



به نام خدا

علم داده در پایتون

مفاهیم و کاربردها

مؤلفان:

محسن یزدی نژاد (کاندیدای دکتری هوش مصنوعی)

محمدرضا پور (استادیار، عضو هیئت علمی وزارت علوم، تحقیقات و فناوری)

معصومه جوادی نیا (کارشناسی مهندسی نرم افزار)

فرهاد استوار نژاد (کارشناسی ارشد حسابداری)



هرگونه چاپ و تکثیر از محتویات این کتاب بدون اجازه کتبی ناشر ممنوع است. متخلفان به موجب قانون حمایت حقوق مؤلفان، مصنفان و هنرمندان تحت پیگرد قانونی قرار می گیرند.

◀ عنوان کتاب: علم داده در پایتون، مفاهیم و کاربردها

◀ مولفان: محسن یزدی نژاد - محمد رضاپور - معصومه جوادی نیا -

فرهاد استوارنژاد

◀ ناشر: موسسه فرهنگی هنری دیباگران تهران

◀ ویراستار: نرگس مهربد

◀ صفحه آرای: نازنین نصیری

◀ نوبت چاپ: اول

◀ تاریخ نشر: ۱۴۰۱

◀ چاپ و صحافی: صدف

◀ تیراژ: ۱۰۰ جلد

◀ قیمت: ۱۱۶۰۰۰۰ ریال

◀ شابک: ۹۷۸-۶۲۲-۲۱۸-۵۲۹-۹

نشانی واحد فروش: تهران، میدان انقلاب،

خ کارگر جنوبی، روبروی پاساژ مهستان،

پلاک ۱۲۵۱-تلفن: ۰۴۶۶۴۱۰۰۶-۲۲۰۸۵۱۱۱

فروشگاههای اینترنتی دیباگران تهران :

WWW.MFTBOOK.IR

www.dibagaran-tehran.com

عنوان و نام پدیدآور: علم داده در پایتون: مفاهیم و کاربردها/

مولفان: محسن یزدی نژاد و...دیگران؛

ویراستار: نرگس مهربد.

مشخصات نشر: تهران: دیباگران تهران: ۱۴۰۱

مشخصات ظاهری: ۲۰۰ ص: مصور، جدول

شابک: ۹۷۸-۶۲۲-۲۱۸-۵۲۹-۹

وضعیت فهرست نویسی: فیبا

یادداشت: مولفان: محسن یزدی نژاد، محمد رضاپور، معصومه جوادی نیا، فرهاد استوارنژاد.

موضوع: پایتون (زبان برنامه نویسی کامپیوتر)

موضوع: Python (computer program language)

موضوع: زبان شناسی کامپیوتری computational linguistics

موضوع: داده کاوی Data mining

شناسه افزوده: یزدی نژاد، محسن، ۱۳۶۲-

رده بندی کنگره: ۷۶/۷۳ QA

رده بندی دیویی: ۰۰۵/۱۳۳

شماره کتابشناسی ملی: ۸۸۰۵۲۴۶

نشانی اینستاگرام دیبا dibagaran_publishing نشانی تلگرام: @mftbook

هر کتاب دیباگران، یک فرصت جدید علمی و شغلی.

هر گوشی همراه، یک فروشگاه کتاب دیباگران تهران.

از طریق سایتهای دیباگران، در هر جای ایران به کتابهای ما دسترسی دارید.

