر پسا ترنسفورمر و ظهور غول‌های زبانی

انتشار مقاله سال ۲۰۱۷، نه یک پایان، بلکه نقطه آغازی بر یک انفجار کامبرین^۱ در اکوسیستم هوش مصنوعی بود. همان‌گونه که در زیست‌شناسی، دوره کامبرین شاهد ظهور سریع و متنوع گونه‌های حیاتی بود، در اینجا نیز معماری ترنسفورمر بستری را فراهم کرد که تمام مدل‌های زبانی پیشرفته امروزی، به عنوان نوادگان مستقیم یا غیرمستقیم آن طبقه‌بندی می‌شوند.

۱.۵) تبارشناسی مدل‌ها: سه شاخه تکاملی

بر مبنای معماری پایه ترنسفورمر، تحقیقات به سه شاخه اصلی تقسیم شدند که هر یک بر تقویت بخش خاصی از این ساختار تمرکز داشتند:

۱.۱.۵) معماری‌های فقط رمزگذار

این دسته از مدل‌ها تنها از پشته «رمزگذار» بهره می‌برند و تمرکز اصلی آن‌ها بر «درک معنایی» و تحلیل روابط درونی متن است.

- کاربرد: ایده‌آل برای وظایفی نظیر طبقه‌بندی متن^۲، تحلیل احساسات و استخراج موجودیت‌های نامدار.
- نمونه شاخص: مدل BERT (معرفی شده توسط گوگل در ۲۰۱۸). این مدل با ایجاد درک «دوطرفه»^۳ از متن، انقلابی در موتورهای جستجو ایجاد کرد و توانست نیت کاربر را با دقتی بی‌سابقه شناسایی کند.

۲.۱.۵) معماری‌های فقط رمزگشا

این خانواده با حذف بخش رمزگذار، تمام توان خود را بر روی پشته «رمزگشا» و قابلیت‌های «تولید متن»^۴ متمرکز کرده‌اند. وظیفه اصلی آن‌ها پیش‌بینی کلمه (یا توکن) بعدی در یک دنباله است.

^۱Cambrian Explosion

^۲Text Classification

^۳Bidirectional

^۴Text Generation

- کاربرد: تولید متن خلاقانه، تکمیل کد و مکالمه.
- نمونه شاخص: سری مدل‌های GPT (توسعه یافته توسط OpenAI). این مدل‌ها اثبات کردند که با افزایش مقیاس مدل و حجم داده‌ها، می‌توان به سطحی از تولید متن دست یافت که تمایز آن از نوشتار انسانی دشوار باشد.

۳.۱.۵ معماری‌های ترکیبی رمزگذار رمزگشا

- این مدل‌ها به معماری اصلی و کامل مقاله ۲۰۱۷ وفادار مانده‌اند و هر دو بخش را حفظ کرده‌اند.
- کاربرد: وظایف دنباله به دنباله (Sequence to Sequence) مانند ترجمه ماشینی و خلاصه‌سازی متون.
 - نمونه شاخص: مدل‌های تی‌فای^۱ (گوگل) و بارت^۲ (متا/فیس‌بوک).

۲.۵ جدول زمانی تکامل: رشد نمایی پارامترها

برای درک شتاب تحولات تکنولوژیک در این حوزه، بررسی روند تکامل مدل‌ها از سال ۲۰۱۷ تا کنون ضروری است. جدول زیر سیر صعودی ابعاد شبکه و قابلیت‌های نوظهور را نمایش می‌دهد.

