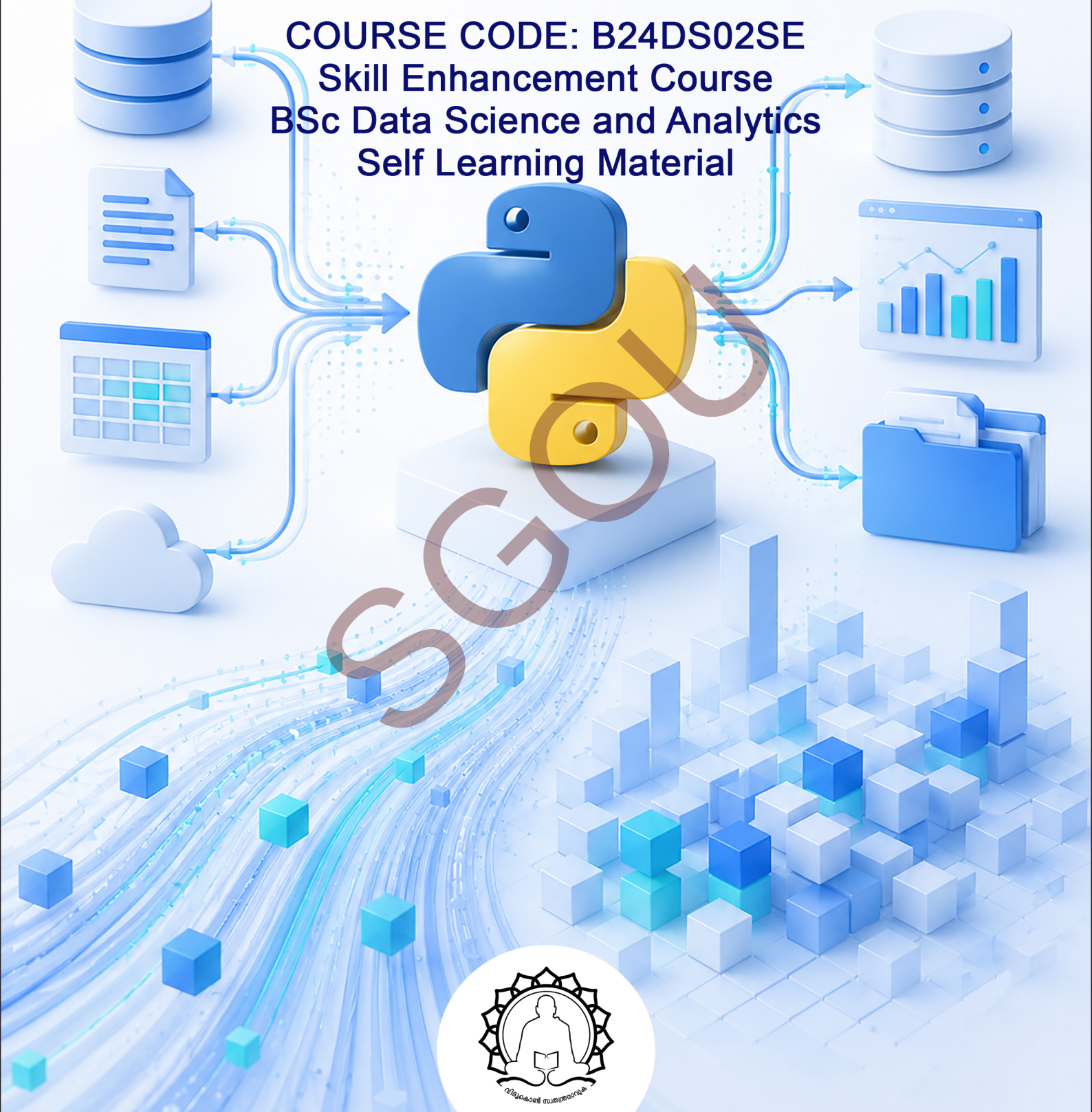


SGOU

Higher
Education
for All

DATA ENGINEERING WITH PYTHON

COURSE CODE: B24DS02SE
Skill Enhancement Course
BSc Data Science and Analytics
Self Learning Material



SREENARAYANAGURU
OPEN UNIVERSITY

SREENARAYANAGURU OPEN UNIVERSITY

The State University for Education, Training and Research in Blended Format, Kerala

SREENARAYANAGURU OPEN UNIVERSITY

Vision

To increase access of potential learners of all categories to higher education, research and training, and ensure equity through delivery of high quality processes and outcomes fostering inclusive educational empowerment for social advancement.

Mission

To be benchmarked as a model for conservation and dissemination of knowledge and skill on blended and virtual mode in education, training and research for normal, continuing, and adult learners.

Pathway

Access and Quality define Equity.

Data Engineering with Python

Course Code: B24DS02SE

Semester- IV

**Skill Enhancement Course
BSc Data Science and Analytics
Self Learning Material
(with model question paper sets)**



SREENARAYANAGURU
OPEN UNIVERSITY

SREENARAYANAGURU OPEN UNIVERSITY

The State University for Education, Training and Research in Blended Format, Kerala



Data Engineering with Python

Course Code: B24DS02SE

Semester- IV

Skill Enhancement Course

BSc Data Science and Analytics

Academic Committee

Dr. T. K. Manoj
Dr. Smitha Dharan
Dr. Satheesh S.
Dr. Vinod Chandra S.S.
Dr. Hari V. S.
Dr. Sharon Susan Jacob
Dr. Ajith Kumar R.
Dr. Smiju I.S.
Dr. Nimitha Aboobaker

Development of the Content

ICT Academy of Kerala

Review and Edit

Riji N. Das

Scrutiny

Greeshma P.P.
Sreerekha V.K.
Shamin S.
Aswathy V.S.
Dr. Kanitha Divakar
Subi Priya Laxmi S.B.N.

Cover Design

Jineshkumar S

Co-ordination

Director, MDDC :

Dr. I.G. Shibi

Asst. Director, MDDC :

Dr. Sajeevkumar G.

Coordinator, Development:

Dr. Anfal M.

Coordinator, Distribution:

Dr. Sanitha K.K.



Scan this QR Code for reading the SLM
on a digital device.

First Edition
June 2026

Copyright
© Sreenarayanaguru Open University

ISBN 978-81-997556-7-3



All rights reserved. No part of this work may be reproduced in any form, by mimeograph or any other means, without permission in writing from Sreenarayanaguru Open University. Printed and published on behalf of Sreenarayanaguru Open University by Registrar, SGOU, Kollam.

www.sgou.ac.in



Visit and Subscribe our Social Media Platforms

MESSAGE FROM VICE CHANCELLOR

Dear learner,

I extend my heartfelt greetings and profound enthusiasm as I warmly welcome you to Sreenarayanaguru Open University. Established in September 2020 as a state-led endeavour to promote higher education through open and distance learning modes, our institution was shaped by the guiding principle that access and quality are the cornerstones of equity. We have firmly resolved to uphold the highest standards of education, setting the benchmark and charting the course.

The courses offered by the Sreenarayanaguru Open University aim to strike a quality balance, ensuring students are equipped for both personal growth and professional excellence. The University embraces the widely acclaimed “blended format,” a practical framework that harmoniously integrates Self-Learning Materials, Classroom Counseling, and Virtual modes, fostering a dynamic and enriching experience for both learners and instructors.

The University is committed to providing an engaging and dynamic educational environment that encourages active learning. The Study and Learning Material (SLM) is specifically designed to offer you a comprehensive and integrated learning experience, fostering a strong interest in exploring advancements in information technology (IT). The curriculum has been carefully structured to ensure a logical progression of topics, allowing you to develop a clear understanding of the evolution of the discipline. It is thoughtfully curated to equip you with the knowledge and skills to navigate current trends in IT, while fostering critical thinking and analytical capabilities. The Self-Learning Material has been meticulously crafted, incorporating relevant examples to facilitate better comprehension.

Rest assured, the university’s student support services will be at your disposal throughout your academic journey, readily available to address any concerns or grievances you may encounter. We encourage you to reach out to us freely regarding any matter about your academic programme. It is our sincere wish that you achieve the utmost success.



Regards,
Dr. Jagathy Raj V. P.

01-06-2026

Contents

BLOCK-1	Introduction to Data Engineering and Working Data in Python	1
UNIT 1	Introduction to Data Engineering	2
UNIT 2	Foundations of Data Engineering and Python for Data Workflows	15
UNIT 3	Data Manipulation with Libraries	28
UNIT 4	File I/O Operations	45
BLOCK-2	Data Processing and Transformation, Data Integration and ETL	56
UNIT 1	Cleaning and Preprocessing Data using Python	57
UNIT 2	Data Normalization and Standardization	71
UNIT 3	Introduction to ETL Processes	82
UNIT 4	Transforming Data using Python, Loading Data, and Security Measures	91
	Model Question Paper Sets	

BLOCK-1

Introduction to Data Engineering and Working Data in Python

UNIT 1

Introduction to Data Engineering

Learning Outcomes

At the end of this unit, the learner will be able to :

- ◆ explain the concept and importance of Data Engineering in modern data-driven environments
- ◆ identify the key components of data engineering systems, including data sources, data ingestion, storage, and processing
- ◆ recognize the major tools and technologies used in data engineering for managing and processing large datasets
- ◆ describe the role of data engineering in supporting data analytics, machine learning, and business decision-making
- ◆ analyze the major challenges faced in data engineering, including scalability, data quality, data integration, and system reliability

Prerequisite

The rapid growth of digital technologies has led to an unprecedented increase in the volume, variety, and speed of data generation. Organizations today collect data from numerous sources such as websites, mobile applications, sensors, and transaction systems. Managing and processing such large volumes of data requires specialized systems and professionals who can design reliable data infrastructures. Before exploring the concepts of Data Engineering in this unit, learners should possess a basic understanding of certain foundational ideas related to data and computing systems.

A fundamental familiarity with computer systems and data handling concepts will help learners better understand how data is generated, stored, and processed in modern applications. Learners should have a general awareness of how computer systems manage information through files, databases, and storage devices. Understanding the

basic structure of data such as records, fields, and tables will also help in appreciating how large datasets are organized and manipulated.

Some exposure to basic programming concepts is also beneficial. Although advanced programming knowledge is not required for this unit, learners should have a general idea of programming logic, algorithms, and problem-solving approaches. This background will help learners understand how data engineers use programming languages and tools to automate data processing tasks.

In addition, learners should possess a general awareness of different types of data that exist in modern computing environments. Data may appear in structured formats such as database tables, semi-structured formats like JSON or XML files, or unstructured formats such as images, videos, and text documents. Recognizing these different data forms will help learners understand why specialized data engineering solutions are required to handle diverse and complex datasets.

Finally, learners should have an interest in understanding how organizations use data for decision-making and innovation. Data engineering plays a crucial role in enabling analytics, artificial intelligence, and business intelligence systems. With this background and curiosity, learners will be better prepared to explore the concepts of data engineering, its components, tools, and the challenges associated with managing large-scale data systems.

Keywords

Data Pipeline, Data Ingestion, Data Processing, ETL, Data Infrastructure, Big Data Technologies, Distributed Systems, Data Integration, Scalability, Data Workflow, Data Analytics.

Discussion

1.1.1 Introduction to Data Engineering

In the digital era, organizations generate massive volumes of data from various sources such as websites, mobile applications, online transactions, sensors, and social media platforms. This rapid growth in data has created new challenges in storing, managing, and processing information efficiently. Traditional data management approaches, which were designed for smaller and structured datasets, are often unable to handle the scale and complexity of modern data environments. As a result, a specialized discipline known as Data Engineering has emerged to address these challenges.

Data Engineering refers to the design, development, and management of systems that collect, store, process, and transform raw data into a format that can be easily used for analysis and decision-making. Data engineers build the infrastructure and data pipelines that allow organizations to move data from various sources to storage systems and analytical platforms. These systems ensure that data is reliable, accessible, and ready for use by data analysts, data scientists, and business intelligence tools.

A key objective of data engineering is to ensure the efficient flow of data across different stages of the data lifecycle. Raw data generated from multiple sources must be collected, cleaned, integrated, and transformed before it can be used for analysis. Data engineers design automated workflows that handle these tasks and ensure that large datasets can be processed without delays or errors. This process often involves techniques such as data ingestion, data transformation, and data storage management.

Modern organizations rely heavily on data to gain insights into customer behavior, optimize operations, and support strategic decision-making. Data engineering plays a crucial role in enabling these capabilities by providing the foundational infrastructure required for advanced analytics and machine learning. Without well-designed data engineering systems, it would be difficult for organizations to extract meaningful insights from large and complex datasets.

As the volume and diversity of data continue to increase, the importance of data engineering has grown significantly. Industries such as finance, healthcare, retail, transportation, and telecommunications increasingly depend on robust data engineering frameworks to manage large-scale data systems. Consequently, data engineering has become a critical component of modern data-driven technologies and digital transformation initiatives.

