

کلان داده، داده کاوی

و فراگیری ماشین

نویسنده: جرد دین

مترجمین:

دکتر محمدرحیم اسفیدانی (هیات علمی دانشکده مدیریت دانشگاه تهران)

امیرعباس دربانی باسمنج

سرشناسه : دین، جرد، ۱۹۷۸ - م.

Dean, Jared

عنوان و نام پدیدآور : کلان‌داده، داده کاوی و فراگیری ماشین: ایجاد ارزش برای رهبران کسب و کار و متخصصان/نویسنده
جرد دین؛ مترجمین محمدرحیم اسفیدانی، امیرعباس درباری باسمنج، مرکز تحقیقات و نوآوری سازمان
اتکا.

مشخصات نشر : تهران: اتکا، انتشارات مرکز تحقیقات و نوآوری سازمان اتکا، ۱۳۹۴.

مشخصات ظاهری : ۳۰۷ ص.: مصور، جدول، نمودار.

شابک : ۲۱۰۰۰۰ ریال ۷-۷۷-۷۶۲۴-۶۰۰-۹۷۸

وضعیت فهرست نویسی : فیبا

یادداشت : عنوان اصلی: Big data, data mining, and machine learning: value creation for business
leaders and practitioners [2014]

یادداشت : کتابنامه.

موضوع : مدیریت - داده پردازي

موضوع : داده کاوی

موضوع : داده‌های کلان

شناسه افزوده : اسفیدانی، محمدرحیم، ۱۳۵۵ -، مترجم

شناسه افزوده : درباری باسمنج، امیرعباس، ۱۳۶۷- مترجم

شناسه افزوده : مرکز تحقیقات و توسعه

رده‌بندی کنگره : HD۳۰/۲/۹۵۰۰۸۱۱۱۴

رده‌بندی دیویی : ۶۵۸/۰۳۱

شماره کتابشناسی ملی : ۶۸۲



انتشارات مرکز تحقیقات و توسعه سازمان اتکا

عنوان : کلان‌داده، داده کاوی و فراگیری ماشین: ایجاد ارزش برای
رهبران کسب و کار و متخصصان

مؤلف : جرد دین

مترجمین : دکتر محمدرحیم اسفیدانی، امیرعباس درباری باسمنج

ویراستار فنی : الیاس خواجه‌کرمی

زمان و نوبت چاپ : آبان ۱۳۹۴، اول

شمارگان : ۲۰۰

طراح جلد : نوید یوسفی

شابک : ۷-۷۷-۷۶۲۴-۶۰۰-۹۷۸

بهاء : ۲۱۰۰۰۰ ریال

کلیه حقوق چاپ این کتاب محفوظ و متعلق به مرکز تحقیقات، توسعه و کیفیت سازمان اتکا است.
هر نوع تکثیر یا نشر تمام یا بخشی از این کتاب منوط به اجازه کتبی سازمان اتکا خواهد بود.

نشانی: تهران، خ امام خمینی (ره)، خ شهید مرادی، ساختمان شهدای مرکزی اتکا، ط ۲

تلفن: ۰۲۱-۶۱۹۱۵۱۲۴، دورنگار: ۰۲۱-۶۱۹۱۵۳۵۵، وبسایت: www.etkac.ir، ایمیل: tec@ERDCenter.com

۳	پیشگفتار
۵	مقدمه
۸	تشکر و قدردانی
۱۰	مقدمه
۱۵	جد ۱: زمانبندی کلان داده:
۱۹	چرا آسون این مبحث مناسب است
۲۱	آیا کردن داده صرفاً یک هیجان زودگذر است؟
۲۴	وقتی استفاده از کلان داده تفاوتی بزرگ ایجاد میکند
۳۵	بخش اول: محیط محاسبه
۳۸	فصل ۱: سخت‌افزار
۳۸	محل ذخیره (دیسک)
۴۱	واحد پردازش مرکزی
۴۳	حافظه
۴۵	شبکه
۴۷	فصل ۲: سامانه‌های توزیع شده
۴۹	محاسبه پایگاه داده
۵۰	محاسبه سامانه فایل
۵۲	ملاحظات
۵۴	فصل ۳: ابزارهای آنالیز
۵۵	و کا

۵۶	زبانهای جاوا و JVM
۵۹	آر (R)
۶۲	پایتون
۶۳	SAS
۶۷	بخش دوم: تبدیل داده به ارزش کسب و کار
۷۰	فصل ۴: مدل‌های پیش‌بینی
۷۴	یک مدل‌های پیش‌بینی برای ساخت مدلها
۷۷	EMMA
۸۱	رده‌بندی باینری
۸۳	رده‌بندی چند سطحی
۸۴	پیش‌بینی بازه‌ای
۸۵	ارزیابی مدل‌های پیش‌بینی
۹۰	فصل ۵: تکنیک‌های معمول برای مدل‌سازی پیش‌بینی
۹۱	RFM
۹۵	رگرسیون
۱۰۶	مدل‌های تعمیم یافته خطی
۱۱۲	شبکه عصبی
۱۲۶	درخت‌های رگرسیون و تصمیم
۱۳۳	ماشین بردار پشتیبانی
۱۴۱	رده‌بندی شبکه متدهای بیزین
۱۵۶	متدهای گروهی
۱۵۹	فصل ۶: بخش‌بندی

