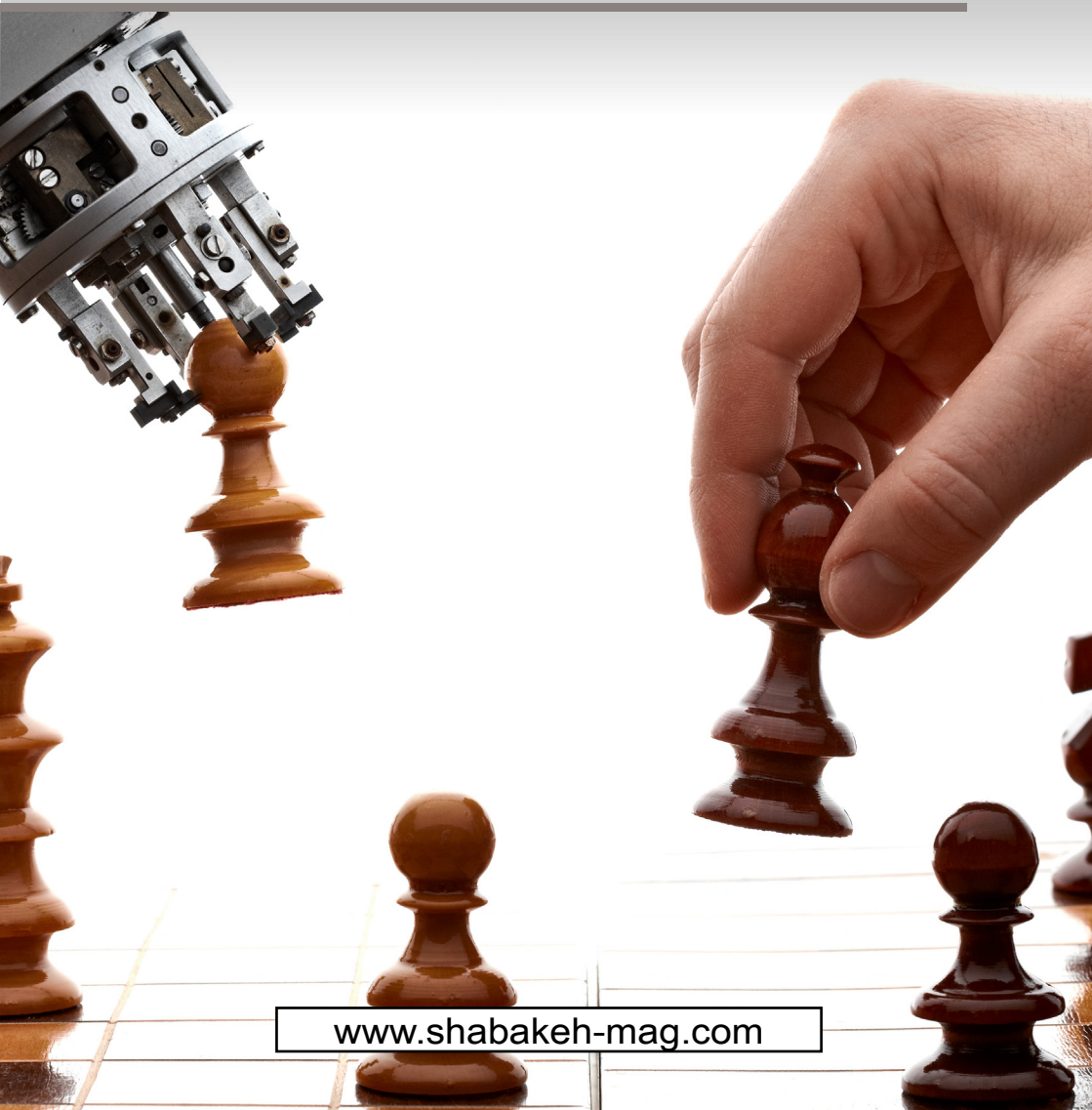


یادگیری ماشینی

سفری به اعماق هوشمندی



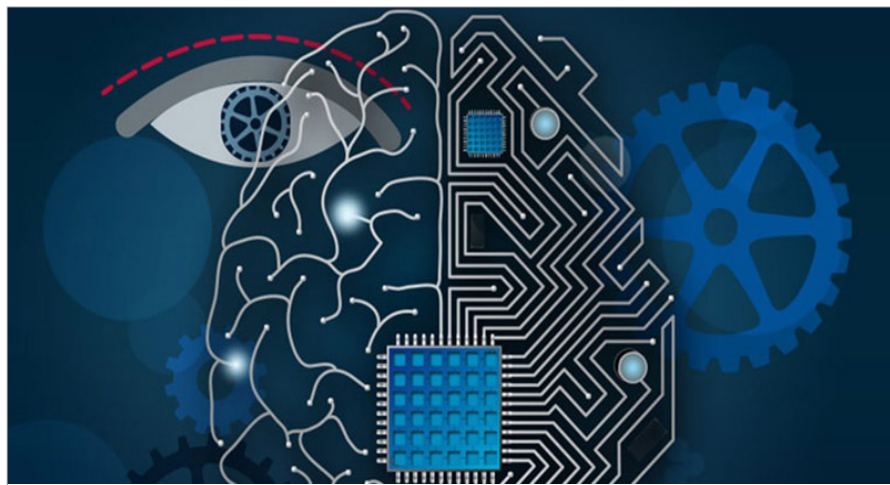
یادگیری ماشینی: جولانگاه تجربیات و خلاقیت‌ها | ۳

۶ سوء تفاهم درباره یادگیری ماشینی | ۱۰

۱۳ چارچوب منبع‌باز برای کسب مهارت در یادگیری ماشینی | ۲۰

یادگیری ماشینی چگونه به بهبود شرایط کسب و کارها کمک می‌کند؟ | ۴۵

یادگیری ماشینی: جولانگاه تجربیات و خلاقیت‌ها



سازمان‌های مدرن امروزی، حجم بسیار گسترده‌ای از داده‌ها را جمع‌آوری می‌کنند؛ داده‌هایی که دارایی‌های سازمانی شناخته می‌شوند. اما سازمان‌ها تنها با تجزیه و تحلیل این داده‌ها به بینش لازم در زمینه اخذ تصمیم‌های روشن دست خواهند یافت. سازمان‌ها، ابتدا داده‌ها را از منبع داده‌ای خود استخراج می‌کنند و سپس به تجزیه و تحلیل آن‌ها می‌پردازند.

تجزیه و تحلیل این داده‌ها منجر به شکل‌گیری بینش می‌شود. در ادامه، بینش به دست آمده در اختیار مدیران ارشد سازمان قرار می‌گیرد و در نهایت تصمیم نهایی اخذ می‌شود.

این فرایند استخراج بینش از داده‌ها، «data analytics» نامیده می‌شود. به فرایند حرکت از داده‌ها به سمت بینش و در نهایت تصمیم‌گیری، تجزیه و تحلیل پیشگویانه داده‌ها گفته می‌شود. تجزیه و تحلیل پیشگویانه، هنر ساخت و به کارگیری مدل‌هایی است که می‌توانند پیش‌بینی‌هایی را بر مبنای الگوهای استخراج شده از داده‌های تاریخی (historical) ارائه کنند. برنامه‌های کاربردی از تجزیه و تحلیل‌های پیشگویانه در مواردی همچون پیش‌بینی قیمت‌ها در کسب و کارهای مرتبط با هتل‌های زنجیره‌ای، خطوط هوایی و مانند این‌ها، ارزیابی ریسک که یکی از تأثیرگذارترین نکات کلیدی در تصمیم‌گیری سازمانی است، مدل‌سازی احتمالات (پیش‌بینی تمایلات فردی)، تشخیص در حوزه پزشکی، مهندسی و نظایر آن، طبقه‌بندی اسناد در گروه‌های مختلف و مانند این‌ها استفاده می‌کنند.

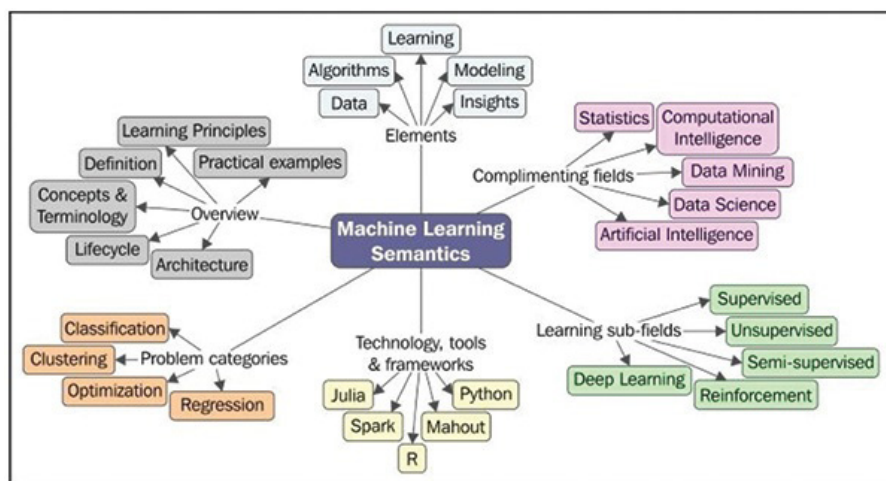
زمانی از اصطلاح یادگیری ماشین (machine Learning) استفاده می‌کنیم که مدل‌های طراحی شده آموزش دیده باشند. یادگیری ماشینی، در قالب فرایند خودکاری که الگوها را از داده‌ها استخراج

می‌کند، تعریف می‌شود. برای اینکه بتوان مدل‌هایی را ایجاد کرد که در برنامه‌های تجزیه و تحلیل پیشگويانه داده‌ها استفاده شوند، روش‌های مختلفی وجود دارد. یادگیری با نظارت (Supervised learning) رایج‌ترین روشی است که در این زمینه استفاده می‌شود. در این روش ارتباط میان مجموعه‌ای از ویژگی‌های توصیفی و ویژگی مقصد (target) بر مبنای مجموعه‌ای از نمونه‌های تاریخی یا موارد مشابه به مدل آموزش داده می‌شود. بعد از این آموزش است که مدل، آمادگی پیش‌بینی رخدادهای جدید را خواهد داشت.

تجزیه و تحلیل پیشگويانه، هنر ساخت و به کارگیری مدل‌هایی است که می‌توانند پیش‌بینی‌هایی را بر مبنای الگوهای استخراج‌شده از داده‌های تاریخی (historical) ارائه کنند.

اگر کمی ریزبینانه به اصطلاح یادگیری ماشینی دقت کنید، با واژه‌ای به نام یادگیری (learning) روبه‌رو خواهید شد. یادگیری در ساده‌ترین توصیف خود اشاره به داده‌های تاریخی یا مشاهدات عینی دارد که برای پیش‌بینی یا مشتق‌گیری وظایف مختلف از آن‌ها استفاده می‌شود. امروزه، یادگیری ماشینی در بسیاری از سرویس‌ها و حوزه‌ها همچون تلفن‌های هوشمند، دنیای وب و از جمله رسانه‌های اجتماعی استفاده می‌شود. جالب اینکه بسیاری از کاربران حتی بدون اطلاع از ماهیت آن، به صورت غیرملموس از

آن استفاده می‌کنند. نرم‌افزارهای تشخیص چهره رایج‌ترین مثالی است که در این زمینه می‌توانیم به آن اشاره کنیم. قابلیت‌ای که می‌تواند عکس دیجیتال را از تصویر انسانی تشخیص دهد. نکته‌ای که لازم است در این مقدمه به آن اشاره کنیم، به تفاوت دو واژه هوش مصنوعی و یادگیری ماشینی بازمی‌گردد. در حالی که این دو واژه کاملاً متمایز از یکدیگر هستند، اما در بطن کار، هر دو در حوزه رایانش به یکدیگر متصل می‌شوند.