1.1.2 Key Components of Data Engineering

Data engineering systems are built around several essential components that work together to collect, process, store, and deliver data for analysis. These components form the foundation of modern data infrastructures and ensure that raw data generated from various sources can be transformed into meaningful information. A well-designed data engineering architecture ensures efficient data flow, scalability, and reliability across the entire data lifecycle.

1. Data Sources

Data sources represent the origin points from which data is generated and collected. In modern digital environments, data is produced from numerous sources such as web applications, mobile devices, sensors, enterprise systems, transaction databases, and social media platforms. These sources generate data in different formats including structured, semi-structured, and unstructured forms. The diversity of data sources increases the complexity of data management and requires specialized systems capable of handling large and heterogeneous datasets.

2. Data Ingestion

Data ingestion refers to the process of collecting data from various sources and transferring it into storage or processing systems. This stage ensures that data generated across different platforms is captured reliably and efficiently. Data ingestion can occur in two primary ways: batch ingestion, where data is collected and processed at regular intervals, and real-time ingestion, where data is continuously streamed into the system. Efficient ingestion mechanisms are essential to prevent data loss and ensure timely availability of information for analysis.

3. Data Storage

Once data is collected, it must be stored in systems capable of handling large volumes of information. Modern data engineering systems use scalable storage solutions such as data warehouses, data lakes, and distributed storage systems. These storage systems are designed to manage large datasets while providing high availability and fault tolerance. Distributed storage allows data to be stored across multiple machines, enabling systems to scale horizontally as data volume grows.

4. Data Processing and Transformation

Data processing involves transforming raw data into structured formats that can be used for analysis and decision-making. During this stage, data engineers perform tasks such as data cleaning, filtering, aggregation, and transformation. Processing frameworks enable large datasets to be processed efficiently through parallel computing techniques. These transformations ensure that the data becomes consistent, accurate, and suitable for analytical applications.

5. Data Pipelines

A data pipeline is a sequence of processes that automate the movement of data from its source to its final destination. Data pipelines integrate multiple components such as ingestion systems, storage platforms, and processing frameworks. Automated pipelines ensure that data flows continuously through the system without manual intervention. This automation improves efficiency, reduces errors, and allows organizations to process large volumes of data in a reliable manner.

Table 1.1.1 Key Components of Data Engineering Systems

Component	Description
Data Sources	Origin points where data is generated
Data Ingestion	Process of collecting data from sources
Data Storage	Systems used to store large datasets
Data Processing	Transformation and cleaning of data
Data Pipelines	Automated workflows for moving data

These components collectively form the backbone of modern data engineering architectures. By integrating data sources, ingestion mechanisms, storage platforms, and processing frameworks into a unified system, organizations can ensure that large volumes of data are efficiently managed and made available for analysis.

1.1.3 Tools and Technologies in Data Engineering

Modern data engineering relies on a wide range of tools and technologies designed to manage large volumes of data efficiently. These technologies support various stages of the data lifecycle, including data collection, storage, processing, integration, and analysis. By combining these tools within a unified ecosystem, organizations can build scalable and reliable data infrastructures capable of supporting advanced analytics and data-driven decision-making.

1. Programming Languages

Programming languages play a crucial role in implementing data engineering workflows. They are used to develop data pipelines, automate data processing tasks, and interact with storage systems and databases. Among the available languages, **Python** has become one of the most widely used languages in data engineering due to its simplicity, readability, and extensive ecosystem of libraries. Other languages such as **Java** and **Scala** are also commonly used, particularly in big data processing frameworks.

Python provides powerful libraries that support data manipulation, workflow automation, and integration with big data platforms. These libraries enable data engineers to perform tasks such as data cleaning, transformation, and large-scale processing efficiently.

2. Data Storage Technologies

Efficient storage systems are essential for handling large datasets generated by modern applications. Data engineers rely on a variety of storage technologies depending on the type and volume of data involved.

Traditional relational databases such as MySQL and PostgreSQL are widely used for structured data storage. However, modern data engineering environments often require more scalable solutions capable of handling massive datasets. This has led to the development of distributed storage systems and NoSQL databases, which can store large volumes of structured and unstructured data across multiple machines.

Data warehouses and data lakes are also widely used storage solutions in data engineering architectures. Data warehouses are optimized for analytical queries, while data lakes store large volumes of raw data in their original formats.

3. Big Data Processing Frameworks

As datasets grow larger, traditional processing techniques become inefficient. Big data processing frameworks allow organizations to process massive datasets using distributed computing. These frameworks divide large computational tasks into smaller units that can be processed simultaneously across multiple machines.

One of the most widely used frameworks is Apache Spark, which supports high-speed data processing and real-time analytics. Another widely used framework is Hadoop, which provides distributed storage and batch processing capabilities. These frameworks allow organizations to process large-scale datasets efficiently while ensuring scalability and fault tolerance.

4. Data Integration and Workflow Tools

Data integration tools are responsible for coordinating the movement of data across different systems. These tools help automate data pipelines and manage complex data workflows.

Workflow management systems allow data engineers to schedule, monitor, and manage data processing tasks. These tools ensure that data pipelines run in the correct sequence and handle failures effectively. By automating data workflows, organizations can maintain consistent and reliable data processing operations.

5. Cloud Platforms and Infrastructure

Cloud computing has significantly transformed modern data engineering practices. Cloud platforms provide scalable infrastructure and managed services that simplify the deployment and management of data systems.

Cloud-based data engineering solutions allow organizations to scale storage and processing resources dynamically according to their needs. These platforms provide services for data storage, data processing, analytics, and machine learning, enabling organizations to build comprehensive data engineering solutions without maintaining complex hardware infrastructures.

Table 1.1.2 Examples of Data Engineering Tools and Technologies

Category	Examples	Purpose
Programming Languages	Python, Java, Scala	Data processing and pipeline development
Storage Technologies	MySQL, PostgreSQL, NoSQL	Data storage and management
Big Data Frameworks	Hadoop, Apache Spark	Large-scale data processing
Workflow Tools	Airflow, Luigi	Pipeline scheduling and automation
Cloud Platforms	AWS, Azure, Google Cloud	Scalable data infrastructure

These tools and technologies collectively enable the development of modern data engineering systems capable of managing large volumes of data efficiently. By integrating programming languages, storage technologies, processing frameworks, and cloud infrastructure, data engineers can design scalable and robust data platforms.

1.1.4 Role of Data Engineering in Data-Driven Organizations

In the modern digital economy, organizations increasingly rely on data-driven decision making to gain competitive advantage, improve operational efficiency, and enhance customer experiences. A data-driven organization uses data as the primary foundation for strategic planning, operational processes, and innovation. However, raw data generated from various sources cannot directly support analytics or decision-making without proper processing and organization. This is where data engineering plays a critical role.

Data engineering focuses on designing and maintaining systems that collect, store, process, and deliver data in a reliable and efficient manner. Data engineers develop data pipelines and infrastructure that transform raw data into structured and usable formats for analytics, reporting, and machine learning applications. In this way, data engineering forms the backbone of modern analytics ecosystems.

Figure 1.1.1 illustrates the role of data engineering within a data-driven organization, showing how data flows from multiple sources through engineering processes before reaching analytical tools and decision-makers.

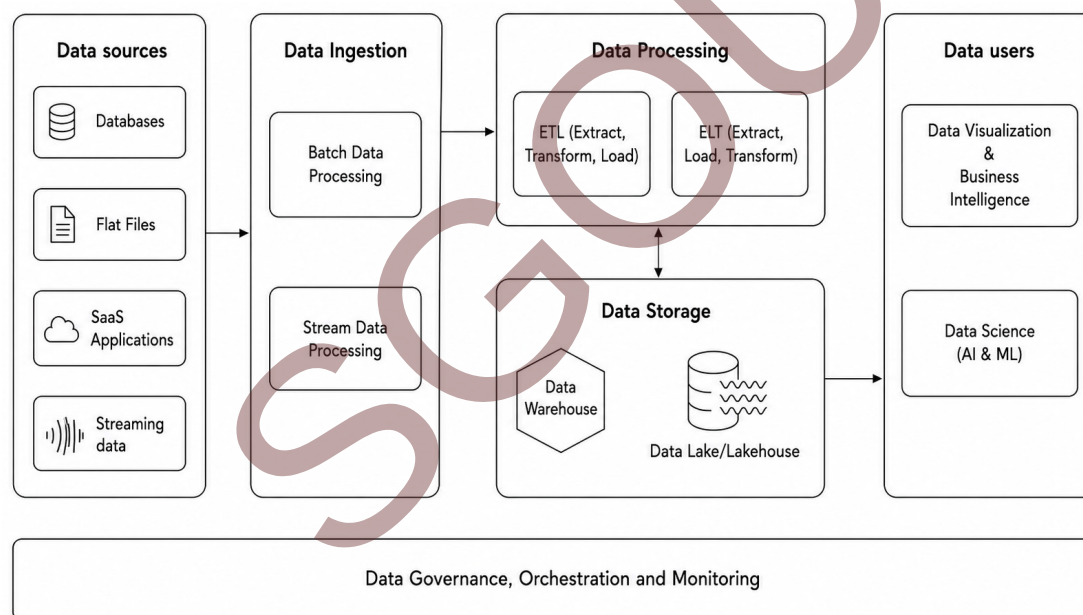


Fig. 1.1.1 Role of data engineering in data-driven organizations

1. Data Integration and Data Collection

Organizations generate data from diverse sources such as transactional systems, websites, mobile applications, sensors, and social media platforms. Data engineers design systems that collect this data efficiently and integrate it into centralized platforms such as data lakes or data warehouses. These systems enable organizations to maintain a unified view of their data.

2. Development of Data Pipelines

A core responsibility of data engineers is the development of data pipelines, which automate the movement and transformation of data from its source to storage and analytical platforms. These pipelines involve processes such as data extraction, transformation, and loading (ETL/ELT). Automated pipelines ensure that data remains continuously updated and available for analytical use.

3. Enabling Reliable Data Storage

Data engineers design scalable storage infrastructures capable of handling large volumes of data. These systems must support both structured and unstructured data while ensuring reliability, fault tolerance, and fast access. Modern data platforms often use distributed storage technologies to achieve these goals.

4. Ensuring Data Quality and Consistency

For organizations to trust their analytical insights, the underlying data must be accurate and consistent. Data engineers implement data validation, cleaning, and transformation processes that improve data quality. These processes remove duplicate records, correct inconsistencies, and standardize data formats before the data is used for analysis.

5. Supporting Analytics and Business Intelligence

Data scientists, analysts, and business intelligence professionals depend on well-prepared datasets for performing analytics and building predictive models. Data engineering systems provide the structured datasets and data infrastructure required for advanced analytics, machine learning, and visualization tools.

Table 1.1.3 Key roles of data engineering in data-driven organizations

Role of Data Engineering	Description	Organizational Benefit
Data Collection and Integration	Gathering data from multiple sources and integrating it into centralized platforms	Unified data access
Data Pipeline Development	Automating the flow of data through ETL or ELT processes	Efficient data processing
Data Storage Management	Designing scalable storage systems such as data lakes and warehouses	Reliable data availability
Data Quality Management	Ensuring data accuracy, consistency, and completeness	Trustworthy analytics
Analytics Enablement	Providing structured datasets for analytics and machine learning	Data-driven decision making

1.1.5 Challenges in Data Engineering

As organizations increasingly rely on data for decision-making and operational processes, the role of data engineering has become more critical. However, managing large-scale data environments presents several technical and organizational challenges. Data engineers must design systems capable of handling massive data volumes, diverse data formats, real-time processing requirements, and strict security regulations. Addressing these challenges requires advanced infrastructure, efficient data pipelines, and reliable data management strategies.

1. Managing Large Data Volumes

One of the primary challenges in data engineering is handling the enormous volume of data generated by modern digital systems. Organizations collect data from numerous sources such as social media, IoT devices, transactional systems, and online platforms. Processing and storing such massive datasets require scalable distributed storage systems and high-performance processing frameworks.

Without efficient infrastructure, large data volumes can cause storage limitations, slow processing speeds, and system failures.

2. Handling Data Variety

Data generated in modern systems exists in multiple formats, including structured, semi-structured, and unstructured data. Integrating these diverse data formats into a unified system can be complex.

For example:

- ◆ Structured data may come from relational databases.
- ◆ Semi-structured data may include JSON or XML files.
- ◆ Unstructured data may include images, audio, video, and text documents.

Managing and integrating these different formats requires flexible data processing systems and advanced data transformation techniques.

3. Ensuring Data Quality

Poor data quality can lead to incorrect insights and unreliable analytical results. Data collected from multiple sources may contain missing values, duplicate records, inconsistent formats, or inaccurate entries.

Data engineers must implement validation mechanisms, cleaning processes, and transformation techniques to maintain high-quality datasets that can be trusted for analysis.