۱۶۶	آنالیز خوشه
۱۶۸	ارزیابی خوشه‌بندی
۱۶۹	تعداد خوشه‌ها
۱۷۲	الگوریتم K-MEANS
۱۷۳	خوشه‌بندی سلسله‌مراتبی
۱۷۴	پایل کردن خوشه‌ها
۱۷۵	فصل ۲: مدل‌سازی پاسخ انگیزه‌ای
۱۷۶	ساختن مدر باست
۱۷۸	اندازه‌گیری پاسخ انگیزش
۱۸۴	فصل ۸: داده‌کاوی
۱۸۴	سری‌های زمانی
۱۸۶	کاهش ابعاد
۱۸۷	یافتن الگوها
۲۰۰	فصل ۹: سامانه‌های توصیه
۲۰۱	سامانه‌های توصیه چه هستند؟
۲۰۲	سامانه‌های توصیه در کجا استفاده میشوند؟
۲۰۳	سامانه‌های توصیه چگونه کار میکنند؟
۲۱۰	ارزیابی کیفیت توصیه و پیشنهاد
۲۱۱	توصیه‌ها در عمل: کتابخانه‌ی SAS
۲۱۸	فصل ۱۰: آنالیز متن
۲۱۹	بازیابی اطلاعات
۲۲۱	طبقه‌بندی محتوا

۲۲۲	متن کاوی
۲۲۴	آنالیز متن در عمل: بیایید "خطرا" بازی کنیم!
۲۴۲	بخش سه: مثال‌های موفق داده‌کاوی در عمل
	فصل ۱۱: بررسی نمونه موردی یک شرکت بزرگ خدمات مالی که دفتر مرکزی آن در ایالات
۲۴۵	متحده است
۲۴۷	فایند کمپین بازاریابی سنتی
۲۵۱	راه حل‌های داده‌ای با عملکرد بالا
۲۵۲	ایجاد ارزش برای معیوب
	فصل ۱۲: بررسی نمونه موردی: یک ارائه دهنده خدماتی بزرگ در حوزه بهداشت و درمان ...
۲۵۶	CAHPS
۲۵۷	HEDIS
۲۵۸	HOS
۲۵۸	IRE
۲۶۶	فصل ۱۳: بررسی نمونه موردی تولید کننده تکنولوژی
۲۶۷	یافتن دستگاه‌های معیوب
۲۷۲	فصل ۱۴: بررسی نمونه موردی مدیریت آنلاین برند
۲۷۷	فصل ۱۵: بررسی نمونه موردی پیشنهاد‌های اپلیکیشن تلفن همراه
۲۸۲	فصل ۱۶: بررسی نمونه موردی یک تولید کننده محصول های-تک
۲۸۳	مدیریت داده گم شده
۲۸۵	کاربری ورای تولید
۲۸۶	فصل ۱۷: نگاهی به آینده
۲۸۷	تحقیقات تجدیدپذیر

۲۸۸	 حریم خصوصی در مورد مجموعه داده عمومی
۲۹۰	 اینترنت اشیاء
۲۹۱	 توسعه نرم افزار در آینده
۲۹۳	 پیشرفت الگوریتمها در آینده
۲۹۶	 در نتیجه
۲۹۸	 فهرست ضمیمه

بسمه تعالی

«کار علمی و تحقیقاتی‌تان را جدی بگیرید. دانش خیلی از نیازهای شما در دانشگاه‌ها

وجود دارد.»

مقام معظم رهبری

حضرت امام خامنه‌ای (مدظله‌العالی)

امروزه با پیشرفت روزافزون علوم رایانه، آمار و ریاضی و همچنین رشد بی‌وقفه فناوری، شاهد گرایش بیش از پیش به استفاده از داده‌ها و اطلاعات در اتخاذ تصمیمات بهینه سازمان‌ها هستیم. با توجه به رشد بی‌سابقه فناوری‌های محاسباتی و بالا رفتن قدرت پردازش رایانه‌ها و همچنین به کارگیری فنون و الگوریتم‌های عظیم محاسباتی، آمار، متخصصان حوزه اطلاعات در دهه‌ی گذشته توانسته‌اند تحلیل‌های بسیار اثربخش را بر روی انبساط داده‌های کلان به منظور استخراج و کشف دانش نهان در آن‌ها انجام دهند.

در یک دهه گذشته رشد نمایی ذخیره‌سازی و بالا رفتن قدرت کشف دانش به گونه‌ای بوده است که بسیاری از کسب‌وکارها به منظور فراهم نمودن زیرساخت‌های مناسب برای ذخیره‌سازی و تحلیل داده، سرمایه‌گذاری‌های کلانی را انجام داده‌اند. امروزه اکثریت شرکت‌ها برای اتخاذ تصمیم مطابق واقعیت، و نه صرفاً یک ظن شخصی، به داده کاوی و تحلیل سطح بالای داده روی آورده و بر آن اساس عمل می‌نمایند. در سال‌های اخیر صنعت خرده‌فروشی از جمله صنایعی بوده است که از این موقعیت و فرصت، بیشترین بهره را برده است. کتاب حاضر به عنوان منبعی کامل و جامع برای امر داده کاوی، توسط پژوهشگر محترم مرکز تحقیقات و نوآوری سازمان اتکا، آقای امیرعباس دریانی ترجمه گردیده است. ضمن سپاسگزاری از مترجم محترم امیدوارم این اثر به جهت ارتقای سطح دانش مدیران و کارشناسان میهن عزیز اسلامی ایران خصوصاً در سازمان اتکا مفید واقع گردد. انشاء...

