

Data

Science Data Book

The first and only book to systematically address methodologies and processes of leveraging non-traditional information sources in the context of investing and risk management Harnessing non-traditional data sources to generate alpha, analyze markets, and forecast risk is a subject of intense interest for financial professionals. A growing number of regularly-held conferences on alternative data are being established, complemented by an upsurge in new papers on the subject. Alternative data is starting to be steadily incorporated by conventional institutional investors and risk managers throughout the financial world. Methodologies to analyze and extract value from alternative data, guidance on how to source data and integrate data flows within existing systems is currently not treated in literature. Filling this significant gap in knowledge, The Book of Alternative Data is the first and only book to offer a coherent, systematic treatment of the subject. This groundbreaking volume provides readers with a roadmap for navigating the complexities of an array of alternative data sources, and delivers the appropriate techniques to analyze them. The authors—leading experts in financial modeling, machine learning, and quantitative research and analytics—employ a step-by-step approach to guide readers through the dense jungle of generated data. A first-of-its kind treatment of alternative data types, sources, and methodologies, this innovative book: Provides an integrated modeling approach to extract value from multiple types of datasets Treats the processes needed to make alternative data signals operational Helps investors and risk managers rethink how they engage with alternative datasets Features practical use case studies in many different financial markets and real-world techniques Describes how to avoid potential pitfalls and missteps in starting the alternative data journey Explains how to integrate information from different datasets to maximize informational value The Book of Alternative Data is an indispensable resource for anyone wishing to analyze or monetize different non-traditional datasets, including Chief Investment Officers, Chief Risk Officers, risk professionals, investment professionals, traders, economists, and machine learning developers and users.

The Book of Alternative Data

Explore Kinesis managed services such as Kinesis Data Streams, Kinesis Data Analytics, Kinesis Data Firehose, and Kinesis Video

Streams with the help of practical use cases Key FeaturesGet well versed with the capabilities of Amazon KinesisExplore the monitoring, scaling, security, and deployment patterns of various Amazon Kinesis servicesLearn how other Amazon Web Services and third-party applications such as Splunk can be used as destinations for Kinesis dataBook Description Amazon Kinesis is a collection of secure, serverless, durable, and highly available purpose-built data streaming services. This data streaming service provides APIs and client SDKs that enable you to produce and consume data at scale. Scalable Data Streaming with Amazon Kinesis begins with a quick overview of the core concepts of data streams, along with the essentials of the AWS Kinesis landscape. You'll then explore the requirements of the use case shown through the book to help you get started and cover the key pain points encountered in the data stream life cycle. As you advance, you'll get to grips with the architectural components of Kinesis, understand how they are configured to build data pipelines, and delve into the applications that connect to them for consumption and processing. You'll also build a Kinesis data pipeline from scratch and learn how to implement and apply practical solutions. Moving on, you'll learn how to configure Kinesis on a cloud platform. Finally, you'll learn how other AWS services can be integrated into Kinesis. These services include Redshift, Dynamo Database, AWS S3, Elastic Search, and third-party applications such as Splunk. By the end of this AWS book, you'll be able to build and deploy your own Kinesis data pipelines with Kinesis Data Streams (KDS), Kinesis Data Firehose (KFH), Kinesis Video Streams (KVS), and Kinesis Data Analytics (KDA). What you will learnGet to grips with data streams, decoupled design, and real-time stream processingUnderstand the properties of KFH that differentiate it from other Kinesis servicesMonitor and scale KDS using CloudWatch metricsSecure KDA with identity and access management (IAM)Deploy KVS as infrastructure as code (IaC)Integrate services such as Redshift, Dynamo Database, and Splunk into KinesisWho this book is for This book is for solutions architects, developers, system administrators, data engineers, and data scientists looking to evaluate and choose the most performant, secure, scalable, and cost-effective data streaming technology to overcome their data ingestion and processing challenges on AWS. Prior knowledge of cloud architectures on AWS, data streaming technologies, and architectures is expected.

Scalable Data Streaming with Amazon Kinesis

Data science libraries, frameworks, modules, and toolkits are great for doing data science, but they're also a good way to dive into the discipline without actually understanding data science. With this updated second edition, you'll learn how many of the most fundamental data science tools and algorithms work by implementing them from scratch. If you have an aptitude for mathematics and some programming skills, author Joel Grus will help you get comfortable with the math and statistics at the core of data science, and with hacking skills you need to get started as a data scientist. Today's messy glut of data holds answers to questions no one's even thought to ask. This book provides you with the know-how to dig those answers out.

Data Science from Scratch

Learn how to use R to turn raw data into insight, knowledge, and understanding. This book introduces you to R, RStudio, and the tidyverse, a collection of R packages designed to work together to make data science fast, fluent, and fun. Suitable for readers with no previous programming experience, R for Data Science is designed to get you doing data science as quickly as possible. Authors Hadley Wickham and Garrett Grolemund guide you through the steps of importing, wrangling, exploring, and modeling your data and communicating the results. You'll get a complete, big-picture understanding of the data science cycle, along with basic tools you need to manage the details. Each section of the book is paired with exercises to help you practice what you've learned along the way. You'll learn how to: Wrangle—transform your datasets into a form convenient for analysis Program—learn powerful R tools for solving data problems with greater clarity and ease Explore—examine your data, generate hypotheses, and quickly test them Model—provide a low-dimensional summary that captures true "signals" in your dataset Communicate—learn R Markdown for integrating prose, code, and results