4. Scalability and System Performance

As organizations grow, the amount of data they generate increases rapidly. Data engineering systems must scale efficiently to handle increasing workloads without affecting performance.

Achieving scalability requires the use of distributed computing technologies, cloud infrastructure, and parallel data processing frameworks.

5. Data Security and Privacy

Modern data systems often store sensitive information such as personal details, financial transactions, and medical records. Protecting this data from unauthorized access is a critical responsibility.

Data engineers must implement security mechanisms such as encryption, authentication, access control, and monitoring systems. In addition, organizations must comply with data protection regulations and privacy policies.

6. Maintaining System Reliability

Data engineering systems typically involve complex architectures composed of multiple components such as data pipelines, storage systems, processing frameworks, and analytics platforms. Failures in any component may disrupt the entire data workflow.

To ensure reliability, data engineering systems must include fault tolerance mechanisms, backup systems, and redundancy strategies that allow the system to continue functioning even when certain components fail.

Table 1.1.4 Common challenges in data engineering

Challenge	Description	Impact
Large Data Volumes	Managing rapidly growing datasets	Storage and processing limitations
Data Variety	Integrating multiple data formats	Increased system complexity
Data Quality Issues	Inconsistent or incomplete data	Inaccurate analysis
Scalability	Handling increasing data workloads	Performance degradation
Data Security	Protecting sensitive information	Risk of data breaches
System Reliability	Ensuring uninterrupted data pipelines	Operational disruptions

In summary, data engineering plays a crucial role in enabling modern organizations to utilize data effectively. However, the growing scale and complexity of data systems introduce significant challenges that must be addressed through robust architectures, advanced technologies, and well-designed data management strategies.

Recap

- ◆ Data engineering plays a crucial role in supporting data-driven organizations by transforming raw data into structured and usable information for analysis.
- ◆ Data engineers design and manage data pipelines that collect, process, and store data from multiple sources efficiently.
- ◆ Modern organizations depend on data engineering systems to enable reliable data storage, integration, and processing across large-scale data environments.
- ◆ Data engineering ensures data quality and consistency, which are essential for accurate analytics and decision-making.
- ◆ Data engineers support analytics, business intelligence, and machine learning applications by preparing well-structured datasets.
- ◆ Despite its importance, data engineering faces several challenges such as large data volumes, diverse data formats, scalability issues, and data security concerns.
- ◆ Maintaining system reliability, fault tolerance, and high performance is essential for ensuring continuous and efficient data processing.

Objective Questions

1. What is meant by a data-driven organization?
2. What is the primary responsibility of a data engineer?
3. What is a data pipeline?
4. Which process is commonly used for data extraction, transformation, and loading?
5. Why is data quality important in data engineering?
6. Name two common storage platforms used in modern data systems.
7. What type of data includes formats such as images, videos, and text documents?
8. What challenge arises due to the rapid growth of data in organizations?
9. Which challenge involves protecting sensitive information from unauthorized access?

10. What mechanism ensures system operation even when some components fail?

Answers

1. An organization that makes decisions based on data analysis
2. Designing and managing data infrastructure and pipelines
3. A system that moves and processes data from sources to storage and analytics platforms
4. ETL (Extract, Transform, Load)
5. It ensures accurate and reliable analysis
6. Data warehouses and data lakes
7. Unstructured data
8. Managing large data volumes
9. Data security and privacy
10. Fault tolerance

Assignments

1. Explain the role of data engineering in supporting data-driven organizations.
2. Discuss the importance of data pipelines in modern data engineering systems.
3. Describe the major challenges faced in data engineering environments.
4. Explain how data engineers ensure data quality and reliability in large-scale systems.
5. Discuss the importance of scalability and fault tolerance in data engineering.

Reference

1. Apache Software Foundation. Apache Hadoop Documentation.
2. Kleppmann, M. (2017). Designing Data-Intensive Applications. O'Reilly Media.
3. Reis, J., & Housley, M. (2022). Fundamentals of Data Engineering. O'Reilly Media.

Suggested Reading

1. Kleppmann, M. (2017). Designing Data-Intensive Applications. O'Reilly Media.
2. Reis, J., & Housley, M. (2022). Fundamentals of Data Engineering. O'Reilly Media.
3. White, T. (2015). Hadoop: The Definitive Guide. O'Reilly Media.

UNIT 2

Foundations of Data Engineering and Python for Data Workflows

Learning Outcomes

At the end of this unit, the learner will be able to :

- ◆ explain the concept and stages of the data engineering lifecycle
- ◆ describe the importance of data engineering in modern data-driven systems
- ◆ understand the role of Python in data engineering workflows
- ◆ identify essential Python libraries used for data processing and analysis\
- ◆ demonstrate awareness of how Python interacts with databases, APIs, and Big Data tools

Prerequisite

As organizations increasingly depend on data to guide business decisions, the importance of efficient data collection, transformation, and storage has become essential. Data engineering serves as a fundamental component in converting raw data into a structured and usable form for analysis. Without a proper understanding of how data moves through systems from its creation to its final use, it becomes challenging to recognize the significance of the tools and technologies employed in modern data environments.

Learners entering this unit are expected to have a basic understanding of data concepts such as structured and unstructured data, as well as familiarity with traditional data processing systems discussed in the previous unit. An introductory exposure to programming concepts will also be helpful, as this unit introduces Python as a key tool for handling data workflows.

This unit builds a bridge between conceptual understanding of Big Data systems and practical tools used to implement data pipelines, thereby preparing learners for more advanced topics in data engineering and analytics.



Keywords

Data Engineering, Data Pipeline, ETL, Data Lifecycle, Python, Pandas, NumPy, APIs, Databases, PySpark, Data Processing

Discussion

1.2.1 Foundations of Data Engineering and Python for Data Workflows

In the modern digital era, organizations generate vast amounts of data from multiple sources such as web applications, mobile devices, sensors, and enterprise systems. However, raw data in its original form is often incomplete, inconsistent, and unsuitable for direct analysis. This creates a need for systematic processes that can transform raw data into meaningful and usable information. This is where the field of data engineering plays a vital role.

Data engineering focuses on designing and managing systems that enable the efficient collection, storage, processing, and delivery of data. It forms the backbone of data-driven decision-making by ensuring that high-quality data is available for analytics, machine learning, and business intelligence applications. Without robust data engineering practices, organizations would struggle to extract value from their data.

At the core of data engineering lies the concept of data workflows, which refer to the sequence of steps involved in moving and transforming data from its source to its final destination. These workflows ensure that data is properly ingested, cleaned, processed, and made available for analysis. As data volumes and complexity increase, managing these workflows efficiently becomes increasingly important.

In recent years, Python has emerged as one of the most widely used programming languages in data engineering. Its simplicity, readability, and extensive ecosystem of libraries make it an ideal choice for building data pipelines and automating data workflows. Python enables data engineers to interact with databases, process large datasets, integrate with APIs, and work with distributed computing frameworks.

The integration of Python into data engineering workflows has significantly improved productivity and flexibility. Instead of relying on multiple specialized tools, data engineers can use Python as a unified platform to perform various tasks such as data extraction, transformation, loading (ETL), and analysis. This has made Python an essential skill for professionals working in data-related fields.

Furthermore, the combination of data engineering principles and Python-based tools supports scalable and efficient data processing. As organizations continue to adopt Big Data technologies, the demand for professionals who understand both the conceptual

foundations of data engineering and practical implementation using Python continues to grow.

This section provides a foundational understanding of how data engineering and Python work together to support modern data workflows. The subsequent sections will explore the data engineering lifecycle in detail, followed by practical aspects of using Python for data-related tasks.

1.2.2 Data Engineering Lifecycle (Overview & Stages)

In modern data-driven environments, data does not become useful immediately after it is generated. Instead, it passes through a structured sequence of processes that transform it from raw, unorganized input into meaningful insights. This structured sequence is known as the data engineering lifecycle.

Understanding this lifecycle is essential because it helps learners visualize how data flows across systems and how different stages contribute to the overall goal of extracting value from data. Each stage in the lifecycle plays a specific role, and failure in any stage can affect the quality and reliability of the final output.

1.2.2.1 Overview of the Data Engineering Lifecycle

The data engineering lifecycle represents a pipeline through which data moves from its source to its final use in analytics or decision-making systems. Unlike traditional systems, where data processing is often isolated, modern data engineering follows a continuous and integrated workflow.

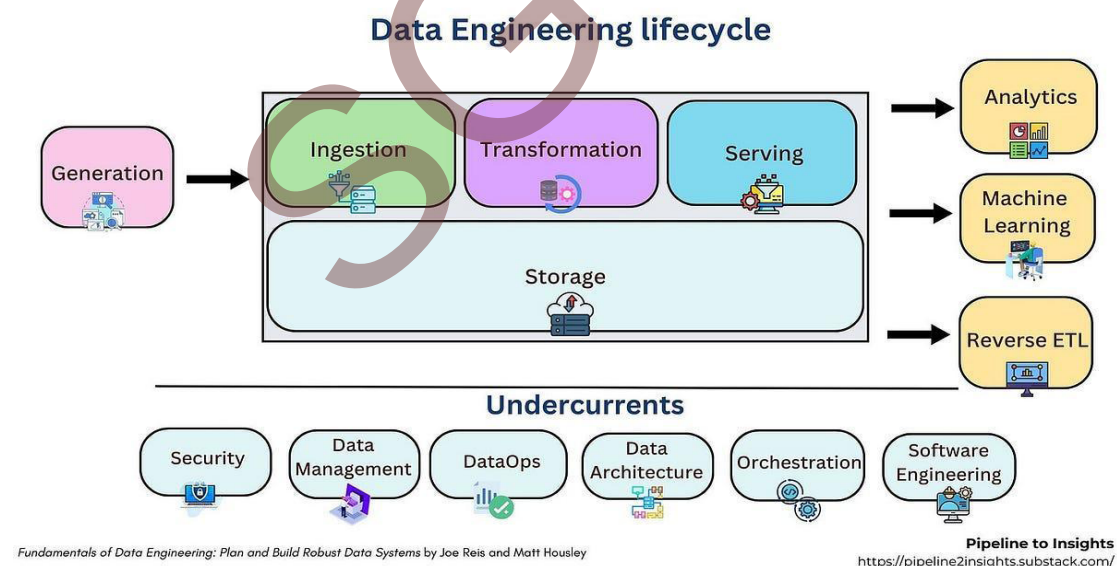


Fig. 1.2.1 Data engineering lifecycle

As illustrated in Figure 1.2.1, the lifecycle typically includes stages such as data generation, ingestion, storage, processing, and serving. These stages are interconnected, forming a pipeline that ensures smooth and efficient data flow.

The lifecycle is not always linear; in many real-world systems, it operates continuously with feedback loops and real-time updates.

1.2.2.2 Stages of the Data Engineering Lifecycle

The data engineering lifecycle can be divided into several key stages. Each stage contributes to transforming raw data into valuable insights.

1. Data Generation

Data generation is the starting point of the lifecycle. Data is produced from multiple sources such as:

- ◆ Social media platforms
- ◆ E-commerce transactions
- ◆ IoT devices and sensors
- ◆ Enterprise applications
- ◆ Log files and system records

This data can be structured, semi-structured, or unstructured, making it complex to handle using traditional systems.

2. Data Ingestion

Data ingestion involves collecting data from various sources and transferring it into storage systems. This stage ensures that data is captured efficiently without loss.

There are two main types of ingestion:

- ◆ **Batch Ingestion** – Data is collected at intervals
- ◆ **Real-time (Streaming) Ingestion** – Data is processed continuously as it is generated

3. Data Storage

Once data is ingested, it must be stored in a scalable and reliable system. Modern data engineering uses:

- ◆ Data lakes
- ◆ Data warehouses
- ◆ Distributed storage systems

These storage systems are designed to handle large volumes of diverse data while ensuring availability and fault tolerance.

4. Data Processing

Data processing is the stage where raw data is transformed into a usable format. This includes:

- ◆ Data cleaning (removing errors and inconsistencies)
- ◆ Data transformation (formatting and structuring data)
- ◆ Data integration (combining multiple data sources)

Processing can be performed using batch processing or real-time processing frameworks.

5. Data Serving

In this stage, processed data is made available for use. It can be accessed through:

- ◆ Query systems
- ◆ Dashboards
- ◆ APIs

Data serving ensures that users and applications can retrieve data efficiently.

6. Data Analysis and Visualization

The final stage involves analyzing the processed data to extract insights. Visualization tools help present data in an understandable format such as charts, graphs, and reports.

This stage supports decision-making in areas such as business strategy, healthcare, finance, and technology.

1.2.2.3 Summary of Lifecycle Stages

To better understand the roles of each stage, Table 1.2.1 provides a concise summary.