جدول ۱.۵: جدول زمانی تکامل مدل‌های زبانی بزرگ و ویژگی‌های بارز آن‌ها

سال	مدل	سازنده	تعداد پارامترها	ویژگی بارز و دستاورد علمی
۲۰۱۷	ترنسفورمر	Google	۶۵ میلیون	معرفی معماری پایه و مکانیزم توجه
۲۰۱۸	GPT-۱	OpenAI	۱۱۷ میلیون	اثبات کارایی یادگیری نظارت‌نشده ^۳ برای تولید متن
۲۰۱۸	BERT	Google	۳۴۰ میلیون	ایجاد درک عمیق و دوطرفه از بافت متن
۲۰۱۹	GPT-۲	OpenAI	۵.۱ میلیارد	تولید متون منسجم طولانی و طرح نگرانی‌های ایمنی
۲۰۲۰	GPT-۳	OpenAI	۱۷۵ میلیارد	ظهور یادگیری چند شات ^۴ ؛ عدم نیاز به بازآموزی برای وظایف جدید
۲۰۲۲	ChatGPT	OpenAI	محرمانه	استفاده از Reinforcement Learning from Human Feedback (RLHF) همسوسازی با قصد انسان و پذیرش عمومی
۲۰۲۳	GPT-۴	OpenAI	(تخمین: تریلیون)	توانایی‌های استدلال پیچیده و درک مدل‌های چندوجهی ^۵

^۱T5

^۲BART

۳.۵) فراتر از متن: تعمیم به بینایی و مدل‌های چندوجهی

قدرت انتزاعی و ریاضیاتی معماری ترنسفورمر به حدی قدرتمند بود که پژوهشگران دریافتند این ساختار محدود به پردازش زبان طبیعی نیست. ایده محوری این بود:

اگر بتوانیم تصاویر را نیز به مانند کلمات در نظر بگیریم، می‌توانیم از مکانیزم توجه بهره ببریم.

این نگرش منجر به تقسیم تصاویر به قطعات کوچک موزاییکی تحت عنوان تکه‌های تصویر^۱ شد. این قطعات دقیقاً مانند کلمات در جمله وارد شبکه می‌شوند. نتیجه این رویکرد، ظهور ترنسفورمرهای بینایی^۲ یا Vision Transformers (ViT) بود که امروزه در حوزه بینایی ماشین جایگزین بسیاری از مدل‌های کانولوشنی شده‌اند.

علاوه بر این، مدل‌های مولد تصویر مانند DALL-E و Midjourney که متن را به تصویر تبدیل می‌کنند، از همین زیرساخت برای ایجاد ارتباط معنایی میان فضای متنی و فضای بصری استفاده می‌کنند که فصل نوبی را در هوش مصنوعی چندوجهی رقم زده است.

^۱ Image Patches

^۲ Vision Transformers

نتیجه‌گیری: میراث ماندگار یک الگو نوین

مقاله جریان‌ساز «توجه تمام آن چیزی است که نیاز دارید» را می‌توان فراتر از یک دستاورد فنی، به عنوان نقطه عطف تاریخی و مرز جداکننده دوران هوش مصنوعی کلاسیک^۱ و عصر نوین هوش مصنوعی مولد^۲ قلمداد کرد. تا پیش از انتشار این اثر، الگو حاکم بر پردازش ماشینی، اسیر محدودیت‌های ذاتی زمان و خطی بودن بود؛ الگوریتم‌ها ناگزیر بودند جهان اطلاعات را به صورت پردازش متوالی^۳، کلمه به کلمه و لحظه به لحظه ادراک کنند. این رویکرد، منجر به خلق مدل‌هایی با حافظه کوتاه‌مدت و ناپایدار می‌شد که در مواجهه با دریای وسیع داده‌ها و ابعاد بالا، دچار سردرگمی و فراموشی می‌گشتند.

شکستن سد زمان و مکان

پژوهشگران گوگل با معرفی معماری ترنسفورمر و برجسته‌سازی نقش محوری مکانیزم توجه، این سد تکنولوژیک را درهم شکستند. آن‌ها اثبات کردند که اگر قیدهای زمانی را حذف کنیم و به ماشین اجازه دهیم نگاهی Global به کل داده‌ها داشته باشد، و همزمان به آن بیاموزیم که چگونه منابع پردازشی خود را مدیریت کند (یا همان تخصیص پویای منابع^۴)، قادر به کشف الگوهای نهفته^۵ خواهد بود که بسیار فراتر از درک سطحی زبان است.