۷..... مقدمه ناشر

۸..... مقدمه نویسندگان

۱۱..... فصل اول: مقدمه‌ای بر داده‌کاوی

۱۲..... ۱-۱- تعریف داده‌کاوی

۱۳..... ۲-۱- چرا داده‌کاوی؟

۱۴..... ۳-۱- هدف از داده‌کاوی

۱۴..... ۴-۱- جایگاه داده‌کاوی، هوش مصنوعی، یادگیری ماشین

۱۶..... ۵-۱- فرآیند داده‌کاوی CRISP

۱۸..... ۶-۱- وظایف داده‌کاوی

۲۰..... فصل دوم: پیش‌پردازش داده‌ها

۲۱..... ۱-۲- مفاهیم اولیه در داده‌کاوی

۲۲..... ۲-۲- کیفیت داده

۲۲..... ۱-۲-۲- نویزها و داده‌های دورافتاده

۲۳..... ۲-۲-۲- روش‌های پیدا کردن نقاط دورافتاده

۲۵..... ۳-۲-۲- مقادیر از دست‌رفته

۲۷..... ۴-۲-۲- داده تکراری

۲۷..... ۵-۲-۲- داده‌های جهت‌دار

۲۷..... ۳-۲-۳- آماده‌سازی داده‌ها

۲۸..... ۱-۳-۲- یکپارچه‌سازی

۲۸..... ۲-۳-۲- تجمیع

۲۸..... ۳-۳-۲- نمونه‌گیری

۳۰..... ۱-۳-۳-۲- نمونه‌گیری تصادفی ساده

۳۰..... ۲-۳-۳-۲- نمونه‌گیری سیستماتیک

۳۱..... ۳-۳-۳-۲- نمونه‌گیری چندبخشی (طبقه‌ای)

۳۲..... ۴-۳-۳-۲- نمونه‌گیری خوشه‌ای

۳۳..... ۴-۳-۲- کاهش بعد

۳۴..... ۱-۴-۳-۲- انتخاب ویژگی

۳۶..... ۲-۴-۳-۲- استخراج ویژگی

۳۷..... ۳-۴-۳-۲- خلق ویژگی

۳۸..... ۴-۴-۳-۲- گسسته‌سازی

۳۸..... ۵-۴-۳-۲- انتقال ویژگی

۴۱..... ۲-۳-۵- متوازن سازی داده ها

فصل سوم: شروع کدنویسی با پایتون ۴۴

۴۵..... ۳-۱- استفاده از محیط های موجود برای کدنویسی

۴۸..... ۳-۲- انواع داده ها

۵۴..... ۳-۳- ساختارهای کنترلی

۵۵..... ۳-۳-۱- جمله if

۵۷..... ۳-۳-۲- حلقه تکرار for

۵۸..... ۳-۳-۳- حلقه While

۵۹..... ۳-۳-۴- دستور Break

۵۹..... ۳-۴- توابع

۶۱..... ۳-۵- کتابخانه ها یا پکیج ها

۶۱..... ۳-۵-۱- کتابخانه numpy

فصل چهارم: روش های پیش پردازش داده ها ۶۵

۶۶..... ۴-۱- کتابخانه PANDAS

۶۶..... ۴-۱-۱- Series

۶۷..... ۴-۱-۲- DataFrame

۷۲..... ۴-۱-۳- فیلتر کردن

۷۴..... ۴-۱-۴- تولید ستون جدید در دیتافریم

۷۴..... ۴-۱-۵- حذف ستون ها و سطر ها در دیتافریم

۷۵..... ۴-۱-۶- مرتب سازی دیتافریم

۷۶..... ۴-۱-۷- تغییر نام ستون های دیتافریم

۷۶..... ۴-۱-۸- یکپارچه سازی داده ها توسط pandas

۷۸..... ۴-۱-۹- تجمیع داده ها در pandas

۷۸..... ۴-۱-۱۰- همبستگی بین متغیر ها

۷۹..... ۴-۲- نحوه شناسایی و مدیریت مقادیر از دست رفته

۸۲..... ۴-۳- مدیریت داده های تکراری

۸۲..... ۴-۴- مدیریت ستون های غیر عددی

۸۲..... ۴-۱-۴- Label Encoding

۸۳..... ۴-۲-۴- One hot Encoding

۸۴..... ۴-۵- نرمال سازی داده ها

فصل پنجم: مصور سازی داده ها در پایتون ۸۶

۸۷..... ۵-۱- کتابخانه MATPLOTLIB

۸۸..... ۵-۱-۱- نمودار خطی

۸۹..... ۵-۱-۲- نمایش داده دو نمودار در یک تصویر

۹۰	۳-۱-۵- نمودار میله‌ای عمودی
۹۱	۴-۱-۵- نمودار میله‌ای
۹۲	۵-۱-۵- نمودار دایره‌ای
۹۳	۶-۱-۵- نمودار پراکندگی
۹۴	۷-۱-۵- نمودار سطح انباشته
۹۴	۸-۱-۵- نمودار هیستوگرام
۹۵	۲-۵- کتابخانه SEARBORN
۹۶	۱-۲-۵- نمودار میله‌ای
۹۶	۲-۲-۵- نمودار پراکندگی
۹۷	۳-۲-۵- نمودار نقشه گرما