Table 1.2.1 Stages of the data engineering lifecycle

Stage	Description	Example
Data Generation	Creation of raw data from sources	Social media posts
Data Ingestion	Collecting and importing data	Streaming logs
Data Storage	Storing data in scalable systems	Data lakes
Data Processing	Cleaning and transforming data	ETL pipelines
Data Serving	Providing access to processed data	APIs, dashboards
Data Analysis	Extracting insights	Business reports

1.2.2.4 Importance of the Data Engineering Lifecycle

The data engineering lifecycle is essential for ensuring that data is reliable, scalable, and ready for analysis. Its importance can be understood through the following aspects:



- ◆ **Improved Data Quality:** Ensures accurate and consistent data
- ◆ **Scalability:** Supports handling of large data volumes
- ◆ **Efficiency:** Automates data workflows and reduces manual effort
- ◆ **Real-time Processing:** Enables immediate insights
- ◆ **Better Decision-Making:** Provides high-quality data for analysis

Without a well-defined lifecycle, organizations may face issues such as data inconsistency, processing delays, and unreliable insights.

1.2.2.5 Role of ETL in the Lifecycle

A key concept within the data engineering lifecycle is ETL (Extract, Transform, Load):

- ◆ **Extract:** Collect data from sources
- ◆ **Transform:** Clean and process data
- ◆ **Load:** Store data into systems

ETL forms the backbone of many data pipelines and ensures smooth data flow across stages.

1.2.3 Introduction to Python for Data Engineering Tasks

In modern data engineering environments, managing large volumes of data efficiently requires not only robust architectures but also powerful programming tools. Among the various programming languages available, **Python** has emerged as a dominant choice for implementing data engineering workflows due to its simplicity, flexibility, and extensive ecosystem.

Python enables data engineers to design, automate, and manage data pipelines that handle tasks such as data extraction, transformation, storage, and integration. Its readability and ease of use make it suitable for both beginners and experienced professionals, while its scalability allows it to be used in complex Big Data environments.

1.2.3.1 Role of Python in Data Engineering

Python plays a central role in almost every stage of the data engineering lifecycle. It serves as a versatile tool that connects data sources, processing systems, and analytical platforms.

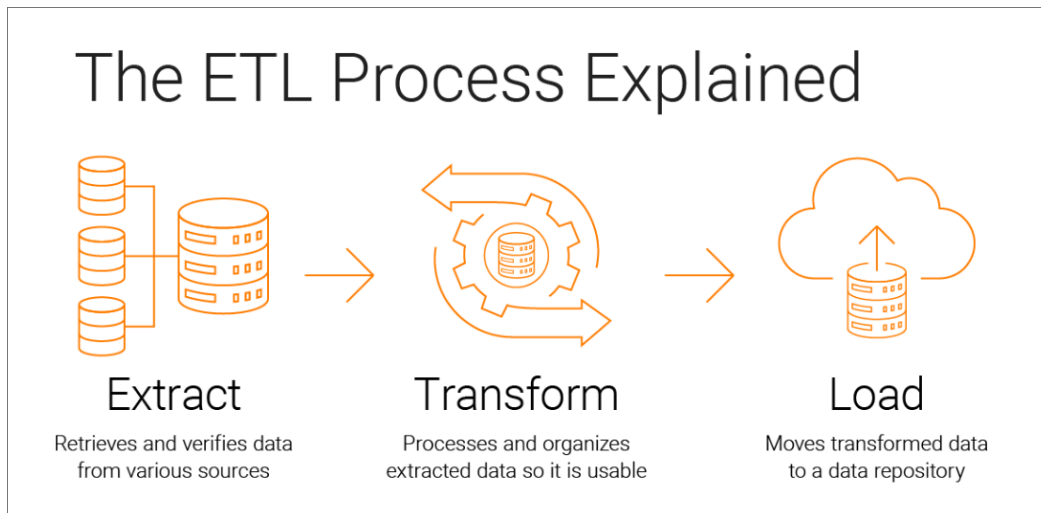


Fig. 1.2.2 Role of python in data engineering workflow

As illustrated in Figure 1.2.2, Python acts as a bridge between different components of the data pipeline. It is used to:

- ◆ Extract data from databases, APIs, and files
- ◆ Transform and clean raw data
- ◆ Load data into storage systems
- ◆ Automate workflows and scheduling tasks
- ◆ Integrate with Big Data frameworks

This flexibility makes Python a preferred language for building end-to-end data solutions.

1.2.3.2 Essential Python Libraries for Data Tasks

One of the key strengths of Python lies in its rich ecosystem of libraries that simplify data engineering tasks. These libraries provide ready-made functionalities for handling different types of data operations.

Table 1.2.2 Essential python libraries for data engineering

Library	Purpose	Example Use
Pandas	Data manipulation and analysis	Cleaning datasets
NumPy	Numerical computations	Handling arrays
Requests	API interaction	Fetching web data
PySpark	Distributed data processing	Big Data analytics
SQLAlchemy	Database interaction	Query execution

These libraries reduce the need to build complex systems from scratch and allow data engineers to focus on solving real-world problems efficiently.

1.2.3.3 Interacting with Databases and APIs

Data engineering workflows often require accessing data from multiple sources, including databases and web services. Python provides powerful tools to interact with these sources seamlessly.

- ◆ **Database Interaction:** Python supports interaction with relational databases such as MySQL, PostgreSQL, and SQLite. Libraries like SQLAlchemy enable querying, inserting, and updating data efficiently.
- ◆ **API Integration:** APIs (Application Programming Interfaces) allow systems to exchange data over the internet. Python libraries such as Requests help in fetching and sending data through APIs.

For example, a data engineer can use Python to extract data from an online service, process it, and store it in a database, all within a single workflow.

1.2.3.4 Working with Big Data Tools (PySpark)

As data size increases, traditional processing tools become insufficient. Python integrates with Big Data frameworks such as **PySpark**, which allows distributed processing of large datasets.

PySpark enables:

- ◆ Parallel processing across multiple nodes
- ◆ Handling of massive datasets
- ◆ Faster computation compared to traditional systems

This makes Python suitable not only for small-scale data tasks but also for enterprise-level Big Data applications.

1.2.3.5 Advantages of Using Python in Data Engineering

Python offers several advantages that make it highly suitable for data engineering:

- ◆ **Ease of Learning:** Simple syntax and readability
- ◆ **Extensive Libraries:** Rich ecosystem for data handling
- ◆ **Flexibility:** Works across different platforms and tools
- ◆ **Scalability:** Supports both small and large-scale data processing
- ◆ **Integration Capability:** Easily connects with databases, APIs, and Big Data tools

1.2.3.6 Python in End-to-End Data Workflows

Python supports the complete data engineering workflow, from data ingestion to data analysis. It can be used to automate pipelines, schedule jobs, and monitor processes.

For example, a typical Python-based workflow may include:

1. Extracting data from an API
2. Cleaning and transforming data using Pandas
3. Storing data in a database
4. Processing large datasets using PySpark
5. Delivering results to analytics tools

This end-to-end capability makes Python an indispensable tool in modern data engineering.

1.2.4 Advantages of Using Python in Data Engineering

In the previous section, we explored how Python supports various stages of the data engineering lifecycle through its libraries and integration capabilities. In this section, we examine why Python has become the preferred programming language for data engineering tasks.

The widespread adoption of Python is not accidental; it is driven by a combination of technical strengths, ease of use, and strong community support. These advantages make Python highly suitable for building scalable and efficient data workflows.

1.2.4.1 Simplicity and Readability

One of the most important advantages of Python is its simple and readable syntax. Unlike many other programming languages, Python code is easy to understand and write, even for beginners.

This simplicity allows data engineers to focus more on solving data-related problems rather than dealing with complex programming constructs. It also improves collaboration among teams, as code can be easily understood and maintained.

1.2.4.2 Extensive Ecosystem of Libraries

Python provides a rich set of libraries that support almost every aspect of data engineering.

These include:

- ◆ Data manipulation libraries (Pandas, NumPy)
- ◆ Database connectivity tools (SQLAlchemy)
- ◆ API handling libraries (Requests)
- ◆ Big Data processing frameworks (PySpark)



This extensive ecosystem eliminates the need to build functionalities from scratch, thereby increasing productivity.

1.2.4.3 Integration with Multiple Systems

Modern data engineering involves working with various systems such as databases, cloud platforms, APIs, and distributed frameworks. Python offers seamless integration with all these components.

1.2.4.4 Scalability and Performance

Although Python is known for its simplicity, it also supports scalable data processing. With frameworks like PySpark and Dask, Python can handle large-scale data processing tasks efficiently.

This allows organizations to use Python for both small-scale data analysis and large-scale Big Data applications.

1.2.4.5 Automation of Data Workflows

Python is widely used for automating repetitive tasks in data engineering. Scripts can be written to:

- ◆ Schedule data ingestion
- ◆ Automate ETL processes
- ◆ Monitor data pipelines
- ◆ Generate reports

Automation reduces manual effort, minimizes errors, and improves efficiency.

1.2.4.6 Strong Community Support

Python has a large and active global community. This provides several benefits:

- ◆ Availability of learning resources
- ◆ Continuous development of libraries
- ◆ Quick resolution of technical issues

Community support ensures that Python remains up-to-date with evolving technologies and industry requirements.

1.2.4.7 Cost-Effectiveness

Python is an open-source programming language, which means it is free to use. Organizations can build powerful data engineering solutions without investing in expensive proprietary software.

1.2.4.8 Summary of Advantages

Table 1.2.3 Advantages of Python in Data Engineering

Advantage	Description
Simplicity	Easy to learn and use
Libraries	Rich ecosystem for data tasks
Integration	Works with multiple systems
Scalability	Supports Big Data processing
Automation	Enables workflow automation
Community Support	Large and active community
Cost	Open-source and free

Recap

- ◆ Data engineering focuses on designing systems for efficient data collection, storage, and processing
- ◆ The data engineering lifecycle includes stages such as generation, ingestion, storage, processing, and analysis
- ◆ Python plays a central role in building and managing data workflows
- ◆ Libraries like Pandas, NumPy, and PySpark simplify data processing tasks
- ◆ Python enables interaction with databases and APIs
- ◆ Integration with Big Data tools allows scalable data processing
- ◆ Python offers advantages such as simplicity, flexibility, and automation capabilities

Objective Questions

1. What is the primary role of data engineering?
2. Which stage of the lifecycle involves collecting data?
3. What does ETL stand for?

4. Which programming language is widely used in data engineering?
5. Name a Python library used for data manipulation.
6. Which tool is used for distributed data processing in Python?
7. What type of ingestion processes data continuously?
8. Which stage involves cleaning and transforming data?
9. What is the purpose of APIs in data engineering?
10. Which feature of Python makes it easy to learn?

Answers

1. Managing and processing data
2. Data ingestion
3. Extract, Transform, Load
4. Python
5. Pandas
6. PySpark
7. Real-time ingestion
8. Data processing
9. Data exchange between systems
10. Simple syntax

Assignments

1. Explain the stages of the data engineering lifecycle in detail.
2. Discuss the role of Python in modern data engineering workflows.
3. Compare batch processing and real-time processing in data pipelines.
4. Describe the advantages of using Python in data engineering.
5. Explain how Python interacts with databases and APIs.

Reference

1. Apache Software Foundation. Python and Big Data Documentation
2. Python Software Foundation. Official Python Documentation

Suggested Reading

1. Kleppmann, M. (2017). *Designing Data-Intensive Applications*. O'Reilly Media.
2. McKinney, W. (2018). *Python for Data Analysis*. O'Reilly Media.
3. Zaharia, M. et al. (2016). *Apache Spark: The Definitive Guide*. O'Reilly Media.

UNIT 3

Data Manipulation with Libraries

Learning Outcomes

At the end of this unit, the learner will be able to :

- ◆ explain the role of NumPy in numerical data processing
- ◆ differentiate between Python lists and NumPy arrays
- ◆ describe the structure and functionality of pandas Series and DataFrames
- ◆ perform data indexing and selection operations using pandas
- ◆ apply common data manipulation, cleaning, and transformation techniques

Prerequisite

Efficient data analysis and processing require tools that can handle large volumes of data with speed and accuracy. While Python provides basic data structures such as lists and dictionaries, these structures are not optimized for large-scale numerical computation or structured data manipulation. This limitation creates the need for specialized libraries that can perform data operations efficiently.

Learners entering this unit are expected to have a basic understanding of Python programming concepts, including variables, lists, and functions. Familiarity with data types and basic operations will help in understanding how advanced libraries like NumPy and pandas extend Python's capabilities.

This unit introduces two of the most powerful Python libraries used in data engineering and analytics. These libraries provide the foundation for handling structured datasets, enabling efficient data manipulation, cleaning, and transformation.

Keywords

NumPy, Array, List, pandas, Series, DataFrame, Indexing, Selection, Data Cleaning, Data Transformation, Data Manipulation

Discussion

1.3.1 Overview of NumPy

In data engineering and analytics, handling numerical data efficiently is a fundamental requirement. While Python provides basic data structures such as lists, these structures are not optimized for high-performance numerical computations. As datasets grow in size and complexity, operations using standard Python lists become slower and less efficient.