گستره تأثیر: از ادبیات تا حیات

ترنسفورمرها نه تنها چالش دیرینه ترجمه ماشینی را با کیفیتی بی‌سابقه مرتفع ساختند، بلکه افق‌های جدیدی را در علوم کامپیوتر گشودند. امروزه این معماری‌ها زیربنای سیستم‌هایی هستند که:

- کدهای پیچیده نرم‌افزاری تولید می‌کنند،
- متون خلاقانه و ادبی می‌نگارند،
- در تشخیص‌های پزشکی دقیق عمل می‌کنند،

^۱ Classical Artificial Intelligence

^۲ Generative Artificial Intelligence

^۳ Sequential Processing

^۴ Dynamic Resource Allocation

^۵ Latent Patterns

- و حتی در حوزه زیست‌شناسی محاسباتی^۱، ساختارهای پیچیده حیات مانند تاخوردگی پروتئین^۲ را رمزگشایی می‌نمایند (اشاره به مدل‌هایی نظیر AlphaFold).

سخن پایانی

امروزه، هر تعاملی که با عامل‌های گفتگوگر^۲ هوشمند صورت می‌گیرد، هر متنی که توسط مترجم‌های عصبی برگردان می‌شود، و هر شگفتی جدیدی که در اکوسیستم هوش مصنوعی رخ‌نمایی می‌کند، در حقیقت پژواکی از آن مقاله هشت صفحه‌ای سال ۲۰۱۷ است. مقاله‌ای که با جسارتی علمی و بیانی‌ای قاطع اعلام کرد برای دستیابی به هوش عمومی، پیچیدگی‌های بازگشتی لازم نیست؛ بلکه توجه، به واقع تمام چیزی است که نیاز دارید.

¹Computational Biology

²Protein Folding

²Conversational Agents

- [1] Wikipedia Contributors. “Attention Is All You Need”. Available at: [Wikipedia](#)
- [2] Wikipedia Contributors. “Cocktail party effect”. Available at: [Wikipedia](#)
- [3] Wikipedia Contributors. “List of large language models”. Available at: [Wikipedia](#)
- [4] DavletD. “Attention Is All You Need”. Available at: [Medium](#)
- [5] Ujwal Tickoo. “Attention is all you need in plain English”. Available at: [Medium](#)
- [6] “History and Evolution of LLMs”. Available at: [GeeksforGeeks](#)
- [7] “From RNNs to Transformers”. Available at: [Baeldung](#)
- [8] “Transformer vs RNN in NLP: A Comparative Analysis”. Available at: [Appinventiv](#)
- [9] Ege Oguz. “RNN vs. Transformer: Why Parallelization Won the AI Race”. Available at: [Medium](#)
- [10] “I Finally Understood Attention is All You Need After So Long”. Available at: [AI Plain English](#)
- [11] Sebastian Raschka. “Understanding and Coding the Self-Attention Mechanism From Scratch”. Available at: [Sebastian Raschka Blog](#)
- [12] Luv Bansal. “Transformer — Attention Is All You Need Easily Explained With Illustrations”. Available at: [Medium](#)
- [13] Debasish Das. “Self-Attention: The Simple Mechanism That Made ChatGPT Possible”. Available at: [Towards AI](#)
- [14] Gunjas Singh. “Positional encoding in transformers: a Visual and Intuitive guide”. Available at: [Medium](#)

- [15] “A Gentle Introduction to Positional Encoding in Transformer Models”. Available at: [MachineLearningMastery](#)
- [16] Animesh Basnet. “A Beginners Guide to Attention is all you need”. Available at: [Medium](#)
- [17] “Deep Dive into Self-Attention by Hand”. Available at: [Towards Data Science](#)
- [18] Mubashir Iqbal. “Self-Attention: Understanding with Easy Analogies”. Available at: [Medium](#)
- [19] Micah Adler. “The Cocktail Party Inside a Neural Network”. Available at: [Medium](#)
- [20] “Attention Is All You Need - A Deep Dive into the Revolutionary Transformer Architecture”. Available at: [Towards AI](#)
- [21] “The history, timeline, and future of LLMs”. Available at: [Toloka AI](#)
- [22] Justin Milner. “A Timeline Recap of LLMs (Part 1)”. Available at: [Medium](#)
- [23] Dr. Alan D. Thompson. “Timeline of AI and language models”. Available at: [LifeArchitect](#)
- [24] Amjad Shaikh. “Attention Is All You Need — And All We’ve Lost”. Available at: [Medium](#)
- [25] Christian Versloot. “From vanilla RNNs to Transformers”. Available at: [GitHub](#)
- [26] “Revolutionizing Vision: A Deep Dive into Attention Is All You Need”. Available at: [ResearchGate](#)
- [27] “Enhancing English-Persian Neural Machine Translation”. Available at: [EasyChair](#)
- [28] “What is an encoder-decoder model?”. Available at: [IBM](#)
- [29] “How Transformers Work: A Detailed Exploration”. Available at: [DataCamp](#)
- [30] “Why were/are LSTMs popular for so long for NLP?”. Available at: [StackExchange](#)
- [31] “Why is it said transformers are more parallelizable than RNN’s?”. Available at: [Stack-Exchange](#)
- [32] “Are Transformers Strictly More Effective Than LSTM RNNs?”. Available at: [Reddit](#)