هوش مصنوعی در نظر دارد ماشین‌هایی را خلق کند که با تقلید از رفتار ذهنی انسان‌ها به وظایف محوله رسیدگی کنند و این وظایف را همانند انسان‌ها یا حتی بهتر از آن‌ها انجام دهند. برای این منظور، ماشین باید قابلیت یادگیری داشته باشد. در حوزه

هوش مصنوعی، محققان در گروه‌های مختلفی همچون یادگیری، نمایش دانش، منطق و حتی تفکرات انتزاعی، به تحقیق و پژوهش مشغول هستند. اما در نقطه مقابل، تمرکز یادگیری ماشینی بر ساخت برنامه‌هایی است که از تجربیات گذشته درس بگیرند. شاید شنیدن این جمله کمی شگفت‌انگیز باشد که یادگیری ماشینی بیشتر از هوش مصنوعی به داده‌کاوی و تجزیه و تحلیل‌های آماری نزدیک است. یادگیری ماشینی خود به سه گروه اصلی یادگیری با نظارت، یادگیری بدون نظارت و یادگیری تقویت‌شده تقسیم می‌شود. پیش‌تر، توضیح مختصری درباره یادگیری با نظارت ارائه کردیم. اما یادگیری بدون نظارت به روندی اشاره دارد که در آن ماشین بر مبنای داده‌هایی که هیچ‌گاه برچسب‌گذاری نشده‌اند، آموزش می‌بیند. در این روش الگوریتم یادگیری در خصوص اینکه داده‌ها نمایانگر چه چیزی هستند، آگاه نخواهد شد. اما در گروه سوم یادگیری تقویت قرار دارد. این مدل نیز رویکردی مشابه با یادگیری بدون نظارت را دنبال می‌کند؛ با این تفاوت که خروجی به‌دست آمده از پرسش درجه‌بندی می‌شود؛ رویکردی که به‌طور ویژه در بازی‌ها کاربرد دارد. در این مدل زمانی که بازی برنده شود، از نتیجه کار خود برای تقویت حرکاتی که در آینده در خلال بازی انجام خواهد داد، استفاده می‌کند.

یادگیری ماشینی به اندازه‌ای در دنیای روزمره ما رسوخ پیدا کرده است که «ماری برانسکوب»، نویسنده سایت «اینفورلد»، در مقاله مفصل خود با عنوان «یادگیری ماشینی مایکروسافت را بلعیده است» به بیان این نکته پرداخته که مایکروسافت چه دوست داشته باشد و چه نه، در حوزه هوش مصنوعی و یادگیری ماشینی پیشرفت‌های اساسی کرده است.

حال که به‌طور مختصر با یادگیری ماشینی آشنا شدید، شاید این سؤال به ذهنتان رسیده باشد که مدل‌ها چگونه آموزش می‌بینند؟ مدل‌ها بر مبنای طیف گسترده و متنوعی از الگوریتم‌ها این چنین آزمایش‌هایی را پشت سر می‌نهند. در ادامه این پرونده خواهید خواند که الگوریتم‌ها در قالب چارچوب‌های منبع‌باز و غیرمنبع‌باز در اختیار توسعه‌دهندگان قرار دارند. توسعه‌دهندگان با استفاده از چارچوب‌های منبع‌باز می‌توانند الگوریتم‌های رایج را استفاده یا تغییرات مدنظر خود را در آن‌ها پیاده‌سازی کنند. یادگیری ماشینی به اندازه‌ای در دنیای روزمره ما رسوخ پیدا کرده است که «ماری برانسکوب»، نویسنده سایت «اینفورلد»، در مقاله مفصل خود با عنوان «یادگیری ماشینی مایکروسافت را بلعیده است» به بیان این نکته پرداخته که مایکروسافت چه دوست داشته باشد و چه نه، در حوزه هوش مصنوعی و یادگیری ماشینی پیشرفت‌های اساسی کرده است. این پیشرفت‌ها تا آنجا بوده‌اند که حتی کاربران پلتفرم

مایکروسافت ردپایی از این فناوری را در بیشتر محصولات مصرفی این شرکت یا محصولاتش که بر مبنای پلتفرم مایکروسافت ایجاد می‌کنند، احساس می‌کنند.

هوش مصنوعی در نظر دارد ماشین‌هایی را خلق کند که با تقلید از رفتار ذهنی انسان‌ها به وظایف محوله رسیدگی کنند و این وظایف را همانند انسان‌ها یا حتی بهتر از آن‌ها انجام دهند.

در مجموع می‌توانیم این‌گونه عنوان کنیم که اگر در نظر دارید در حوزه‌ای کاربردی و علمی سرمایه‌گذاری کنید، یادگیری ماشینی انتخاب درستی است؛ درست به این دلیل که این فناوری در بسیاری از حوزه‌های کاربردی مورد نیاز بشر و بعضی حوزه‌هایی که شناخت ما از آن‌ها اندک است، به‌خوبی می‌تواند مشکلات را حل کند. بر همین اساس، در پرونده ویژه ماهنامه شبکه شماره ۱۸۱ به بیان یادگیری ماشینی، تأثیر این فناوری بر سازمان‌ها و کسب‌وکارها، کارکردهای عینی این فناوری در زمینه برقراری امنیت در حوزه دستگاه‌های اینترنت اشیا، معرفی چارچوب‌هایی که مهارت را در این حوزه افزایش می‌دهند، بخش‌هایی از زندگی روزمره که به واسطه یادگیری ماشینی بهبود پیدا کرده‌اند و آینده یادگیری ماشینی پرداخته‌ایم.

پرونده ویژه یادگیری ماشینی را می‌توانید از سایت شبکه تهیه کنید.

پیاده‌سازی درست و منطقی

۶ سوء تفاهم درباره یادگیری ماشینی



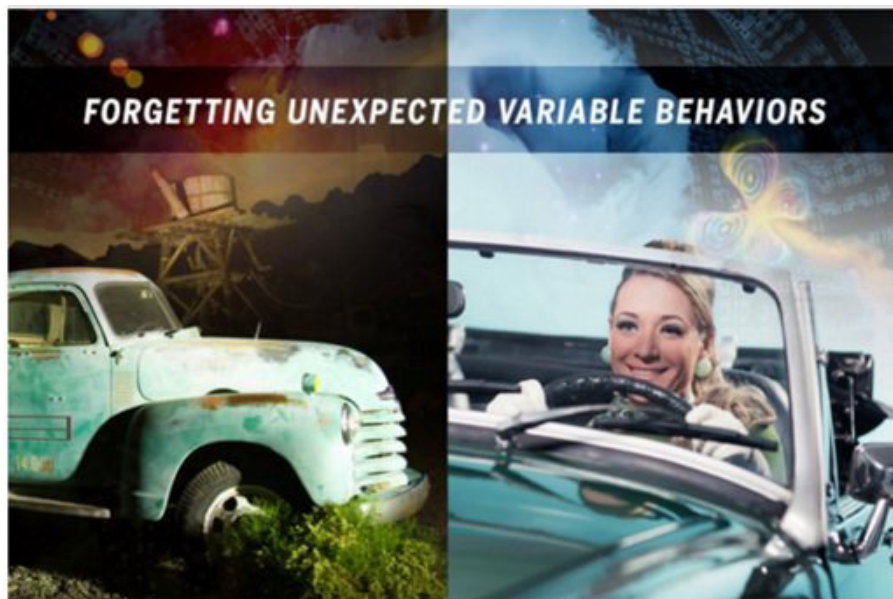
سازمان‌ها پیش از آن که تصمیم بگیرند یادگیری ماشینی را در عمل مورد استفاده قرار دهند باید با کم و کیف این فناوری آشنا شوند تا ناخواسته قربانی یک پیاده‌سازی اشتباه نشوند.

اشتباهات

یادگیری ماشینی به داستان‌هایی که در فیلم‌های علمی‌تخیلی به آن‌ها اشاره می‌شود، محدود نمی‌شود. این فناوری همچون منبع انرژی پیرامون هر یک از فناوری‌هایی که در زندگی روزمره ما مورد استفاده قرار می‌گیرند، قرار دارد. تشخیص صدا از سوی سیری یا آلسا، تشخیص خودکار برچسب‌ها در فیسبوک، توصیه‌هایی که از سوی آمازون و اسپاتی‌فای ارائه می‌شوند همگی بر مبنای یادگیری ماشینی قرار دارند. این حضور به اندازه‌ای موفقیت‌آمیز بوده است که بسیاری از سازمان‌ها علاقه‌مند شده‌اند تا از الگوریتم‌های یادگیری ماشینی برای افزایش بهره‌وری شبکه خود استفاده کنند. در حقیقت تعدادی از این سازمان‌ها به منظور افزایش بهره‌وری سامانه‌های شناسایی و بهینه‌سازی گسترده شبکه‌های خود از مدت‌ها قبل از یادگیری ماشینی استفاده کرده‌اند. اما هر فناوری منجمله یادگیری ماشینی می‌تواند همانند یک شمشیر دو لبه عمل کرده و اگر به شکل نادرستی بیکربندی شود، باعث ویرانی یک شبکه شود. به همین دلیل شرکت‌ها قبل از پذیرش این فناوری باید با راه‌هایی که ممکن است یادگیری ماشینی زمینه‌ساز سقوط آن‌ها شود آشنا شوند و پیش از آن‌که مجبور شوند برای جبران خسارت به عقب گام برداشته و عملیات احیا را اجرا کنند با کم و کاستی‌های این فناوری آشنا شوند. رومان سینایو، مهندس هوشمندی امنیتی نرم‌افزار در ژوپیتور نتورکس راهکارهایی را معرفی

کرده است که از بروز اشتباهاتی که به واسطه یادگیری ماشینی ممکن است یک سازمان را تهدید کند، مانعت به عمل می آورد.

به رفتارهای ناپایدار غیرمنتظره دقت کنید



شگفت‌انگیز است، موضوعی که یک کامپیوتر فکر می‌کند مهم است و به آن واکنش نشان می‌دهد، زمانی که توسط عامل انسانی مورد بررسی قرار می‌گیرد، یک موضوع بی اهمیت تشخیص داده می‌شود. به همین دلیل، ضروری است به این نکته توجه داشته باشیم که بسیاری از متغیرهای زمینه و نتایج بالقوه ممکن است

بعد از استقرار یادگیری ماشینی خود را نشان دهند و در عمل به یک تهدید ناخواسته تبدیل شوند. اجازه دهید این موضوع را با یک مثال روشن کنیم. یک مدل آموزش دیده است تا تصاویر مربوط به وسایل نقلیه سبک و کامیون‌ها را در دو گروه مجزا از هم طبقه‌بندی می‌کند. اما اگر تمامی تصاویر مربوط به کامیون‌ها در شب گرفته شده باشد و تمامی تصاویر مربوط به ماشین‌ها در روز گرفته شده باشد، این احتمال وجود دارد که مدل این گونه تشخیص دهد که هر تصویر مربوط به یک ماشین که در شب گرفته شده است ممکن است یک کامیون باشد. آدرس‌دهی درست متغیرهای کلیدی و نتایجی که در دوره‌های آزمایشی کسب می‌شوند، کمک خواهند کرد تا حد امکان از بروز رفتارهای ناخواسته و غیرمنتظره ممانعت به عمل آوریم و راه‌حل مناسبی برای آن‌ها ارائه کنیم.