محمد مهدی کربلایی

مدیرعامل سازمان اتکا

پیشگفتار

من شیفته آنالیز پیش‌بینی هستم و تمام عمر کاری خود را در این حوزه گذرانده‌ام. ریاضیات لذت‌بخش است (حداقل برای من)؛ اما تبدیل آنچه الگوریتم‌ها از آن پرده برمی‌دارند به راه‌حلی که یک مجموعه بتواند از آن استفاده و با آن تولید سود کند ریاضیات را ارزشمندتر هم می‌کند. در این رابطه من و جارد دین به نحوی باهم هم‌عقیده هستیم. ما هر دو دوست داریم شاهد چگونگی تأثیر راه‌حل‌ها بر سازمان‌هایی که با آن‌ها کار می‌کنیم باشیم. گرچه چیزی که ما را شگفت‌زده می‌کند این است که حوزه‌ای که ما پیش‌تر و به‌صورت گام‌به‌گام به کار گرفته‌ایم اکنون به یکی از جذاب‌ترین مشاغل قرن بیست و یکم بدل شده. چگونه چنین اتفاقی افتاده است؟

ما در دنیایی زندگی می‌کنیم که داده‌ها در ابعاد همواره رو به رشد در حال جمع‌آوری شدن هستند. این داده‌ها فعالان سازمان با ماشین‌ها را خلاصه کرده و دانه‌بندی بهتری از رفتارهای آن‌ها ارائه می‌دهند. این روش در بعضی موارد توسط شاخص‌های حجم، تنوع و سرعت توضیح داده می‌شوند و این یعنی ریج کلان داده. این سه شاخص برای درک ارزش در داده‌ای جمع‌آوری می‌شوند که حتی ما هم نمی‌دانیم به چه کار می‌آیند. پیش‌ازاین بسیاری از سازمان‌ها داده‌ها را جمع‌آوری کرده و سمود رویکرد هوش کسب‌وکار که رایج شده است خلاصه‌ای از آن را ارائه می‌دهند.

اما در سال‌های اخیر یک تغییر نمونه‌ای رخ داده است. سازمان‌ها دریافته‌اند که آنالیز پیش‌بینی^۱ مسیر تصمیم‌گیری آن‌ها را تغییر دگرگون می‌کند. الگوریتم‌ها رویکردهای مربوط به مدل‌سازی پیش‌بینی که در این کتاب توضیح داده شده‌اند در بیشتر موارد نیز جدیدی نیستند. خود جارد هم عنوان کرده است که مسئله کلان داده به‌هیچ‌وجه چیز جدیدی نیست. الگوریتم‌هایی که او در مورد آن‌ها توضیح داده است حداقل ۱۵ سال است که

وجود دارند و این گواهی است بر این که اساساً به الگوریتم‌های تازه نیازی نیست. با این وجود مدل‌سازی پیش‌بینی در واقع برای بسیاری از سازمان‌هایی که می‌خواهند تصمیم‌گیری بر اساس داده را گسترش دهند ناشناخته است. این سازمان‌ها نه تنها باید به درک درستی از دانش و اصول مدل‌سازی پیش‌بینی دست یابند، بلکه باید بدانند چگونه می‌توانند آن‌ها را در مورد

^۱ Predictive Analytics

مسائلی که بر ضد استانداردها و پاسخهای این اصول هستند به کارگیرند. اما مدل سازی پیش بینی چیزی بیش از صرفاً ساختن مدل های پیش بینی است. جنبه های عملکردی پروژه های مدل سازی پیش بینی عموماً نادیده گرفته می شوند و کمتر پیش می آید که در کتابها پوشش داده شوند. این امر در ابتدا شامل مشخص کردن سخت افزارها و نرم افزارهای مورد نیاز برای مدل سازی پیش بینی است. همان گونه که جارد توضیح می دهد این موضوع بستگی به سازمان، داده و تحلیل گرانی دارد که بر روی پروژه کار می کنند. اگر منابع مناسب در اختیار تحلیل گران قرار نگیرد پروژه با مشکل برخورد مواجه شده و معمولاً شکست می خورد. به شخصه در زمانی که با در اختیار داشتن سخت افزاری نامناسب بر روی پروژه ای بود شاهد این مسئله بودم. این مشکل باعث شد مجبور شوم زمان قابل توجهی را صرف کاربر روی محدودیت های رم و سرعت پردازش کنم.

در نهایت موفقیت پروژه های مدل سازی پیش بینی بر اساس معیارهای سازمانی که از آن استفاده می کند سنجیده می شود. سائیل^۱ مانند افزایش بهره وری، بازگشت سرمایه، ارزش طول عمر مشتری^۲ و یا معیارهای نرم مانند اعتبار شرکت. من به نمونه های موردی مطرح شده در این کتاب که به این موضوع اشاره دارند بسیار علاقه مند. تعداد زیادی از این نمونه ها در اینجا مطرح شده اند که می توانند توجه شما را جلب کنند. این موضوع به خصوص برای مدیرانی که سعی در درک چگونگی تأثیر مدل سازی پیش بینی بر دستاورد نهایی شان دارند بسیار حائز اهمیت است.

