

هوش مصنوعی: دنیای نوین فناوری

بهار ۱۴۰۴

شماره

۱

HIGH
HEDIUN
LOW

انیاک

نشریه انجمن علمی علوم کامپیوتر دانشگاه سمنان



شناسنامه

فصل نامه علمی، فرهنگی و خبری انیاک

سال اول - شماره اول - بهار ۱۴۰۴

صاحب امتیاز: انجمن علمی علوم کامپیوتر دانشگاه سمنان

استاد مشاور: کیمیا پیوندی

مدیر مسئول: علیرضا شجاعی بهاءآباد

سردبیر: فاطمه جدیدی

صفحه آرا: رامتین طریک

ویراستار علمی: امیرسجاد سعدی

نویسندگان: شکیبای عربی - زهرا انارکی - امیرسجاد سعدی - امیرسپهر - رامتین طریک - کیوان صحافی

آدرس: سمنان - روبروی پارک سوکان - پردیس شماره ۱ - دانشکده ریاضی، آمار و علوم کامپیوتر

ایمیل: computersciencesemnan@gmail.com | تلگرام: @CSSSU

فهرست

- | | |
|----|--|
| ۴ | مقدمه ای بر مدل های مولد: تولد یک آفریننده |
| ۶ | فشرده سازی فایل های صوتی |
| ۸ | طبقه بندی متون کوتاه بر اساس CNN و بسط معنایی |
| ۱۰ | طبقه بندی تصویر با استفاده از SVM و CNN |
| ۱۲ | بهینه سازی و تشخیص تومورهای مغزی با هوش مصنوعی |
| ۱۴ | معرفی کتاب های جذاب و کاربردی در حوزه علوم کامپیوتر |
| ۱۶ | برگزاری ایونت لینوکسی Distro Talk دانشکده علوم کامپیوتر |
| ۱۷ | افتخار آفرینی دانشجویان علوم کامپیوتر در آیین ناب و نماد |



دانشگاه سمنان



دانشگاه سمنان
معاونت دانشجویی و فرهنگی



انیاک



Semnan Computer Science
انجمن علمی علوم کامپیوتر سمنان

سخن استاد مشاور

با توجه به گستردگی و محبوبیت رشته علوم کامپیوتر در سال های اخیر، به نظر می رسد داشتن یک نشریه برای اطلاع رسانی و آگاهی بیشتر می تواند کمک شایانی به علاقه مندان در این حوزه داشته باشد. از طرف دیگر، امروزه سرعت پیشرفت فناوری با رشد سریع و تصاعدی همراه است و این امر مستلزم بروز رسانی مداوم دانش و آگاهی برای دانشجویان و متخصصان این حوزه می باشد. در همین راستا انجمن علوم کامپیوتر با جمع آوری مطالب کاربردی و جدید و انتشار این مطالب تلاش می کند اطلاعات مختصر و مفیدی را برای دانشجویان این رشته فراهم نماید. اینجانب به عنوان استاد مشاور از تلاش های کلیه اعضای انجمن علوم کامپیوتر کمال تشکر را دارم و امیدوارم این شماره نشریه مورد توجه خوانندگان واقع شود.

کیمیا پیوندی

استاد مشاور انجمن علمی علوم کامپیوتر



سخن سر دبیر

هر صبح منوریم و هر شام خوشیم
مایم که بی پیچ سرانجام خوشیم

مایم که از باده ی بی جام خوشیم،
گویند سرانجام ندارد شما...

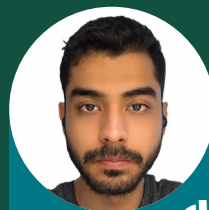
در جهانی که هر چیز را باید در قالبی مشخص، عددی دقیق، و چارچوبی منطقی گنجاند، ما تصمیم گرفتیم بار دیگر، خوشی بی دلیل فهمیدن را تجربه کنیم. نه به سرانجام محتوم فکر کردیم و نه به جداول پایان بندی. در میان صفر و یک، در میانه بیت ها و بایت ها، لبخند زدیم و خلق کردیم.

این شماره از ENIAC، حاصل تلاشی ست برای نگاه کردن به علم، نه فقط از لنز تکنیک، که از پنجره ی اندیشه. از مدل های مولد تا تشخیص بیماری با تصویر، از ساختار صدا تا عمق داده ها، ما به کنجکاوی وفادار ماندیم و با پرسش جلو رفتیم. اگر این شماره، حتی لحظه ای ذهن شما را روشن تر، و فکر تان را عمیق تر کرد، رسالت ما تمام است.

بی نیاز از سرانجام.

فاطمه جدیدی





نویسنده: امیر سپهر

مقدمه ای بر مدل های مولد: تولد یک آفریننده

هدف یک مدل مولد یافتن الگوهای پنهان و شکل داده با استفاده از هوش مصنوعی است. به دیگر سخن، مدل های مولد نوعی از فرایندهای یادگیری ماشین هستند که به ماشین ها این امکان را میدهند که مانند انسان رویا ببینند، و بر اساس آنچه از قبل تجربه کرده اند داده های جدید خلق کنند. اهمیت این نوآوری بشری در قدرت آفرینش است. قدرت خلق کردن چیزهای جدید که کاربرد های فراوانی در علوم، هنر و سایر زمینه ها دارد.

خیال کنید که در حال آموزش نقاشی حیوانات به یک کودک هستید. پس از مدتی، کودک مشخصه های کلی حیوانات را فرا می گیرد و بنا به آنچه از شما یاد گرفته است، شروع به کشیدن حیوانات جدید می کند. این فرایند دقیقاً مشابه آنچه هست که در مدل های مولد اتفاق می افتد.

فهم تفاوت مدل های مولد با مدل های تمایزی از امور اساسی در یادگیری ماشین است:

مدل های مولد بر روی این تمرکز می کنند که چطور داده ورودی یا یک المان از دنیای واقعی را تقلید کنند. یک مدل مولد مانند کودکی است که ابتدا سعی می کند بفهمد صورت یک گربه به چه شکل است. پس از دیدن تصاویر متعددی از گربه ها، کودک سعی می کند گربه خودش را نقاشی کند. مدل های تمایزی از سوی دیگر تلاش نمی کنند که گربه یا سگ بکشند. آنها فقط این وظیفه را بر عهده دارند که بتوانند بین داده ها تمایز قایل شوند و تفاوت آنها را درک کنند. مدل تمایزی سعی می کند بتواند به صراحت تشخیص دهد که چه موقع در حال نگاه کردن به سگ، و چه زمانی در حال نگاه کردن به یک گربه است.

انواع متفاوتی از مدل های مولد وجود دارد که هر یک شیوه و روش خود را در فهم و تولید داده دارند:

شبکه های بیضوی (Bayesian Networks):

این شبکه ها مدل های گرافیکی هستند که روابط آماری را میان داده ها پیدا میکنند. این شبکه ها در تشخیص علائم بیماری کاربرد دارند.

مدل های انتشاری (Diffusion Models):

این مدل می تواند بفهمد که چطور داده ها در گذر زمان دچار تغییر می شوند و تکامل پیدا میکنند. به طور مثال شناسایی

روند شیوع یک ویروس به این مدل واگذار میگردد.

شبکه های عصبی مولد (Generative Adversarial Networks (GANs):

این شبکه ها از دو شبکه عصبی تشکیل شده اند. مولد (generator) و تفکیک کننده (discriminator). این دو شبکه عصبی که با هم آموزش میبینند با هم رقابت می کنند. شبکه مولد سعی بر این دارد تا داده ای را خلق کند، که شبکه عصبی تفکیک کننده نتواند تفاوت آن را با داده اصلی و واقعی تشخیص دهد. حاصل این نبرد یک داده مصنوعی با کیفیت هست که بسیار نزدیک به دنیای واقعی است. این شبکه ها در زمینه هایی همچون خلق تصویر و تقلید فرم صورت انسان کاربرد دارند.

خود رمزنگار ها (Variational Autoencoders or VAEs):

یک فرم فشرده از داده ورودی تولید می کنند. سپس از این



جعبه سیاه:

بشتر مولد ها به خصوص آنهایی که مبتنی بر یادگیری عمیق هستند مانند یک جعبه سیاه عمل می کنند، به این معنی که به طور دقیق نمی توان رفتار آنها را درک کرد. نمی دانیم آنها چگونه تصمیم می گیرند و به سوالات ما پاسخ میدهند. این امر باعث می شود تا در رابطه با استفاده از مولد ها در حوزه های پزشکی و سلامت نگران باشیم!