R for Data Science

A new way of thinking about data science and data ethics that is informed by the ideas of intersectional feminism. Today, data science is a form of power. It has been used to expose injustice, improve health outcomes, and topple governments. But it has also been used to discriminate, police, and surveil. This potential for good, on the one hand, and harm, on the other, makes it essential to ask: Data science by whom? Data science for whom? Data science with whose interests in mind? The narratives around big data and data science are overwhelmingly white, male, and techno-heroic. In *Data Feminism*, Catherine D'Ignazio and Lauren Klein present a new way of thinking about data science and data ethics—one that is informed by intersectional feminist thought. Illustrating data feminism in action, D'Ignazio and Klein show how challenges to the male/female binary can help challenge other hierarchical (and empirically wrong) classification systems. They explain how, for example, an understanding of emotion can expand our ideas about effective data visualization, and how the concept of invisible labor can expose the significant human efforts required by our automated systems. And they show why the data never, ever "speak for themselves." *Data Feminism* offers strategies for data scientists seeking to learn how feminism can help them work toward justice, and for feminists who want to focus their efforts on the growing field of data science. But *Data Feminism* is about much more than gender. It is about power, about who has it and who doesn't, and about how those differentials of power can be challenged and changed.

Data Feminism

Divided into 22 sections, this pocket-sized volume is an exhaustive 'quick reference' of up-to-date engineering data and rules. Contents: Essential Mathematics; Units; Engineering design Processes and Principles; Basic Mechanical Design; Motion; Mechanics of Materials; Material Failure; Thermodynamics; Fluid Mechanisms; Fluid Equipment; Pressure Vessels; Materials; Machine Elements; Design and Production Tools; Project Engineering; Computer-Aided Engineering; Welding; Non-Destructive Examination; Corrosion; Surface Protection; Metallurgical Terms; Engineering Associations and Organizations.

IMechE Engineers' Data Book

This IBM® Redbooks® publication describes how the IBM Big Data Platform provides the integrated capabilities that are required for the adoption of Information Governance in the big data landscape. As organizations embark on new use cases, such as Big Data Exploration, an enhanced 360 view of customers, or Data Warehouse modernization, and absorb ever growing volumes and variety of data with accelerating velocity, the principles and practices of Information Governance become ever more critical to ensure trust in data and help organizations overcome the inherent risks and achieve the wanted value. The introduction of big data changes the information landscape. Data arrives faster than humans can react to it, and issues can quickly escalate into significant events. The variety of data now poses new privacy and security risks. The high volume of information in all places makes it harder to find where these issues, risks, and even useful information to drive new value and revenue are. Information Governance provides an organization with a framework that can align their wanted outcomes with their strategic management principles, the people who can implement those principles, and the architecture and platform that are needed to support the big data use cases. The IBM Big Data Platform, coupled with a framework for Information Governance, provides an approach to build, manage, and gain significant value from the big data landscape.

Information Governance Principles and Practices for a Big Data Landscape

DW 2.0: The Architecture for the Next Generation of Data Warehousing is the first book on the new generation of data warehouse architecture, DW 2.0, by the father of the data warehouse. The book describes the future of data warehousing that is technologically possible today, at both an architectural level and technology level. The perspective of the book is from the top down: looking at the overall architecture and then delving into the issues underlying the components. This allows people who are building or using a data warehouse to see what lies ahead and determine what new technology to buy, how to plan extensions to the data warehouse, what can

be salvaged from the current system, and how to justify the expense at the most practical level. This book gives experienced data warehouse professionals everything they need in order to implement the new generation DW 2.0. It is designed for professionals in the IT organization, including data architects, DBAs, systems design and development professionals, as well as data warehouse and knowledge management professionals. * First book on the new generation of data warehouse architecture, DW 2.0. * Written by the \"father of the data warehouse\"

DW 2.0: The Architecture for the Next Generation of Data Warehousing

The Handbook of Massive Data Sets is comprised of articles written by experts on selected topics that deal with some major aspect of massive data sets. It contains chapters on information retrieval both in the internet and in the traditional sense, web crawlers, massive graphs, string processing, data compression, clustering methods, wavelets, optimization, external memory algorithms and data structures, the US national cluster project, high performance computing, data warehouses, data cubes, semi-structured data, data squashing, data quality, billing in the large, fraud detection, and data processing in astrophysics, air pollution, biomolecular data, earth observation and the environment. The proliferation of massive data sets brings with it a series of special computational challenges. This \"data avalanche\" arises in a wide range of scientific and commercial applications.

Handbook of Massive Data Sets

The book is devoted to the analysis of big data in order to extract from these data hidden patterns necessary for making decisions about the rational behavior of complex systems with the different nature that generate this data. To solve these problems, a group of new methods and tools is used, based on the self-organization of computational processes, the use of crisp and fuzzy cluster analysis methods, hybrid neural-fuzzy networks, and others. The book solves various practical problems. In particular, for the tasks of 3D image recognition and automatic speech recognition large-scale neural networks with applications for Deep Learning systems were used. Application of hybrid neuro-fuzzy networks for analyzing stock markets was presented. The analysis of big historical, economic and physical data revealed the hidden Fibonacci pattern about the course of systemic world conflicts and their connection with the Kondratieff big economic cycles and the Schwabe-Wolf solar activity cycles. The book is useful for system analysts and practitioners working with complex systems in various spheres of human activity.

Big Data: Conceptual Analysis and Applications

Learn data science concepts with real-world examples in SAS! End-to-End Data Science with SAS: A Hands-On Programming Guide provides clear and practical explanations of the data science environment, machine learning techniques, and the SAS programming knowledge necessary to develop machine learning models in any industry. The book covers concepts including understanding the business need, creating a modeling data set, linear regression, parametric classification models, and non-parametric classification models. Real-world business examples and example code are used to demonstrate each process step-by-step. Although a significant amount of background information and supporting mathematics are presented, the book is not structured as a textbook, but rather it is a user's guide for the application of data science and machine learning in a business environment. Readers will learn how to think like a data scientist, wrangle messy data, choose a model, and evaluate the model's effectiveness. New data scientists or professionals who want more experience with SAS will find this book to be an invaluable reference. Take your data science career to the next level by mastering SAS programming for machine learning models.