فصل ششم: روش‌های پیش‌بینی..... ۹۹

۱۰۰	۱-۶- دسته‌بندی الگوریتم‌های یادگیری ماشین
۱۰۱	۲-۶- انواع الگوریتم‌های پیش‌بینی
۱۰۱	۱-۲-۶- الگوریتم‌های طبقه‌بندی
۱۰۲	۱-۱-۲-۶- بیز ساده
۱۰۶	۲-۲-۶- الگوریتم درخت تصمیم
۱۰۹	۱-۲-۲-۶- انتخاب صفت تقسیم‌کننده
۱۱۰	۱-۱-۲-۲-۶- معیار Entropy
۱۱۳	۲-۱-۲-۱-۶- معیار Gini Index
۱۱۴	۲-۲-۲-۶- کم‌برازش و بیش‌برازش
۱۱۸	۳-۲-۶- الگوریتم K نزدیک‌ترین همسایگی
۱۲۰	۴-۲-۶- شبکه عصبی پرسپترون چندلایه
۱۲۶	۵-۲-۶- الگوریتم ماشین بردار پشتیبان
۱۲۸	۶-۲-۶- رگرسیون لجستیک
۱۳۰	۳-۶- الگوریتم‌های رگرسیون
۱۳۰	۱-۳-۶- الگوریتم نزدیک‌ترین همسایه
۱۳۰	۲-۳-۶- الگوریتم درخت تصمیم
۱۳۱	۳-۳-۶- الگوریتم شبکه عصبی
۱۳۱	۴-۳-۶- الگوریتم ماشین بردار پشتیبان
۱۳۱	۵-۳-۶- رگرسیون خطی

فصل هفتم: روش‌های توصیفی..... ۱۳۴

۱۳۵	۱-۷- خوشه‌بندی
۱۳۶	۱-۱-۷- روش k-means
۱۴۰	۲-۱-۷- خوشه‌بندی مشتریان به روش RFM
۱۴۴	۳-۱-۷- روش DBSCAN

۱۵۰ ۲-۷- الگوهای مکرر و قواعد انجمنی

فصل هشتم: ارزیابی ۱۵۵

۱۵۶ ۱-۸- ارزیابی روش‌های پیش‌بینی
۱۵۶ ۱-۱-۸- معیارهای ارزیابی در مدل‌های طبقه‌بندی
۱۵۸ ۲-۱-۸- معیارهای ارزیابی در مدل‌های رگرسیون
۱۶۱ ۲-۸- استراتژی‌های ارزیابی
۱۶۱ ۱-۲-۸- روش train test split
۱۶۱ ۲-۲-۸- روش K-Fold Cross Validation
۱۶۳ ۳-۲-۸- روش Leave one out
۱۶۳ ۴-۲-۸- روش stratified k-fold cross validation
۱۶۴ ۵-۲-۸- روش group k-fold cross validation
۱۶۴ ۳-۸- شناسایی مدل و پارامترهای بهینه

فصل نهم: روش‌های ترکیبی ۱۶۶

۱۶۸ ۱-۹- تکنیک‌های ترکیبی پایه
۱۶۸ ۱-۱-۹- رأی‌گیری
۱۶۹ ۲-۱-۹- میانگین‌گیری
۱۷۰ ۲-۹- تکنیک‌های ترکیبی پیشرفته
۱۷۰ ۱-۲-۹- Bagging
۱۷۰ ۲-۲-۹- Random Forest
۱۷۱ ۳-۲-۹- Stacking
۱۷۲ ۳-۹- تکنیک‌های پیشرفته بر پایه BOOSTING
۱۷۲ ۱-۳-۹- AdaBoost
۱۷۳ ۲-۳-۹- Gradient Boosting Machine
۱۷۵ ۳-۳-۹- Xgboost
۱۷۶ ۴-۳-۹- lightgbm
۱۷۷ ۵-۳-۹- Catboost

فصل دهم: مهندسی ویژگی‌ها ۱۷۸

۱۷۹ ۱-۱۰- مهندسی ویژگی‌ها چیست؟
۱۸۲ ۲-۱۰- رابطه متقابل ویژگی‌ها با متغیر هدف
۱۸۵ ۳-۱۰- خلق ویژگی
۱۸۸ ۴-۱۰- خوشه‌بندی با الگوریتم K-MEANS
۱۸۹ ۵-۱۰- خلق ویژگی با PCA
۱۹۳ ۶-۱۰- کدگذاری هدف
۱۹۵ ۷-۱۰- انتخاب زیرمجموعه‌ای از ویژگی