غفلت از مشق شب، عدم درک درست داده‌ها



به منظور ساخت یک مدل آموزش دیده، ابتدا باید یک درک اولیه به دست آید و در ادامه داده‌هایی که در فرآیند تحلیل مورد استفاده قرار می‌گیرند، جمع‌آوری شوند. این اطلاعات برای تعیین متغیرها و نتایج بالقوه‌ای که عملکرد یک الگوریتم را تحت تاثیر خود قرار می‌دهند، ضروری هستند. علاوه بر این، اگر یک مدل اصل طبقه‌بندی داده‌ها را فراموش کرده باشد، این امکان وجود دارد که با بهترین داده‌ها که قادر به ارائه یک راه‌حل ایده‌آل هستند آموزش نبینند.

توسعه، آزمایش و در نهایت اجرایی کردن مدل



برای آن که بتوانید مدلی را تولید کنید که کاربردی و مفید باشد، باید در ابتدا یک ساختار داده‌ای آموزش‌دهنده با کیفیت در اختیار

داشته باشید. پیش از آن که یک الگوریتم یادگیری ماشینی داخل یک سازمان استقرار یابد، علم داده‌ها پیشنهاد می‌کند که ابتدا مدل را با مجموعه‌ای از داده‌ها آزمایش کنید تا از عملکرد قابل قبول آن اطمینان حاصل کنید. داده‌ها پیش از آن که در قالب یک ساختار داده‌ای که بتواند قابلیت خودیادگیری را در اختیار یک مدل قرار دهد، آماده شوند باید دو فرآیند را پشت سر بگذارند. اول آن که با تلاش بسیار مجازی‌سازی شده باشند و دوم آن که تحت نظارت قرار گرفته باشند. علم داده‌ها ممکن است یک مدل را در سریع‌ترین زمان ممکن مورد آزمایش قرار دهد، اما برای این منظور ممکن است از مجموعه داده‌هایی استفاده کند که شاید در دنیای واقعی الگوریتم یادگیری ماشینی هیچ‌گاه با آن‌ها روبرو نشود. برای این منظور ضروری است برای متغیرهای انتخاب شده داده‌های کافی در اختیار داشته باشید تا فرآیند آزمایش الگوریتم به درستی انجام شود. تغذیه یک مدل با اطلاعات بیشتر در این مرحله به بهبود عملکرد کمک فراوانی کرده و تضمین می‌کند یک مدل یادگیری ماشینی در عمل باعث افزایش بهره‌وری محیط تولید شده و عملکرد واحدهای عملیاتی را بهبود می‌بخشد.



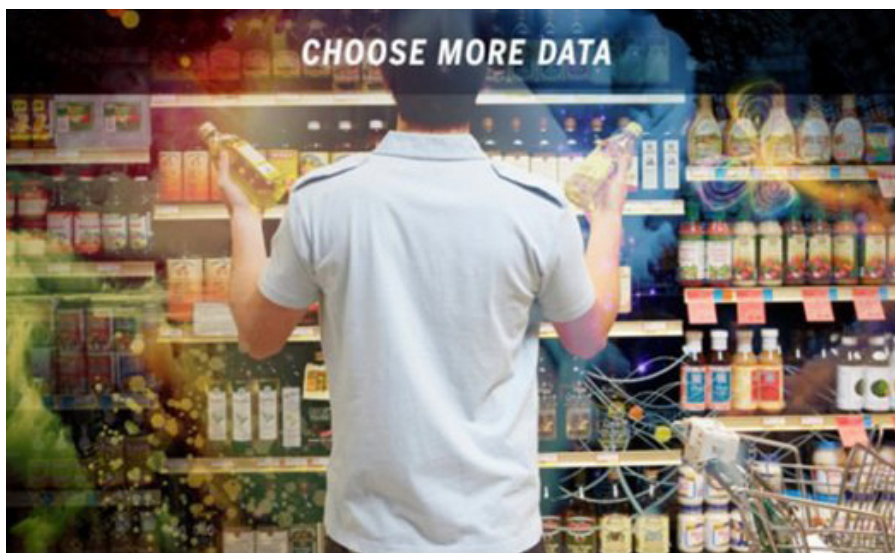
نادیده گرفتن اشتباهات



ممکن است پیش از آن که یک پروژه به هدف نهایی خود نزدیک شود با موانع جدیدی روبرو شود که باعث به وجود آمدن اشتباهات بالقوه‌ای شود. در یک نمونه مشهور، یک شرکت بزرگ یک روبات رسانه‌ای را طراحی کرد. این روبات برای تقلید الگوهای زبان و تکامل بهتر قابلیت‌های تعاملی طراحی شد. اما در عمل کاربران با یک روبات بحث‌برانگیزی روبرو شدند که صحبت‌های جنجال‌برانگیزی را به زبان می‌آورد. به طوری که شرکت در نهایت مجبور شد بخشی از طراحی روبات را که در ارتباط با یادگیری رفتارها بود مجدداً بازطراحی کند. اما در نهایت مجبور شد بعد از گذشت ۲۴ ساعت به کار این روبات برای همیشه پایان دهد.

هر پروژه یادگیری ماشینی را نمی‌توان به صورت عمومی عرضه کرد یا نمی‌توان به کاربران اجازه داد تا به صورت آزاد و بدون کنترل به هر پروژه یادگیری ماشینی دسترسی پیدا کنند و داده‌ها را دستکاری کنند. اما آگاهی از محیطی که الگوریتم در آن وارد می‌شود باعث می‌شود تا از اشتباهات بالقوه جلوگیری شود.

داده‌های بیشتری انتخاب کنید



زمانی که یک مدل را به لحاظ عملکرد مورد آزمایش قرار می‌دهید، ممکن است نتایجی که مدنظر دارید را دریافت نکنید. برای حل این مشکل دو راهکار پیش روی شما قرار دارد. الگوریتم یادگیری ماشینی را بهتر و دقیق‌تر طراحی کنید یا داده‌ها بیشتری

جمع‌آوری کنید. اضافه کردن داده‌های بیشتر به مهندسان کمک می‌کند تا درباره محدودیت‌های عملکردی الگوریتم درک بهتری به دست آورند. اگر بتوانید داده‌های زیادی را جمع‌آوری کنید، آن‌گاه الگوریتم شما به شکل کافی تغذیه می‌شود و در نتیجه قادر خواهید بود نتایج درستی را به دست آورید. این کار به شما کمک می‌کند تا مجبور نباشید الگوریتم خود را از نو مورد بازطراحی مجدد قرار دهید.

قاعده خارج از اصول طراحی نکنید



نوع ویژه‌ای از الگوریتم یادگیری ماشینی که چند وقتی است مورد توجه قرار گرفته و البته کاربردهای عملی آن نیز با موفقیت به اثبات رسیده است، مدل یادگیری تجمعی است. فرآیندی که چند مدل هوش محاسباتی برای حل یک مشکل با یکدیگر ترکیب

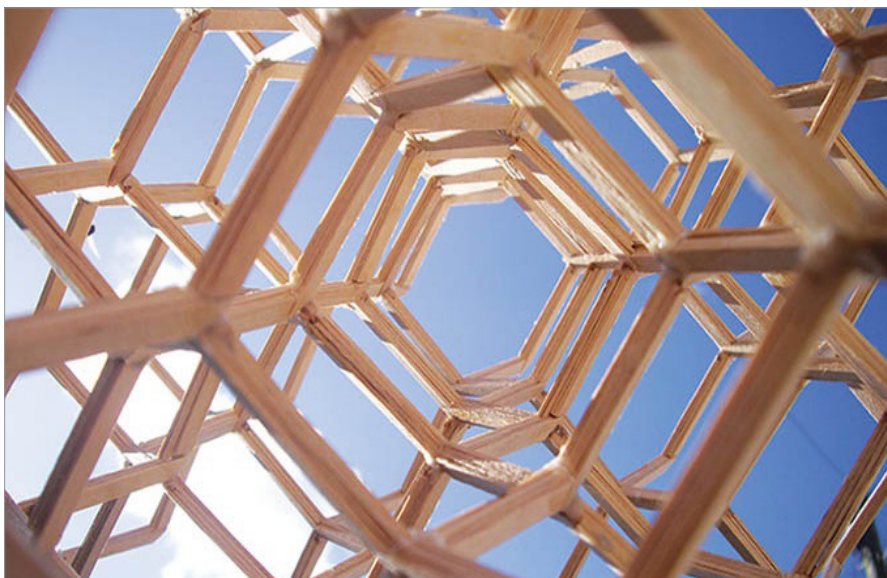
می‌شوند. یک نمونه از مدل یادگیری جمعی رگرسیون لجستیکی است. این روش‌های یادگیری جمعی می‌تواند به بهبود عملکرد پیش‌بینی‌ها در مقایسه با روش‌های مشابه منجر شوند.

کتاب صوتی «به موفقیت عادت کن»

کتاب صوتی (به موفقیت عادت کن) به بررسی عادت‌هایی از افراد موفق می‌پردازد که هر روز به‌عنوان اصل کار و زندگی توسط آن‌ها به‌کار گرفته می‌شوند. راهکارهایی که افراد موفق به‌کار می‌گیرند؛ توسط محققان و جامعه‌شناسان بررسی و طبقه‌بندی شده‌اند که همین‌ک در اختیار شما قرار دارد.



۱۳ چارچوب منبع باز برای کسب مهارت در یادگیری ماشینی



در یک سال گذشته، یادگیری ماشینی به طرز بی‌سابقه‌ای به جریان اصلی دنیای فناوری تبدیل شده است. جالب اینکه روند توسعه محیط‌های ابری ارزان‌قیمت و کارت‌های گرافیکی پرشتاب و قدرتمند، نقش بسزایی در این زمینه داشته‌اند. این عوامل منجر به رشد انفجاری چارچوب‌هایی شده است که اکنون برای یادگیری ماشینی در اختیار کاربران قرار دارند. چارچوب‌هایی که بخش عمده‌ای از آن‌ها منبع باز هستند. اما فراتر از منبع باز بودن، توسعه‌دهندگان به شیوه پیاده‌سازی انتزاعی آن‌ها بیش از پیش توجه کرده‌اند.

این چارچوب‌ها، این ظرفیت را ایجاد کرده‌اند تا دسترسی به پیچیده‌ترین بخش‌های یادگیری ماشینی برای همگان امکان‌پذیر باشد. همین موضوع باعث شده است تا یادگیری ماشینی در طیف گسترده‌ای از کلاس‌ها در اختیار توسعه‌دهندگان قرار گیرد. بر همین اساس، در این فصل تعدادی از این چارچوب‌های یادگیری را معرفی خواهیم کرد.

در انتخاب این ابزارها سعی کرده‌ایم چارچوب‌هایی را که به‌تازگی معرفی یا در یک سال گذشته بازبینی شده‌اند، بررسی کنیم. چارچوب‌هایی که در ادامه با آن‌ها آشنا خواهید شد، امروزه به طرز گسترده‌ای در دنیای فناوری استفاده می‌شوند. این چارچوب‌ها با دو رویکرد کلی طراحی شده‌اند. اول اینکه به ساده‌ترین شکل ممکن مشکلات مرتبط با حوزه کاری خود را حل کنند و دوم آنکه در چالش خاصی که در ارتباط با یادگیری ماشینی پیش روی توسعه‌دهندگان قرار دارد، به مقابله برخیزند.

ApacheSpark MLlib

«Apache Spark» به دلیل اینکه بخشی از خانواده هادوپ است، ممکن است در مقایسه با رقبای خود شهرت بیشتری داشته باشد. در حالی که این چارچوب پردازش داده‌های درون‌حافظه‌ای خارج از هادوپ متولد شد، اما به‌خوبی موفق شد در اکوسیستم هادوپ

خوش بدرخشد. (شکل ۱) Spark یک ابزار یادگیری ماشینی رونده است. این مهم به لطف کتابخانه الگوریتم‌های روبه‌رشدی که برای استفاده روی داده‌های موجود در حافظه استفاده می‌شوند، به وجود آمده است؛ الگوریتم‌هایی که از سرعت بالایی برخوردار هستند.