مدل سازی پیش بینی یک دانش است؛ اما به کارگیری موفق راه حل های آن نیازمند مربوط کردن مدل ها به کسب و کار است. تجربه امری ضروری برای تشخیص این رابطه است و این کتاب گنجی از تجاربی است که می توانند شما را در سفر مدل سازی پیش بینی یاریت کنند.

دین ابوت^۲

شرکت تجزیه و تحلیل ابوت^۳

مارس ۲۰۱۴

^۱ Customer lifetime value

^۲ Dean Abbott

^۳ Abbott Analytics

مقدمه

پروژه این کتاب برای اولین بار در اولین هفته کاری من در شغل کنونی به عنوان مدیر داده کاوی اس ای اس^۱ به من پیشنهاد داده شد. نوشتن یک کتاب همواره یکی از آرزوهای من بوده است پس درگیر شدن با آن برای من هیجان زیادی در پی داشت. من متوجه شدم که چرا بسیاری از مردم می خواهند که یک کتاب بنویسند اما تنها اندکی از آن ها آن قدری خوش شانس هستند که بتوانند عقاید و افکارشان را منتشر کنند.

من این فرصت را داشته ام که در طول تحصیل و فعالیت حرفه ایم پیشگام و یا در مرکز پیشرفت های علمی در زمینه داده کاوی باشم و پژوهش هایم را تحت نظارت ذهن های برجسته ای به انجام برسانم. این تجارب به من کمک کردند تا مهارت و تجربه لازم برای خلق این اثر را کسب کنم.

داده کاوی حوزه ای است که به آن عشق می ورزم. من از زمان کودکی می خواستم درک کنم که چیزها چگونه با بازدهی معمول و حتی خیلی بالا کار می کنند. از دوران دبستان تا دبیرستان فکر می کردم مهندسی می تواند شغلی باشد که کنج کاوی و علاقه من به توضیح دنیای پیرامونم را به صورت هم زمان ارضا کند. با این وجود در دوران کارشناسی آمار را کشف کردم و به دام افتادم. در بخش اول کتاب به سال ها سخت افزار و معماری سامانه پرداخته ام. این علاقه ای است که والدینم مرا با مهربانی به آن تشویق کردند. آن هم زمانی که کامپیوترها قیمتی خیلی خیلی بیشتر از ۲۹۹ دلار داشتند. اولین کامپیوتر من یک اپل آی سی با دو فلاپی درایو ۵.۲۵ بود و هارد درایو هم نداشت. چند سال بعد با یک کیت یک کامپیوتر اینتل ۳۸۶ ساختم و به خوبی فشردن دکمه توربو را برای بالا بردن سرعت سی پی یو از ۸ مگاهرتز به ۱۶ مگاهرتز را به یاد دارم. من قانون مور را به صورت مستقیم مشاهده کردم و هنوز این مسئله که تلفن هوشمندم قدرت محاسبه ای بیش از کامپیوترهای برنامه نویسی مشتری، برنامه فضایی آپولو و برنامه شاتل فضایی مدارگرد به صورتی ترکیبی داشتند من را شگفت زده می کند. پس از پایان دوره لیسانس در آمار کار خود را در اداره سرشماری دولت فدرال ایالات متحده آغاز کردم. آنجا من برای اولین بار با کلان داده مواجه شدم. پیش از پیوستن به اداره سرشماری من هرگز یک برنامه کامپیوتری که اجرای آن به زمانی بیش از یک دقیقه نیاز داشت ننوشته بودم، مگر این که هدف نوشتن برنامه ای بود که اجرای بیش از یک دقیقه زمان احتیاج

^۱ SAS

داشت. یکی از اولین پروژه‌های من کار با پرونده اصلی نشانی‌ها بود، یعنی یک فهرست نشانی که توسط اداره سرشماری نگهداری می‌شد. این فهرست نشانی‌ها همچنین شامل نظرسنجی‌های اولیه‌ای است که چارچوب نظرسنجی‌های امروزی است که اداره سرشماری انجام می‌دهد و برای این امر در نه سال بعدی کارهای زیاد دیگری صورت گرفت. این فهرست ۳۰۰ میلیون مورد ثبت شده دارد که اطلاعات مربوط به آدرس، طول جغرافیایی و اطلاعات جغرافیایی در آن‌ها باهم ترکیب شده‌اند. صدها جنبه مختلف در رابطه با هر مسکن در این فهرست موجود است. من در زمان کار کردن با چنین مجموعه‌ای از اطلاعات بهره‌وری، مقیاس پذیری و بهینه سازی، منت‌افزاری در برنامه را آموختم.