اخلاقیات:

توانایی مولد ها در تولید محتوای بصری که از واقعیت قابل تمیز نیستند، سبب نگرانی جوامع بشری شده است. مولد های هوش مصنوعی به راحتی می توانند با شبیه سازی سناریو های غیر واقعی سبب گسترش اطلاعات غلط شوند که ممکن است عواقب فاجعه باری به همراه داشته باشد.

وابستگی به داده ورودی:

کیفیت داده های جدید وابسته به داده هایی است که برای آموزش مولد استفاده شده اند. اگر داده های آموزشی معیوب باشند خروجی مولد نیز معیوب خواهد بود.

فروپاشی (Mode Collapse):

فروپاشی یک امر شایع در شبکه های عصبی مولد (GAN) است. فروپاشی زمانی رخ می دهد که داده هایی که توسط مدل تولید میشوند تنوع پایینی دارند.

مدل های مولد تصویر:

مولد های تصویر به طور کل به دو دسته VAE ها و GAN ها تقسیم بندی می شوند که هر کدام پیش تر در این مقاله توضیح داده شد.

از جمله مولد های تصویر قدرتمند و مشهور می توان به Midjourney و DALL-E اشاره کرد که روزانه میزان شمار زیادی کاربر از سراسر جهان هستند.

سخن نهایی:

مدل های مولد به سبب توانایی خود در خلق آثار جدید، پیشبینی و تحلیل داده ها دنیای کنونی ما را متحول کرده و به جزء جدایی ناپذیری از زندگی بدل گشته اند. تکنولوژی که صنعت را دگرگون کرده و سبب گسترش نوآوری و خلاقیت انسان ها میشود.

فرم فشرده برای تولید داده های جدید استفاده می کنند. معمولاً در خلق کردن و حذف کردن نویز تصاویر به کار می روند. ماشین بولتزمن محدود شده (Restricted Boltzmann Machines (RBMs):

این شبکه عصبی دو لایه در سیستم های پیشنهادگر استفاده می شوند. اینستاگرام را باز می کنید و پست هایی را می بینید که به شما پیشنهاد شده است. در پشت پرده یک سیستم پیشنهادگر یا RBM در حال کار است!

شبکه عصبی بازگشتی پیکسلی (Pixel Recurrent Neural Networks (PRNNs):

این شبکه های عصبی هر عنصر یا پیکسل از تصویر را با استفاده از محتوای پیکسل های قبلی تولید می کنند. این شبکه ها نیز در تولید تصاویر کاربرد دارند.

زنجیره مارکوف (Markov Chains):

این ها مدل هایی هستند که آینده را از آنچه در حال حاضر وجود دارد پیش بینی می کنند، نه آنچه که در گذشته وجود داشته است. این مدل ها معمولاً در تولید متن به کار برده میشوند که در آن هر کلمه طبق کلمه قبل از خودش حدس زده می شود.

محدودیت ها و چالش های مدل های مولد:

هزینه یادگیری:

مدل های مولد پیشرفته به زمان زیاد و سخت افزار های قدرتمند نیاز دارند تا بتوانند به خوبی آموزش ببینند.

کیفیت داده تولید شده:

مدل های مولد در تولید داده های جدید عملکرد مطلوبی ارایه می کنند. باید توجه داشت که کیفیت و مطابقت داده تولید شده با دنیای واقعی از اهمیت فراوانی برخوردار است. سنجش کیفیت داده ها امری چالش برانگیز است. شاید در نگاه نخست همه چیز بی نقص به نظر برسد اما اگر با دقت بیشتری بنگرید متوجه برخی نواقص خواهید شد!

بیش برازش (Overfitting):

این امکان وجود دارد که مولد هوش مصنوعی داده هایی را تولید کند که بسیار شبیه به داده هایی است که پیش تر برای آموزش آن استفاده شده است. این اختلال بیش برازش نام دارد که سبب پایین آمدن تنوع در داده های جدید تولید شده می شود.



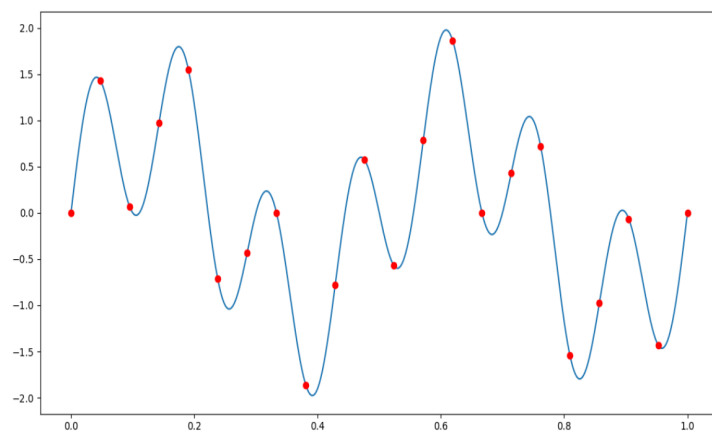
نویسنده: کیوان صحافی

فشرده سازی فایل های صوتی

یه فایل صوتی چجوری فشرده میشه؟

یه آهنگ سه دقیقه ای رو چطوری توی چند مگابایت جا می دن و همچنان کیفیتش اون قدر بد نیست؟ اصلاً مگه صدا رو میشه کوچیک کرد؟

صدا در اصل یه سیگنال آنالوگه، یعنی یه موجی که با گذشت زمان تغییر می کنه. وقتی این سیگنال رو دیجیتال می کنیم (مثلاً با میکروفرن ضبطش می کنیم)، تبدیل می شه به کلی عدد که هر کدوم بلندی صدا رو توی یه لحظه خاص نشون می دن. حالا تصور کنید برای هر ثانیه از یه آهنگ، هزاران عدد داریم! طبیعتاً این حجم زیادی از داده ست. ما باید یه راهی پیدا کنیم که فقط چیزهای مهم رو نگه داریم و بقیه رو دور بریزیم. درست مثل وقتی که یه عکس رو می خوایم کوچیک کنیم ولی همچنان واضح بمونه.



۱: نمونه برداری از سیگنال آنالوگ برای ذخیره سازی دیجیتال

تبدیل از حوزه زمان به حوزه فرکانس

وقتی به یه موج صوتی نگاه می کنیم، در حوزه زمان (نموداری که نشون می ده بلندی صدا چطور در طول زمان تغییر می کنه)، پتانسیلی برای فشرده سازی نمی بینیم. ولی اگه همین سیگنال رو ببریم توی فضای فرکانس، ناگهان نظم ظاهر می شه. فضای فرکانسی چیه؟ فرض کنید یه تابع عادی رو به صورت جمع چند تابع سینوسی بنویسیم. فضای ضرایب این توابع سینوسی و خودشون رو میتونیم فضای فرکانسی حساب کنیم. این تبدیل فضا میتونه با کمک تبدیل هایی مثل تبدیل فوریه، DCT و ... شکل بگیره. در فضای فرکانس می بینیم که فقط تعداد خیلی کمی از فرکانس ها هستن که واقعاً انرژی زیادی دارن، و بقیه تقریباً صفر یا خیلی خیلی کوچیکن.

به این حالت می گیم اسپارس بودن (Sparsity) - یعنی اطلاعات اصلی فقط توی تعداد کمی از نقاط متمرکزن، بقیه تقریباً بی ارزشن.

این فقط مخصوص صدا نیست. بیشتر سیگنال های طبیعی (مثل عکس، ویدیو، حتی موج مغزی!) در فضای فرکانس، اسپارس می شن. دلیلش هم ساده ست: طبیعت معمولاً منظم تر از اونه که فکر

می کنیم. مثلاً توی گفتار، بیشتر انرژی صوت توی فرکانس های خاصی متمرکز - مثلاً فرکانس های مربوط به هارمونیک ها و رزونانس های گلو و دهان. یا توی عکس ها، بیشتر انرژی توی لبه ها و کانتراست هاست، نه توی هر پیکسلی.