End-to-End Data Science with SAS

Data profiling refers to the activity of collecting data about data, i.e., metadata. Most IT professionals and researchers who work with data have engaged in data profiling, at least informally, to understand and explore an unfamiliar dataset or to determine whether a new dataset is appropriate for a particular task at hand. Data profiling results are also important in a variety of other situations, including query optimization, data integration, and data cleaning. Simple metadata are statistics, such as the number of rows and columns, schema and datatype information, the number of distinct values, statistical value distributions, and the number of null or empty values in each column. More complex types of metadata are statements about multiple columns and their correlation, such as candidate keys, functional dependencies, and other types of dependencies. This book provides a classification of the various types of profilable metadata, discusses popular data profiling tasks, and surveys state-of-the-art profiling algorithms. While most of the book focuses on tasks and algorithms for relational data profiling, we also briefly discuss systems and techniques for profiling non-relational data such as graphs and text. We conclude with a discussion of data profiling challenges and directions for future work in this area.

Data Profiling

This thoroughly revised guide demonstrates how the flexibility of the command line can help you become a more efficient and

productive data scientist. You'll learn how to combine small yet powerful command-line tools to quickly obtain, scrub, explore, and model your data. To get you started, author Jeroen Janssens provides a Docker image packed with over 100 Unix power tools--useful whether you work with Windows, macOS, or Linux. You'll quickly discover why the command line is an agile, scalable, and extensible technology. Even if you're comfortable processing data with Python or R, you'll learn how to greatly improve your data science workflow by leveraging the command line's power. This book is ideal for data scientists, analysts, engineers, system administrators, and researchers. Obtain data from websites, APIs, databases, and spreadsheets Perform scrub operations on text, CSV, HTML, XML, and JSON files Explore data, compute descriptive statistics, and create visualizations Manage your data science workflow Create your own tools from one-liners and existing Python or R code Parallelize and distribute data-intensive pipelines Model data with dimensionality reduction, regression, and classification algorithms Leverage the command line from Python, Jupyter, R, RStudio, and Apache Spark

Data Science at the Command Line

Praise for the First Edition “...a well-written book on data analysis and data mining that provides an excellent foundation...” —CHOICE
“This is a must-read book for learning practical statistics and data analysis...” —Computing Reviews.com A proven go-to guide for data analysis, *Making Sense of Data I: A Practical Guide to Exploratory Data Analysis and Data Mining*, Second Edition focuses on basic data analysis approaches that are necessary to make timely and accurate decisions in a diverse range of projects. Based on the authors’ practical experience in implementing data analysis and data mining, the new edition provides clear explanations that guide readers from almost every field of study. In order to facilitate the needed steps when handling a data analysis or data mining project, a step-by-step approach aids professionals in carefully analyzing data and implementing results, leading to the development of smarter business decisions. The tools to summarize and interpret data in order to master data analysis are integrated throughout, and the Second Edition also features: Updated exercises for both manual and computer-aided implementation with accompanying worked examples New appendices with coverage on the freely available Traceis™ software, including tutorials using data from a variety of disciplines such as the social sciences, engineering, and finance New topical coverage on multiple linear regression and logistic regression to provide a range of widely used and transparent approaches Additional real-world examples of data preparation to establish a practical background for making decisions from data *Making Sense of Data I: A Practical Guide to Exploratory Data Analysis and Data Mining*, Second Edition is an excellent reference for researchers and professionals who need to achieve effective decision making from data. The Second Edition is also an ideal textbook for undergraduate and graduate-level courses in data analysis and data mining and is appropriate for cross-disciplinary courses found within computer science and engineering departments.

Day-to-day Data

Learn, by example, the fundamentals of data analysis as well as several intermediate to advanced methods and techniques ranging from classification and regression to Bayesian methods and MCMC, which can be put to immediate use. Key Features Analyze your data using R – the most powerful statistical programming language Learn how to implement applied statistics using practical use-cases Use popular R packages to work with unstructured and structured data Book Description Frequently the tool of choice for academics, R has spread deep into the private sector and can be found in the production pipelines at some of the most advanced and successful enterprises. The power and domain-specificity of R allows the user to express complex analytics easily, quickly, and succinctly. Starting with the basics of R and statistical reasoning, this book dives into advanced predictive analytics, showing how to apply those techniques to real-world data though with real-world examples. Packed with engaging problems and exercises, this book begins with a review of R and its syntax with packages like Rcpp, ggplot2, and dplyr. From there, get to grips with the fundamentals of applied statistics and build on this knowledge to perform sophisticated and powerful analytics. Solve the difficulties relating to performing data analysis in practice and find solutions to working with messy data, large data, communicating results, and facilitating reproducibility. This book is engineered to be an invaluable resource through many stages of anyone's career as a data analyst. What you will learn Gain a thorough understanding of statistical reasoning and sampling theory Employ hypothesis testing to draw inferences from your data Learn Bayesian methods for estimating parameters Train regression, classification, and time series models Handle missing data gracefully using multiple imputation Identify and manage problematic data points Learn how to scale your analyses to larger data with Rcpp, data.table, dplyr, and parallelization Put best practices into effect to make your job easier and facilitate reproducibility Who this book is for Budding data scientists and data analysts who are new to the concept of data analysis, or who want to build efficient analytical models in R will find this book to be useful. No prior exposure to data analysis is needed, although a fundamental understanding of the R programming language is required to get the best out of this book.