- [33] “Why is it said that the transformer is more parallelizable than RNN’s?”. Available at: [Reddit](#)
- [34] “How We Use Selective Attention to Filter Information and Focus”. Available at: [Very-well Mind](#)
- [35] “The cocktail party explained”. Available at: [Speech Neurolab](#)
- [36] “LLM Timeline”. Available at: [LLM Timeline App](#)
- [37] “Attention is all you need (Transformer) - Model explanation”. Available at: [YouTube](#)
- [38] “Transformer and Attention Mechanism”. Available at: [YouTube](#)

واژه‌نامه انگلیسی به فارسی

A

Aggregation	تلفیق
Auto-Regressive	خودرگرسیو
Auto-regressive	خودرگرسیو
Artificial Intelligence	هوش مصنوعی
Attention Head	سرِ توجه
Attention Is All You Need	توجه تمام آن چیزی است که نیاز دارید
Attention Mechanism	مکانیزم توجه

B

Background Noise	نویز پس‌زمینه
Backpropagation	پس‌انتشار
BART	بارت
Bidirectional	دوطرفه

C

Cambrian Explosion	انفجار کامبرین
Classical Artificial Intelligence	هوش مصنوعی کلاسیک
Cluster	کلاستر
Computational Biology	زیست‌شناسی محاسباتی
Computational Complexity	پیچیدگی محاسباتی
Computational Instability	ناپایداری محاسباتی
Concatenate	متصل کردن
Conversational Agents	عامل‌های گفتگوگر
Context Vector	بردار زمینه
Continuous Vector Space	فضای برداری پیوسته

D

Decoder	رمزگشا
Decoder Only	فقط رمزگشا
Deep Learning	یادگیری عمیق
Dynamic Resource Allocation	تخصیص پویای منابع
Dimension	بعد
Dimensional Expansion	انبساط ابعادی
Discrete Probability Distribution	توزیع احتمال گسسته

E

Euclidean Distance	فاصله اقلیدسی
Embedding	تعبیه
Encoder	رمزگذار
Encoder Decoder	رمزگذار رمزگشا
Encoder Only	فقط رمزگذار
Encoder-Decoder Architecture	معماری رمزگذار-رمزگشا

F

Feed-Forward	شبکه پیش‌خور
Few-Shot Learning	یادگیری چند شات

G

Gate	دروازه
Generative Artificial Intelligence	هوش مصنوعی مولد
Geometric Representation	نمایش هندسی
Global Access	دسترسی سراسری
Gradient	گرادیان
Greedy Decoding	رمزگشایی حریصانه

H

Hidden State	حالت پنهان
--------------	------------

I

ID	شناسه
Image Patches	تکه‌های تصویر
Information Bottleneck	گلوگاه اطلاعاتی
Information Distribution	توزیع اطلاعات
Input Embedding	تعبیه ورودی

K

Key	کلید
-----------	------

L

Layer Normalization	نرمال‌سازی لایه
Large Language Models	مدل‌های زبانی بزرگ
Latent Patterns	الگوهای نهفته
Learning Rate	نرخ یادگیری
Linear Combination	ترکیب خطی
Linear Layer	لایه خطی
Linear Projection	نگاشت خطی
Linear Transformation	تبدیل خطی
Logit	امتیاز خام
Logits	امتیازات خام
Long-Range Dependencies	وابستگی‌های طولانی‌برد
Lookup Table	جدول جستجو