خب حالا این چه کمکی بهمون می کنه؟

اگه بدونیم فقط چند تا فرکانس مهم ان، می تونیم:

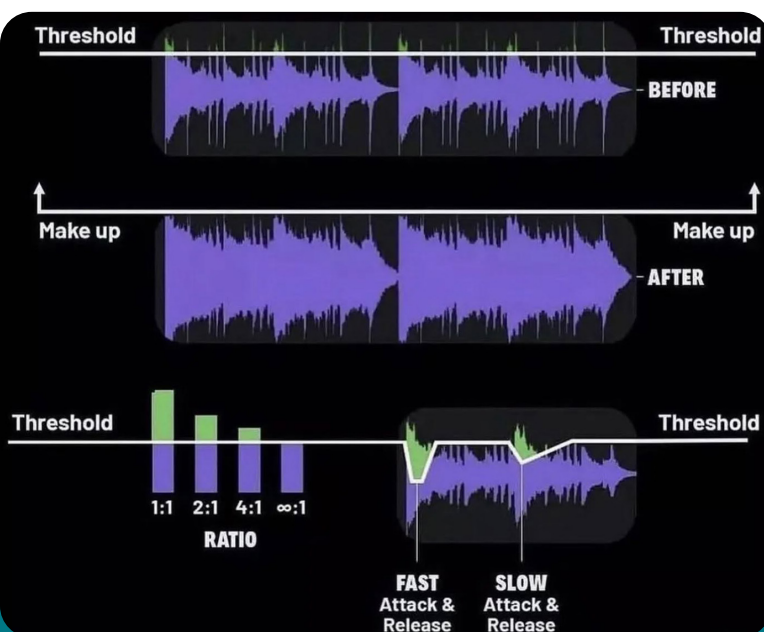
. فقط همونا رو ذخیره کنیم و بقیه رو دور بریزیم؛

. یا ضرایب کوچیک رو با دقت کمتر (quantized)

ذخیره کنیم چون کم اهمیت ترن؛

. و این یعنی حجم خیلی کمتر، بدون افت محسوس کیفیت!

در واقع، اسپارس بودن همون نقطه طلاییه که فشرده سازی رو ممکن و مؤثر می کنه. اگه



سیگنال‌ها توی همه فرکانس‌ها انرژی یکنواخت داشتن، دیگه تقریباً هیچ کاری نمی‌شد کرد! گوش انسان هم خودش یه فشرده‌سازه!

جالبه بدونید مغز ما همه فرکانس‌ها رو با یه دقت نمی‌شنوه. مثلاً اگه دو صدای نزدیک به هم هم‌زمان پخش بشن، فقط صدای قوی‌تر رو می‌شنویم. این ویژگی رو بهش می‌گن پوشش‌دهی یا **masking**.

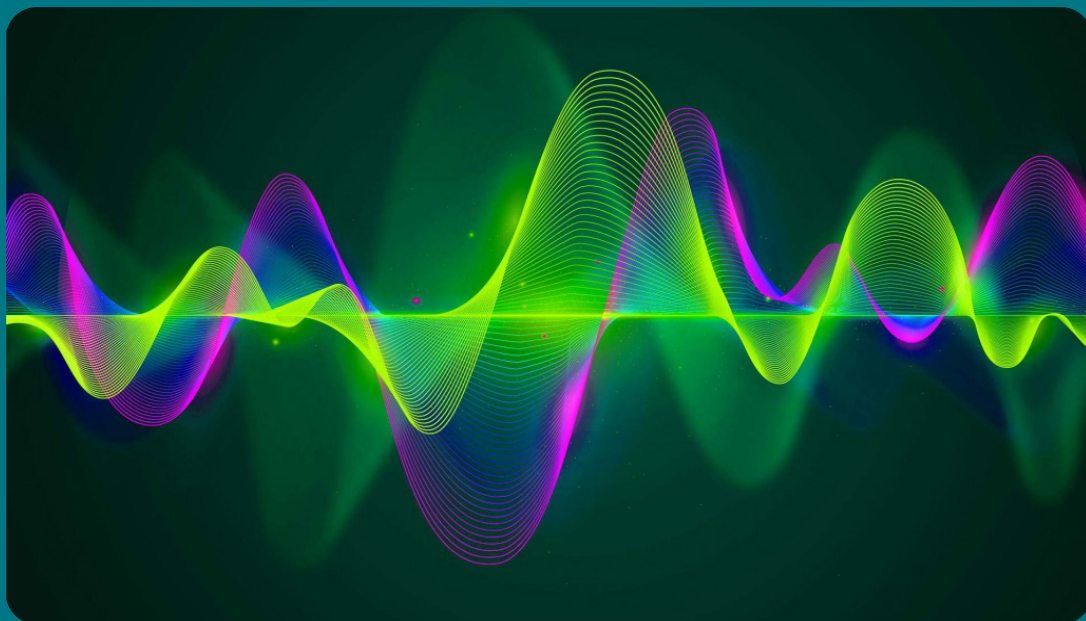
فرمت‌های فشرده‌سازی مثل MP3 دقیقاً از همین ترفند استفاده می‌کنن: اون چیزایی رو حذف می‌کنن که مغز ما احتمالاً نمی‌شنوه. خلاصه، می‌ذارن «واقعیت شبیه واقعیت» بمونه، نه «واقعیت کامل». MP3 چیکار می‌کنه؟

MP3 که همه باهاش آشنایم، یکی از محبوب‌ترین فرمت‌های فشرده‌سازی با اتلاف (Lossy) هست. یعنی بعضی از اطلاعات صدا حذف می‌شن، ولی طوری که شنونده متوجه نشه. مراحلش تقریباً این‌طوره: ۱. تقسیم صدا به بخش‌های کوچیک (مثلاً هر ۲۰ میلی ثانیه)

۲. اعمال تبدیل DCT برای رفتن به فضای فرکانس

۳. استفاده از مدل شنوایی انسان برای تشخیص بخش‌های قابل حذف

۴. کوانتیزاسیون و کدگذاری اطلاعات باقی‌مونده





نویسنده: رامتین طریک

طبقه بندی متون کوتاه بر اساس CNN و بسط معنایی

چالش‌های طبقه‌بندی متون کوتاه
متون کوتاه، مانند توییت‌ها یا نظرات آنلاین، به دلیل طول محدود، معمولاً فاقد اطلاعات زمینه‌ای کافی هستند و ممکن است شامل خطاهای املائی یا چندمعنایی باشند. روش‌های سنتی یادگیری ماشین، مانند نایو بیس (Naive Bayes)، ماشین بردار پشتیبان (SVM) و کان نزدیک‌ترین همسایه (KNN)، از مدل کیسه کلمات (Bag-of-Words) و وزن دهی TF-IDF استفاده می‌کنند. این روش‌ها روابط زمینه‌ای بین کلمات را نادیده می‌گیرند و با مشکلاتی مانند ابعاد بالا و پراکندگی داده‌ها مواجه هستند، که منجر به استخراج ناکافی اطلاعات معنایی می‌شود.

با ظهور یادگیری عمیق، مدل‌های مبتنی بر شبکه‌های عصبی، مانند شبکه‌های عصبی بازگشتی (RNN) و کانولوشنی (CNN)، عملکرد بهتری در NLP نشان داده‌اند. با این حال، وابستگی این مدل‌ها به تعداد لایه‌های شبکه برای استخراج ویژگی‌های معنایی، باعث افزایش پیچیدگی محاسباتی و زمان آموزش می‌شود. برای رفع این مشکلات، مقاله پیشنهاد می‌کند که با استفاده از پایگاه دانش خارجی و بسط معنایی، اطلاعات معنایی متون کوتاه غنی‌سازی شود.

روش پیشنهادی: SECNN

روش SECNN ترکیبی از پیش‌پردازش پیشرفته، مکانیزم توجه، بسط معنایی و شبکه‌های عصبی کانولوشن است که در چهار مرحله اصلی اجرا می‌شود:

۱. پیش‌پردازش متن با شباهت جارو-وینکلر بهبودیافته

متون کوتاه اغلب شامل خطاهای املائی هستند که با کلمات موجود در جدول بردار کلمات (مانند Word2Vec) مطابقت ندارند. مقاله یک نسخه بهبودیافته از معیار شباهت جارو-وینکلر را معرفی می‌کند که هم تطابق پیشوند و هم پسوند کلمات را در نظر می‌گیرد. این روش با محاسبه شباهت بین کلمات متون کوتاه و کلمات جدول بردار، خطاهای املائی را شناسایی و اصلاح می‌کند. فرمول بهبودیافته به شرح زیر است:

این روش پوشش جدول بردار کلمات را افزایش داده و دقت استخراج ویژگی‌ها را بهبود می‌بخشد.