خط‌مشی انتشارات مؤسسه فرهنگی هنری دیباگران تهران در عرصه کتاب‌هایی با کیفیت عالی است که بتواند
خواسته‌های به‌روز جامعه فرهنگی و علمی کشور را تا حد امکان پوشش دهد.
هر کتاب دیباگران تهران، یک فرصت جدید شغلی و علمی

حمد و سپاس ایزد منان را که با الطاف بی‌کران خود این توفیق را به ما ارزانی داشت تا بتوانیم در راه ارتقای دانش عمومی و فرهنگی این مرز و بوم در زمینه چاپ و نشر کتب علمی و آموزشی گام‌هایی هرچند کوچک برداشته و در انجام رسالتی که بر عهده داریم، مؤثر واقع شویم.

گسترده‌گی علوم و سرعت توسعه روزافزون آن، شرایطی را به وجود آورده که هر روز شاهد تحولات اساسی چشمگیری در سطح جهان هستیم. این گسترش و توسعه، نیاز به منابع مختلف از جمله کتاب را به عنوان قدیمی‌ترین و راحت‌ترین راه دستیابی به اطلاعات و اطلاع‌رسانی، بیش از پیش برجسته نموده است.

در این راستا، واحد انتشارات مؤسسه فرهنگی هنری دیباگران تهران با همکاری اساتید، مؤلفان، مترجمان، متخصصان، پژوهشگران و محققان در زمینه‌های گوناگون و مورد نیاز جامعه تلاش نموده برای رفع کمبودها و نیازهای موجود، منابعی پُر بار، معتبر و با کیفیت مناسب در اختیار علاقمندان قرار دهد.

کتابی که در دست دارید تألیف " دکتر محسن یزدی نژاد، دکتر محمدرضا پور، مهندس معصومه جوادی نیا، مهندس فرهاد استوارنژاد" است که با تلاش همکاران ما در نشر دیباگران تهران منتشر گشته و شایسته است از یکایک این گرامیان تشکر و قدردانی کنیم.

با نظرات خود مشوق و راهنمای ما باشید

با ارائه نظرات و پیشنهادات و خواسته‌های خود، به ما کمک کنید تا بهتر و دقیق‌تر در جهت رفع نیازهای علمی و آموزشی کشورمان قدم برداریم. برای رساندن پیام‌هایتان به ما از رسانه‌های دیباگران تهران شامل سایتهای فروشگاهی و صفحه اینستاگرام و شماره‌های تماس که در صفحه شناسنامه کتاب آمده استفاده نمایید.

مدیر انتشارات

مؤسسه فرهنگی هنری دیباگران تهران
dibagaran@mftplus.com

مقدمه نویسندگان

جایگاه «هوش مصنوعی» در دنیا از سطح متخصصین و مهندسان فراتر رفته و مورد توجه رهبران برجسته‌ی جهان قرار گرفته است. این که در کشور ما در چه وضعیتی قرار داریم مشخص است که از نظر «نیروی انسانی» توانمند چیزی کم از کشورهای پیشرفته در این عرصه کم نداریم؛ همچنین علاوه بر جوانان پای کار و پرتلاش، اکنون می‌توان در اغلب فراخوان‌های مجلات علمی و همایش‌ها دید که در هر رشته‌ای که باشند جویای کاربرد هوش مصنوعی در آن حوزه بوده و خیلی مشتاق‌اند که بدانند در زمینه‌ی تخصصی‌شان چه توصیفی وجود دارد و اگر هم بشود از وضعیت آینده‌شان باخبر شوند.

در طرح‌های پژوهشی که از سوی سازمان‌ها مورد حمایت قرار می‌گیرند نیز محورهایی از این زیرشاخه‌های هوش مصنوعی برجسته می‌شوند و بنابراین استقبال دانشگاه و صنعت از هوش مصنوعی قابل ملاحظه می‌باشد. هر چند سهم کمتر از یک درصدی شرکت‌های دانش‌بنیان از تولید ناخالص داخلی، مطلوب نبوده و رهبر عزیز^(حفظه الله تعالی) نیز تأکید کرده‌اند: «این سهم باید حتماً افزایش یابد و در طول ۳ تا ۴ سال، حداقل به ۵ درصد برسد.»^۱