To address these limitations, the NumPy (Numerical Python) library was developed. NumPy provides powerful data structures and functions that enable efficient storage and processing of numerical data. It forms the foundation for many other data science and data engineering libraries, including pandas.

NumPy introduces the concept of arrays, which are more efficient than traditional lists in terms of both memory usage and computational speed. These arrays allow operations to be performed on entire datasets simultaneously, rather than element by element. This approach, known as vectorization, significantly improves performance.

In modern data workflows, NumPy is widely used for:

- ◆ Performing mathematical and statistical operations
- ◆ Handling multi-dimensional datasets
- ◆ Supporting data preprocessing tasks
- ◆ Acting as the underlying structure for advanced libraries

Thus, understanding NumPy is essential for learners who wish to work with data efficiently in Python-based environments.

1.3.1.1 NumPy Array and Python List

A key feature that distinguishes NumPy from standard Python is its use of arrays instead of lists. While both arrays and lists are used to store collections of data, they differ significantly in performance, structure, and functionality.

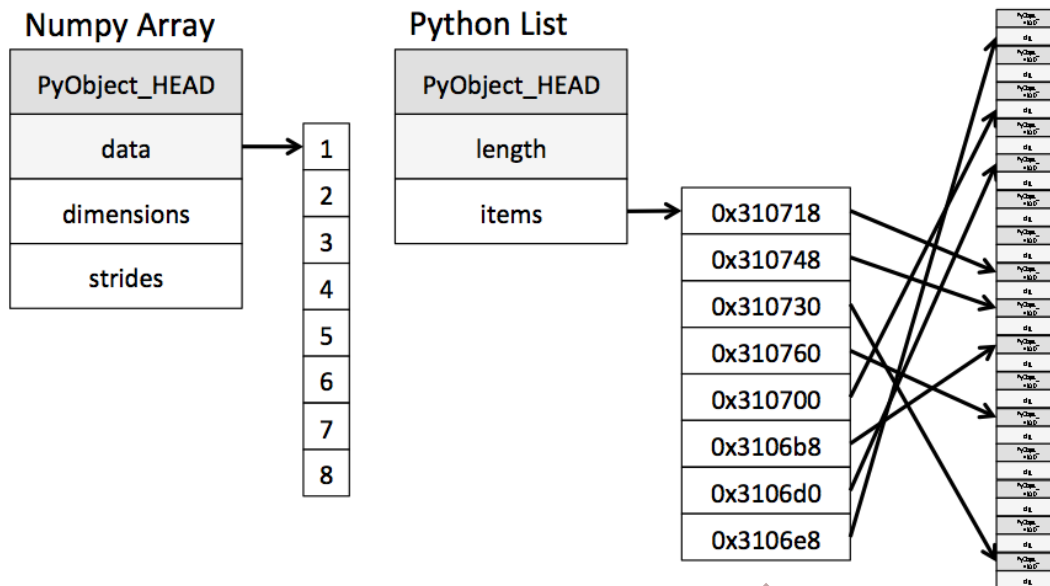


Fig. 1.3.1 NumPy Array vs Python List

A Python list is a general-purpose data structure that can store elements of different data types. While this flexibility is useful, it reduces efficiency when performing numerical computations.

In contrast, a NumPy array is designed specifically for numerical operations. It stores elements of the same data type and allows operations to be performed on all elements simultaneously.

Key Differences

Table 1.3.1 Python List vs NumPy Array

Feature	Python List	NumPy Array
Data Type	Heterogeneous	Homogeneous
Performance	Slower	Faster
Memory Usage	High	Efficient
Operations	Loop-based	Vectorized
Use Case	General purpose	Numerical computing

Explanation

- ◆ **Homogeneity:** NumPy arrays store elements of the same type, enabling optimized memory usage.
- ◆ **Vectorization:** Operations such as addition or multiplication are applied to entire arrays at once.
- ◆ **Performance:** NumPy uses optimized C-based implementations, making it significantly faster.

Consider the example of adding two datasets:

- ◆ Using lists → requires looping through each element
- ◆ Using NumPy → single operation applied to entire array

This difference becomes significant when dealing with large datasets.

1.3.2 Overview of pandas

While NumPy provides powerful support for numerical computations, real-world data is often more complex and structured in tabular formats such as spreadsheets, databases, and CSV files. Handling such structured data efficiently requires more advanced tools that can manage rows, columns, labels, and heterogeneous data types. This need is addressed by the pandas library.

pandas is a widely used Python library designed for data manipulation and analysis. It builds upon NumPy and extends its capabilities by introducing flexible and user-friendly data structures that simplify working with structured datasets.

In modern data engineering workflows, pandas plays a crucial role in:

- ◆ Organizing data into meaningful structures
- ◆ Performing data cleaning and preprocessing
- ◆ Transforming and reshaping datasets
- ◆ Supporting exploratory data analysis

Its intuitive design allows users to perform complex operations with minimal code, making it an essential tool for both beginners and professionals.

1.3.2.1 Introduction to pandas Data Structures

The core strength of pandas lies in its two primary data structures: **Series** and **DataFrame**. These structures allow users to represent and manipulate data efficiently.

Series			Series			DataFrame	
	apples			oranges			
0	3	+	0	0	=	0	3
1	2		1	3		1	2
2	0		2	7		2	0
3	1		3	2		3	1
							oranges
						0	3
						7	2

Fig.1.3.2 pandas Data Structures (Series and Data Frame)

Series

A Series is a one-dimensional labeled array capable of holding data of any type (integer, float, string, etc.). It can be thought of as a single column of data with an associated index.

Key characteristics:

- ◆ One-dimensional structure
- ◆ Contains labels (index)
- ◆ Supports various data types
- ◆ Suitable for representing a single variable

DataFrame

A DataFrame is a two-dimensional labeled data structure, similar to a table in a database or a spreadsheet. It consists of rows and columns, where each column can contain different data types.

Key characteristics:

- ◆ Two-dimensional structure
- ◆ Rows and columns with labels
- ◆ Supports heterogeneous data
- ◆ Most commonly used pandas structure

Comparison of Series and DataFrame

Table 1.3.2 Series vs DataFrame

Feature	Series	DataFrame
Dimension	One-dimensional	Two-dimensional
Structure	Single column	Multiple columns
Data Type	Homogeneous or mixed	Column-wise heterogeneous
Use Case	Individual variable	Complete dataset

Why pandas is Important

pandas simplifies data handling by providing:

- ◆ Efficient data alignment and indexing
- ◆ Built-in functions for data cleaning
- ◆ Easy handling of missing data
- ◆ Integration with NumPy and other libraries

For example, instead of writing complex loops to filter or transform data, pandas allows users to perform these operations using simple and readable commands.

Conceptual Understanding

Consider a dataset containing student records:

ID	Name	Marks
1	Asha	85
2	Ravi	90

- ◆ Each column → Series
- ◆ Entire table → DataFrame

This structure makes it easier to manipulate and analyze data.

1.3.3 Indexing and Selection Techniques

In data analysis and engineering, working with large datasets requires efficient methods to access, retrieve, and manipulate specific portions of data. Simply storing data in a structured format is not sufficient; users must be able to quickly locate and extract relevant information. This is achieved through indexing and selection techniques in pandas.

Indexing allows users to identify and access data using labels or positions, while selection enables filtering and retrieving subsets of data based on conditions. These operations are essential for performing data analysis, cleaning, and transformation tasks efficiently.

1.3.3.1 Concept of Indexing in pandas

In pandas, every data structure (Series or DataFrame) is associated with an index, which acts as a label for identifying data elements. In a DataFrame, rows are labeled using an index, while columns are identified by their names.

As shown in Figure 1.3.3, a DataFrame consists of rows and columns, where each element can be accessed using its corresponding row and column labels.

1. Label-Based Indexing (.loc)

Label-based indexing uses row and column labels to access data. This method is particularly useful when the dataset contains meaningful labels.

For example, a row can be selected using its label, or a column can be accessed using its name. This method improves readability and clarity in data operations.

2. Integer-Based Indexing (.iloc)

Integer-based indexing uses numerical positions of rows and columns. As seen in **Figure 1.3.3**, positions start from zero and increase sequentially.

This method is useful when the exact position of data is known rather than its label.

3. Boolean Indexing

Boolean indexing allows selection of data based on conditions. Instead of specifying positions or labels, users define a condition, and only the rows satisfying that condition are returned.

This method is widely used in real-world data analysis for filtering datasets.

1.3.3.3 Comparison of Indexing Methods

Table 1.3.3 Comparison of pandas Indexing Techniques

Method	Type	Use Case
.loc	Label-based	Access using names
.W	Position-based	Access using index numbers
Boolean Indexing	Condition-based	Filtering data

1.3.3.4 Data Selection Operations

Data selection involves extracting specific parts of a dataset for analysis. Using the indexing techniques illustrated in Figure 1.3.3, users can perform operations such as:

- ◆ Selecting a single column
- ◆ Selecting multiple columns
- ◆ Extracting specific rows
- ◆ Filtering data based on conditions

These operations allow users to focus only on relevant data, improving both efficiency and clarity.

1.3.3.5 Importance of Indexing and Selection

Efficient indexing and selection provide several advantages:

- ◆ Faster data access
- ◆ Improved performance
- ◆ Simplified data manipulation
- ◆ Enhanced readability

Without proper indexing techniques, handling large datasets would become time-consuming and inefficient.

1.3.4 Common Data Manipulation Tasks in pandas

In real-world data engineering and analytics, raw data is rarely in a form that can be directly used for analysis. It often requires modification, restructuring, and summarization. These operations are collectively known as data manipulation.

pandas provides a rich set of functions that allow users to efficiently manipulate datasets. These operations help transform raw data into a structured and meaningful format, enabling better analysis and decision-making.

As illustrated in Figure 1.3.4, data manipulation involves multiple operations such as filtering, sorting, grouping, and merging, which are applied sequentially as part of a data workflow.



Fig. 1.3.4 Data manipulation workflow in pandas

1.3.4.1 Sorting Data

Sorting is one of the most basic yet essential data manipulation tasks. It involves arranging data in a specific order, either ascending or descending.

Using pandas, datasets can be sorted based on one or more columns. As part of the workflow shown in Figure 1.3.4, sorting helps organize data before further processing.

For example:

- ◆ Sorting student records based on marks
- ◆ Ordering transactions by date

Sorting improves readability and helps identify patterns in data.

1.3.4.2 Filtering Data

Filtering involves selecting a subset of data based on specific conditions. This is closely related to Boolean indexing discussed earlier.

As illustrated in Figure 1.3.4, filtering allows users to extract only relevant data from large datasets.

Examples include:

- ◆ Selecting records where marks are greater than a threshold
- ◆ Extracting data for a specific category

Filtering reduces data size and focuses analysis on important information.

1.3.4.3 Grouping Data

Grouping is used to organize data into categories based on one or more attributes. After grouping, aggregation operations can be performed on each group.

For example:

- ◆ Grouping students by department
- ◆ Grouping sales data by region

As part of the workflow in Figure 1.3.4, grouping enables summarization and comparative analysis across categories.

1.3.4.4 Aggregation Operations

Aggregation involves performing calculations on grouped data. Common aggregation functions include:

- ◆ Sum
- ◆ Mean

- ◆ Count
- ◆ Minimum and Maximum

These operations help summarize large datasets into meaningful statistics.

For instance:

- ◆ Calculating average marks of students
- ◆ Finding total sales per region

Aggregation plays a key role in data analysis and reporting.

1.3.4.5 Merging and Joining Data

In many real-world scenarios, data is stored across multiple datasets. pandas allows combining these datasets using merging and joining operations.

As shown in Figure 1.3.4, merging integrates data from different sources into a single dataset.

Examples include:

- ◆ Combining customer data with transaction records
- ◆ Joining employee details with department information

This operation is essential in building comprehensive datasets.

1.3.4.6 Reshaping Data

Reshaping involves changing the structure of data without altering its content. This includes:

- ◆ Pivoting data
- ◆ Transforming rows into columns
- ◆ Flattening hierarchical data

Reshaping helps in organizing data for better visualization and analysis.

1.3.4.7 Summary of Data Manipulation Tasks

Table 1.3.4 Common Data Manipulation Tasks in pandas

Operation	Purpose	Example
Sorting	Arrange data	Sort by marks
Filtering	Select subset	Marks > 80

Grouping	Categorize data	Group by department
Aggregation	Summarize data	Average marks
Merging	Combine datasets	Join tables
Reshaping	Change structure	Pivot table

1.3.4.8 Importance of Data Manipulation

Data manipulation is essential because:

- ◆ Raw data is often incomplete or unorganized
- ◆ Analysis requires structured and clean datasets
- ◆ Efficient manipulation improves performance
- ◆ It enables meaningful insights from data

Without proper data manipulation, even large datasets may fail to provide useful information.