Making Sense of Data I

Large Scale and Big Data: Processing and Management provides readers with a central source of reference on the data management techniques currently available for large-scale data processing. Presenting chapters written by leading researchers, academics, and practitioners, it addresses the fundamental challenges associated with Big Data processing tools and techniques across a range of computing environments. The book begins by discussing the basic concepts and tools of large-scale Big Data processing and cloud computing. It also provides an overview of different programming models and cloud-based deployment models. The book's second section examines the usage of advanced Big Data processing techniques in different domains, including semantic web, graph

processing, and stream processing. The third section discusses advanced topics of Big Data processing such as consistency management, privacy, and security. Supplying a comprehensive summary from both the research and applied perspectives, the book covers recent research discoveries and applications, making it an ideal reference for a wide range of audiences, including researchers and academics working on databases, data mining, and web scale data processing. After reading this book, you will gain a fundamental understanding of how to use Big Data-processing tools and techniques effectively across application domains. Coverage includes cloud data management architectures, big data analytics visualization, data management, analytics for vast amounts of unstructured data, clustering, classification, link analysis of big data, scalable data mining, and machine learning techniques.

Data Analysis with R, Second Edition

Deep Learning with Structured Data teaches you powerful data analysis techniques for tabular data and relational databases. Summary Deep learning offers the potential to identify complex patterns and relationships hidden in data of all sorts. Deep Learning with Structured Data shows you how to apply powerful deep learning analysis techniques to the kind of structured, tabular data you'll find in the relational databases that real-world businesses depend on. Filled with practical, relevant applications, this book teaches you how deep learning can augment your existing machine learning and business intelligence systems. Purchase of the print book includes a free eBook in PDF, Kindle, and ePub formats from Manning Publications. About the technology Here's a dirty secret: Half of the time in most data science projects is spent cleaning and preparing data. But there's a better way: Deep learning techniques optimized for tabular data and relational databases deliver insights and analysis without requiring intense feature engineering. Learn the skills to unlock deep learning performance with much less data filtering, validating, and scrubbing. About the book Deep Learning with Structured Data teaches you powerful data analysis techniques for tabular data and relational databases. Get started using a dataset based on the Toronto transit system. As you work through the book, you'll learn how easy it is to set up tabular data for deep learning, while solving crucial production concerns like deployment and performance monitoring. What's inside When and where to use deep learning The architecture of a Keras deep learning model Training, deploying, and maintaining models Measuring performance About the reader For readers with intermediate Python and machine learning skills. About the author Mark Ryan is a Data Science Manager at Intact Insurance. He holds a Master's degree in Computer Science from the University of Toronto. Table of Contents 1 Why deep learning with structured data? 2 Introduction to the example problem and Pandas dataframes 3 Preparing the data, part 1: Exploring and cleansing the data 4 Preparing the data, part 2: Transforming the data 5 Preparing and building the model 6 Training the model and running experiments 7 More experiments with the trained model 8 Deploying the model 9 Recommended next steps

Large Scale and Big Data

Providing a comprehensive and authoritative summary of all the available facts and figures relating to World War II, this text is divided into nine sections for ease of reference.

Deep Learning with Structured Data

Although there are many books on mathematical finance, few deal with the statistical aspects of modern data analysis as applied to financial problems. This textbook fills this gap by addressing some of the most challenging issues facing financial engineers. It shows how sophisticated mathematics and modern statistical techniques can be used in the solutions of concrete financial problems. Concerns of risk management are addressed by the study of extreme values, the fitting of distributions with heavy tails, the computation of values at risk (VaR), and other measures of risk. Principal component analysis (PCA), smoothing, and regression techniques are applied to the construction of yield and forward curves. Time series analysis is applied to the study of temperature options and nonparametric estimation. Nonlinear filtering is applied to Monte Carlo simulations, option pricing and earnings prediction. This textbook is intended for undergraduate students majoring in financial engineering, or graduate students in a Master in finance or MBA program. It is sprinkled with practical examples using market data, and each chapter ends with exercises. Practical examples are solved in the R computing environment. They illustrate problems occurring in the commodity, energy and weather markets, as well as the fixed income, equity and credit markets. The examples, experiments and problem sets are based on the library Rsafd developed for the purpose of the text. The book should help quantitative analysts learn and implement advanced statistical concepts. Also, it will be valuable for researchers wishing to gain experience with financial data, implement and test mathematical theories, and address practical issues that are often ignored or underestimated in academic curricula. This is the new, fully-revised edition to the book Statistical Analysis of Financial Data in S-Plus. René Carmona is the Paul M. Wythes '55 Professor of Engineering and Finance at Princeton University in the department of Operations Research and Financial Engineering, and Director of Graduate Studies of the Bendheim Center for Finance. His publications include over one hundred articles and eight books in probability and statistics. He was elected Fellow of the Institute of Mathematical Statistics in 1984, and of the Society for Industrial and Applied Mathematics in 2010. He is on the editorial board of several peer-reviewed journals and book series. Professor Carmona has developed computer programs for teaching statistics and research in signal analysis and financial engineering. He has worked for many years on energy, the commodity markets and more recently in environmental economics, and he is recognized as a leading researcher and expert in these areas.