۲. شناسایی کلمات مرتبط با مکانیزم توجه

با گسترش فناوری‌های اطلاعات مانند رایانش ابری و داده‌های کلان، حجم عظیمی از داده‌های متنی، به‌ویژه متون کوتاه نظیر نظرات کاربران و پست‌های شبکه‌های اجتماعی تولید می‌شود. این متون کوتاه ویژگی‌هایی مانند حجم زیاد و محدودیت‌هایی نظیر کمبود اطلاعات زمینه‌ای، چندمعنایی و خطاهای املائی دارند. استخراج اطلاعات ارزشمند از این داده‌ها چالش بزرگی در پردازش زبان طبیعی (NLP) است. مقاله‌ای که توسط وانگ و همکاران در سال ۲۰۲۰ منتشر شد، روشی نوآورانه برای طبقه‌بندی با استفاده از شبکه‌های عصبی کانولوشن (CNN) و بسط معنایی پیشنهاد می‌کند. این روش با عنوان SECNN (Semantic Extension Convolutional Neural Network) معرفی شده و با ترکیب پیش‌پردازش پیشرفته، مکانیزم توجه و پایگاه دانش خارجی، عملکرد طبقه‌بندی را بهبود می‌بخشد. این خلاصه مروری جامع بر روش پیشنهادی، اهمیت آن و نتایج تجربی ارائه می‌دهد.



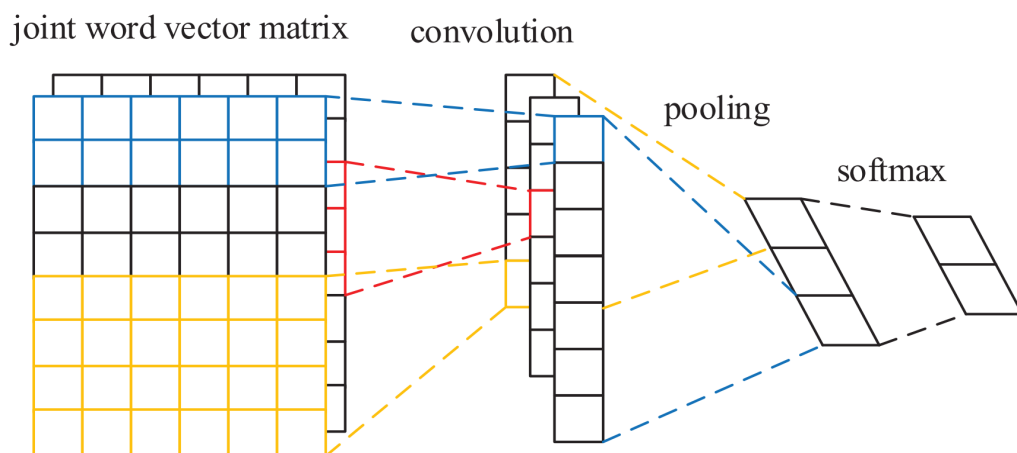


Figure 4 | Convolutional neural network (CNN) short text classification structure.

در متون کوتاه، تنها تعداد محدودی از کلمات معنای اصلی جمله را منتقل می‌کنند. مقاله یک مدل CNN مبتنی بر مکانیزم توجه طراحی کرده است که کلمات کلیدی (مرتبط) را شناسایی می‌کند. این مدل شامل لایه‌های پیچشی و تجمیع است که ویژگی‌های معنایی سطح بالا را استخراج می‌کنند. وزن‌های توجه با استفاده از تابع تانژانت هیپربولیک (tanh) محاسبه شده و کلمات مرتبط با مقایسه فاصله اقلیدسی در

فضای بردار کلمات شناسایی می‌شوند.

۳. بسط معنایی با پایگاه دانش Probase

برای غلبه بر کمبود اطلاعات، از پایگاه دانش Probase برای مفهوم‌سازی متن کوتاه و کلمات مرتبط استفاده می‌شود. این فرآیند شامل تولید بردارهای مفهومی برای متن و کلمات مرتبط است که با استفاده از مدل Word2Vec محاسبه می‌شوند. سه ماتریس بردار (متن کوتاه، مفاهیم متن، و مفاهیم کلمات مرتبط) ترکیب شده و ماتریس ورودی نهایی را تشکیل می‌دهند.

۴. طبقه‌بندی با CNN کلاسیک

ماتریس ورودی نهایی به یک مدل CNN کلاسیک وارد می‌شود که شامل لایه‌های پیچشی، تجمیع بیشینه (Max Pooling) و کاملاً متصل (Fully Connected) است. این مدل ویژگی‌های معنایی را استخراج کرده و با استفاده از لایه Softmax، احتمال تعلق متن به هر دسته را محاسبه می‌کند.

داده‌های مجموعه:

روش SECNN روی پنج مجموعه داده استاندارد آزمایش شد:

- MR: نظرات فیلم با ۱۰,۶۶۲ نمونه، ۲ دسته (مثبت/منفی)، میانگین طول ۲۰ کلمه
- TREC: پرس‌وجوها با ۶,۴۵۲ نمونه، ۶ دسته، میانگین طول ۱۰ کلمه
- AG News: مقالات خبری با ۱۲۷,۶۰۰ نمونه، ۴ دسته، میانگین طول ۷ کلمه

Twitter: توییت‌ها با ۱۱,۲۰۹ نمونه، ۳ دسته (مثبت/خنثی/منفی)، میانگین طول ۱۹ کلمه
SST-۲: نسخه توسعه‌یافته MR با ۹,۶۱۳ نمونه، ۲ دسته، میانگین طول ۱۹ کلمه

روش SECNN پیشنهادی توسط وانگ و همکاران، رویکردی نوآورانه برای طبقه‌بندی متون کوتاه ارائه می‌دهد که با ترکیب پیش‌پردازش پیشرفته، مکانیزم توجه، بسط معنایی شبکه‌های عصبی کانولوشن، محدودیت‌های روش‌های سنتی را برطرف می‌کند. این روش با اصلاح خطاهای املایی، شناسایی کلمات کلیدی و غنی‌سازی اطلاعات معنایی، دقت و پایداری بالایی در مجموعه داده‌های متنوع نشان داده است. با وجود پیچیدگی محاسباتی، SECNN گامی مهم در بهبود پردازش متون کوتاه در NLP است و برای کاربردهایی مانند تحلیل احساسات، فیلتر هرزنامه و دسته‌بندی اخبار ارزشمند است. این خلاصه برای خوانندگان نشریه Eniac مروری جامع و قابل فهم ارائه می‌دهد که پیشرفت‌های اخیر در یادگیری ماشین و NLP را برجسته می‌کند.



نویسنده: امیرسجاد سعدی

طبقه‌بندی تصویر با استفاده از SVM و CNN

این سند خلاصه‌ای از مقاله «طبقه‌بندی تصویر با استفاده از SVM و CNN» است که در سال ۲۰۲۰ در *Procedia Computer Science* منتشر شده است. مقاله روی تشخیص اعداد دست‌نویس با استفاده از یک مدل ترکیبی به نام CNN-SVM تمرکز دارد که ویژگی‌های تصاویر را با شبکه‌های کانولوشنی (CNN) استخراج کرده و با ماشین بردار پشتیبان (SVM) طبقه‌بندی می‌کند. هدف این خلاصه، ارائه توضیحات ساده، دقیق و جامع از الگوریتم‌ها و اصطلاحات کلیدی، همراه با حفظ ساختار اصلی مقاله است. برای توضیح SVM، از مثال «میز و توپ» استفاده می‌کنیم تا مفهوم ترفند هسته (Kernel Trick) به خوبی جا بیفتد.

مقاله عملکرد یک مدل ترکیبی CNN-SVM را برای طبقه‌بندی تصاویر، به‌خصوص تشخیص اعداد دست‌نویس، بررسی می‌کند. این مدل با SVM تنها مقایسه شده و معیارهایی مثل دقت (Accuracy)، صحت (Precision) و یادآوری (Recall) برای ارزیابی استفاده شده‌اند. داده‌ها

شامل اعداد دست‌نویس یا اشیایی مثل توپ فوتبال، گل آفتاب‌گردون و پیتزا هستند. نتایج نشان می‌دهند که مدل ترکیبی با دقت تست تا ۹۸,۹۵ درصد، از SVM تنها (تا ۹۷,۸۵ درصد) بهتر عمل می‌کند. این بهبود به دلیل توانایی CNN در استخراج ویژگی‌های پیچیده و قدرت SVM در طبقه‌بندی دقیق است.