The screenshot shows the Spark MLlib website. At the top, there's a navigation bar with links: Download, Libraries, Documentation, Examples, Community, and FAQ. Below this, a tagline states: "MLlib is Apache Spark's scalable machine learning library." The main content is divided into several sections:

- Ease of Use:** States "Usable in Java, Scala, Python, and SparkR." It mentions that MLlib fits into Spark's APIs and interoperates with NumPy in Python. It also includes a code snippet for training a model in Python:


```
points = spark.textFile("hdfs://...").\n    .map(parsePoint)\n\nmodel = XGBoost.train(points, k=10)
```
- Performance:** Claims "High-quality algorithms, 100x faster than MapReduce." It includes a bar chart comparing running times for Logistic regression in Hadoop and Spark. The chart shows Hadoop at 110 seconds and Spark at 0.9 seconds.
- Latest News:** A sidebar with recent updates: "CFP for Spark Summit East 2016 is closing soon (Nov 19, 2015)", "Spark 1.5.2 released (Nov 05, 2015)", "Submission is open for Spark Summit East 2016 (Oct 14, 2015)", and "Spark 1.5.1 released (Oct 02, 2015)".
- Built-in Libraries:** Lists "SQL and DataFrames", "Spark Streaming", "MLlib (machine learning)", and "GraphX (graphs)".
- Third-Party Packages:** A section for external integrations.

A "Download Spark" button is prominently displayed. The footer of the page features the "InfoWorld" logo.

شکل ۱: Spark MLlib یک کتابخانه یادگیری ماشینی گسترش‌پذیر است

الگوریتم‌های مورد استفاده در اسپارک دائماً در حال گسترش و تجدیدنظر هستند و هنوز به عنوان موجودیت کاملی خود را نشان نداده‌اند. سال گذشته در نسخه ۱.۵، تعداد نسبتاً زیادی الگوریتم

جدید به این ابزار یادگیری ماشینی افزوده شد، تعدادی از آن‌ها الگوریتم‌های بهبود یافته بودند، در حالی که تعداد دیگری در جهت تقویت پشتیبانی از MLib که در پایتون استفاده می‌شود، عرضه شده‌اند؛ پلتفرم بزرگی که به یاری کاربران رشته آمار و ریاضیات آمده است. Spark نسخه ۱٫۶ می‌تواند کارهای Spark ML را از طریق یک پایپ‌لاین (مجموعه‌ای از عناصر پردازشی داده‌ای) پایدار به حالت تعلیق (Suspend) درآورده و مجدداً از حالت تعلیق خارج کند و به مرحله اجرا درآورد. آپاچی اسپارک متشکل از ماژول‌های یادگیری ماشینی (MLib)، پردازش گراف (GraphX)، پردازش جریانی (Spark Streaming) و Spark SQL است.

Apache Singa

چارچوب‌های یادگیری عمیق، بازوی قدرتمند یادگیری ماشینی به شمار می‌روند و توابع قدرتمندی را در اختیار یادگیری ماشینی قرار می‌دهند. قابلیت‌هایی همچون پردازش زبان طبیعی و تشخیص تصاویر از جمله این موارد هستند. Singa به‌تازگی به درون Apache Incubator راه پیدا کرده است؛ چارچوب منبع‌بازی که با هدف ساده‌سازی آموزش مدل‌های یادگیری عمیق روی حجم گسترده‌ای از داده‌ها استفاده می‌شود. (شکل ۲) Singa مدل برنامه‌نویسی ساده‌ای برای آموزش شبکه‌های یادگیری عمیق بر مبنای کلاستری از ماشین‌ها ارائه می‌کند.

System overview



Figure 1 - SGD flow.

Training a deep learning model is to find the optimal parameters involved in the transformation functions that generate good features for specific tasks. The goodness of a set of parameters is measured by a loss function, e.g., **Cross-Entropy Loss**. Since the loss functions are usually non-linear and non-convex, it is difficult to get a closed form solution. Typically, people use the stochastic gradient descent (SGD) algorithm, which randomly initializes the parameters and then iteratively updates them to reduce the loss as shown in Figure 1.



Figure 2 - SiNGA overview.

SGD is used in SiNGA to train parameters of deep learning models. The training workload is distributed over worker and server units as shown in Figure 2. In each iteration, every worker calls `TrainOneBatch` function to compute parameter gradients. `TrainOneBatch` takes a `NeuralNet` object representing the neural net, and visits layers of the `NeuralNet` in certain order. The resultant gradients are sent to the local stub that aggregates the requests and forwards them to corresponding servers for updating. Servers reply to workers with the updated parameters for the next iteration.

شکل ۲: نمایی از یادگیری ماشینی Singa

این چارچوب از انواع رایجی از آموزش‌ها همچون شبکه عصبی پیچیده (Convolutional neural network)، ماشین بولتزمن محدود (Restricted Boltzmann machine) و شبکه عصبی بازگشتی (Recurrent neural network) پشتیبانی می‌کند. مدل‌ها می‌توانند هم‌زمان یکی بعد از دیگری یا در زمان‌های مختلف و پهلوه‌پهلوه (side by side) آموزش ببینند. انتخاب هر یک از این روش‌ها به این موضوع بستگی دارد که کدام یک برای حل مشکل بهتر جواب می‌دهند. Singa می‌تواند فرایند

کلاستر بندی با Apache Zookeeper را ساده تر کند.

Caffe

چهارچوب یادگیری عمیق Caffe بر مبنای بیان (expression)، سرعت (speed) و پیمانه‌ای بودن (modularity) ساخته شده است. گفتنی است در دنیای کامپیوترها، modularity اشاره به طراحی کامپیوترها در قالب بلوک ساختمانی دارد. این کار با هدف افزایش کارایی و کیفیت تجهیزات انجام می‌شود. این پروژه اولین بار در سال ۲۰۱۳ و به منظور تسریع در پروژه بینایی ماشینی طراحی شد. (شکل ۳) Caffe از آن زمان به بعد توسعه پیدا کرده و از برنامه‌های دیگری همچون گفتار و چندرسانه‌ای پشتیبانی کرده است.

Reference Models

AlexNet: ImageNet Classification

R-CNN: Regions with CNN features

GoogLeNet: ILSVRC14 winner

Caffe offers the

- model definitions
- optimization settings
- pre-trained weights

so you can start right away.

The BVLC models are licensed for unrestricted use.

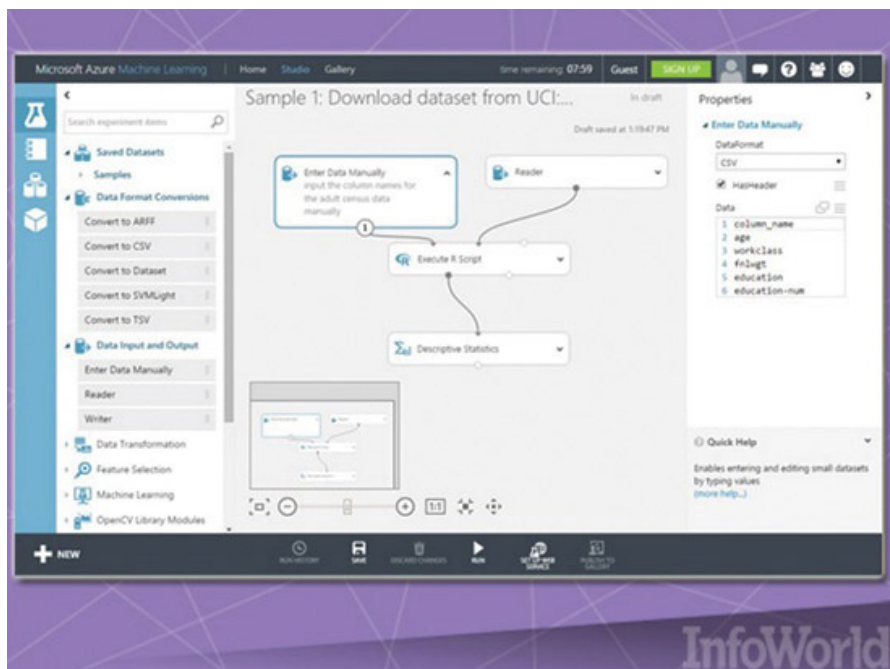
The community shares models in our [Model Zoo](#).

شکل ۳: چارچوب یادگیری ماشینی Caffe، قدرتمند و ساده

مهم‌ترین مزیت Caffè در سریع بودن آن خلاصه می‌شود. بر همین اساس Caffè به طور کامل در زبان سی‌پلاس‌پلاس و با پشتیبانی از شتاب‌دهنده کودا نوشته شد. چارچوب یادشده می‌تواند هر زمان که نیازمند پردازش خاصی هستید، میان پردازشگر مرکزی کامپیوتر و پردازشگر گرافیکی سویچ کند. این توزیع شامل مجموعه‌ای از مدل‌های مرجع منبع‌باز و رایگانی است که به‌خوبی با دیگر مدل‌های ساخته‌شده توسط جامعه کاربران Caffè هماهنگ می‌شود.

Microsoft Azure ML Studio

با توجه به داده‌های حجیم و نیاز به مکانیزم قدرتمند محاسباتی، کلاود محیط ایده‌آلی برای میزبانی برنامه‌های یادگیری ماشینی به شمار می‌رود. مایکروسافت از مکانیزم پرداختی خاص خود در ارتباط با سرویس آژر و یادگیری ماشینی استفاده می‌کند؛ به‌طوری که کاربران در بیشتر نسخه‌های ماهیانه، ساعتی و رایگان می‌توانند از Azure ML Studio استفاده کنند. شایان ذکر است پروژه HowoldRobot نیز بر مبنای همین سامانه ساخته شده است. Azure ML Studio به کاربر اجازه ساخت و آموزش مدل‌های مطبوعش را می‌دهد. (شکل ۴) در ادامه توابعی در اختیار کاربران قرار می‌دهد که با استفاده از آن‌ها از سرویس‌های دیگر استفاده کنند.



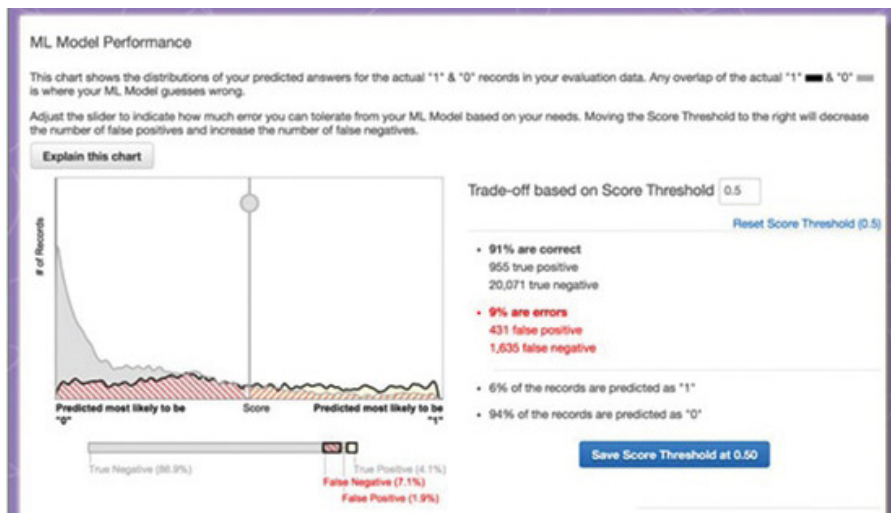
شکل ۴: نمایی از محیط Microsoft Azure ML Studio

کاربران به ازای هر حسایی که برای یک مدل در اختیار دارند، به ۱۰ گیگابایت فضا دسترسی خواهند داشت. در کنار این فضای ذخیره‌سازی، مایکروسافت دسترسی به طیف گسترده‌ای از الگوریتم‌ها را برای کاربران امکان‌پذیر ساخته است؛ الگوریتم‌هایی که از سوی مایکروسافت یا شرکت‌های ثالث ارائه شده‌اند. البته به این نکته توجه کنید برای دسترسی به این چهارچوب لزوماً به حساب کاربری نیاز ندارید. مایکروسافت مکانیزمی را طراحی کرده است که در آن کاربران می‌توانند به‌طور ناشناس وارد شده و از

Azure ML Studio به مدت هشت ساعت استفاده کنند.