1.3.5 Data Cleaning and Transformation

In real-world data environments, raw data is often incomplete, inconsistent, and noisy. Such data cannot be directly used for analysis or decision-making. Therefore, it is essential to preprocess the data to improve its quality and structure. This process involves data cleaning and data transformation.

Data cleaning focuses on identifying and correcting errors, while data transformation involves converting data into a suitable format for analysis. Together, these processes ensure that datasets are accurate, consistent, and ready for further processing.

As illustrated in Figure 1.3.5, data cleaning and transformation are key stages in the data preparation pipeline, bridging raw data and meaningful insights.

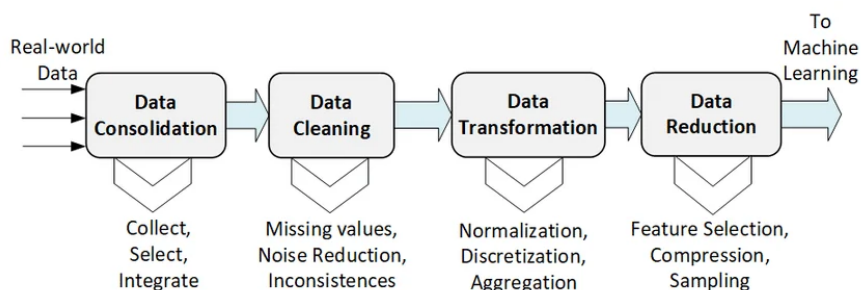


Fig. 1.3.5 Data cleaning and transformation process

1.3.5.1 Data Cleaning

Data cleaning is the process of detecting and correcting errors or inconsistencies in data. As shown in Figure 1.3.5, cleaning is the first step in preparing data for analysis.

Common data quality issues include:

- ◆ Missing values
- ◆ Duplicate records
- ◆ Incorrect or inconsistent data
- ◆ Outliers and noise

Handling Missing Values

Missing data is a common problem in datasets. pandas provides methods to:

- ◆ Remove missing values
- ◆ Replace missing values with default or calculated values

For example:

- ◆ Replacing missing marks with the average value
- ◆ Removing rows with incomplete data

Removing Duplicates

Duplicate records can lead to incorrect analysis. Data cleaning ensures that repeated entries are identified and removed.

Correcting Inconsistent Data

Data may contain inconsistencies such as different formats or incorrect entries. For example:

- ◆ Different date formats
- ◆ Spelling variations
- ◆ Incorrect numerical values

Cleaning ensures uniformity across the dataset.

1.3.5.2 Data Transformation

After cleaning, data must often be transformed into a suitable format for analysis. As depicted in Figure 1.3.5, transformation reshapes and standardizes data.

Data Type Conversion

Data may need to be converted into appropriate types, such as:

- ◆ String to integer
- ◆ Integer to float
- ◆ Date strings to datetime format

This ensures correct processing and analysis.

Normalization and Scaling

Normalization adjusts values to a common scale without distorting relationships. This is important in:

- ◆ Machine learning models
- ◆ Comparative analysis

Feature Creation

New features can be created from existing data to enhance analysis.

Examples:

- ◆ Calculating total marks from subject-wise scores
- ◆ Extracting year from date

Encoding Categorical Data

Categorical data (e.g., gender, category) must often be converted into numerical format for processing.

Table 1.3.5 Data cleaning and transformation tasks

Task	Description	Example
Missing Value Handling	Fill or remove missing data	Replace null values
Duplicate Removal	Eliminate repeated records	Remove duplicate rows
Data Correction	Fix inconsistencies	Standardize formats
Type Conversion	Change data types	String to integer
Normalization	Scale values	Normalize marks
Feature Creation	Derive new attributes	Total score

1.3.5.4 Importance of Data Cleaning and Transformation

Data cleaning and transformation are essential for the following reasons:

- ◆ **Improves Data Quality:** Ensures accuracy and consistency
- ◆ **Enhances Reliability:** Reduces errors in analysis
- ◆ **Supports Better Decision-Making:** Provides trustworthy insights
- ◆ **Prepares Data for Advanced Processing:** Essential for machine learning and analytics

Without proper data preparation, even advanced analytical tools may produce incorrect or misleading results.

Recap

- ◆ NumPy is a fundamental library for numerical computing in Python
- ◆ NumPy arrays are more efficient than Python lists in terms of speed and memory
- ◆ pandas provides powerful data structures such as Series and DataFrame
- ◆ Indexing techniques like .loc, .iloc, and Boolean indexing enable efficient data access
- ◆ Data manipulation tasks include sorting, filtering, grouping, aggregation, and merging
- ◆ Data cleaning ensures removal of errors, duplicates, and inconsistencies
- ◆ Data transformation converts data into suitable formats for analysis
- ◆ Proper data preparation is essential for reliable insights and decision-making

Objective Questions

1. What is NumPy primarily used for?
2. What type of data structure does NumPy provide?
3. What is the main purpose of pandas?
4. What is a pandas DataFrame?
5. Which method is used for label-based indexing?

6. Which method is used for position-based indexing?
7. What is Boolean indexing used for?
8. Name one data manipulation operation in pandas.
9. What is the purpose of data cleaning?
10. What is data transformation?
11. What does normalization do?
12. Which library is used for tabular data manipulation in Python?
13. What is aggregation in pandas?
14. What is the purpose of merging datasets?
15. What type of data does a Series represent?

Answers

1. Numerical computing
2. Array
3. Data manipulation and analysis
4. Two-dimensional data structure
5. .loc
6. .iloc
7. Filtering data
8. Sorting
9. Removing errors and inconsistencies
10. Converting data into suitable format
11. Scales data to a common range
12. pandas
13. Performing calculations on data
14. Combining datasets
15. One-dimensional data



Assignments

1. Explain the advantages of NumPy arrays over Python lists.
2. Describe the structure and functionality of pandas DataFrame.
3. Discuss different indexing techniques used in pandas with examples.
4. Explain common data manipulation tasks in pandas.
5. Describe the importance of data cleaning and transformation in data analysis.

Reference

1. Python Software Foundation. NumPy Documentation
2. Python Software Foundation. pandas Documentation

Suggested Reading

1. McKinney, W. (2018). *Python for Data Analysis*. O'Reilly Media.
2. VanderPlas, J. (2016). *Python Data Science Handbook*. O'Reilly Media.
3. Grus, J. (2019). *Data Science from Scratch*. O'Reilly Media.

UNIT 4

File I/O Operations

Learning Outcomes

At the end of this unit, the learner will be able to :

- ◆ explain the concept of file input and output operations
- ◆ describe different file modes used in Python
- ◆ demonstrate how to read data from files
- ◆ explain how to write and append data to files
- ◆ understand practical applications of file handling in data workflows

Prerequisite

In data engineering and analytics, data is often stored externally in files such as text files, CSV files, and logs. While previous units focused on processing data within Python using libraries like NumPy and pandas, it is equally important to understand how data is read from and written to external sources.

Learners entering this unit should have a basic understanding of Python programming, including variables, data types, and functions. Familiarity with structured data formats such as tables and datasets will also help in understanding file operations.

This unit introduces file handling concepts that act as a bridge between raw data storage and data processing, enabling learners to build complete data workflows.

Keywords

File I/O, File Handling, Read, Write, Append, File Modes, Text File, CSV File, Data Storage, Python File Operations



Discussion

1.4.1 Basics of File I/O Operations

In modern data engineering and analytics, data is rarely stored directly within programs. Instead, it is stored externally in files such as text files, CSV files, logs, and datasets. To work with such data, programs must be able to read data from files and write data back to them. This process is known as File Input/Output (I/O).

File I/O operations form a fundamental part of data workflows because they enable interaction between a program and external storage systems. Without file handling, it would be impossible to persist data, share data between systems, or process large datasets efficiently.

As illustrated in Figure 1.4.1, file I/O acts as a bridge between data storage and data processing systems. Data is read from files into the program, processed, and then written back to files for storage or further use.

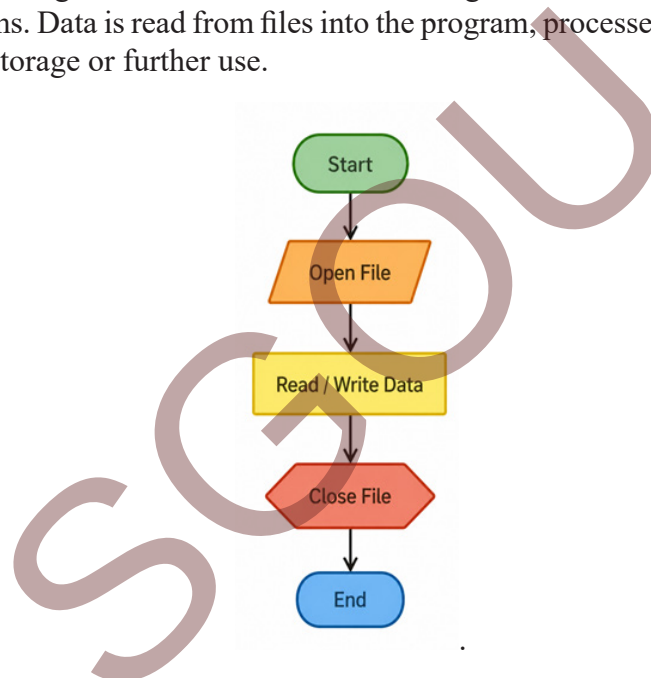


Fig. 1.4.1 File input and output process

Understanding File I/O

File I/O operations involve two main activities:

- ◆ **Input (Reading):** Retrieving data from a file into the program
- ◆ **Output (Writing):** Saving data from the program into a file

These operations allow data to move between memory and storage, ensuring that information is not lost when a program ends.

Importance of File I/O

File I/O operations are essential for several reasons:

- ◆ **Data Persistence:** Data can be stored permanently for future use
- ◆ **Handling Large Data:** Large datasets can be processed without storing them in memory
- ◆ **Data Sharing:** Files allow data exchange between different systems
- ◆ **Automation:** Enables automated data pipelines and workflows

Real-World Context

Consider an e-commerce system:

- ◆ Customer data is stored in files
- ◆ Orders are read from transaction files
- ◆ Reports are written back to files

Without file I/O operations, such systems would not function effectively.

1.4.1.1 Types of Files and File Modes

In file handling, it is important to understand both the types of files used for storing data and the modes in which these files are accessed. These concepts determine how data is read, written, and managed within a program.

As introduced in Figure 1.4.1, file I/O operations involve interaction between storage and processing systems. The type of file and the mode of access together define how this interaction takes place.

Types of Files

Files used in data engineering and programming can be broadly classified into two main categories:

1. Text Files

Text files store data in a human-readable format. Each piece of data is stored as a sequence of characters.

Common examples include:

- ◆ **.txt** files
- ◆ **.csv** (Comma-Separated Values) files
- ◆ **.log** files

As commonly used in data workflows, text files allow easy inspection and modification of data. For instance, CSV files are widely used to store tabular datasets.



2. Binary Files

Binary files store data in a non-readable format, typically as sequences of bytes. These files are not meant to be interpreted directly by humans.

Examples include:

- ◆ Image files (.jpg, .png)
- ◆ Audio files (.mp3)
- ◆ Compiled program files

Binary files are more efficient for storing complex data but require specific programs to interpret them.

Comparison of File Types

Table 1.4.1 Text files vs binary files

Feature	Text Files	Binary Files
Format	Human-readable	Machine-readable
Examples	.txt, .csv	.jpg, .mp3
Ease of Use	Easy to edit	Requires special tools
Storage Efficiency	Less efficient	More efficient

File Modes in Python

When opening a file in Python, a file mode must be specified. This mode determines how the file will be accessed—whether for reading, writing, or appending.

Table 1.4.2 File modes in python

Mode	Description
r	Read mode (default)
w	Write mode (creates/overwrites file)
a	Append mode (adds data to existing file)
rb	Read binary mode
wb	Write binary mode

Understanding File Modes

- ◆ **Read Mode (r):** Used to read data from an existing file. If the file does not exist, an error occurs.

- ◆ **Write Mode (w):** Used to create a new file or overwrite an existing file.
- ◆ **Append Mode (a):** Used to add data to the end of an existing file without deleting previous content.
- ◆ **Binary Modes (rb, wb):** Used when working with non-text files such as images or audio.

Conceptual Flow

As shown earlier in **Figure 1.4.1**, when a file is opened:

- ◆ The **type of file** determines how data is stored
- ◆ The **mode** determines how data is accessed

Together, they control the behavior of file operations.

Real-World Example

Consider a student database stored in a CSV file:

- ◆ Using **read mode (r)** → load student data
- ◆ Using **write mode (w)** → create a new report file
- ◆ Using **append mode (a)** → add new student records

1.4.2 Reading Data from Files

In data engineering and analytics, data is often stored externally in files such as text files, CSV files, and logs. To process this data, it must first be read into the program. Reading data from files is therefore a fundamental operation that enables data analysis, transformation, and decision-making.