قدرت‌های مدیر صبورم "ماری آن"^۱ هستم که به من فرصت آموختن داد و پروژه‌های جالب و ارزشمندی را در اختیارم گذاشت که به من تجربه عملی و امکان نوآوری دادند. این شغل بسیار مهم بود چرا که من توانستم تکنیک‌ها و رویکردهای جدیدی را به کار گیرم که پیش از آن در آن دپارتمان در مورد آن تحقیق نشده بود. درست مثل تمام پروژه‌های جدید دیگر بعضی نظرات عملی شدند و بعضی شکست مواجه گشتند. یکی از پروژه‌های مشخصی که من در آن مشارکت داشتم مشخص کردن این موضوع بود که کدام بلوک در سرشماری سال ۲۰۰۰ بیش شماری و یا کم شماری شده است (داره سرشماری ایالات متحده به مناطق جغرافیایی مخصوصی شامل ایالت، شهرستان، منطقه، مجموعه بلوک‌ها و بلوک تقسیم شده که تقریباً ۸,۲ میلیون بلوک در ایالات متحده وجود دارد). ما راهی برای مشخص کردن دقت مدلمان برای پیش‌بینی انحراف تعداد واقعی واحدهای مسکونی نسبت به تعداد واحدهای مسکونی گزارش شده از طریق داده‌های موجود نداشتیم. خوشبختانه پروژه برای انجام تحقیقات میدانی برای تهیه بازخورد و اعتبار مدل از بودجه مجلس بهره‌مند بود. این ولس‌باری بود که من عبارت "داده‌کاوی" را می‌شنیدم و برای بار اول بود که با سامانه SAS™ Enterprise Miner and CART by Salford مواجه می‌شدم. بعد از مدتی کار در اداره سرشماری ریافتیم که برای رسیدن به اهداف کاریم نیاز به تحصیلات بیشتری دارم و در نتیجه در دپارتمان آمار دانشگاه جرج میسون در فایرفاکس، وی ای^۲ ثبت نام کردم.

در طول دوره کارشناسی ارشد جزئیات بیشتری در مورد الگوریتم‌های معمول در حوزه داده‌کاوی، یادگیری ماشین و آمار آموختم. تجزیه و تحلیل بقا، نمونه‌گیری در پژوهش و آمار محاسباتی بخشی از این موضوع بودند. در زمان تحصیل در دوره کارشناسی ارشد توانستم

^۱ Maryann

^۲ George Mason University in Fairfax, VA.

درس‌هایی که در کلاس آموزش داده می‌شدند را با تجزیه و تحلیل داده‌ای که به صورت عملی در دفتر کارم انجام می‌شد ترکیب، و نوآوری‌های لازم را به آن اضافه کنم. من توانستم تئوری و قوت و ضعف‌های رویکردهای مختلف تجزیه و تحلیل داده و آنالیز پیش‌بینی مرتبط با آن تئوری‌ها را درک کنم.

پس از پایان دوره کارشناسی ارشد مسیر شغلیم را از تجزیه و تحلیل داده به تولید نرم افزار تغییر دادم. پس شروع به کار در موسسه اس‌ای‌اس کردم که خالق همان نرم افزاری بود که من قبلاً از آن استفاده می‌کردم. من از مرحله استفاده از نرم افزار به تولید آن تغییر مکان داده بودم. این مسئله چالش‌ها و فرصت‌های جدیدی برای رشد در بر داشت چرا که متوجه صحت و دقت محدودی که اس‌ای‌اس در نرم افزارها در کنار مستندات همه جانبه‌شان، تلاش خستگی ناپذیر آن، برای رسیدن به پیشرفت‌های جدید نرم افزاری قابل انطباق با نرم افزارهای موجود و ارائه خصوصیات نرم افزاری جدید مورد نیاز مشتریان شدم.

در سال‌هایی که در اس‌ای‌اس کار می‌کردم کاملاً فهمیدم که نرم افزار چگونه تولید می‌شود و مشتریان چگونه از آن بسته ده می‌کنند. من معمولاً بخت آن را دارم که با مشتریان دیدار کنم، چالش‌های شغلی آن‌ها را بشنوم و در روش‌ها یا فرآیندهایی را به آن‌ها پیشنهاد بدهم که به موفقیت رهنمونشان کند و برای سازمانشان ایجاد بهره کند.

این کتاب در نتیجه این مجموعه تجارب من در کنار کمک‌های کارمندان و همکاران بی نظیرم هم در موسسه اس‌ای‌اس و هم در جاهای دیگر نوشته شده است.

مقدمه

درون این پشته داده دانشی نهفته است که می‌تواند زندگی یک بیمار و یا جهان را دگرگون کند.

اتول برت، دانشگاه استنفورد^۱

سرطان نامی است که به استهوان از بیماری‌ها داده شده که در آن‌ها سلول‌های غیرمعمول به شکلی غیر قابل واپایش تهاجم می‌شود و به بافت‌های بدن حمله می‌کنند.