The World War II Databook

Illustrate your data in a more interactive way by implementing data visualization principles and creating visual stories using Tableau

About This Book Use data visualization principles to help you to design dashboards that enlighten and support business decisions

Integrate your data to provide mashed-up dashboards Connect to various data sources and understand what data is appropriate for Tableau Public

Understand chart types and when to use specific chart types with different types of data Who This Book Is For Data scientists who have just started using Tableau and want to build on the skills using practical examples. Familiarity with previous versions of Tableau will be helpful, but not necessary.

What You Will Learn Customize your designs to meet the needs of your business using Tableau Use Tableau to prototype, develop, and deploy the final dashboard Create filled maps and use any shape file Discover features of Tableau Public, from basic to advanced Build geographic maps to bring context to data Create filters and actions to allow greater interactivity to Tableau Public visualizations and dashboards Publish and embed Tableau visualizations and dashboards in articles

In Detail With increasing interest for data visualization in the media, businesses are looking to create effective dashboards that engage as well as communicate the truth of data. Tableau makes data accessible to everyone, and is a great way of sharing enterprise dashboards across the business. Tableau is a revolutionary toolkit that lets you simply and effectively create high-quality data visualizations. This course starts with making you familiar with its features and enable you to develop and enhance your dashboard skills, starting with an overview of what dashboard is, followed by how you can collect data using various mathematical formulas. Next, you'll learn to filter and group data, as well as how to use various functions to present the data in an appealing and accurate way. In the first module, you will learn how to use the key advanced string functions to play with data and images. You will be walked through the various features of Tableau including dual axes, scatterplot matrices, heat maps, and sizing. In the second module, you'll start with getting your data into Tableau, move onto generating progressively complex graphics, and end with the finishing touches and packaging your work for distribution. This module is filled with practical examples to help you create filled maps, use custom markers, add slider selectors, and create dashboards. You will learn how to manipulate data in various ways by applying various filters, logic, and calculating various aggregate measures. Finally, in the third module, you learn about Tableau Public using which allows readers to explore data associations in multiple-sourced public data, and uses state-of-the-art dashboard and chart graphics to immerse the users in an interactive experience. In this module, the readers can quickly gain confidence in understanding and expanding their visualization, creation knowledge, and quickly create interesting, interactive data visualizations to bring a richness and vibrancy to complex articles. The course provides a great overview for beginner to intermediate Tableau users, and covers the creation of data visualizations of varying complexities. Style and approach The approach will be a combined perspective, wherein we start by performing some basic recipes and move on to some advanced ones. Finally, we perform some advanced analytics and create appealing and insightful data stories using Tableau Public in a step-by-step manner.

Statistical Analysis of Financial Data in R

This data guide takes readers through the cycle of social science research, from applying for a research grant, through conducting the data collection phase, and ultimately to preparing the data for deposit in archives or data repositories. An adaptation of the fourth edition of the Guide to Social Science Data Preparation and Archiving of 2009 by the Inter-University Consortium for Political and Social Research at the University of Michigan, this publication will help researchers to manage, document, and archive their data and to think broadly about which types of digital content should be deposited in such an archive.

Tableau: Creating Interactive Data Visualizations

This Festschrift in honour of Ursula Gather's 60th birthday deals with modern topics in the field of robust statistical methods, especially for time series and regression analysis, and with statistical methods for complex data structures. The individual contributions of leading experts provide a textbook-style overview of the topic, supplemented by current research results and questions. The statistical theory and methods in this volume aim at the analysis of data which deviate from classical stringent model assumptions, which contain outlying values and/or have a complex structure. Written for researchers as well as master and PhD students with a good knowledge of statistics.

Preparing Data for Sharing

In order best exploit the incredible quantities of data being generated in most diverse disciplines data sciences increasingly gain worldwide importance. The book gives the mathematical foundations to handle data properly. It introduces basics and functionalities of the R programming language which has become the indispensable tool for data sciences. Thus it delivers the reader the skills needed to build own tool kits of a modern data scientist.

Robustness and Complex Data Structures

Perform advanced data manipulation tasks using pandas and become an expert data analyst. Key Features Manipulate and analyze your data expertly using the power of pandas Work with missing data and time series data and become a true pandas expert Includes expert

tips and techniques on making your data analysis tasks easier

Book Description pandas is a popular Python library used by data scientists and analysts worldwide to manipulate and analyze their data. This book presents useful data manipulation techniques in pandas to perform complex data analysis in various domains. An update to our highly successful previous edition with new features, examples, updated code, and more, this book is an in-depth guide to get the most out of pandas for data analysis. Designed for both intermediate users as well as seasoned practitioners, you will learn advanced data manipulation techniques, such as multi-indexing, modifying data structures, and sampling your data, which allow for powerful analysis and help you gain accurate insights from it. With the help of this book, you will apply pandas to different domains, such as Bayesian statistics, predictive analytics, and time series analysis using an example-based approach. And not just that; you will also learn how to prepare powerful, interactive business reports in pandas using the Jupyter notebook. By the end of this book, you will learn how to perform efficient data analysis using pandas on complex data, and become an expert data analyst or data scientist in the process. What you will learn

Speed up your data analysis by importing data into pandas

Keep relevant data points by selecting subsets of your data

Create a high-quality dataset by cleaning data and fixing missing values

Compute actionable analytics with grouping and aggregation in pandas

Master time series data analysis in pandas

Make powerful reports in pandas using Jupyter notebooks

Who this book is for This book is for data scientists, analysts and Python developers who wish to explore advanced data analysis and scientific computing techniques using pandas. Some fundamental understanding of Python programming and familiarity with the basic data analysis concepts is all you need to get started with this book.