As illustrated earlier in Figure 1.4.1, reading involves transferring data from storage (files) into memory, where it can be processed by a program. This step forms the starting point of most data workflows.

1.4.2.1 Process of Reading Files

Reading a file in Python typically involves three steps:

1. Opening the file in read mode (r)
2. Accessing the contents of the file
3. Closing the file after reading

This process ensures that system resources are managed efficiently and that data is accessed correctly.

1.4.2.2 Methods for Reading Data

Python provides several methods to read data from files. These methods offer flexibility depending on how the data needs to be processed.

1. Reading the Entire File

The entire content of the file is read at once. This method is suitable for small files.

Conceptual Example:

Name,Marks

Asha,85

Ravi,90

This approach loads all data into memory in a single operation.

2. Reading Line by Line

Instead of loading the entire file, data can be read one line at a time. This is useful for large files.

- ◆ Reduces memory usage
- ◆ Allows processing of data incrementally

3. Reading Specific Portions

Data can also be read in smaller chunks or a specific number of characters.

- ◆ Useful for very large datasets
- ◆ Improves performance in streaming applications

1.4.2.3 Reading Structured Data Files

In data engineering, files often contain structured data such as tables. For example, CSV files store data in rows and columns.

As part of the workflow introduced in Figure 1.4.1, reading structured files allows:

- ◆ Loading datasets into memory
- ◆ Preparing data for analysis
- ◆ Integrating data into pipelines

1.4.2.4 Applications of File Reading

Reading data from files is widely used in real-world scenarios:

- ◆ Data Analysis: Loading datasets for processing

- ◆ Log Processing: Reading system logs for monitoring
- ◆ Configuration Management: Reading application settings
- ◆ Data Pipelines: Ingesting data from external sources

1.4.2.5 Challenges in Reading Files

While reading files is straightforward, certain challenges may arise:

- ◆ Large file sizes leading to memory issues
- ◆ Missing or inconsistent data
- ◆ File format incompatibility
- ◆ Encoding issues

Proper handling techniques are required to address these challenges effectively.

1.4.3 Writing Data to Files

In data engineering workflows, processing data is only one part of the task. Once data has been processed, analyzed, or transformed, it must often be stored for future use. This is achieved through writing data to files.

Writing data to files involves transferring data from the program (memory) to external storage. As illustrated earlier in Figure 1.4.1, this represents the output stage of file I/O operations, where processed data is saved for persistence, sharing, or further processing.

1.4.3.1 Process of Writing to Files

Writing data to a file in Python generally involves the following steps:

1. Opening the file in a suitable mode (write or append)
2. Writing data into the file
3. Closing the file after writing

This process ensures that data is correctly stored and system resources are properly managed.

1.4.3.2 Writing Modes

When writing data, Python provides specific modes that determine how the file behaves:

1. Write Mode (w)

- ◆ Creates a new file if it does not exist
- ◆ Overwrites the file if it already exists

This mode is useful when creating new output files or replacing old data.

2. Append Mode (a)

- ◆ Adds new data to the end of an existing file
- ◆ Does not remove previous content

This mode is useful for:

- ◆ Maintaining logs
- ◆ Updating datasets
- ◆ Adding new records

Conceptual Illustration

Consider a program that processes student data and generates results:

Average Marks: 87.5

Top Performer: Asha

Using write or append operations, this output can be stored in a file for future reference.

1.4.3.3 Writing Structured Data

In real-world applications, data is often written in structured formats such as CSV or text files.

As part of the workflow shown in Figure 1.4.1, writing structured data enables:

- ◆ Storing processed datasets
- ◆ Exporting results for reporting
- ◆ Sharing data across systems

For example:

- ◆ Writing processed sales data into a CSV file
- ◆ Saving analysis results into a report file

1.4.3.4 Applications of File Writing

Writing data to files is widely used in:

- ◆ Data Storage: Saving processed data
- ◆ Report Generation: Creating summaries and reports
- ◆ Logging Systems: Recording system activities
- ◆ Data Pipelines: Output stage of ETL processes

1.4.3.5 Best Practices in File Writing

To ensure efficient and reliable file writing:

- ◆ Use appropriate file modes (w or a)
- ◆ Avoid overwriting important data unintentionally
- ◆ Ensure files are properly closed after writing
- ◆ Handle errors during file operations

These practices help maintain data integrity and system performance.

1.4.3.6 Summary of Writing Operations

Table 1.4.3 Writing Operations in Python

Mode	Purpose	Behavior
w	Write	Creates or overwrites file
a	Append	Adds data without deleting existing content

Recap

- ◆ File I/O operations involve reading from and writing to files
- ◆ Files can be text or binary
- ◆ File modes determine how files are accessed
- ◆ Reading allows data to be loaded into programs
- ◆ Writing enables storing processed data
- ◆ File handling is essential for real-world data workflows

Objective Questions

1. What does file I/O stand for?
2. Which mode is used for reading a file?
3. Which mode overwrites a file?
4. What is append mode used for?



5. What type of file is CSV?
6. What is the purpose of file reading?
7. What is the purpose of file writing?
8. Which mode is used for binary reading?
9. Name one application of file I/O.
10. Why is file handling important?

Answers

1. Input/Output
2. r
3. w
4. Adding data to file
5. Text file
6. Load data
7. Store data
8. rb
9. Data storage
10. Enables data persistence

Assignments

1. Explain the concept of file I/O operations in Python.
2. Describe different file modes with examples.
3. Explain how data is read from files in Python.
4. Discuss the importance of file writing in data workflows.
5. Compare reading and writing operations in file handling.

Reference

1. Python Software Foundation. File Handling Documentation

Suggested Reading

1. McKinney, W. (2018). *Python for Data Analysis*. O'Reilly Media.
2. Downey, A. (2015). *Think Python*. O'Reilly Media.

SGOU

BLOCK-2

**Data Processing and
Transformation , Data
Integration and ETL**

UNIT 1

Cleaning and Preprocessing Data using Python

Learning Outcomes

At the end of this unit, the learner will be able to :

- ◆ explain the importance of data cleaning and preprocessing
- ◆ identify common data quality issues such as missing values and outliers
- ◆ apply techniques for handling missing data using Python
- ◆ detect and treat outliers in datasets
- ◆ understand the role of preprocessing in data analysis workflows

Prerequisite

In real-world data environments, raw data is rarely perfect. It often contains missing values, inconsistencies, and irregular patterns that can affect analysis. Without proper preparation, such data may lead to incorrect conclusions and unreliable results. Therefore, understanding how to clean and preprocess data is essential for any data-related task.

Learners entering this unit are expected to have a basic understanding of Python programming and familiarity with libraries such as pandas and NumPy. Knowledge of data structures like DataFrames and basic operations such as filtering and selection will help in applying preprocessing techniques effectively.

This unit builds upon earlier concepts and introduces systematic approaches to improving data quality, preparing learners for more advanced topics such as feature engineering and ETL processes.

Keywords

Data Cleaning, Data Preprocessing, Missing Values, Outliers, Data Quality, Data Transformation, pandas, Data Integrity, Noise, Data Preparation



Discussion

2.1.1 Introduction to Data Cleaning and Preprocessing

In data engineering and analytics, the quality of data directly influences the accuracy of results. Raw data collected from various sources is often incomplete, inconsistent, and noisy. Before performing any analysis, it is essential to clean and preprocess the data to ensure its reliability.

Data cleaning refers to the process of identifying and correcting errors in the dataset, while data preprocessing involves transforming data into a suitable format for analysis. These steps are fundamental in preparing data for tasks such as statistical analysis, machine learning, and reporting.

As illustrated in Figure 2.1.1, raw data undergoes several preprocessing steps before becoming suitable for analysis.

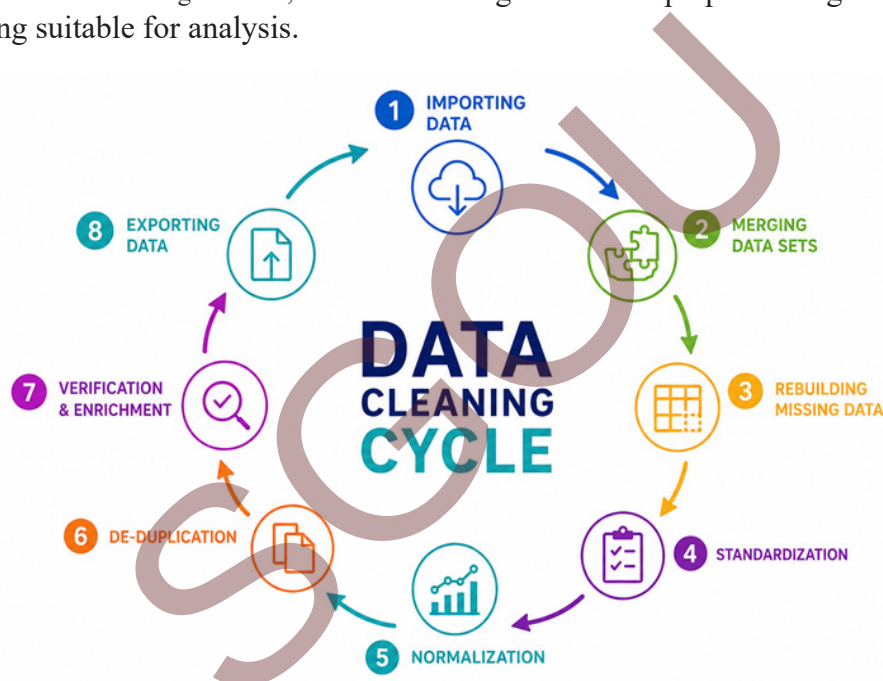


Fig. 2.1.1 Data cleaning and preprocessing workflow

Understanding the Need for Data Cleaning

Raw data may contain several issues:

- ◆ Missing values
- ◆ Duplicate records
- ◆ Inconsistent formats
- ◆ Outliers and noise

These issues can lead to inaccurate analysis if not handled properly.

Role of Preprocessing in Data Workflows

As shown in Figure 2.1.1, preprocessing acts as an intermediate stage between raw data and analysis. It ensures that:

- ◆ Data is consistent and structured
- ◆ Errors are minimized
- ◆ Analytical models perform better

Real-World Context

Consider a healthcare dataset where data is collected from hospital records. The table below shows common data quality issues:

Table 2.1.1 Patient data table

Patient ID	Name	Age	Gender	Diagnosis	Issue Identified
P101	Anu	29	Female	Fever	No issue
P102	Ravi Kumar	–	Male	Diabetes	Missing Age
P103	Meera	150	Female	Hypertension	Incorrect Age
P104	Kumar	45	Male	Diabetes	Duplicate Record
P105	Ravi Kumar	45	Male	Diabetes	Duplicate Record
P106	Sinu	32	–	Asthma	Missing Gender

In this dataset:

- ◆ **Missing Values :** Patient ID P102 has no age value, and P106 has missing gender information. Missing values can reduce the completeness and reliability of the dataset.
- ◆ **Incorrect Values :** The age of 150 for Patient ID P103 is unrealistic and indicates a data entry error, which can affect analysis results.
- ◆ **Duplicate Records :** P104 and P105 refer to the same patient (Ravi Kumar) but are recorded twice with slight variation in name. This leads to redundancy and biased outcomes.
- ◆ **Overall Impact :** If such issues are not handled through data cleaning and preprocessing, the final analysis may produce incorrect or misleading conclusions.

2.1.2 Key Concepts in Data Cleaning and Preprocessing

Before applying specific techniques for cleaning and preprocessing data, it is important to understand the key concepts that define data quality and influence how data should be prepared. These concepts help data engineers identify issues and apply appropriate corrective measures.

Data cleaning and preprocessing are not random processes; they are guided by well-defined principles that ensure data becomes reliable, consistent, and suitable for analysis.

2.1.2.1 Data Quality

Data quality refers to the condition of a dataset in terms of its accuracy, completeness, consistency, and reliability. High-quality data leads to meaningful insights, while poor-quality data can produce misleading results.

As illustrated in Figure 2.1.1, preprocessing improves data quality by removing errors and inconsistencies before analysis.

Key dimensions of data quality include:

- ◆ **Accuracy:** Data correctly represents real-world values
- ◆ **Completeness:** No important data is missing
- ◆ **Consistency:** Data is uniform across the dataset
- ◆ **Validity:** Data follows defined formats and rules

2.1.2.2 Missing Values

Missing values occur when no data is stored for a particular variable. This is one of the most common issues in real-world datasets.

For example:

Name	Age	Marks
Asha	20	85
Ravi	—	90

In the above example, Ravi's age is missing.

Missing values can arise due to:

- ◆ Data entry errors
- ◆ System failures
- ◆ Incomplete data collection

Handling missing values is crucial because they can affect calculations and model performance.

2.1.2.3 Outliers

Outliers are data points that significantly differ from the rest of the dataset. They may occur due to errors or represent rare but valid observations.