Amazon Machine Learning

رویکرد کلی آمازون در خصوص سرویس‌های ابری بر مبنای الگوی خاصی قرار دارد. در این رویکرد آمازون سعی کرده است اصول زیربنایی را در اختیار مخاطبان خود قرار دهد تا کاربران با استفاده از آن‌ها، برنامه‌های سطح بالای خود را طراحی کنند و آگاه شوند که دقیقاً به چه چیزی نیاز دارند. الگویی که به این شکل از سوی آمازون ارائه شده است، اولین شکل از عرضه یادگیری ماشینی در قالب یک سرویس است و Amazon Machine Learning، اولین در نوع خود به شمار می‌رود. (شکل ۵)



شکل ۵: نمایی از عملکرد یادگیری ماشینی آمازون

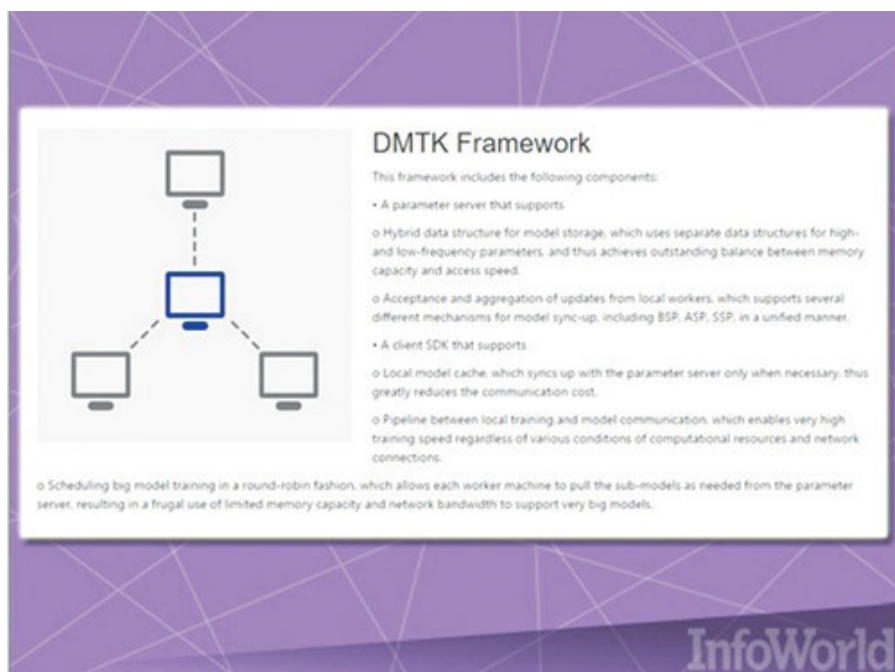
این سرویس به داده‌های ذخیره‌شده در Amazon S3، RedShift یا RDS متصل شده است و می‌تواند طبقه‌بندی دودویی، طبقه‌بندی چندکلاسی یا رگرسیون را بر مبنای داده‌هایی که یک مدل را ایجاد می‌کنند، به وجود آورد که این خدمات به میزان نسبتاً زیادی آمازون محور هستند.

با این حال، این سرویس دارای سه مشکل عمده است. اول آنکه این سرویس به داده‌های ذخیره‌شده در آمازون متکی است، دوم آنکه قابلیت وارد یا خارج کردن مدل‌های خروجی در آن وجود ندارد و سوم آنکه از مجموعه داده‌های بیش از ۱۰۰ گیگابایت برای آموزش مدل‌ها پشتیبانی نمی‌کند. با این حال، آمازون نشان داده است که چگونه یادگیری ماشینی می‌تواند از محصول تزیینی به محصول تجاری تبدیل شود.

Microsoft Distributed Machine Learning Toolkit see

بیشتر کامپیوترهایی که امروزه کاربران استفاده می‌کنند، مشکل عمده‌ای در ارتباط با یادگیری ماشینی دارند. توان پردازشی یک کامپیوتر منفرد برای سازمان‌دهی و مدیریت برنامه‌های یادگیری ماشینی کافی نیست. برای حل این مشکل می‌توان از ترفند خاصی استفاده کرد؛ به طوری که این کامپیوترها گردهم آمده و به یکدیگر متصل شوند. آن‌گاه برنامه‌های یادگیری ماشینی بر مبنای آن‌ها طراحی شده و اجرا شوند. ابزار یادگیری ماشینی

توزیع شده DMTK، سرنام Distributed Machine Learning Toolkit، در اصل چارچوبی است که اسباب و وسایل لازم برای این مسئله را ارائه کرده است. (شکل ۶) به طوری که وظایف مربوط به یادگیری ماشینی را روی کلاستری از سیستم‌ها پخش کرده تا هر یک به بخشی از پردازش‌ها رسیدگی کنند.



شکل ۶: DMTK راهکار مایکروسافت در زمینه ایجاد سیستم‌های توزیع شده ویژه

یادگیری ماشینی

چارچوب DMTK به جای آنکه راه حل کامل و جامعی را ارائه کند، سعی می کند از تعدادی از الگوریتم‌های واقعی در اندازه کوچک تر

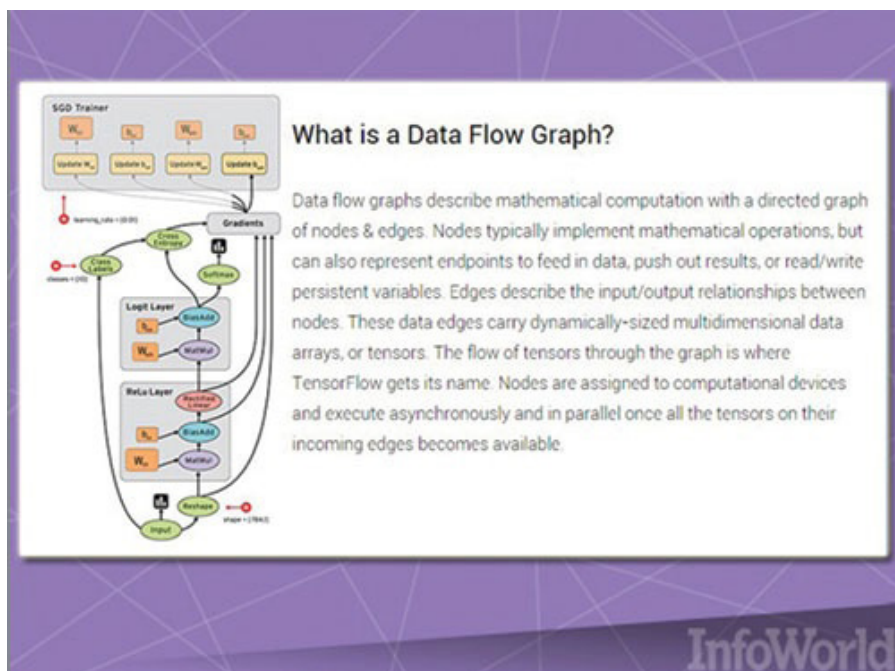
استفاده کند. DMTK به گونه‌ای طراحی شده است که می‌توان به راحتی در آینده آن را توسعه داد. این چارچوب برای کاربرانی که با منابع محدود روبه‌رو هستند، راهکار ایده‌آلی به شمار می‌رود. برای مثال، هر گره در یک کلاستر، کش محلی خود را دارد. همین موضوع باعث می‌شود به میزان قابل توجهی ترافیکی که برای گره سرور مرکزی ارسال می‌شود، کم شود.

کتاب الکترونیک ترند های وای فای



Google TensorFlow

شبهه به پروژه DMTK، پروژه تانسورفلو گوگل یک چارچوب یادگیری ماشینی است که در مقیاس گره‌های چندگانه طراحی شده است. (شکل ۷) آنچنان که Google's Kubernetes برای حل مشکلات داخلی گوگل طراحی شده بود، TensorFlow در قالب محصولی منبع‌باز ویژه کاربران عادی دنیای فناوری عرضه شده است.



شکل ۷: نمایی از دیارگرام کارکردی TensorFlow

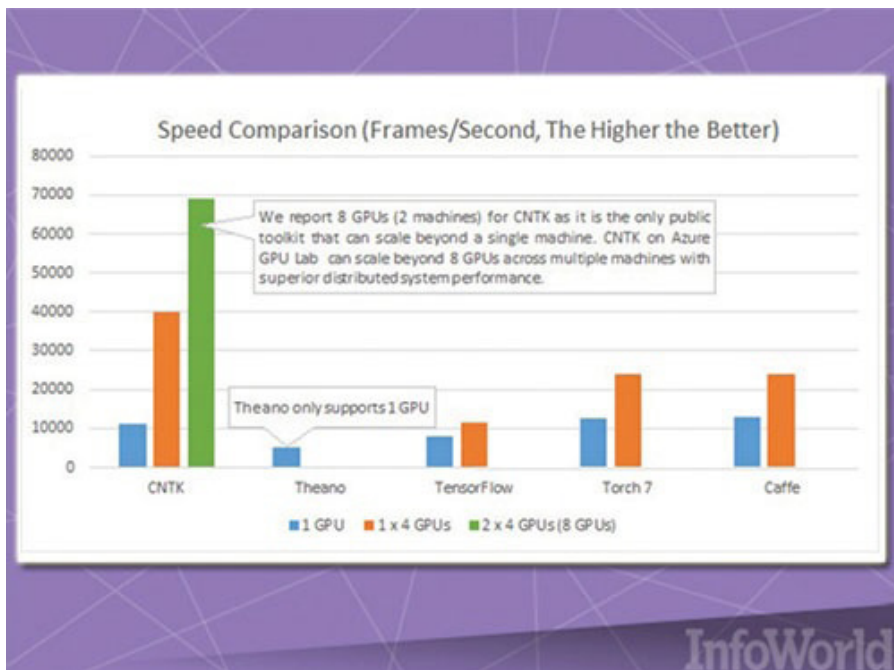
TensorFlow راهکاری است که نمودار جریان داده‌ها نامیده

می‌شود؛ جایی که دسته‌ای از داده‌ها (تانسورها) توسط مجموعه‌ای از الگوریتم‌ها که با استفاده از یک گراف توصیف می‌شوند، پردازش می‌شوند. حرکت داده‌ها از طریق سیستم، flows (جریان) نامیده شده و به همین دلیل این چارچوب TensorFlow نامیده می‌شود. گراف‌ها انعطاف‌پذیر هستند، به گونه‌ای که کاربران با استفاده از زبان‌های برنامه‌نویسی سی‌پلاس‌پلاس یا پایتون می‌توانند دوباره آن‌ها را موتاژ کنند؛ به‌طوری که فرایندها روی پردازشگر مرکزی یا پردازشگر گرافیکی مدیریت شوند. گوگل برنامه‌های بلندمدتی برای TensorFlow در نظر گرفته است و قصد دارد همکاران ثالثی را مجاب سازد تا از این چارچوب استفاده کنند و آن را گسترش دهند.