M

Maximum Path Length	طول مسیر بیشینه
Masked Multi-Head Attention	توجه چندسر پوشیده
Matrix Parallelization	موازی‌سازی ماتریسی
Multi-Head Attention	توجه چند سر
Multimodal Models	مدل‌های چندوجهی
Moore's Law	قانون مور

N

Natural Language Processing (NLP)	پردازش زبان طبیعی
---	-------------------

O

Output Probabilities احتمالات خروجی

P

Parallel Processing پردازش موازی
Performance Saturation اشباع عملکرد
Positional Encoding رمزگذاری موقعیتی
Position-wise Feed-Forward Network شبکه‌ی پیش‌خور نقطه-به-نقطه
Probability Simplex سیمپلکس احتمالی
Projection نگاشت
Protein Folding تاخوردگی پروتئین

Q

Query پرس‌وجو

R

Recurrent Neural Network (RNN) شبکه‌های عصبی بازگشتی
Relative Distances فواصل نسبی
ReLU تابع فعال‌ساز خطی اصلاح‌شده
Residual Connection اتصال میان‌بر
Residual Connections اتصالات میان‌بر

S

Saturation اشباع شدن
Scaled Dot-Product Attention توجه ضرب داخلی مقیاس‌شده
Scaling هم‌مقیاس شدن
Scaling Laws قوانین مقیاس‌پذیری
Self-Attention توجه خودکار
Sequential Processing پردازش متوالی
Sub-word زیر-واحد
Softmax سافت‌مکس
Spatial Information اطلاعات مکانی
Stack پشته

T

T5	تی‌فای
Text Classification	طبقه‌بندی متن
Text Generation	تولید متن
Temporal Dependency	وابستگی زمانی
The Cocktail Party Effect	اثر مهمانی کوکتل
Time Complexity	پیچیدگی زمانی
Token	توکن
Tokenization	توکن‌گذاری
Transformer	ترنسفورمر

U

Unsupervised Learning	یادگیری نظارت نشده
-----------------------	--------------------

V

Value	مقدار
Vanishing Gradient	محو شدن گرادیان
Vanishing Gradient Effect	تأثیر محوشدگی گرادیان
Vector	بردار
Vector Subspace	زیرفضای برداری
Vision Transformers	ترنسفورمرهای بینایی
Vocabulary	فضای دایره لغات

W

Weight Allocation	تخصیص وزن
Wikipedia	ویکی‌پدیا
Word Embedding	تعبیه کلمات
Word Embeddings	تعبیه کلمات

واژه‌نامه فارسی به انگلیسی

ا

Residual Connection	اتصال میان‌بر
Residual Connections	اتصالات میان‌بر
The Cocktail Party Effect	اثر مهمانی کوکتل
Output Probabilities	احتمالات خروجی
Saturation	اشباع شدن
Performance Saturation	اشباع عملکرد
Spatial Information	اطلاعات مکانی
Latent Patterns	الگوهای نهفته
Logit	امتیاز خام
Logits	امتیازات خام
Dimensional Expansion	انبساط ابعادی
Cambrian Explosion	انفجار کامبرین

ب

BART	بارت
Vector	بردار
Context Vector	بردار زمینه
Dimension	بعد

پ

Natural Language Processing (NLP)	پردازش زبان طبیعی
Sequential Processing	پردازش متوالی
Parallel Processing	پردازش موازی
Query	پرس‌وجو
Backpropagation	پس‌انتشار
Stack	پشته

Time Complexity	پیچیدگی زمانی
Computational Complexity	پیچیدگی محاسباتی

ت

Vanishing Gradient Effect	تأثیر محو شدن گرادیان
ReLU	تابع فعال ساز خطی اصلاح شده
Protein Folding	تا خوردگی پروتئین
Linear Transformation	تبدیل خطی
Dynamic Resource Allocation	تخصیص پویای منابع
Weight Allocation	تخصیص وزن
Linear Combination	ترکیب خطی
Transformer	ترنسفورمر
Vision Transformers	ترنسفورمرهای بینایی
Embedding	تعبیه
Word Embeddings	تعبیه کلمات
Input Embedding	تعبیه ورودی
Image Patches	تکه های تصویر
Aggregation	تلفیق
Attention Is All You Need	توجه تمام آن چیزی است که نیاز دارید
Multi-Head Attention	توجه چند سر
Masked Multi-Head Attention	توجه چند سر پوشیده
Self-Attention	توجه خود کار
Scaled Dot-Product Attention	توجه ضرب داخلی مقیاس شده
Discrete Probability Distribution	توزیع احتمال گسسته
Information Distribution	توزیع اطلاعات
Token	توکن
Tokenization	توکن گذاری
Text Generation	تولید متن
T5	تی فای