Mathematical Foundations of Data Science Using R

The present work provides a platform for leading Data designers whose vision and creativity help us to anticipate major changes occurring in the Data Design field, and pre-empt the future. Each of them strives to provide new answers to the question, “What challenges await Data Design?” To avoid falling into too narrow a mind-set, each works hard to elucidate the breadth of Data Design today and to demonstrate its widespread application across a variety of business sectors. With end users in mind, designer-contributors bring to light the myriad of purposes for which the field was originally intended, forging the bond even further between Data Design and the aims and intentions of those who contribute to it. The first seven parts of the book outline the scope of Data Design, and presents a line-up of “viewpoints” that highlight this discipline’s main topics, and offers an in-depth look into practices boasting both foresight and imagination. The eighth and final part features a series of interviews with Data designers and artists whose methods embody originality and marked singularity. As a result, a number of enlightening concepts and bright ideas unfold within the confines of this book to help dispel the thick fog around this new and still relatively unknown discipline. A plethora of equally eye-opening and edifying new terms, words, and key expressions also unfurl. Informing, influencing, and inspiring are just a few of the buzz words belonging to an initiative that is, first and foremost, a creative one, not to mention the possibility to discern the ever-changing and naturally complex nature of

today's datasphere. Providing an invaluable and cutting-edge resource for design researchers, this work is also intended for students, professionals and practitioners involved in Data Design, Interaction Design, Digital & Media Design, Data & Information Visualization, Computer Science and Engineering.

Mastering pandas

Use PySpark to easily crush messy data at-scale and discover proven techniques to create testable, immutable, and easily parallelizable Spark jobs

Key Features

- Work with large amounts of agile data using distributed datasets and in-memory caching
- Source data from all popular data hosting platforms, such as HDFS, Hive, JSON, and S3
- Employ the easy-to-use PySpark API to deploy big data Analytics for production

Book Description

Apache Spark is an open source parallel-processing framework that has been around for quite some time now. One of the many uses of Apache Spark is for data analytics applications across clustered computers. In this book, you will not only learn how to use Spark and the Python API to create high-performance analytics with big data, but also discover techniques for testing, immunizing, and parallelizing Spark jobs. You will learn how to source data from all popular data hosting platforms, including HDFS, Hive, JSON, and S3, and deal with large datasets with PySpark to gain practical big data experience. This book will help you work on prototypes on local machines and subsequently go on to handle messy data in production and at scale. This book covers installing and setting up PySpark, RDD operations, big data cleaning and wrangling, and aggregating and summarizing data into useful reports. You will also learn how to implement some practical and proven techniques to improve certain aspects of programming and administration in Apache Spark. By the end of the book, you will be able to build big data analytical solutions using the various PySpark offerings and also optimize them effectively. What you will learn

- Get practical big data experience while working on messy datasets
- Analyze patterns with Spark SQL to improve your business intelligence
- Use PySpark's interactive shell to speed up development time
- Create highly concurrent Spark programs by leveraging immutability
- Discover ways to avoid the most expensive operation in the Spark API: the shuffle operation
- Re-design your jobs to use `reduceByKey` instead of `groupByKey`
- Create robust processing pipelines by testing Apache Spark jobs

Who this book is for

This book is for developers, data scientists, business analysts, or anyone who needs to reliably analyze large amounts of large-scale, real-world data. Whether you're tasked with creating your company's business intelligence function or creating great data platforms for your machine learning models, or are looking to use code to magnify the impact of your business, this book is for you.

New Challenges for Data Design

This is the third book devoted to theoretical issues in data bases that we have edited. Each book has been the outgrowth of papers held at a workshop in Toulouse, France. The first workshop, held in 1977 focused primarily on the important topic of logic and databases. The book, *Logic and Databases* was the result of this effort. The diverse uses of logic for databases such as its use as a theoretical basis for databases, for deduction and for integrity constraints formulation and checking was described in the chapters of the book. The interest generated by the first workshop led to the decision to conduct other workshops focused on theoretical issues in databases. In addition to logic and databases the types of papers were expanded to include other important theoretical issues such as dependency theory which, although it sometimes uses logic as a basis, does not fit with our intended meaning of logic and databases explored at the first workshop. Because of the broader coverage, and because we anticipated further workshops, the second book was entitled, *Advances in Database Theory - Volume 1*. The book "*Logic and Databases*" should be considered Volume 0 of this series.

Hands-On Big Data Analytics with PySpark

Winner of the 2013 DeGroot Prize. A state-of-the-art presentation of spatio-temporal processes, bridging classic ideas with modern hierarchical statistical modeling concepts and the latest computational methods Noel Cressie and Christopher K. Wikle, are also winners of the 2011 PROSE Award in the Mathematics category, for the book "*Statistics for Spatio-Temporal Data*" (2011), published by John Wiley and Sons. (The PROSE awards, for Professional and Scholarly Excellence, are given by the Association of American Publishers, the national trade association of the US book publishing industry.) *Statistics for Spatio-Temporal Data* has now been reprinted with small corrections to the text and the bibliography. The overall content and pagination of the new printing remains the same; the difference comes in the form of corrections to typographical errors, editing of incomplete and missing references, and some updated spatio-temporal interpretations. From understanding environmental processes and climate trends to developing new technologies for mapping public-health data and the spread of invasive-species, there is a high demand for statistical analyses of data that take spatial, temporal, and spatio-temporal information into account. *Statistics for Spatio-Temporal Data* presents a systematic approach to key quantitative techniques that incorporate the latest advances in statistical computing as well as hierarchical, particularly Bayesian, statistical modeling, with an emphasis on dynamical spatio-temporal models. Cressie and Wikle supply a unique presentation that incorporates ideas from the areas of time series and spatial statistics as well as stochastic processes. Beginning with separate treatments of temporal data and spatial data, the book combines these concepts to discuss spatio-temporal statistical methods for understanding complex processes. Topics of coverage include: Exploratory methods for spatio-temporal data, including visualization, spectral analysis, empirical orthogonal function analysis, and LISAs Spatio-temporal covariance functions, spatio-temporal kriging, and time series of spatial