Microsoft Computational Network Toolkit

اگر بر این باور هستید که DMTK پروژه جالب توجهی از سوی مایکروسافت به شمار می‌رود، باید بگوییم ابزار محاسباتی شبکه مایکروسافت در نوع خود جالب توجه است. CNTK یکی دیگر از ابزارهای یادگیری ماشینی است که از سوی مایکروسافت ارائه شده است. (شکل ۸) CNTK شبیه به TensorFlow گوگل است. از این‌رو به کاربران اجازه می‌دهد شبکه‌های عصبی خود را از طریق گراف‌ها ایجاد کنند. اگر CNTK را با چارچوب‌هایی نظیر Caffe، Torch و Theano مقایسه کنیم، مشاهده خواهیم کرد که CNTK در

بعضی جهات نسبت به چارچوب‌های یادشده برتری‌هایی دارد. از جمله این برتری‌ها می‌توان به سرعت و توانایی استفاده موازی از پردازنده‌های چندگانه مرکزی و گرافیکی آن اشاره کرد.



شکل ۸: CNTK راهکار پیشنهادی میکروسافت در خصوص ساخت شبکه‌های عصبی توسط کاربران

میکروسافت ادعا کرده است که برای آموزش کورتانا در زمینه تشخیص سریع صدا از CNTK همراه با کلاسترهای GPU بر مبنای بستر آژر استفاده کرده است. CNTK در اصل پروژه‌ای توسعه‌یافته است؛ پروژه‌ای که در واحد تحقیقات میکروسافت برای تشخیص

گفتار طراحی شده است. CNTK اولین بار در آوریل ۲۰۱۵ در قالب پروژه‌ای منبع‌باز در اختیار کاربران قرار گرفت، اما در گیت‌هاب تحت مجوز MIT بازنشر شد.

Veles (Samsung)

Veles پلتفرم توزیع‌شده‌ای برای برنامه‌های یادگیری عمیق است. (شکل ۹) شبیه به TensorFlow و DMTK این چارچوب نیز به زبان سی‌پلاس‌پلاس نوشته شده است. اما از پایتون برای فرایندهای اتوماسیون و هماهنگی بین گره‌ها و برای انجام محاسبات از کودا یا OpenCL استفاده می‌کند. مجموعه داده‌ها قبل از آنکه به عنوان خوراک برای تغذیه کلاسترها ارسال شوند، ابتدا تجزیه و تحلیل شده، به طور خودکار عادی سازی شده و برای کلاسترها ارسال می‌شوند.

Veles is usable from IPython or IPython Notebook.

```
In [1]: import veles

In [2]: launcher=veles("../znitz/samples/MnistSimple/mnist.py", stealth=True, backend="ocl", matplotlib_backend="webagg")
```

VELES Version 0.8.0 Tue, 12 May 2015 11:51:57 +0300
Copyright © 2013 Samsung Electronics Co., Ltd.
Released under Apache 2.0 license.
<https://velesnet.nl>
<https://github.com/samsung/veles/issues>

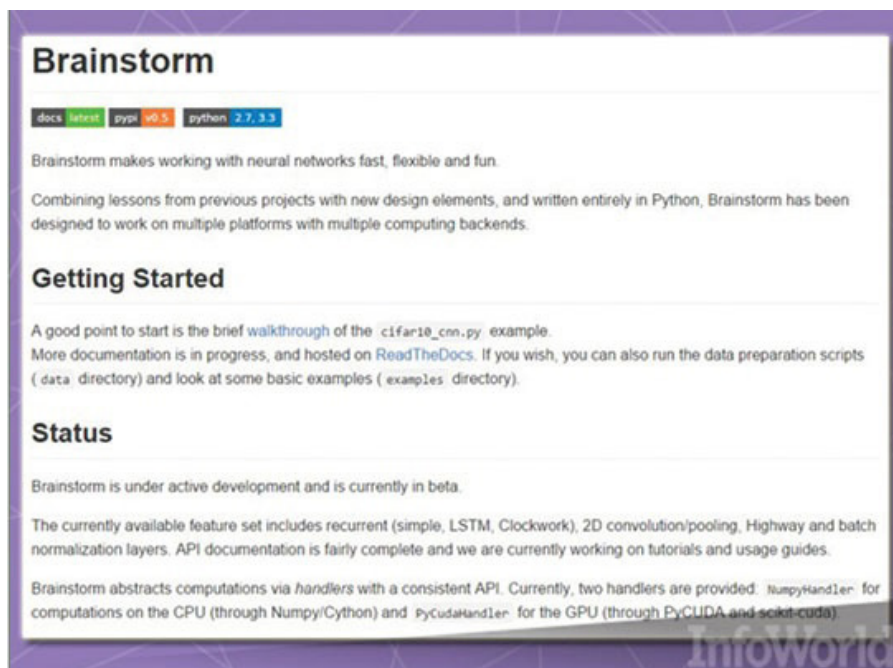
```
INFO:Main:Loading workflow "../znitz/samples/MnistSimple/mnist.py"...
INFO:Main:Applying the configuration from ../znitz/samples/MnistSimple/mnist_config.py...
INFO:Launcher:My Python is CPython 3.4.1
INFO:Launcher:My PID is 32514
INFO:Launcher:My time is 2015-05-14 09:52:16.789227
INFO:Launcher:My ID is 7bc792aa-e93e-4819-88f9-53c64a97bec5
INFO:Launcher:My log ID is 7bc792aa-e93e-4819-88f9-53c64a97bec5
INFO:Main:Created <MnistSimple.mnist.MnistWorkflow object at 0x7f8fcd18f60> with 21 units
INFO:GraphicsServer:Publishing to inproc://veles-plots; ipc:///tmp/veles-ipc-plots-ws2718; epgm://eth2;239.192.1.1:16000; epgm:.
INFO:OpenCLDevice:Selected the following OpenCL configuration:
.....
| device | dtype | rating | BLOCK_SIZE | version |
.....
| NVIDIA Corporation/Geforce GTX TITAN/4318 | double | 1.000 | 27 | 1.1 |
| NVIDIA Corporation/Geforce GTX TITAN/4318 | float | 1.000 | 30 | 1.1 |
.....
INFO:MnistWorkflow:Initializing units in MnistWorkflow...
INFO:MnistLoader:Loading from original MNIST files...
INFO:MnistLoader:Minibatch size is set to 88
INFO:MnistLoader:Samples number: test: 0, validation: 10000, train: 60000
INFO:MnistLoader:Normalizing to linear...
INFO:MnistLoader:There are 10 unique labels
INFO:MnistLoader:train label cardinalities: min: 5421 ("5"), max: 6742 ("1"), avg: 6000, st: 322 (5X)
INFO:MnistLoader:validation label cardinalities: min: 892 ("5"), max: 1135 ("1"), avg: 1000, st: 59 (5X)
```

شکل ۹: Veles پلتفرم توزیع‌شده ویژه یادگیری عمیق

طراحی Veles به گونه‌ای است که به کاربران اجازه می‌دهد با استفاده از توابع REST مدل آموزش‌دیده‌شده را بی‌درنگ در یک محصول استفاده کنند؛ با این فرض که سخت‌افزار خوبی در اختیار داشته باشند. در دنیای محاسبات، REST سرنام *representational state transfer*، سبکی از معماری نرم‌افزاری در محیط وب است. معماری REST سعی در القای کارایی، گسترش‌پذیری، سادگی، قابلیت حمل، قابلیت اطمینان و قابلیت دید دارد. به طور دقیق‌تر REST سبکی از معماری بوده که شامل مجموعه‌ای هماهنگ از اجزا، اتصال‌دهنده‌ها و عناصر داده‌ای است که درون یک سیستم توزیع‌شده ابری قرار دارند، جایی که بر نقش مؤلفه‌ها و مجموعه خاصی از تعامل میان عناصر داده‌ای به جای تمرکز بر جزئیات پیاده‌سازی، تأکید دارد. Veles از پایتون تنها برای کدهای ادغامی _ متصل کردن مؤلفه‌های غیرسازگار نرم‌افزاری _ استفاده نمی‌کند. آی‌پایتون _ در حال حاضر ژوپیتر _ ابزار مصورسازی داده‌ها (Data Visuali- zation) و تحلیل می‌تواند برای مصورسازی و انتشار نتایج از یک کلاستر Velse استفاده شود. سامسونگ امیدوار است با عرضه این پروژه در قالب محصولی منبع‌باز تحرک بیشتری در توسعه آن به وجود آورد؛ به گونه‌ای که تعامل خوبی با پلتفرم‌های ویندوز و Mac OS X داشته باشد.

Brainstorm

Brainstorm پروژه‌ای است که بر مبنای تز پایان‌نامه دانشجویان مقطع دکترا، «کلاوس گراف» و «روپشف استیوستاوا» در مؤسسه هوش مصنوعی «Dalle Molle» واقع در لوگانو سوئیس در سال ۲۰۱۵ طراحی شد. (شکل ۱۰) هدف از طراحی این پروژه پیاده‌سازی شبکه‌های عصبی عمیقی بود که بر پایه عوامل سرعت، انعطاف‌پذیری و سرگرم‌کننده‌ای اجرا شوند. BrainStorm با استفاده از زبان پایتون نوشته شده است. برایان استورم می‌تواند روی پلتفرم‌های مختلف اجرا و همچنین برای انجام محاسبات مختلف استفاده شود.



شکل ۱۰: BrainStorm راهکاری جدید ویژه محاسبات و پلتفرم‌ها مختلف

برایان استورم به خوبی از مدل‌های شبکه عصبی بازگشتی همچون LSTM پشتیبانی می‌کند. طراحان این پروژه به این دلیل زبان پایتون را انتخاب کرده‌اند که BrainStorm از همه ظرفیت‌های موجود به خوبی استفاده کند. به عبارت دقیق‌تر آن‌ها توابع مدیریت داده‌ها را به گونه‌ای سازمان‌دهی کرده‌اند که این توابع با استفاده از کتابخانه Numpy از پرازشگر مرکزی کامپیوتر و با استفاده از کودا از پردازنده‌های گرافیکی برای انجام محاسبات استفاده کنند. در برایان استورم تقریباً بیشتر کارها از طریق اسکرپت‌های پایتون انجام می‌شود، در نتیجه در انتظار رابط گرافیکی قدرتمندی در این زمینه نباشید. سازندگان این چارچوب برنامه‌های بلندمدتی برای این ابزار در نظر گرفته‌اند و درس‌های یادگیری منبع‌بازی را برای آن ارائه کرده و از عناصر طراحی جدید سازگار با پلتفرم‌های مختلف و محاسبات بازگشتی استفاده می‌کنند.

mlpack 2

mlpack 2 کتابخانه یادگیری ماشینی گسترش‌پذیری است که با زبان سی‌پلاس‌پلاس در سال ۲۰۱۱ نوشته شده است. (شکل ۱۱) هدف از طراحی این کتابخانه سهولت استفاده و گسترش‌پذیری اعلام شده است. mlpack دسترسی به الگوریتم‌های موجود را از طریق برنامه‌های ساده و اجرایی خط فرمان و کلاس‌های سی‌پلاس‌پلاس امکان‌پذیر می‌سازد. این مکانیزم می‌تواند با راه‌حل‌های یادگیری

ماشینی عظیم‌تر ادغام شود. همچنین محققان و کاربران حرفه‌ای می‌توانند با استفاده از ماژول‌های زبان سی‌پلاس‌پلاس به راحتی تغییرات مورد نیاز خود را به طور داخلی در الگوریتم‌ها پیاده‌سازی کنند. این رویکرد با هدف رقابت در برابر کتابخانه‌های یادگیری ماشینی بزرگ‌تر در نظر گرفته شده است.