ماشین بردار پشتیبان (SVM)

ماشین بردار پشتیبان (SVM) یک الگوریتم یادگیری ماشین است که برای طبقه‌بندی داده‌ها استفاده می‌شود. هدفش پیدا کردن یک خط یا صفحه (به نام هپرپلین) است که داده‌های دو دسته را با بیشترین حاشیه (فاصله) ممکن از هم جدا کند.

مثال میز و توپ

فرض کنید روی یک میز، توپ‌های قرمز و آبی پخش شده‌اند. می‌خواهید یک خط بکشید که قرمزها را از آبی‌ها جدا کند. SVM خطی را انتخاب می‌کند که فاصله‌اش با

نزدیک‌ترین توپ‌ها از هر رنگ (به نام بردارهای پشتیبان) حداکثر باشد. این حاشیه بزرگ باعث می‌شود مدل در برابر داده‌های جدید مقاوم‌تر باشد.

مشکل داده‌های غیرخطی: حالا اگر توپ‌ها جوری پخش باشند که هیچ خط راستی نتواند آنها را جدا کند (مثلاً قرمزها در مرکز و آبی‌ها دورشان)، چه؟ اینجا SVM از ترفند هسته استفاده می‌کند.

ترفند هسته: تبدیل میز به تپه

تصور کنید میز یک پارچه صاف است و توپ‌ها به صورت یک دایره قرمز در مرکز و آبی‌ها دورش پخش شده‌اند. هیچ خطی نمی‌تواند آنها را جدا کند. حالا وسط میز را بکشید بالا تا یک تپه سه‌بعدی تشکیل شود. در این فضای جدید، قرمزها بالای تپه و آبی‌ها پایین آن هستند. حالا یک صفحه صاف می‌تواند آنها را جدا کند. ترفند هسته همین است: داده‌ها را به یک فضای با ابعاد بالاتر می‌برد (بدون محاسبات سنگین) تا جداسازی آسان‌تر شود.

در مقاله، از هسته‌های پلی‌نومیل (با درجه ۳ یا ۵) و پارامتر گاما (۰,۱ یا ۱) استفاده شده است. گاما تعیین می‌کند که مدل چقدر به داده‌های نزدیک حساس باشد (گامای بالا مدل را پیچیده‌تر می‌کند). همچنین، دو روش تصمیم‌گیری تست شده‌اند:

یکی در برابر یکی: هر جفت کلاس را جداگانه مقایسه می‌کند.

یکی در برابر بقیه: یک کلاس را با همه کلاس‌های دیگر مقایسه می‌کند.

شبکه‌های کانولوشنی (CNN)

شبکه‌های کانولوشنی (CNN) مدل‌های یادگیری عمیقی هستند که برای پردازش تصاویر طراحی شده‌اند. آنها با استخراج ویژگی‌های تصویری مثل لبه‌ها، شکل‌ها و بافت‌ها، تصاویر را تحلیل می‌کنند.

مثال کارآگاه

فکر کنید یک کارآگاه با ذره‌بین یک عدد دست‌نویس را بررسی می‌کند. او ابتدا تکه‌های کوچک تصویر را نگاه می‌کند و سرنخ‌های ساده مثل خطوط افقی یا عمودی پیدا می‌کند (لایه‌های کانولوشنی). بعد، این سرنخ‌ها را ترکیب می‌کند تا الگوهای پیچیده‌تر مثل حلقه‌ها یا تقاطع‌ها را تشخیص

نتایج عملکرد

مقاله یک جدول مقایسه‌ای ارائه می‌دهد که دقت CNN و SVM را نشان می‌دهد:

گاما	درجه	روش تصمیم‌گیری	دقت آموزش SVM (%)	دقت تست SVM (%)	دقت تست CNN-SVM (%)
1	3	یکی در برابر یکی	97.80	97.52	98.82
1	3	یکی در برابر بقیه	97.38	97.74	98.74
0.1	3	یکی در برابر یکی	97.95	97.85	98.95
0.1	5	یکی در برابر یکی	98.35	97.68	98.88

مدل CNN-SVM در همه موارد دقت تست بالاتری (تا ۹۸,۹۵ درصد) نسبت به SVM تنها (تا ۹۷,۸۵ درصد) دارد. دلیل اصلی، توانایی CNN در ارائه ویژگی‌های باکیفیت‌تر به SVM است که کار طبقه‌بندی را آسان‌تر می‌کند.

معیارهای ارزیابی

برای ارزیابی، از معیارهای زیر استفاده شده است:

صحت (Precision) = (مثبت‌های درست) / (مثبت‌های درست + مثبت‌های غلط)

یادآوری (Recall) = (مثبت‌های درست) / (مثبت‌های درست + منفی‌های غلط)

یک نمودار پراکندگی نشان می‌دهد که صحت و یادآوری برای کلاس‌های مختلف (مثل اعداد یا اشیاء) بین ۰,۷۸ تا ۰,۸۸ است، با میانگین ۰,۸۸.

برای درک بهتر، چند اصطلاح مهم توضیح داده می‌شود:

ویژگی (Feature): یک خصوصیت قابل اندازه‌گیری از داده، مثل شکل یک حلقه در عدد ۸.

هپرلین: صفحه‌ای در فضای چندبعدی که داده‌ها را جدا می‌کند.

حاشیه (Margin): فاصله بین هپرلین و نزدیک‌ترین داده‌ها؛ حاشیه بزرگ‌تر یعنی مدل قوی‌تر.

Overfitting: وقتی مدل بیش‌ازحد به داده‌های آموزشی وابسته می‌شود و روی داده‌های جدید خوب عمل نمی‌کند.

نتیجه‌گیری

مدل ترکیبی CNN-SVM یک رویکرد قدرتمند برای طبقه‌بندی تصویر، به‌خصوص تشخیص اعداد دست‌نویس، ارائه می‌دهد. با ترکیب استخراج ویژگی پیشرفته CNN و طبقه‌بندی مقاوم SVM، این مدل به دقت تست ۹۸,۹۵ درصد می‌رسد، که از SVM تنها (۹۷,۸۵ درصد) بهتر است. مثال «میز و توپ» نشان داد که ترنند هسته در SVM چگونه داده‌های پیچیده را قابل جداسازی می‌کند. این مدل برای کاربردهای پردازش تصویر، از تشخیص اعداد تا شناسایی اشیاء، بسیار مناسب است.

دهد (لایه‌های عمیق‌تر). در نهایت، تصمیم می‌گیرد که تصویر چه عددی است (لایه‌های تمام‌متصل).

اجزای اصلی CNN شامل:

لایه‌های کانولوشنی: فیلترهایی که ویژگی‌های محلی (مثل لبه‌ها) را استخراج می‌کنند.

لایه‌های پولینگ: اندازه داده‌ها را کم می‌کنند (مثلاً با گرفتن حداکثر یا میانگین) تا محاسبات ساده‌تر شود.

لایه‌های تمام‌متصل: ویژگی‌های استخراج‌شده را به یک طبقه‌بندی نهایی تبدیل می‌کنند.

در مدل ترکیبی، CNN فقط ویژگی‌ها را استخراج می‌کند و طبقه‌بندی به SVM سپرده می‌شود.

مدل ترکیبی CNN-SVM

مدل ترکیبی CNN-SVM از نقاط قوت هر دو الگوریتم استفاده می‌کند. مراحل کار:

ورودی تصویر: یک تصویر عدد دست‌نویس به CNN داده می‌شود.

استخراج ویژگی: CNN با لایه‌های کانولوشنی و پولینگ، یک بردار ویژگی (شامل الگوهای کلیدی مثل شکل حلقه‌ها) تولید می‌کند.

طبقه‌بندی با SVM: این بردار به SVM داده می‌شود که با یک هپرلین، عدد را طبقه‌بندی می‌کند.

این ترکیب به CNN اجازه می‌دهد ویژگی‌های پیچیده را از تصاویر استخراج کند، در حالی که SVM با حاشیه بزرگ خود، طبقه‌بندی دقیق و مقاوم‌تری ارائه می‌دهد. این روش به‌خصوص برای مجموعه داده‌های کوچک‌تر که CNN به‌تنهایی ممکن است بیش‌ازحد یاد بگیرد (Overfitting)، بسیار مؤثر است.