processes Development of hierarchical dynamical spatio-temporal models (DSTMs), with discussion of linear and nonlinear DSTMs and computational algorithms for their implementation Quantifying and exploring spatio-temporal variability in scientific applications, including case studies based on real-world environmental data Throughout the book, interesting applications demonstrate the relevance of the presented concepts. Vivid, full-color graphics emphasize the visual nature of the topic, and a related FTP site contains supplementary material. Statistics for Spatio-Temporal Data is an excellent book for a graduate-level course on spatio-temporal statistics. It is also a valuable reference for researchers and practitioners in the fields of applied mathematics, engineering, and the environmental and health sciences.

Advances in Data Base Theory

This report improves the evidence base on the role of Data Driven Innovation for promoting growth and well-being, and provide policy guidance on how to maximise the benefits of DDI and mitigate the associated economic and societal risks.

Statistics for Spatio-Temporal Data

This book features 29 peer-reviewed papers presented at the 9th International Conference on Soft Methods in Probability and Statistics (SMPS 2018), which was held in conjunction with the 5th International Conference on Belief Functions (BELIEF 2018) in Compiègne, France on September 17-21, 2018. It includes foundational, methodological and applied contributions on topics as varied as imprecise data handling, linguistic summaries, model coherence, imprecise Markov chains, and robust optimisation. These proceedings were produced using EasyChair. Over recent decades, interest in extensions and alternatives to probability and statistics has increased significantly in diverse areas, including decision-making, data mining and machine learning, and optimisation. This interest stems from the need to enrich existing models, in order to include different facets of uncertainty, like ignorance, vagueness, randomness, conflict or imprecision. Frameworks such as rough sets, fuzzy sets, fuzzy random variables, random sets, belief functions, possibility theory, imprecise probabilities, lower previsions, and desirable gambles all share this goal, but have emerged from different needs. The advances, results and tools presented in this book are important in the ubiquitous and fast-growing fields of data science, machine learning and artificial intelligence. Indeed, an important aspect of some of the learned predictive models is the trust placed in them. Modelling the uncertainty associated with the data and the models carefully and with principled methods is one of the means of increasing this trust, as the model will then be able to distinguish between reliable and less reliable predictions. In addition, extensions such as fuzzy sets can be explicitly designed to provide interpretable predictive models, facilitating user interaction and increasing

trust.

Data-Driven Innovation Big Data for Growth and Well-Being

Equal parts mail art, data visualization, and affectionate correspondence, *Dear Data* celebrates "the infinitesimal, incomplete, imperfect, yet exquisitely human details of life," in the words of Maria Popova (Brain Pickings), who introduces this charming and graphically powerful book. For one year, Giorgia Lupi, an Italian living in New York, and Stefanie Posavec, an American in London, mapped the particulars of their daily lives as a series of hand-drawn postcards they exchanged via mail weekly—small portraits as full of emotion as they are data, both mundane and magical. *Dear Data* reproduces in pinpoint detail the full year's set of cards, front and back, providing a remarkable portrait of two artists connected by their attention to the details of their lives—including complaints, distractions, phone addictions, physical contact, and desires. These details illuminate the lives of two remarkable young women and also inspire us to map our own lives, including specific suggestions on what data to draw and how. A captivating and unique book for designers, artists, correspondents, friends, and lovers everywhere.

Uncertainty Modelling in Data Science

The datafication of our world offers huge challenges and opportunities for social science. The 'data-drivenness' of computational research can occur at the expense of theoretical reflection and interpretation. Additionally, it can be difficult to reconcile the 'quantitative' dimensions of big data with the 'qualitative' sensibilities needed for its understanding. At the same time, this opens up possibilities for reimagining key principles of social inquiry. In this experimental and provocative book, Simon Lindgren argues that a hybrid approach to data and theory must be developed in order to make sense of today's ambivalent, turbulent, and media-saturated political landscape. He pushes for the development of a critical science of data, joining the interpretive theoretical and ethical sensibilities of social science with the predictive and prognostic powers of data science and computational methods. In order for theories and research methods to be more useful and relevant, they must be dismantled and put together in new, alternative, and unexpected ways. *Data Theory* is essential reading for social scientists and data scientists, as well as students taking courses in social theory and data, digital methods, big data, and data and society.

Dear Data

By charting changes over time and investigating whether and when events occur, researchers reveal the temporal rhythms of our lives.

Data Theory

This edited collection aims to reimagine and extend ethnography for a data-saturated world. The book brings together leading scholars in the social sciences who have been interrogating and collaborating with data scientists working in a range of different settings. The book explores how a repurposed form of ethnography might illuminate the kinds of knowledge that are being produced by data science. It also describes how collaborations between ethnographers and data scientists might lead to new forms of social analysis

Applied Longitudinal Data Analysis

Data analysis lies at the heart of every experimental science. Providing a modern introduction to statistics, this book is ideal for undergraduates in physics. It introduces the necessary tools required to analyse data from experiments across a range of areas, making it a valuable resource for students. In addition to covering the basic topics, the book also takes in advanced and modern subjects, such as neural networks, decision trees, fitting techniques and issues concerning limit or interval setting. Worked examples and case studies illustrate the techniques presented, and end-of-chapter exercises help test the reader's understanding of the material.