بیش از ۱۰۰ نوع منحصر به فرد و متفاوت سرطان وجود دارند که بیشترشان بر اساس جایی (معمولاً یک عضو) که بیماری از آن آغاز شده نام‌گذاری شده‌اند. سرطان از سلول‌های بدن آغاز می‌شود. در شرایط طبیعی بدن میزان تولید و واپاشی جدید جهت جایگزینی با سلول‌های فرسوده یا آسیب دیده را کنترل می‌کند. سرطان طبعاً نیست. در بیماران سرطانی سلول‌ها در زمان مقرر نمی‌میرند و سلول‌های جدید در زمانی که به آن‌ها نیاجی نیست تولید می‌شوند (درست مثل موقعی که من از بچه‌هایم می‌خواهم برایم چیزی با استفاده از روگرفت چاپ کنند و آن‌ها به جای یک نسخه ده نسخه از چیزی که خواسته بودم، با من می‌آورند). سلول‌های اضافی تشکیل یک توده بافت را می‌دهند که به آن تومور گفته می‌شود. تومورها به دو دسته تقسیم می‌شوند: تومورهای خوش خیم که سرطانی نیستند و تومورهای بدخیم که سرطانی هستند. تومورهای بدخیم در سراسر بدن پخش می‌شوند و به بافت حمله می‌کنند. خانواده من، مانند بیشتر خانواده‌هایی که می‌شناسم، اعضای زیادی دارد که این بیماری را دارند. در سال ۲۰۱۳ در ایالات متحده ۱٫۶ میلیون مورد سرطانی جدید تخمین زده شده است و بیش از ۵۸۰۰۰۰ مرگ در اثر این بیماری رخ داده است.

¹ Atul Butte, Stanford University

در حدود ۲۳۵۰۰۰ نفر در ایالات متحده در سال ۲۰۱۴ مبتلا به سرطان پستان تشخیص داده شده‌اند و در حدود ۴۰۰۰۰ نفر در سال ۲۰۱۴ در نتیجه ای بیماری در خواهند گذشت. شایع‌ترین نوع سرطان پستان کارسینوم داکتال^۱ است که از پوشش داخلی مجاری شیری آغاز می‌شود. دومین گونه شایع سرطان پستان نوع لوبولار^۲ است. درمان‌های متفاوتی از جمله عمل جراحی، شیمی درمانی، پرتو درمانی، ایمن درمانی و واکسن درمانی برای سرطان پستان وجود دارند. معمولاً یک یا بیش از یک گزینه درمانی برای اطمینان از دست یابی به بهترین نتیجه برای بیمار استفاده می‌شوند. سازمان غذا و دارو^۳ در حدود ۶۰ داروی مختلف را برای درمان سرطان پستان تأیید کرده است. تصمیمات در زمینه دوره درمان و انتخاب پروتکل درمانی مورد استفاده در مشوره‌های بین پزشک و بیمار صورت می‌گیرند و برخی فاکتورها به همین تصمیمات بر می‌گردند.

یکی از داروهای مهم، داکتال سازمان غذا و دارو برای درمان سرطان تاموکسیفن سترات^۴ است. این دارو که تحت نام تجاری نولوا^۵ به فروش می‌رسد اولین بار در سال ۱۹۶۹ در انگلستان تجویز و در سال ۱۹۹۸ توسط سازمان غذا و دارو تأیید شده است. تاموکسیفن معمولاً به صورت قرص‌های روزانه و در دوزهای ۱۰، ۲۰ و ۴۰ میلی گرم مصرف می‌شود. دارو دارای اثراتی جانبی از جمله تهوع، سوء هاضمه و سرگیجه^۶ است. ساق پا است. تاموکسیفن برای درمان میلیون‌ها زن و مرد مبتلا به سرطان پستان وابسته به هورمون به کار گرفته شده است. معمولاً تاموکسیفن به دلیل میزان موفقیت بالا، یعنی ۸۰٪ در حدود ۸۰٪ اولین دارویی است که برای درمان سرطان پستان تجویز می‌شود. علم به این موضوع که یک دارو ۸۰٪ موفق بوده ما را به تأثیر آن بر روی بیماران امیدوار می‌کند. اما یک نکته مهم در مورد دارو تا پیش از دوران کلان داده ناشناخته بود: این که تاموکسیفن، ۸۰٪ بر روی بیماران تأثیر انداخته است، بلکه در ۸۰٪ بیماران ۱۰۰٪ تأثیر گذار و در باقی موارد بی تأثیر بوده است. این یافته سالانه زندگی هزاران نفر را تغییر داده است. با استفاده از تکنیک‌ها و ایده‌های مطرح شده در این کتاب دانشمندان توانسته‌اند نشانگرهای ژنتیکی را شناسایی کنند که در همان ابتدا مشخص می‌کنند که آیا تاموکسیفن شخص مبتلا به سرطان پستان را به طور موثر درمان خواهد کرد یا

^۱ ductal carcinoma

^۲ lobular

^۳ Food and Drug Administration (FDA)

^۴ tamoxifen citrate

^۵ Nolvadex

^۶ Nolvadex

خیر. این گونه از تجزیه و تحلیل پیش از دوران کلان داده ممکن نبود. چرا چنین چیزی امکان پذیر نبود؟ چرا که حجم و دانه بندی داده وجود نداشت. حجم از ادغام نتایج مربوط به بیماران و دانه بندی از ترتیب گذاری دی ان ای به دست می آیند. علاوه بر داده، منابع محاسباتی برای حل مسائلی از این دست برای بیشتر دانشمندانی که بیرون از آزمایشگاه محاسبات عالی بودند به سادگی در دسترس نبود. در آخر، سومین عنصر یعنی الگوریتم ها و یا تکنیک های مدل سازی برای درک این رابطه در سال های اخیر پیشرفت قابل توجهی کرده اند.