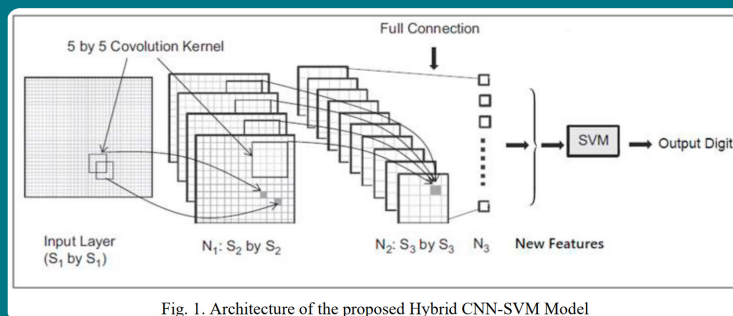


Fig. 1. Architecture of the proposed Hybrid CNN-SVM Model



بهینه سازی و تشخیص تومور های مغزی با استفاده از هوش مصنوعی و تصاویر MRI

چهار حالت مختلف بررسی خواهد شد که عبارت اند از بدون استفاده از روش خاص، استفاده از یادگیری انتقالی، استفاده از افزایش داده ها، ترکیب یادگیری انتقالی و افزایش داده ها عملکرد مدل ها با معیارهایی ارزیابی شده و نتایج با استفاده از آزمون Kruskal-Wallis مقایسه شده اند. در جدول پایین هفت شبکه با معیار های ارزیابی مقایسه و نشان داده شده.

Maximum scores achieved by the seven DL neural networks on the test data. BTD data.

Network	F1_score	Accuracy	Sensitivity	Specificity	Precision
InceptionResNetV2	95,39	97,22	96,17	98,15	97,67
InceptionV3	94,85	96,89	97,81	97,43	94,80
DenseNet121	94,82	96,89	96,72	97,66	96,55
Xception	94,59	96,73	96,72	97,87	94,14
ResNet50V2	93,11	95,91	96,17	98,72	93,89
VGG19	74,04	76,76	82,17	100,00	89,29
EfficientNetB7	68,50	76,92	96,15	100,00	100,00

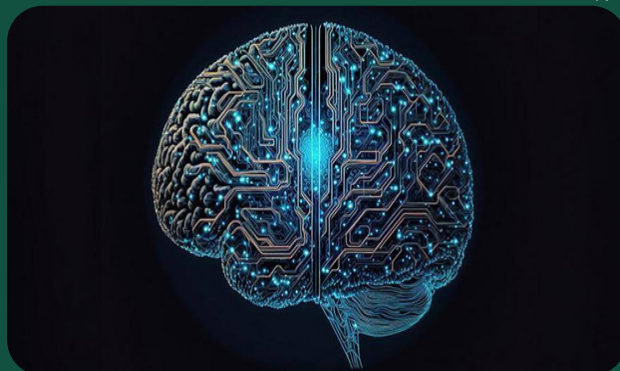
در این جدول سه نوع تومور مغزی با معیار های ارزیابی قابل نمایش است.

Maximum scores achieved by the seven DL neural networks as a function of the three classes.

Class	F1_score	Accuracy	Sensitivity	Specificity	Precision
Pituitary	95,39	97,22	97,81, 8	100,00	100,00
Glioma	93,59	93,94	96,15	100,00	97,67
Meningioma	82,71	92,14	85,92	100,00	100,00

در آموزش و اعتبارسنجی، دقت حدود ۰/۹۵ و خطا نزدیک به ۰/۱ بود و همچنین در ۳۵ دور (Epoch) نشان داد که مدل میتواند با تعداد کمتری از دور ها به عملکرد بالایی دست یابد. تحلیل دو کلاس عدم وجود یا وجود تومور: نتایج نشان داد که توالی های مناسب مانند FLAIR و ترکیب یادگیری انتقالی و افزایش داده می تواند عملکرد تشخیص تومورها را به طور چشمگیری بهبود بدهد. تشخیص تومور در مجموعه داده ها: نتایج نشان می دهند که همه شبکه ها عملکرد بالایی (بیش از ۹۰٪) دارند: - نمره F1 برای آموزش از صفر ۳/۴٪ بهبود یافت. - استفاده از افزایش داده ۰/۰۱٪ بهبود ایجاد کرد. - یادگیری انتقالی دقت را بین کلاس ها تا ۶٪ افزایش داد. - ترکیب یادگیری انتقالی و افزایش داده، دقت را تا ۹٪ نسبت به حالت آموزش از صفر افزایش داد.

تشخیص و شناسایی زودهنگام تومور های مغزی همانند سایر سرطان ها برای اتخاذ اقدامات پیشگیرانه مناسب، بسیار حیاتی است. هوش مصنوعی در سال های اخیر رشد چشمگیری داشته و در محیط های پیچیده ای مانند پزشکی نیز مورد استفاده قرار گرفته است. هوش مصنوعی به ویژه روش های مبتنی بر یادگیری عمیق (DL)، پیشرفت چشمگیری در تشخیص، طبقه بندی و ارزیابی تومور های مغزی داشته اند که این روش ها دقتی قابل مقایسه با رادیولوژیست های متخصص نشان داده اند. توسعه های مختلف در زمینه AI منجر به ایجاد مدل و معماری های جدید برای پردازش تصاویر طبیعی شده اند. تحقیقات اخیر نشان داده که شبکه های یادگیری عمیق برای طبقه بندی و شناسایی تومور های مغزی نتایج امیدوار کننده ای داشته اند، معماری جدیدی به نام cross-Transformer معرفی شده که مبتنی بر مدل های توجه (Attention) است و هفت شبکه یادگیری عمیق دیگر نیز برای طبقه بندی و شناسایی تومورها مقایسه شده اند.



دیتاست ها شامل BTD، MRI-D و TCGA-LGG و سه نوع داده MRI شامل T1، T1W1 و Gd-T1 و FLAIR تحلیل شده است. معیار های متفاوتی برای ارزیابی عملکرد وجود دارند که ما در اینجا از Accuracy (دقت)، Sencitivity (حساسیت)، Specificity (ویژگی)، Precision (دقت پیش بینی)، F1 score استفاده می کنیم. طراحی آزمایش: داده های اولیه به نسبت ۸۰ درصد برای آموزش و ۲۰ درصد برای تست تقسیم شده است. طبقه بندی تومورها به سه نوع (مننژیوما، گلیوما، هیپوفیز) انجام شده است و وجود یا عدم وجود تومور نیز بررسی می شود که در

زمان آموزش مدل ها :

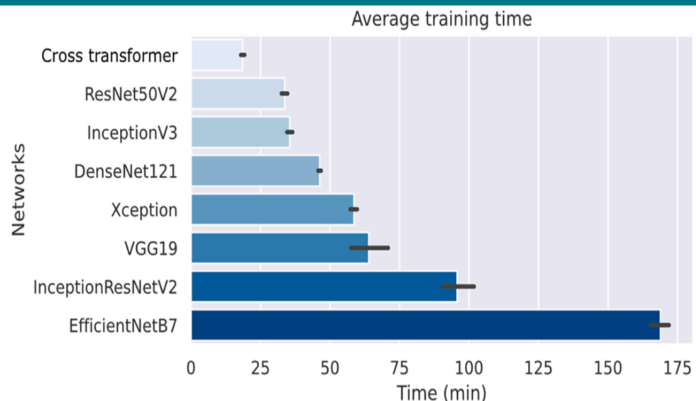


Fig. 10. The average training time of the 8 architectures with 95 % confidence interval (black lines). Training.

مدل Cross-Transformer کارآمدترین مدل از نظر زمان آموزش بود. این مدل به طور میانگین ۱۸ دقیقه برای آموزش ۵۰ دوره با حدود ۳۰۰۰ تصویر زمان نیاز داشت. مدل InceptionResNetV2، پنج برابر و EfficientNetB7 نه برابر بیشتر از cross-Transformer زمان نیاز داشت در نتیجه :

– مدل Cross-Transformer از نظر زمان آموزش و کارایی عملکرد بهتری داشت.

– توالی FLAIR به عنوان بهترین گزینه برای تشخیص تومورهای مغزی تأیید شد.

– مدل های پیشنهادی می توانند به طور مؤثر در تشخیص پزشکی استفاده شوند و بهره وری را افزایش دهند.