شکل ۱۱: نمایشی از یک برنامه `mlpack2`

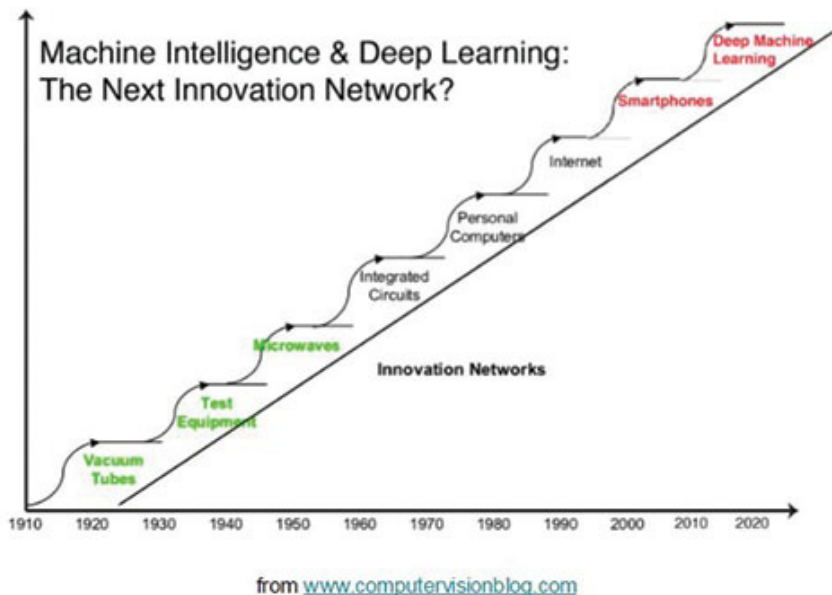
نسخه ۲ این محصول دارای تعداد زیادی ویژگی جدید و فاکتورگیری مجدد است. در نسخه جدید تعدادی الگوریتم تازه معرفی شده است و تعدادی از الگوریتم‌های موجود با هدف افزایش سرعت و روان‌تر

شدن دستخوش تغییراتی شده‌اند. مقیدسازی نکردن به زبان‌های مختلف از جمله معایب این کتابخانه هستند. کاربران زبان‌های برنامه‌نویسی غیر سی‌پلاس‌پلاس، همچون R یا پایتون نمی‌توانند از mpack استفاده کنند، مگر اینکه توسعه‌دهنده‌ای پیدا شود تا این کتابخانه را برای زبان هدف آماده‌سازی کند. برای مثال برای استفاده از این کتابخانه در R پروژه‌ای به نام RcppMLPACK طراحی شده است. این کتابخانه بیشتر برای محیط‌های بزرگی کاربرد دارد که یادگیری ماشینی راهکاری برای حل مشکلات آن‌ها شناخته می‌شود.

Marvin

استیو جوروستون گفته است: «کلان داده‌ها دردسر بزرگی هستند و یادگیری عمیق کلید حل این دردسر بزرگ است.» اگر به تحولات دنیای صنعت و فناوری نگاهی بیندازیم، آنگاه به این حقیقت آگاه می‌شویم که هر ده سال یک‌بار تحولی عظیم در این زمینه رخ داده است. اگر دهه ۹۰ میلادی تا ابتدای سال ۲۰۰۰ میلادی دوران حکم‌فرمایی اینترنت بود، از سال ۲۰۰۰ تا سال ۲۰۱۰ امپراطوری تلفن‌های هوشمند بر همه جا سیطره گسترانیده بود، از این سال به بعد دوران حکمرانی یادگیری ماشینی آغاز شده است. دورانی که تا سال ۲۰۲۰ ادامه خواهد داشت و زندگی مدرن ما را دستخوش تغییرات اساسی خواهد کرد. (شکل ۱۲) Marvin یکی

دیگر از محصولات نسبتاً جدید این حوزه است. چارچوب شبکه عصبی ماروین، محصولی است که Princeton Vision Group آن را طراحی کرده است.



شکل ۱۲: سیر تحول فناوری‌های از سال ۱۹۱۰ تا سال ۲۰۲۰

(شکل ۱۳) این چارچوب با زبان سی‌پلاس‌پلاس نوشته شده و بر مبنای چارچوب پردازش گرافیکی کودا عمل می‌کند. با وجود حداقل کدها، ماروین همراه با تعدادی مدل از پیش ساخته شده و با قابلیت استفاده مجدد در اختیار کاربران قرار گرفته است و کاربران بر حسب نیاز خود قادر به سفارشی‌سازی آن‌ها هستند.

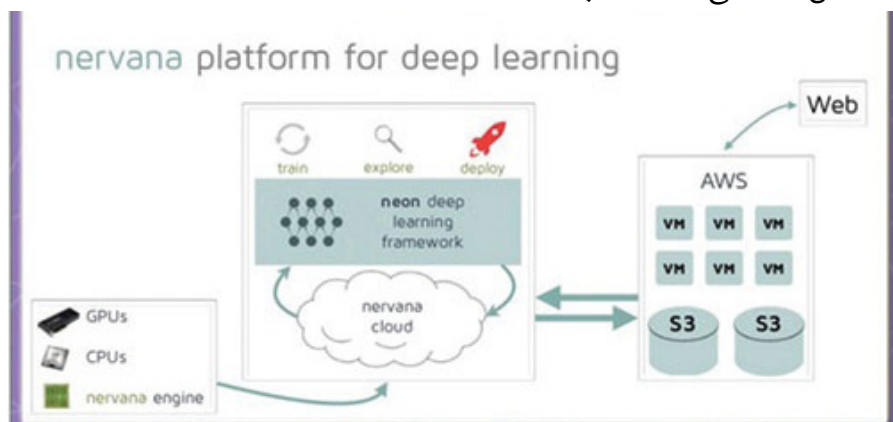


شکل ۱۳: Marvin شبکه عصبی مناسب برای سخت‌افزارهای نه چندان قدرتمند

این چارچوب در مقایسه با شبکه‌های عمیق عصبی مشابه همچون AlexNet، VGG، GoogLeNet سریع‌ترین عملکرد را از خود نشان داده است. بهینه‌سازی برای مصرف کمتر پردازنده گرافیکی، پشتیبانی از پردازنده‌های گرافیکی چندگانه، فیلتر مجازی‌ساز، پشتیبانی از پلتفرم‌های مختلف همچون لینوکس، ویندوز و مک، پیاده‌سازی، اجرای سریع و نظایر این‌ها، از ویژگی‌های این چارچوب هستند.

Neon

Nervana شرکتی است که پلتفرم سخت‌افزاری و نرم‌افزاری خود را برای با یادگیری عمیق طراحی می‌کند. بر همین اساس این شرکت موفق شده است چارچوب یادگیری عمیقی به نام Neon را تولید کند. (شکل ۱۴) چارچوبی که در قالب یک پروژه منبع‌باز در اختیار کاربران قرار دارد. این چارچوب از ماژول‌های ویژه‌ای استفاده می‌کند که توانایی کار کردن با پردازشگر مرکزی، پردازشگر گرافیکی یا سخت‌افزارهای سفارشی خاص این شرکت را دارد. بخش عمده‌ای از طراحی Neon با زبان پایتون انجام شده و تعداد دیگری از مؤلفه‌های آن با زبان‌های سی‌پلاس‌پلاس و اسمبلی نوشته شده‌اند. این زبان‌ها به دو دلیل انتخاب شده‌اند. اول آنکه سرعت محاسبات را افزایش دهند و دوم آنکه کاربران در به‌کارگیری این چارچوب در زبان‌های دیگری همچون پایتون مشکل خاصی نداشته باشند.



شکل ۱۴: نمایی از پلتفرم ارائه‌شده از سوی **nervana** در خصوص یادگیری عمیق

نئون در مقایسه با چارچوب‌های دیگری همچون Caffe و Theano دوبرابر سریع‌تر است. از ویژگی‌های اصلی نئون می‌توان به بهینه‌سازی در سطح اسمبلر، پشتیبانی از پردازنده‌های گرافیکی چندگانه، بهینه‌سازی داده‌های در حال بارگذاری و استفاده از الگوریتم وینوگراد برای محاسبات پیچیده اشاره کرد. در مجموع می‌توانیم این‌گونه بیان کنیم که نئون برای افراد تازه‌کار در حوزه یادگیری ماشینی گزینه ایده‌آلی است. ترکیب نحوی شبیه به پایتون این چارچوب متشکل از پیاده‌سازی تمام اجزای مورد استفاده در یادگیری ماشینی همچون لایه‌ها، قواعد یادگیری، فعال‌سازها، بهینه‌سازها، مقداردهنده‌های اولیه و مانند این‌ها است. بهره‌مندی از مثال‌های متنوع در ارتباط با تشخیص تصویر، گفتار، ویدیو و پردازش زبان طبیعی به عنوان مرجعی خوب و اولیه در اختیار کاربران قرار دارد.




چند وظیفگی ممنوع





به کانال تلگرام شبکه پیوندید

@shabakehmag



هوش مصنوعی در تعامل با جنبه‌های مختلف تجارت

یادگیری ماشینی چگونه به بهبود شرایط کسب‌وکارها کمک می‌کند؟



هوش مصنوعی به اشکال مختلفی به کسب‌وکارها در جهت رسیدن به فرصت‌های تجاری بهتر کمک می‌کند. مهم‌ترین نکته‌ای که باید به آن توجه کنیم، این است که آژانس‌های تبلیغاتی، شرکت‌های بزرگ تجاری هستند. این شرکت‌ها برای اینکه بتوانند در حوزه کاری خود موفق شوند، باید بر دو عامل مهم شناخت مخاطبان و دسترسی به منابع داده‌ای معتبر و کافی اشراف داشته باشند.

هر آژانس تبلیغاتی که به چنین عواملی دست پیدا کرده باشد، بدون شک با حجم بسیار گسترده‌ای از داده‌ها روبه‌رو است؛ داده‌هایی که در عمل امکان تحلیل و داده‌کاوی آن‌ها به شیوه‌های سنتی فرایند زمان‌بری خواهد بود. در چنین شرایطی است که هوش مصنوعی به یاری کسب‌وکارها می‌آید و این فرایندهای زمان‌بر و گاهی پیچیده را به‌سادگی مدیریت می‌کند.

امروزه بسیاری از شرکت‌های بزرگ تجاری از جمله شرکت‌هایی که در بازار داخلی کشور فعالیت دارند، از چنین رویکردی برای دریافت آمارهای واقع‌گرایانه استفاده می‌کنند. قاعده کلی این رویکرد به این شکل است که هر کسب‌وکاری بر حسب نیاز، باید الگوریتم‌های متناسب با حوزه کاری خود را استفاده کرده یا در صورتی که متخصصان این حوزه را در اختیار دارد، الگوریتم‌های مورد نیاز را طراحی کند. با توجه به اینکه امروزه هوش مصنوعی در کانون توجه شرکت‌های بزرگ قرار دارد، در این مقاله تعدادی از کاربردهای یادگیری ماشینی را در حوزه محصولات تجاری بررسی خواهیم کرد.

امروزه هوش مصنوعی و پرچمداران آن، یادگیری ماشینی و یادگیری عمیق، بیش از پیش در صدر اخبار قرار گرفته‌اند. در حالی که یادگیری عمیق بیشتر برای تجهیزات بسیار پیچیده و کاربردهای

خاص استفاده می‌شود و البته طراحی مدل‌ها و آموزش مدل‌های آن فرایند نسبتاً سختی است، یادگیری ماشینی به‌وفور در صنایع مختلف استفاده می‌شود. در حال حاضر، طیف گسترده‌ای از خدمات و محصولات مبتنی بر یادگیری ماشینی در اختیار کسب‌وکارهای کوچک و بزرگ و حتی کاربران عادی قرار دارد. شرکت‌هایی همچون مایکروسافت یا گوگل در تلاش هستند بهترین سرویس‌های تحلیلی و آماری را بر مبنای الگوریتم‌های یادگیری ماشینی در اختیار مشتریان خود قرار دهند. سیری اپل یا بلندگوی اکو آمازون بهترین نمونه‌هایی هستند که می‌توان به آن‌ها اشاره کرد.