داستان تاموکسیفن فرصت های هیجان انگیزی را که داشتن داده های بیشتر و بیشتر در کنار منابع ناساتی و الگوریتم ها برای ما فراهم می کنند را مشخص تر می کند. این موضوع به ما در طبقه بندی و پیش بینی نیز یاری می رساند. با چنین دانشی که دستاورد دانشمندانی است که بر روی تار نسیم تحقیق کرده اند ما می توانیم تغییر درمان را آغاز کنیم و بسیاری از حوزه های دیگر نیز تسلیما را نیز به صورت مثبت متحول کنیم. ما می توانیم با استفاده از این پیشرفت ها از درمان ها و وسایلی که ما را از جلوگیری کنیم و به جای آن مشخص کنیم که کدام درمان خاص برای هر شخص مفیدتر است. دیگر یک دارو ۵٪ تأثیر گذار نخواهد بود، بلکه ما اکنون می توانیم بگوییم دارو به کدام ۱٪ از این بیماران کمک خواهد کرد. مفهوم داروی شخصی سال ها است که مطرح شده است. با پیشرفت های که در کار با کلان داده و آنالیز پیش بینی صورت گرفته است این مفهوم بیش از هر زمان دیگری به حقیقت نزدیک شده است. دارویی با ۲٪ نرخ موفقیت تا زمانی که مشخص نشود برای کدام دسته از بیماران مفید است توسط تولیدکنندگان دارو دنبال و از جانب سازمان غذا و دارو تأیید نخواهد شد. در صورت وجود این اطلاعات می توان زندگی ها را نجات داد. تاموکسیفن یکی از مثال های بسیاری است که پتانسیل هایی که در بهره گیری از منابع محاسباتی و صوری در وقت از دست دادن داده های اطراف ما وجود دارد را به ما نشان می دهد.

ما در حال حاضر در دوران کلان داده زندگی می کنیم. عبارت "کلان داده" در زمانی نزدیک به شروع دوره کلان داده ایجاد شد. همان قدر که من تاریخ شروع دوره کلان داده را سال ۲۰۰۱ در نظر می گیرم، این تاریخ منبع مناظرات و بحث های مهیج در وبلاگ ها و حتی نیویورک تایمز هم بوده است. به نظر می رسد که عبارت "کلان داده" به مفهوم امروزی آن برای اولین بار در اواخر دهه نود به کار گرفته شده است. اولین مقاله در این رابطه با عنوان "مدل های کلان داده

با فاکتور دینامیک برای سنجش اقتصاد کلان و پیش‌بینی^۱ در سال ۲۰۰۳ توسط فرانسیس اکس دایبولت^۲ منتشر شد، گرچه اعتبار این موضوع بیشتر به جان مشی^۳، دانشمند ارشد اس جی آی^۴، به‌عنوان اولین کسی که عبارت "کلان داده" را به کار برده است داده می‌شود. در اواخر دهه نود مشی برای گروه کوچکی از افراد در مورد موج جزر و مدی کلان داده که در حال وقوع بود یک مجموعه سخنرانی را ترتیب داد. دوره کلان داده دوره‌ای است که با حجم به سرعت در حال گسترش داده به‌صورتی بزرگ‌تر از آنچه بیشتر افراد حتی قادر به تصور وقوع آن نبودند تعریف می‌شود.

حجم بزرگ داده به تنهایی مشخص‌کننده دوره کلان داده نیست، چرا که همواره حجم‌هایی از داده وجه داشته‌اند که بیشتر از حد توان ما برای کار کردن به‌صورت موثر بر روی این داده‌های موجرت بودند. چیزی که این دوران داده را از دوره‌های دیگر متمایز می‌کند این است که شرکت‌ها، دولت‌ها، سازمان‌های غیر انتفاعی دستخوش تغییر در رفتار شده‌اند. در این دوره آن‌ها می‌خوانند، شروع به جمع‌آوری تمام داده‌های ممکن برای اهداف ناشناخته کنونی و یا آتی کنند تا بتوانند به‌کار خود را پیشرفت دهند. بسیاری بر این باورند که سازمان‌هایی که در تصمیم‌گیری‌هایشان داده به همراه پشتیبانی پژوهش و نمونه‌موردی‌های مربوط به آن استفاده می‌کنند در طولانی مدت تصمیمات بهتری می‌گیرند. تصمیماتی که به کسب‌وکاری قدرتمندتر و ماندگارتر منتهی خواهد شد. شرکت‌ها در پاسخ به رشد پر سرعت داده‌های تولید شده شروع به جمع‌آوری داده در بالاترین حد ممکن و ارزش دادن بیش از پیش به پتانسل‌های آن برای موارد آتی کرده‌اند. ما چه اندازه داده شخصی تولید می‌کنیم؟ اولین سوال این است: داده شخصی چیست؟ در سال ۱۹۹۵ اتحادیه اروپایی در حیطه قوانین اشخاص، داده شخصی را هرگونه اطلاعاتی عنوان کرد که می‌توانند کسی را به‌صورت مستقیم و یا غیر مستقیم از دیگران مشخص کنند. شرکت بین‌المللی داده تخمین زده است که ۲٫۸ زتابایت داده در سال ۲۰۱۲ ایجاد شده است و میزان داده تولید شده در هر سال ۲۰۱۵ دو برابر خواهد شد. با چنین عدد بزرگی درک این‌که چه میزان از این داده‌ها درواقع مربوط به شما هستند مشکل است. این عدد به چیزی در حدود ۵ گیگابایت در روز برای هر کارمند