نتایج این تحقیق نشان داد که شبکه های عصبی پیشرفته می توانند به طور مؤثری در تشخیص و طبقه بندی تومورهای مغزی استفاده شوند. توالی FLAIR برای تشخیص تومور بسیار مناسب است و مدل Cross-Transformer علاوه بر دقت بالا، زمان آموزش کوتاه تری دارد که آن را به گزینه ای مناسب برای کاربردهای پزشکی تبدیل می کند. این دستاوردها مسیر روشنی برای توسعه سیستم های خودکار دقیق تر و کاربردهای بالینی فراهم می کند.

آزمون های آماری :

Table 5

P-value evaluated between the seven different neural networks using the Kruskal-Wallis test. Classification of the three types of tumors.

p-value							
Networks	1	2	3	4	5	6	7
InceptionResNetV2	1	1,00	0,00	0,77	0,01	0,00	0,00
InceptionV3	2	0,00	1,00	0,00	0,00	0,19	0,00
DenseNet121	3	0,77	0,00	1,00	0,00	0,00	0,00
Xception	4	0,01	0,00	0,00	1,00	0,00	0,00
ResNet50V2	5	0,00	0,19	0,00	0,00	1,00	0,00
VGG19	6	0,00	0,00	0,00	0,00	0,00	0,11
EfficientNetB7	7	0,00	0,00	0,00	0,00	0,00	0,11

بیشترین تعداد توزیعات مشابه با سایر شبکه ها را ResNet50V2 نشان داد.

با بررسی و مقایسه هشت شبکه عصبی از جمله Cross-transformer با سه نوع توالی تصویر به این نتیجه رسیدیم که مدل InceptionResNetV2 با توالی FLAIR بهترین عملکرد را داشت.

– توالی FLAIR در تشخیص تومورها برای بیشتر شبکه ها موفق ترین بود و این روند برای متریک های دقت و نمره F1 نیز صادق بود.

تحلیل مدل با آزمون kruskal wallis :

مدل InceptionResNetV2 بالاترین کارایی را در تشخیص تومورهای مغزی داشت. این مدل در مقایسه با اکثر شبکه ها تفاوت آماری معنی داری نشان داد، اما با مدل Xception شباهت بیشتری داشت (p-value برابر ۰/۰۷)

توالی FLAIR برای اکثر شبکه ها تفاوت آماری معنا داری داشت، به جز مدل های EfficientNetB7 و VGG19.

عملکرد مدل InceptionResNetV2 :

– دقت آموزش و اعتبارسنجی با گذر زمان افزایش یافت و به بالای ۰/۹ رسید.

– خطا برای آموزش و اعتبارسنجی کاهش یافتند و به حدود ۰/۵ همگرا شدند.

– اعتبارسنجی بهتر از آموزش بود، اما برای بهبود عملکرد مدل ممکن است به تعداد بیشتری از دوره ها (epochs) نیاز باشد.

عملکرد مدل Cross-Transformer :

– این مدل دقت بالایی نشان داد، که پس از افزایش تعداد دوره ها به بیش از ۰/۹ رسید.

– با این حال، پس از ۴۰ دوره، دقت اعتبارسنجی کمتر از دقت آموزش شد که نشان دهنده آموزش بیش از حد است.

– خطا نیز مقدار کمی داشت، اما از دوره ۲۵ به بعد، اعتبارسنجی مقادیر بیشتری نسبت به آموزش نشان داد که باز هم نشان دهنده آموزش بیش از حد است.



نویسنده: زهرا انارکی

معرفی کتاب جذاب و کاربردی در حوزه علوم کامپیوتر

اولین کتابی که به ذهنم رسید تا به شما علاقه مندان حوزه علوم کامپیوتر معرفی کنم کتابی در مورد میزان اهمیت مهارت‌های نرم برای مهندسان نرم افزار بود. تقریباً همه‌ی ما می‌دانیم که علاوه بر مهارت‌های برنامه نویسی که لازم و ضروری است همه تا حدودی از آن برخوردار باشیم، داشتن روحیه‌ی کار تیمی، اشتیاق به یادگیری مستمر و همچنین ارتباط گیری موثر با افراد از ارزش بالایی برخوردار است. در ادامه نام و توضیحاتی در مورد کتاب مورد نظر به شما خواهم داد:

Soft Skills: The Software Developer's Life Manual

نویسنده: John Z. Sonmez

سونمز، با تجربه‌ی سال‌ها فعالیت در حوزه‌ی برنامه نویسی در این کتاب سعی کرده است از کد زدن و الگوریتم نوشتن کمی فراتر رود و به سوالاتی مانند اینکه: چطور یک توسعه دهنده‌ی موفق باشیم؟ یا چطور رشد کنیم، تاثیرگذار باشیم و در عین حال نیز تعادل را حفظ کنیم؟ می‌پردازد.

ساختار کتاب این گونه است که در ۷ بخش اصلی به موضوعات زیر می‌پردازد:

۱- Career (شغل):

چطور یک مسیر شغلی موفق بسازیم، چگونه در مصاحبه‌ها بدرخشیم و شغل مناسب خود را پیدا کنیم؟

۲- Marketing yourself (بازاریابی خودمان):

اهمیت برندسازی شخصی برای توسعه‌دهندگان و راهکارهایی برای دیده‌شدن در جامعه‌ی فناوری.

۳- Productivity (بهره‌وری):

تکنیک‌هایی برای مدیریت زمان، تمرکز، کار عمیق و غلبه بر اهمال‌کاری

۴- Learning (یادگیری):

چه روش‌های یادگیری موثر، غلبه بر فراموشی و حفظ مهارت‌ها در دنیای متغیر فناوری وجود دارد؟

۵- Financial (مالی):

مفاهیم ابتدایی سرمایه‌گذاری، پس‌انداز و استقلال مالی برای متخصصان فناوری.

۶- Fitness (تناسب اندام):

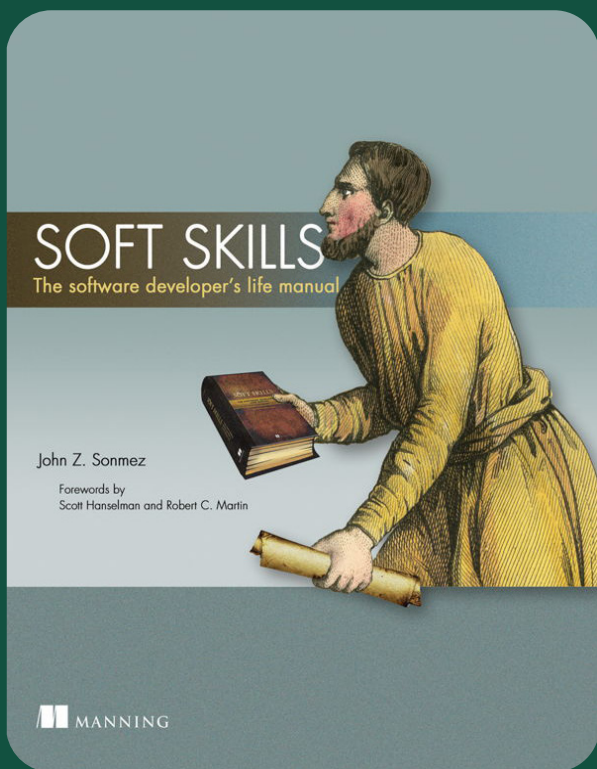
چرا سلامت جسمی برای برنامه‌نویسان مهم است و چطور آن را حفظ کنیم؟

۷- Spirit (ذهنیت):

نقش طرز فکر یک توسعه دهنده در پیشرفت شغلی، یادگیری مداوم و غلبه در چالش‌های حرفه‌ای.

نثر سونمز ساده، خودمانی و بر پایه‌ی تجربیات واقعی خودش است. او به عنوان یک توسعه‌دهنده، مربی و کارآفرین، از چالش‌های شخصی و حرفه‌ای‌اش می‌گوید و با زبانی صادقانه و جذاب، مخاطب را همراهی می‌کند.

همچنین اگر علاقه مند به خواندن ترجمه‌ی این کتاب به زبان فارسی هستید و یا مانند من به مطالعه از روی PDF اعتقادی ندارید هم نگران نباشید، این کتاب توسط آقای ابراهیم نقیب‌زاده‌مشایخ ترجمه گردیده و می‌توانید نسخه‌ی فیزیکی آن را خریداری کنید.