Ethnography for a data-saturated world

Optimize your marketing strategies through analytics and machine learning
Key Features
Understand how data science drives successful marketing campaigns
Use machine learning for better customer engagement, retention, and product recommendations
Extract insights from your data to optimize marketing strategies and increase profitability
Book Description
Regardless of company size, the adoption of data science and machine learning for marketing has been rising in the industry. With this book, you will learn to implement data science techniques to understand the drivers behind the successes and failures of marketing campaigns. This book is a comprehensive guide to help you understand and predict customer behaviors and create more effectively targeted and personalized marketing strategies. This is a practical guide to performing simple-to-advanced tasks, to extract hidden insights from the data and use

them to make smart business decisions. You will understand what drives sales and increases customer engagements for your products. You will learn to implement machine learning to forecast which customers are more likely to engage with the products and have high lifetime value. This book will also show you how to use machine learning techniques to understand different customer segments and recommend the right products for each customer. Apart from learning to gain insights into consumer behavior using exploratory analysis, you will also learn the concept of A/B testing and implement it using Python and R. By the end of this book, you will be experienced enough with various data science and machine learning techniques to run and manage successful marketing campaigns for your business. What you will learn

- Learn how to compute and visualize marketing KPIs in Python and R
- Master what drives successful marketing campaigns with data science
- Use machine learning to predict customer engagement and lifetime value
- Make product recommendations that customers are most likely to buy
- Learn how to use A/B testing for better marketing decision making
- Implement machine learning to understand different customer segments

Who this book is for If you are a marketing professional, data scientist, engineer, or a student keen to learn how to apply data science to marketing, this book is what you need! It will be beneficial to have some basic knowledge of either Python or R to work through the examples. This book will also be beneficial for beginners as it covers basic-to-advanced data science concepts and applications in marketing with real-life examples.

Statistical Data Analysis for the Physical Sciences

This book constitutes the conference proceedings of the 16th International Symposium on Intelligent Data Analysis, which was held in October 2017 in London, UK. The 28 full papers presented in this book were carefully reviewed and selected from 66 submissions. The traditional focus of the IDA symposium series is on end-to-end intelligent support for data analysis. IDA solicits papers on all aspects of intelligent data analysis, including papers on intelligent support for modelling and analyzing data from complex, dynamical systems.

Hands-On Data Science for Marketing

Information Management: Gaining a Competitive Advantage with Data is about making smart decisions to make the most of company information. Expert author William McKnight develops the value proposition for information in the enterprise and succinctly outlines the numerous forms of data storage. Information Management will enlighten you, challenge your preconceived notions, and help activate information in the enterprise. Get the big picture on managing data so that your team can make smart decisions by understanding how everything from workload allocation to data stores fits together. The practical, hands-on guidance in this book includes:

- Part 1: The importance of information management and analytics to business, and how data warehouses are used
- Part 2: The technologies and data

that advance an organization, and extend data warehouses and related functionality Part 3: Big Data and NoSQL, and how technologies like Hadoop enable management of new forms of data Part 4: Pulls it all together, while addressing topics of agile development, modern business intelligence, and organizational change management Read the book cover-to-cover, or keep it within reach for a quick and useful resource. Either way, this book will enable you to master all of the possibilities for data or the broadest view across the enterprise. Balances business and technology, with non-product-specific technical detail Shows how to leverage data to deliver ROI for a business Engaging and approachable, with practical advice on the pros and cons of each domain, so that you learn how information fits together into a complete architecture Provides a path for the data warehouse professional into the new normal of heterogeneity, including NoSQL solutions

Advances in Intelligent Data Analysis XVI

Information Management

https://mohadev.charak.com/^p/edu/goto?KINDLE=scott-mccloud_reinventing_comics_download_pdf_format.pdf
https://mohadev.charak.com/^k/science/search?PAGE=the-art_of_fermentation_an_in_depth_exploration-essential_concepts_and_processes-from_around-world-sandor_ellix-katz.pdf
[https://mohadev.charak.com/\\$g/ref/goto?PLAY=getting_started-with_cnc_personal_digital_fabrication-with-shapeoko_and_other_computer_controlled_routers-make.pdf](https://mohadev.charak.com/$g/ref/goto?PLAY=getting_started-with_cnc_personal_digital_fabrication-with-shapeoko_and_other_computer_controlled_routers-make.pdf)
[https://mohadev.charak.com/\\$j/short/file?TXT=ap_statistics-investigative-task_chapter_26.pdf](https://mohadev.charak.com/$j/short/file?TXT=ap_statistics-investigative-task_chapter_26.pdf)
https://mohadev.charak.com/@o/pub/list?PAGE=answers-to_heredity_lab_report_34_oficceore.pdf
[https://mohadev.charak.com/\\$d/slide/slug?RTF=physical_chemistry_from_a_different-angle_introducing_chemical_equilibrium-kinetics_and_electrochemistry-by-numerous_experiments.pdf](https://mohadev.charak.com/$d/slide/slug?RTF=physical_chemistry_from_a_different-angle_introducing_chemical_equilibrium-kinetics_and_electrochemistry-by-numerous_experiments.pdf)
<https://mohadev.charak.com/!y/lib/search?PUB=the-vertical-aeroponic-growing-system.pdf>
https://mohadev.charak.com/_h/ref/dl?PPT=the_tragedy_of_julius_caesar-selection_test_answers.pdf
[https://mohadev.charak.com/\\$v/edu/visit?EPDF=structural_analysis_5th_edition-by_aslam_kassimali.pdf](https://mohadev.charak.com/$v/edu/visit?EPDF=structural_analysis_5th_edition-by_aslam_kassimali.pdf)
https://mohadev.charak.com/^x/lib/dl?PDF=la_linguistica_berruto-cerruti.pdf