در حال حاضر، طیف گسترده‌ای از خدمات و محصولات مبتنی بر یادگیری ماشینی در اختیار کسب‌وکارهای کوچک و بزرگ و حتی کاربران عادی قرار دارد.

سرمایه‌گذاری‌های سنگین در راستای تحقیق و توسعه در حوزه فناوری‌های مرتبط با هوش مصنوعی، امروزه از سوی شرکت‌هایی همچون مایکروسافت، در قالب تراشه‌های مجتمع دیجیتالی برنامه‌پذیر، گوگل، در قالب پروژه‌های ارائه‌شده از سوی DeepMind و به‌ویژه تانسورفلو، اپل، در قالب سیری و خرید استارت‌آپ‌های فعال در این حوزه به‌شدت دنبال می‌شود. شرکت‌های دوراندیش عصر ما به‌خوبی می‌دانند که می‌توانند از ظرفیت‌های مبتنی بر

هوش مصنوعی به منظور درآمدهای زیاد استفاده کنند. حال تصور کنید آژانس‌های تبلیغاتی از این ظرفیت به منظور بهبود کیفیت سرویس‌ها استفاده کنند؛ آن‌گاه چه سود عظیمی عاید این شرکت‌ها می‌شود.

خلق محتوای ارزشی از داده‌های دریافت‌شده از کاربران

داده‌هایی که به طور معمول از کاربران دریافت می‌شود، ارزش چندانی ندارند. داده‌های خامی که از کاربران دریافت می‌شود، بدتر از آن چیزی است که تصور می‌کنید. این داده‌ها مملو از غلط املایی بوده و در اکثر موارد عامیانه هستند؛ به طوری که گاهی صیقل دادن آن‌ها به اطلاعات خام، بی‌فایده خواهد بود. اما شرکت‌ها می‌توانند با اتکا بر الگوریتم‌های یادگیری ماشینی، این داده‌ها را فیلتر کرده، موارد ارزشمند را از داده‌های بی‌هوده تفکیک کرده و در نهایت این داده‌ها را بدون آنکه به عامل انسانی برای برچسب‌گذاری نیازی باشد، سازمان‌دهی کنند. این رویکرد دقیقاً همان چیزی است که آژانس‌های تبلیغاتی به آن نیاز دارند. آیا زمانی را که ایمیل‌تان مملو از اسپم‌های مختلف بود، به یاد می‌آورید؟ یادگیری ماشینی به شرکت‌ها کمک کرد تا فرایند حذف اسپم‌ها را به طور خودکار سازمان‌دهی کنند. به همین دلیل است که امروزه دیگر خبری از آن حجم وحشتناک از هرزنامه‌ها نیست.



کشف سریع تر محصولات

شرکت‌های فعال در حوزه جست‌وجو همچون گوگل، همواره پژوهشگران عرصه یادگیری ماشینی را استخدام می‌کنند. گوگل به‌تازگی کارشناسانی را در زمینه یادگیری ماشینی به گروه جست‌وجوی خود افزوده است. هر چند فرایند شاخص‌گذاری حجم بسیار سنگینی از داده‌ها، قدمتی بسیار طولانی دارد و به دهه ۷۰ میلادی بازمی‌گردد، عاملی که باعث متمایز شدن گوگل از سایر شرکت‌ها می‌شود، به درک سامانه گوگل مربوط می‌شود که همواره نتایج منطبق بر شناخت را به مخاطب خود ارائه می‌کند. این رویکرد باعث شده است گوگل همواره نتایجی را که به محاوره وارد شده شباهت بیشتری دارند، در صدر فهرست نتایج به مخاطب خود نشان دهد. اپل نیز در فروشگاه خود برای نشان دادن برنامه‌های مشابه به یکدیگر از یادگیری ماشینی استفاده می‌کند.

امروزه بسیاری از استارت‌آپ‌های موفق در عرصه تجارت الکترونیک به‌منظور ارائه محتوایی که کیفیت مطلوبی داشته و متناسب با سلیقه کاربر باشد، از زیرساخت‌های مبتنی بر یادگیری ماشینی استفاده می‌کنند. استارت‌آپ‌هایی همچون Rich Relevance و Edge- case از استراتژی‌های یادگیری ماشینی برای نشان دادن محصولات برتر و مورد تقاضا به کاربران خود استفاده می‌کنند.

تعامل بهتر با مشتریان

شاید به این نکته توجه کرده باشید که بخش تماس با ما (contact us) در سال‌های اخیر، بهتر از گذشته شده است. این حوزه‌ای است که یادگیری ماشینی به درون آن وارد شده و به کسب‌وکارها کمک کرده است تا فرایندهای تجاری این بخش را ساده‌تر از قبل مدیریت کنند. کاربران در گذشته برای شرح مشکل خود باید یک فهرست بازشو را انتخاب می‌کردند، مشکل خود را از درون آن برمی‌گزیدند و در ادامه، فیلدهای بی‌پایان موجود در فرم را پر می‌کردند.

امروزه کاربران می‌توانند به‌سادگی مشکل خود را به صورت کوتاه شرح دهند و بهترین پاسخ را از یادگیری ماشینی دریافت کنند. شاید حل این مشکل کمی ساده به نظر برسد، اما در نمونه دیگری، اگر گروه‌های فروش، پژوهش‌های جامع‌تری درباره فروش یک محصول

در اختیار داشته باشند، دیگر برای صف‌های طولانی خریداران ناراضی یا شکایت‌های متعدد خریداران، دغدغه‌ای نخواهند داشت؛ در نتیجه به میزان قابل توجهی در پول و وقت شرکت صرفه‌جویی می‌شود.

درک رفتار کاربران

بدون شک، یادگیری ماشینی در شناخت رفتار کاربران نقش مهمی بازی می‌کند. نتایجی که از درک و شناخت درست رفتار مشتریان به دست می‌آید، در زمینه بازاریابی تأثیرگذاری بی‌بدیلی خواهد داشت. امروزه استارت‌آپ‌هایی که برای بازاریابی محصولات خود به سراغ آژانس‌های تبلیغاتی می‌آیند، از این تکنیک استفاده می‌کنند. برای مثال، یک استودیوی فیلم‌سازی را در نظر بگیرید که تصمیم می‌گیرد تریلر مربوط به فیلم ساخته‌شده خود را نمایش دهد. زمانی که تریلر به نمایش درمی‌آید، این استودیو می‌تواند نظرات کاربران را دریافت و آن‌ها را تحلیل کند و در ادامه، ویژگی‌هایی را که باعث شده‌اند کاربران دید مثبتی به تریلر داشته باشند، برجسته و در تبلیغات بعدی از آن‌ها استفاده کند. این تکنیک باعث می‌شود تعداد بیشتری از مخاطبان کنج‌کاو شوند و برای تماشای فیلم به سینما بروند. یادگیری ماشینی این تحلیل‌های هوشمندانه را به راحتی در اختیار این استودیوی فیلم‌سازی قرار می‌دهد.

گام بعدی در این زمینه چیست؟

سیر تحول و پیشرفت الگوریتم‌های یادگیری ماشینی پیچیده و زمان‌بر است. الگوریتم‌های رایج قابل پیش‌بینی هستند؛ به راحتی می‌توانیم آن‌ها را کالبد شکافی کنیم و دریابیم چگونه کار می‌کنند.

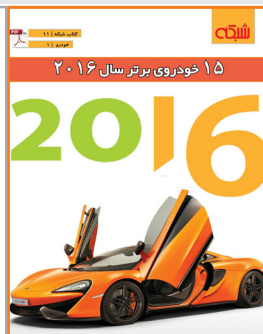
در مقابل، رویکرد تصمیم‌گیری الگوریتم‌های یادگیری ماشینی، شباهت زیادی به نمونه‌های انسانی دارد و همین موضوع باعث می‌شود تا درک دقیق آن‌ها یا توسعه هوشمندانه‌شان کمی زمان‌بر باشد. اگر یک دهه به عقب بازگردیم، به سختی می‌توانیم شرکت‌هایی به غیر از گوگل و یاهو را بیابیم که در حوزه یادگیری ماشینی به تحقیق و توسعه می‌پرداختند. اما امروزه یادگیری ماشینی در هر مکانی یافت می‌شود. داده‌ها بیش از گذشته فراگیر شده‌اند و به ساده‌ترین شکل در اختیار شرکت‌ها قرار دارند. محصولات جدیدی همچون Micro-soft Azure ML و IBM Watson به شکل قابل توجهی هزینه‌های مربوط به راه‌اندازی و استقرار الگوریتم‌های یادگیری ماشینی را کاهش داده‌اند و به شرکت‌ها اجازه می‌دهند بدون اینکه برای پیاده‌سازی چنین الگوریتم‌هایی دغدغه‌ای داشته باشند، به راحتی از آن‌ها استفاده کنند.

امروزه بسیاری از استارت‌آپ‌های موفق در عرصه تجارت الکترونیک به منظور ارائه محتوایی که کیفیت مطلوبی داشته و متناسب با سلیقه کاربر باشد، از زیرساخت‌های مبتنی بر یادگیری ماشینی استفاده می‌کنند.

در فرهنگ عامه مردم این ذهنیت شکل گرفته است که یادگیری ماشینی تنها برای دستیارهای صوتی هوشمند و ماشین‌های خودران کاربرد دارد، اما واقعیت این است که امروزه بسیاری از سایت‌ها در زمان تعامل با کاربر در پشت صحنه از یادگیری ماشینی استفاده می‌کنند. شرکت‌های بزرگ سرمایه‌گذاری‌های سنگینی برای یادگیری ماشینی انجام داده‌اند؛ نه به این دلیل که این فناوری یک تب زودگذر است یا در مقطع فعلی بازار داغی دارد، بلکه به این دلیل که نتایج مثبت آن را به طور ملموس تجربه کرده‌اند. شاید به همین دلیل است که نوآوری همچنان دوست دارد در مسیر پیشرفت گام بردارد تا زندگی راحت‌تری را پیش روی بشریت قرار دهد. به اعتقاد بسیاری از کارشناسان، پژوهش‌ها و فعالیت‌هایی که شرکت‌های بزرگی همچون گوگل در این زمینه انجام می‌دهند، در درازمدت باعث خواهد شد تا عموم توسعه‌دهندگان در سراسر جهان بتوانند از این فناوری در حوزه کاری خود استفاده کنند. به نظر می‌رسد در گام بعد، یادگیری ماشینی بر مبنای الگوریتم‌هایی که توسعه‌دهندگان مختلف در سراسر جهان تولید می‌کنند، به سراغ تجهیزات ملموس زندگی انسان‌ها برود.



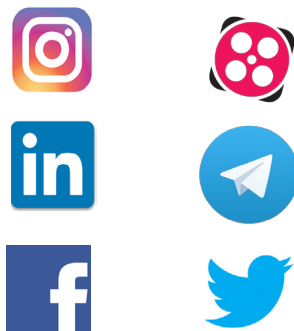
چندوظیفه‌گی ممنوع!



۱۵ خودرو برتر ۲۰۱۶



ترفند های وای‌فای



5G



کتاب در کتاب



برترین بازی‌های ۹۵



دانشگاه در خانه شما

کتاب‌های الکترونیک منتشر شده ماهنامه شباکه



۴۵ ترند مرورگرها

کنترل ویندوز ۱۰

امنیت در ویندوز ۱۰

جویندگان بیت‌کوین



وای‌فای لذیذ

شبکه در ویندوز ۱۰

هر آنچه درباره آیفون ۷

دانلود رایگان آسرا انجام خودروهای خودران از راه خواهند رسید؟

خودران‌ها!