¹ Big Data Dynamic Factor Models for Macroeconomic Measurement and Forecasting

² Francis X. Diebolt

³ John Mashey

⁴ SGI

متوسط ایالات متحده تقسیم می‌شود. این داده‌ها شامل ای میل، دانلود فیلم، فایل‌های شنیداری در جریان، صفحات گسترده اکسل و غیره می‌شود. این داده‌ها همچنین شامل داده‌هایی می‌شود که به عنوان اطلاعات در اینترنت ایجاد می‌شوند. بسیاری از این داده‌های تولید شده مستقیماً توسط شما یا من دیده نمی‌شوند بلکه در مورد ما ذخیره می‌شوند. طول چراغ‌های راهنمایی و رانندگی، جی پی اس که از طریق تلفن‌های ما مختصات را مشخص می‌کند یا میزان دریافت عوارضی زمانی که ما به سرعت از میان باجه‌های الکترونیکی دریافت عوارض می‌گذریم مثالی برای این نوع از داده‌ها هستند.

پیش از شروع دوران کلان داده کسب‌وکارها برای داده‌هایی که ارزش آئی نداشتند به نسبت ارزش کمی قائل بودند. وقتی دوران کلان داده آغاز شد سرمایه‌گذاری بر روی داده‌هایی که پتانسیل مفید واقع شدن در آینده را داشتند تغییر کرد و سازمان‌ها شروع به تلاش آگاهانه برای نگهداری داده‌های ارزشمند با پتانسیل مربوطه کردند. این تغییر رفتاری دایره بزرگی را ایجاد کرد که در آن داده ذخیره می‌شد و سپس به دلیل ارزشمند بودن داده اشخاص بر روی یافتن ارزش در داده برای سازمان سرمایه‌گذاری می‌کردند. موفقیت در یافتن ارزش در داده به جمع‌آوری بیشتر داده منتهی شد. بسیاری از این داده‌ها به جایی رهنمون نمی‌گشتند اما در بسیاری از موارد نتیجه ثابت کرد که هر چه داده بیشتری داشته باشیم امکان رسیدن به مطلوب برایمان بیشتر است. تغییر بزرگ دیگری که در ابتدای دوران کلان داده رخ داد رشد و تولید سریع و پختگی در تکنولوژی‌های ذخیره، اداره و تبادل این داده‌ها به صورت موثر بود.

حال که در دوران کلان داده هستیم چالش ما جمع‌آوری داده نیست بلکه استفاده از کامپیوترها برای افزودن به دامنه دانشمان و تشخیص الگوهایی است که پیش از این قادر به دیدن و یا یافتن آن‌ها نبوده‌ایم.

برخی از تکنولوژی‌های کلیدی و افت‌های پیش آمده در بازار منجر به افزایش جمع‌آوری، ذخیره و استفاده از داده در کارهای مربوط به تجزیه و تحلیل، به میزان بسیار زیاد شد. این مسئله ناشی از فاکتورهای زیادی از جمله پروتکل اینترنتی نسخه ۶، وسائل پیشرفته ارتباط از راه دور، تکنولوژی‌هایی مانند آر اف آی دی، حس‌گرهای کامپیوتری-مخابراتی، کاهش هزینه تولید هر واحد الکترونیکی، رسانه‌های گروهی و اینترنت است.

در ادامه جدول زمانی‌ای را مشاهده خواهید کرد که در آن برخی از اتفاقات کلیدی که به دوران کلان داده منجر شده‌اند و وقایعی که به شکل‌گیری و استفاده از کلان داده در آینده کمک کرده‌اند، مشخص شده‌اند.

جدول زمان‌بندی کلان داده:

در اینجا تعدادی از مواردی که نشان‌گر اتفاقات تأثیر گذار بر آماده شدن مسیر برای دوران کلان داده و نقاط عطف در طول این دوره هستند آورده شده‌اند.

۱۹۹۱

- اینترنت با ایجاد شبکه جهانی متولد شد. پروتکل انتقال ابرمتن (HTTP) به وسیله‌ای استاندارد برای به اشتراک گذاشتن اطلاعات در این رسانه نوین تبدیل شد.

۱۹۹۵

- سان^۱ پلت فورم جاوا^۲ را معرفی کرد، جاوا که در سال ۱۹۹۱ ابداع شد تبدیل به دومین زبان رایج پس از C شد. جاوا فضای برنامه‌های کاربردی وب غالب می‌شود و استاندارد بالفعال برای برنامه‌های کاربردی رده متوسط است. این برنامه‌ها منبعی برای ضبط و ذخیره شدآمد وب هستند.
- سیستم جهانی تعیین موقعیت (GPS) به صورت کامل عرضه شد. جی پی اس در اصل توسط دارپا (آژانس پروژه‌های پژوهشی پیشرفته) برای برنامه‌های کاربردی نظامی در ابتدای دهه ۷۰ تولید شد. این تکنولوژی سپس در برنامه‌های مربوط اتوموبیل، ناوبری هواپیما و یافتن آی فون حضور یافت.

۱۹۹۸

^۱ Sun

^۲ Java platform

^۳ DARPA (Defense Advanced Research Projects Agency)