تا اینجا بحث، به دو عامل «حمایت حاکمیتی» و «نیروی انسانی» اشاره کردیم اما محقق شدن این چشم‌انداز آیا تا همین جای کار مسلّم خواهد بود و به عامل زمینه‌ساز دیگری نیاز ندارد؟ باید خیلی صریح پاسخ بگوئیم که: خیر! اجازه دهید علت را با مثالی به عرض شما خوانندگان عزیز برسانیم:

در سال ۱۹۹۸ میلادی دو دانشجوی دکترای دانشگاه استنفورد، ابزاری برای وبگردی ساختند که کم‌کم نام این ابزار شد Google! مایکروسافت هم در فضای رقابتی موجود، بدش نمی‌آمد گوی و میدان را از گوگل برآید و مانند سیستم‌عامل کدبسته‌ی ویندوز، محصول خودش در این عرصه را نیز به جهانیان دیکته کند! بنابراین، پس از یک سرمایه‌گذاری عظیم و البته کارهای زرنگ‌مآبانه و عادت غربی‌ها در جذب نخبگان، و این بار با جذب نخبگانی از تیم فنی گوگل توسط مایکروسافت، بالاخره موتور جستجوی مایکروسافتی با نام Bing در سال ۲۰۰۹، یعنی چیزی در حدود ۱۳ سال بعد از گوگل، رونمایی شد.

اما این بار دیگر عوامل غیرمعمول – شاید شانس و شاید حمایت‌های ویژه‌ی دولتی – منجر به پیشتازی بیل گیتس نشد و موتور بینگ نتوانست در زمره‌ی برگ زرین دیگری در خاطرات این مرد پولدار نقش ببندد و با حضور در بین مهندسين، خیلی سربلند بگوید که من اخراجی دانشگاهم ولی از همه شما موفق‌تر و در واقع همان پولدارترم!

«گیتس» هر چه کرد نشد که نشد و طبق آخرین برآوردها، بیش از ۸۰٪ بازار موتورهای جستجو همچنان به گوگل اختصاص یافت و سهم بینگ در خوشبینانه‌ترین حالت به ۵٪ هم نمی‌رسد! این بار نه سرمایه‌گذاری

^۱ دیدار نخبگان و استعدادهای برتر علمی با رهبر انقلاب، ۲۶ آذر ۱۴۰۰

مایکروسافت از گوگل کمتر بود و نه تیم فنی ضعیف‌تری داشت زیرا اگر هر کدام از این علّت‌ها باعث این ناکامی شده بود، قطعاً «بیل» وارد چنین رقابتی نمی‌شد.

البته در این زمینه چندان نمی‌توان بیل‌گیتس را سرزنش کرد زیرا علت این اتفاق در آن زمان قابل پیش‌بینی نبود، و شاید تصور چنین بوده که ۱۳ سال تأخیر از گوگل را با سرمایه‌گذاری بیشتر بتوان جبران نمود. اما در این حوزه، یعنی «هوش مصنوعی» و هسته‌ی آن «یادگیری ماشین»، داده‌ها اهمیت بالایی دارند و به‌ویژه، جمع‌آوری رفتارهای کاربران توسط الگوریتم‌های هوشمند جستجو، خزیدن (کروالینگ) در بین وبلاگ‌ها، اطلاع از نیازهای مشتریان و نهایتاً برچسب زنی دادگان و وجود «داده‌های برچسب‌خورده»، نقش تعیین‌کننده‌ای در «پایدار ماندن مزیت رقابتی» یک فناوری دارند که همه‌ی اینها را مدت‌ها قبل گوگل شروع کرده بود و آن‌قدر فاصله گرفته بود که مایکروسافت هم نتوانست از وی سبقت بگیرد.