ج

Lookup Table	جدول جستجو
--------------------	------------

ح

Hidden State	حالت پنهان
--------------------	------------

خ

خودرگرسیو Auto-regressive

د

دروازه Gate

دسترسی سراسری Global Access

دوطرفه Bidirectional

ر

رمزگذار Encoder

رمزگذار رمزگشا Encoder Decoder

رمزگذاری موقعیتی Positional Encoding

رمزگشا Decoder

رمزگشایی حریصانه Greedy Decoding

ز

زیرفضای برداری Vector Subspace

زیر-واحد Sub-word

زیست‌شناسی محاسباتی Computational Biology

س

سافت‌مکس Softmax

سرِ توجه Attention Head

سیمپلکس احتمالی Probability Simplex

ش

شبکه پیش‌خور Feed-Forward

شبکه‌های عصبی بازگشتی Recurrent Neural Network (RNN)

شبکه‌ی پیش‌خور نقطه-به-نقطه Position-wise Feed-Forward Network

شناسه ID

ط

Text Classification طبقه‌بندی متن
Maximum Path Length طول مسیر بیشینه

ع

Conversational Agents عامل‌های گفتگوگر

ف

Euclidean Distance فاصله اقلیدسی
Continuous Vector Space فضای برداری پیوسته
Vocabulary فضای دایره لغات
Encoder Only فقط رمزگذار
Decoder Only فقط رمزگشا
Relative Distances فواصل نسبی

ق

Moore's Law قانون مور
Scaling Laws قوانین مقیاس‌پذیری

ک

Cluster کلاستر
Key کلید

گ

Gradient گرادیان
Information Bottleneck گلوگاه اطلاعاتی

ل

Linear Layer لایه خطی

م

Concatenate	متصل کردن
Vanishing Gradient	محو شدن گرادیان
Multimodal Models	مدل‌های چندوجهی
Large Language Models	مدل‌های زبانی بزرگ
Encoder-Decoder Architecture	معماری رمزگذار-رمزگشا
Value	مقدار
Attention Mechanism	مکانیزم توجه
Matrix Parallelization	موازی‌سازی ماتریسی

ن

Computational Instability	ناپایداری محاسباتی
Learning Rate	نرخ یادگیری
Layer Normalization	نرمال‌سازی لایه
Projection	نگاشت
Linear Projection	نگاشت خطی
Geometric Representation	نمایش هندسی
Background Noise	نویز پس‌زمینه

و

Temporal Dependency	وابستگی زمانی
Long-Range Dependencies	وابستگی‌های طولانی‌برد
Wikipedia	ویکی‌پدیا

ه

Scaling	هم‌مقیاس شدن
Artificial Intelligence	هوش مصنوعی
Classical Artificial Intelligence	هوش مصنوعی کلاسیک
Generative Artificial Intelligence	هوش مصنوعی مولد

ی

Few-Shot Learning	یادگیری چند شات
Deep Learning	یادگیری عمیق
Unsupervised Learning	یادگیری نظارت‌نشده

فهرست اختصارات

B

BERT Bidirectional Encoder Representations from Transformers

C

CNN Convolutional Neural Networks

F

FFN Feed-Forward Network

G

GPT Generative Pre-trained Transformer

GPU Graphics Processing Unit

GRU Gated Recurrent Unit

L

LLM Large Language Models

LSTM Long Short-Term Memory

N

NLP Natural Language Processing

R

RLHF Reinforcement Learning from Human Feedback

RNN Recurrent Neural Network

S

Seq2Seq Sequence to Sequence

T

TPU Tensor Processing Units

V

ViT Vision Transformers