روایتی ساده از دنیای کامپیوترها

کتاب دوم:

The Pattern on the Stone: The Simple Ideas That Make Computers Work

نوشته‌ی W. Daniel Hillis

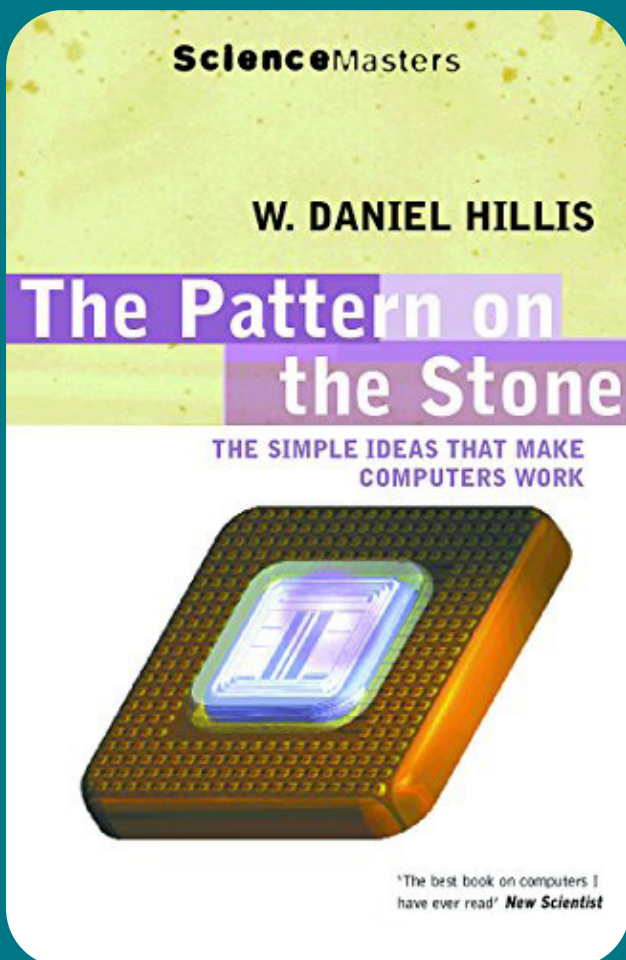
دومین کتاب که به نظرم روایت بسیار ساده و جذاب برای فهم مفاهیم ماشین‌ها را دارد و مناسب دانشجویان تازه ورود به این رشته یا عموم علاقه‌مندان رشته‌مان است کتابی ست که بالاترنامش را آوردم. این کتاب در سال ۱۹۹۸ به زبان انگلیسی نوشته شد. نویسنده، W. Daniel Hillis بنیان‌گذار شرکت Thinking Machines بود که در دهه ۸۰ و ۹۰ از پیشروترین شرکت‌های سازنده‌ی کامپیوترهای موازی بود. او همچنین سابقه همکاری با MIT، Disney و The Long Now Foundation را دارد. دانش او صرفاً نظری نیست، بلکه در عمل سیستم‌هایی ساخته که دیدگاهش را اثبات کرده‌اند.

اگر بخواهم از دلایلی که این کتاب را برای من خاص کرده صحبت کنم اولین و مهم‌ترین دلیل خوشخوان بودن و نترسیدن از پیچیدگی‌ها در این کتاب است. کتاب تفکر عمیق را پرورش می‌دهد و حتی اگر برنامه‌نویس نباشید و یا چیزی از آن ندانید، تصویری دقیق‌تر از چپستی آن در ذهن خواهید داشت.

کتاب ساختار منسجم و کوتاهی داشته و شاید اگر آدم متوسط یا حتی کند خوانی هم باشید حدوداً یک روزه آن را تمام کنید. فصل‌های ابتدایی کتاب در مورد اینکه دکمه‌های روشن و خاموش چگونه کار می‌کنند و منطق دیجیتال پشت آنها چیست می‌گوید: نویسنده با چراغ قوه و کلید برق شروع می‌کند. از همین سادگی، وارد بحث منطق بولی می‌شود: چگونه «روشن» و «خاموش» می‌توانند تبدیل به صفر و یک شوند و این دوتایی‌ها چه کارهای بزرگی می‌توانند انجام دهند.

در میانه‌ی کتاب در مورد مدارهای منطقی و حافظه و همچنین ماشین تورینگ و ایده‌ی محاسبه‌پذیری می‌خوانیم. او با یک معرفی ساده از ماشین تورینگ، این ایده را مطرح می‌کند که چرا بعضی مسائل را می‌توان حل کرد و بعضی‌ها ذاتاً حل‌ناپذیرند. اینجا مرز علوم کامپیوتر با فلسفه روشن می‌شود.

در پایان کتاب کم‌کم متوجه می‌شویم که در واقع گفتن قدم به قدم کارها به یک ماشین ساده و اعمال دستورات روی داده‌ها چیزی جز برنامه‌نویسی و الگوریتم نیست. همچنین Hillis وارد دنیای پیچیده‌تری می‌شود مثلاً اینکه: چرا پردازنده‌ها باید چندوظیفه‌ای شوند، چالش‌های تفکر موازی چیست، و در نهایت، آیا ماشین‌ها می‌توانند واقعاً یاد بگیرند؟



دیستروتاک

اولین ایونت لینوکسی دانشکده علوم کامپیوتر



Semnan Computer Science
انجمن علمی علوم کامپیوتر سمنان

آن و همچنین معرفی Ubuntu توسط سید مهدی اسماعیل زاده رستمی مطرح گردید. در ادامه آرش جعفرپور فوزی به معرفی مفهوم GUI و نقش آن‌ها در شخصی‌سازی محیط کاربری پرداختند. کیوان صحافی توزیع Nix را معرفی کرده و ویژگی‌های منحصر به فرد آن را برای شرکت‌کنندگان توضیح دادند و در پایان، نوید مافی تجربه‌های خود را در استفاده از توزیع Arch Linux به اشتراک گذاشتند و ویژگی‌های حرفه‌ای این توزیع قدرتمند را معرفی کردند.

در انتهای این رویداد، مسابقه‌ی (Capture The Flag) با عنوان SCS CTF: Day Zero نیز برگزار شد که با استقبال دانشجویان همراه بود. شرکت‌کنندگان با حل چالش‌های متنوع امنیتی در فضایی رقابتی و آموزنده در ۱۱ مرحله، مهارت‌های خود را در زمینه امنیت سایبری محک زدند.

برگزاری اولین رویداد لینوکسی دانشکده علوم کامپیوتر «Distro Talk»

اولین ایونت تخصصی «Distro Talk» با محوریت آشنایی با توزیع‌های لینوکس و گفت‌وگو پیرامون آنها، روز دوشنبه ۱ اردیبهشت‌ماه ۱۴۰۴ توسط انجمن علمی علوم کامپیوتر دانشگاه سمنان در تالار خیام دانشکده علوم پایه برگزار شد. این رویداد با استقبال چشمگیر دانشجویان علاقه‌مند به حوزه نرم‌افزارهای متن‌باز و سیستم‌عامل لینوکس همراه بود.

در این برنامه، ابتدا دکتر کیمیا پیوندی؛ استاد مشاور انجمن علمی علوم کامپیوتر درمورد دیستروتاک، برنامه‌های انجمن علمی و اهداف آنها توضیحاتی دادند. سپس سخنرانی دانشجویان علوم کامپیوتر آغاز شد. شکل‌گیری لینوکس و تاریخچه

آیین ناب و نماد تقدیر از دانشجویان دانشکده علوم کامپیوتر

انجمن علمی علوم کامپیوتر دانشگاه سمنان به دانشجویان این رشته که در مراسم آیین ناب و نماد، در روز چهارشنبه ۱۰ اردیبهشت ۱۴۰۴ در تالار خیام دانشکده ریاضی، آمار و علوم کامپیوتر، مورد تقدیر قرار گرفتند، صمیمانه تبریک می‌گوید.

آقای نوید مافی به پاس قبولی در المپیاد علمی دانشجویی، آقای رامتین طریک به دلیل منتخب شدن در طرح راد، آقای محمدمهدی عباسی به عنوان سرآمد فرهنگی و آموزشی، و خانم فاطمه جدیدی به عنوان دانشجوی نمونه سال ۱۴۰۳، افتخارآفرین شدند.



انیاک

نشریه انجمن علمی علوم کامپیوتر دانشگاه سمنان



اولین نسخه از نشریه انیاک، را بخوانید...