البته فناوری‌های از این نوع، یک ویژگی غیرملموس هم دارند و آن اینکه علاوه بر «خود ساختن» که در پاراگراف قبل اشاره شد، از «تکمیل‌شدنی توسط کاربران» نیز برخوردارند؛ یعنی شما وقتی از یک محصول خودرویی استفاده می‌کنید، به‌عنوان یک مصرف‌کننده نقشی در آموزش دادن خودرو از رفتار تان و از انتظاری که از خودرو دارید نخواهید داشت؛ به بیان ملموس‌تر، انسان‌ها در تعامل با همدیگر و در زندگی، همچنان که از خدمات هم بهره‌مند می‌شوند، به مرور باعث تکامل طرف مقابل نیز بایستی بشوند (البته وقتی با دوست خوب مواجهیم!!) و سامانه‌های مبتنی بر «هوش مصنوعی» در این ویژگی‌هاست که به انسان‌ها شبیه می‌شوند، زیرا این‌گونه سامانه‌ها اساساً حیات‌شان وابسته به «آموزش دیدن» و «تکمیل شدن مداوم» است؛ آری، گوگل این ثانیه، با گوگل ثانیه‌ی قبل فرق دارد، و هر ثانیه‌اش آموزش‌دیده‌تر از ثانیه‌ی قبلش شده و این روند ۱۳ سال قبل از بینگ شروع شده بود و کاربران هم در آن ۱۳ سال به گوگل خودشان کمک کرده بودند تا بفهمد از وی چه می‌خواهند و در ادامه نیز وقتی به نیاز همدیگر پی بُردند، طبیعی بود که از وی به ابزار دیگری که تازه راه رفتن را داشت یاد می‌گرفت، روی نیاورند. در واقع، کودک نحیفِ بینگ، حتماً نمی‌توانست جای نوجوان فربه‌ی گوگل را در این کارزار بستاند هر چند که مایکروسافت، لباس‌های قشنگ و گران بر تنش بکشانند!

از مثال فوق – بدون آنکه خواسته باشیم روی لوگوی خاصی تأکید کرده باشیم – بگذریم و به عامل سوم موفقیت هوش مصنوعی پردازیم و آن هم وابستگی شدیدش به «داده‌های تمیز و برچسب‌خورده‌ای» است که کاربران نیز نقش مؤثری در این آماده‌سازی داده‌ها ایفا می‌کنند.

وضعیت فعلی در کشور، که نمی‌توان اهمیت یافتن «هوش مصنوعی» را در آن کتمان کرد، مرهون تأکید رهبر عزیز می‌باشد که در فراراهبرد مد نظر ایشان «این مسئله در اداره آینده دنیا نقش دارد و باید به‌گونه‌ای عمل کنیم که ایران جزو ۱۰ کشور برتر هوش مصنوعی در دنیا قرار بگیرد.»^۱

^۱ همان مأخذ قبل

بنابراین کتاب حاضر در راستای چند کتابی که در عرصه‌ی هوش مصنوعی با عنوان «هم‌اکنون، داده‌کاوی»، «ارائه راهکار برای تمام مسائل یادگیری ماشین» و «آمار کاربردی در علم داده» ارائه کرده‌ایم، به مبحث جدی و سازنده‌ی «علم داده‌ها» می‌پردازد و ابزار انتخاب شده هم پایتون است زیرا هم کدباز است هم یادگیری آن برای کاربران با سطوح مختلف قابل شروع کردن است. کدباز بودن برنامه فقط منجر به راحتی و آماده‌بودن قطعه‌کدها توسط دیگر برنامه‌نویسان نبوده، بلکه علاوه بر این هم‌افزایی، قابلیت این نرم‌افزار و چند هزار کتابخانه‌ی آن در پردازش انواع مختلف داده‌ها هم مطلب دیگر برای انگیزه بهره‌برداری از آن بوده و به‌علاوه، رعایت حق کپی مسأله‌ای است که شما اگر مقاله‌ای چاپ کنید و به نرم‌افزاری تجاری برای تحلیل داده‌هایتان اشاره کنید که کدباز نباشد، حتما لازم است اجازه‌ی شرکت مربوطه و مسائل مالی آنرا اثبات نمائید که در پایتون با چنین مشکلی هم مواجه نخواهید بود. نکته‌ی پایانی اینکه کدهای کتاب در آدرس مندرج در آخرین منبع کتاب نیز در دسترس عزیزان علاقمند قرار داده شده است.

متن کتاب حاضر در ایام بزرگداشت حکیم نظامی تفرشی به اتمام رسید و دیباچه را نیز با بیتی از آن حکیم پایان می‌بریم:

نظامی را ره تحقیق بنمای

خداوند او را توفیق بگشای

و در این حد اشاره نماییم که «تحقیق» به دنبال حقیقت بودن است، و علم داده‌ها نیز به دنبال چنین معرفتی است...

فصل اول: مقدمه‌ای بر داده‌کاوی



python