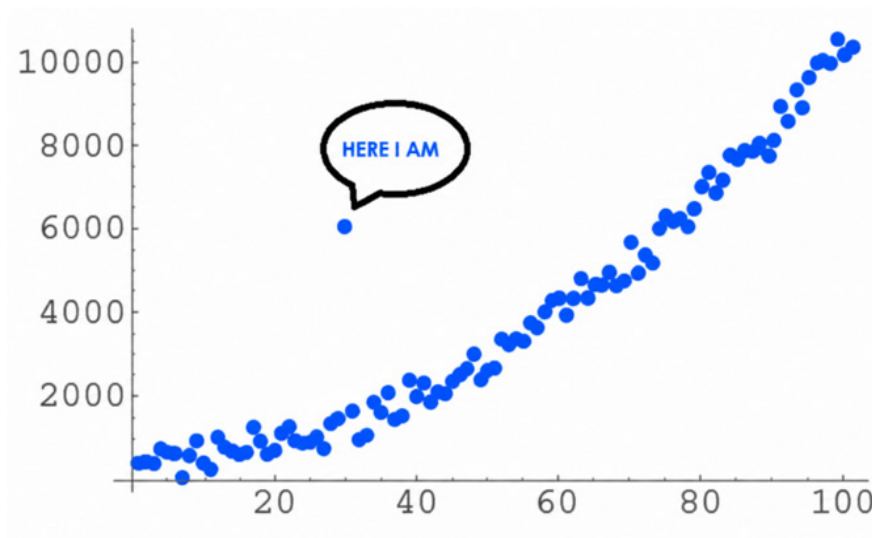


Fig. 2.1.2 Outliers in a dataset

As shown in Figure 2.1.2, outliers appear as points that lie far from the general pattern of data.

Outliers can:

- ◆ Distort statistical measures such as mean and standard deviation
- ◆ Affect model performance
- ◆ Indicate data errors or unusual events

2.1.2.4 Noise in Data

Noise refers to random errors or irrelevant data that do not contribute meaningful information.

Examples include:

- ◆ Typographical errors
- ◆ Sensor errors
- ◆ Random fluctuations

Noise reduces data quality and must be minimized during preprocessing.

2.1.2.5 Data Consistency

Consistency ensures that data values are uniform across the dataset.

Examples of inconsistency:

- ◆ Different formats for dates (DD/MM/YYYY vs MM/DD/YYYY)
- ◆ Different representations (Male vs M)

Ensuring consistency improves data reliability and usability.

2.1.2.6 Data Transformation

Data transformation involves converting data into a suitable format for analysis.

This includes:

- ◆ Changing data types
- ◆ Scaling values
- ◆ Standardizing formats

Transformation is a key part of preprocessing and prepares data for advanced analysis.

2.1.2.7 Summary of Key Concepts

Table 2.1.2 Key concepts in data cleaning and preprocessing

Concept	Description	Example
Data Quality	Accuracy and reliability of data	Correct values
Missing Values	Absence of data	Empty fields
Outliers	Extreme values	Very high marks
Noise	Random errors	Typographical errors
Consistency	Uniform data format	Standard date format
Transformation	Data conversion	Scaling values

2.1.3 Techniques for Handling Missing Values

Missing values are one of the most common data quality issues encountered in real-world datasets. As discussed in the previous section, missing data can arise due to incomplete data collection, system errors, or data entry problems. If not handled properly, missing values can lead to inaccurate analysis and unreliable results.

Therefore, it is essential to apply appropriate techniques to handle missing values before performing any data analysis or modeling.

As illustrated in Figure 2.1.1, handling missing values is a key step in the data preprocessing pipeline, ensuring that the dataset becomes complete and usable.

2.1.3.1 Identifying Missing Values

Before handling missing data, it is important to identify where missing values exist in the dataset.

In Python (using pandas), missing values are typically represented as:

- ◆ NaN (Not a Number)
- ◆ None

Identification methods include:

- ◆ Checking for null values
- ◆ Counting missing entries in each column
- ◆ Visual inspection of datasets

Recognizing missing values is the first step toward effective data cleaning.

2.1.3.2 Removing Missing Values

One of the simplest approaches is to remove missing data.

1. Deleting Rows

Rows containing missing values can be removed entirely.

- ◆ Suitable when missing values are few
- ◆ Helps maintain data integrity

2. Deleting Columns

Columns with a large number of missing values can be removed.

- ◆ Useful when a column is mostly incomplete
- ◆ Prevents noise in analysis

Advantages

- ◆ Simple and easy to apply
- ◆ No assumptions required

Limitations

- ◆ Loss of potentially useful data
- ◆ Not suitable for large-scale missing data

2.1.3.3 Replacing Missing Values (Imputation)

Instead of removing data, missing values can be replaced using appropriate techniques.

1. Mean/Median/Mode Imputation

Missing values are replaced with statistical measures:

- ◆ **Mean:** For numerical data
- ◆ **Median:** For skewed data
- ◆ **Mode:** For categorical data

Example:

- ◆ Replacing missing marks with average marks

2. Forward Fill and Backward Fill

- ◆ **Forward Fill (ffill):** Replaces missing values with the previous value
- ◆ **Backward Fill (bfill):** Replaces missing values with the next value

These methods are useful in time-series data.

3. Constant Value Replacement

Missing values can be replaced with a fixed value such as:

- ◆ 0
- ◆ “Unknown”

This method is simple but may not always be appropriate.

2.1.3.4 Advanced Imputation Techniques

In more complex scenarios, advanced methods can be used:

- ◆ Predictive models to estimate missing values
- ◆ Interpolation techniques
- ◆ Machine learning-based imputation

These techniques improve accuracy but require more computational effort.

2.1.3.5 Comparison of Missing Value Handling Techniques

Table 2.1.3 Techniques for handling missing values

Technique	Description	Use Case
Deletion	Remove rows/columns	Small missing data
Mean/Median/Mode	Replace with statistical value	Numerical/categorical data
Forward/Backward Fill	Use neighboring values	Time-series data
Constant Replacement	Replace with fixed value	Simple datasets
Advanced Methods	Predict missing values	Complex datasets

2.1.3.6 Impact of Handling Missing Values

Proper handling of missing values:

- ◆ Improves data quality
- ◆ Enhances model accuracy
- ◆ Prevents bias in analysis
- ◆ Ensures consistency in datasets

Improper handling, however, can distort results and reduce reliability.

2.1.4 Techniques for Handling Outliers

In addition to missing values, another critical issue in data preprocessing is the presence of outliers. Outliers are data points that significantly deviate from the general pattern of the dataset. These values may arise due to measurement errors, data entry mistakes, or genuine rare events.

If not handled properly, outliers can distort statistical measures such as mean and standard deviation, and negatively impact analytical models. Therefore, identifying and managing outliers is an essential step in data cleaning.

As illustrated in Figure 2.1.2, outliers appear as extreme values that lie far from the majority of data points.

2.1.4.1 Identifying Outliers

Before handling outliers, it is important to detect them accurately.

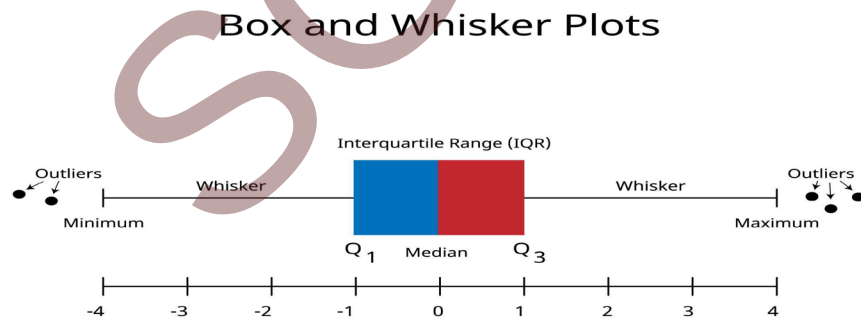


Fig. 2.1.3 Outlier detection using visualization techniques

As shown in Figure 2.1.3, outliers can be identified using visual and statistical methods.

1. Visual Methods

- ◆ Box Plot : Highlights outliers as points outside whiskers
- ◆ Scatter Plot : Shows distant data points
- ◆ Histogram : Reveals unusual distributions

2. Statistical Methods

- ◆ Z-Score Method: Values beyond a certain threshold (e.g., ± 3) are considered outliers
- ◆ Interquartile Range (IQR): Outliers are values below $Q1 - 1.5 \times IQR$ or above $Q3 + 1.5 \times IQR$

2.1.4.2 Techniques for Handling Outliers

Once identified, outliers can be handled using different techniques depending on the context.

1. Removing Outliers

Outliers can be removed from the dataset.

- ◆ Suitable when outliers are errors
- ◆ Improves data consistency

Limitation : May result in loss of important information if outliers are valid observations

2. Capping (Winsorization)

Outliers are replaced with a maximum or minimum threshold value.

- ◆ Preserves dataset size
- ◆ Reduces extreme variation

Example: Values above a threshold are replaced with the maximum allowed value

3. Transformation Techniques

Mathematical transformations can reduce the impact of outliers:

- ◆ Log transformation
- ◆ Square root transformation

These techniques compress large values and reduce skewness.

4. Using Robust Statistical Methods

Instead of removing outliers, robust methods can be used:

- ◆ Median instead of mean
- ◆ Interquartile range instead of standard deviation

These methods reduce sensitivity to extreme values.

5. Treating Outliers as Special Cases

In some scenarios, outliers represent meaningful information:

- ◆ Fraud detection
- ◆ Rare events in healthcare

Instead of removing them, they can be analyzed separately.

2.1.4.3 Comparison of Outlier Handling Techniques

Table 2.1.4 Techniques for handling outliers

Technique	Description	Use Case
Removal	Delete extreme values	Erroneous data
Capping	Limit extreme values	Preserve dataset
Transformation	Reduce skewness	Large-scale data
Robust Methods	Use resistant statistics	Sensitive analysis
Special Handling	Analyze separately	Rare events

2.1.4.4 Impact of Outlier Handling

Proper handling of outliers:

- ◆ Improves accuracy of analysis
- ◆ Enhances model performance
- ◆ Reduces bias in datasets

Improper handling may lead to:

- ◆ Loss of critical information
- ◆ Distorted results

Recap

- ◆ Data cleaning and preprocessing are essential steps to improve data quality before analysis
- ◆ Raw data often contains issues such as missing values, outliers, noise, and inconsistencies
- ◆ Data quality depends on accuracy, completeness, consistency, and validity
- ◆ Missing values can be handled using deletion or imputation techniques such as mean, median, or forward fill

- ◆ Outliers are extreme values that can distort analysis and must be carefully identified
- ◆ Techniques such as IQR, Z-score, capping, and transformation are used to handle outliers
- ◆ Proper preprocessing ensures reliable results and improves model performance
- ◆ Data preprocessing is a foundational step in data engineering and analytics workflows

Objective Questions

1. What is data cleaning?
2. What is data preprocessing?
3. What is meant by data quality?
4. What are missing values?
5. Name one cause of missing data.
6. What is an outlier?
7. Which method replaces missing values with average values?
8. What does IQR stand for?
9. Which technique removes extreme values?
10. What is data consistency?
11. Which method uses previous values to fill missing data?
12. What is noise in data?
13. Which statistical measure is less affected by outliers?
14. What is the purpose of preprocessing?
15. Which method caps extreme values?

Answers

1. Process of correcting data errors
2. Preparing data for analysis
3. Accuracy and reliability of data
4. Absence of data
5. Data entry error
6. Extreme value
7. Mean imputation
8. Interquartile Range
9. Removal method
10. Uniform data format
11. Forward fill
12. Random error
13. Median
14. Improve data quality
15. Capping (Winsorization)

Assignments

1. Explain the importance of data cleaning and preprocessing in data engineering.
2. Describe different types of data quality issues with examples.
3. Discuss various techniques for handling missing values.
4. Explain how outliers can affect data analysis and how they can be handled.
5. Compare different methods of handling missing values and outliers.

Reference

1. Python Software Foundation. pandas Documentation
2. Python Software Foundation. NumPy Documentation

Suggested Reading

1. McKinney, W. (2018). *Python for Data Analysis*. O'Reilly Media.
2. VanderPlas, J. (2016). *Python Data Science Handbook*. O'Reilly Media.
3. Kleppmann, M. (2017). *Designing Data-Intensive Applications*. O'Reilly Media.

SGOU

UNIT 2

Data Normalization and Standardization

Learning Outcomes

At the end of this unit, the learner will be able to :

- ◆ explain the concept of data normalization and standardization
- ◆ understand the importance of feature scaling in data preprocessing
- ◆ describe different normalization techniques
- ◆ apply Min-Max scaling and Z-score normalization
- ◆ evaluate the impact of scaling on data analysis and models

Prerequisite

In the previous unit, learners explored how to clean and preprocess data by handling missing values and outliers. Once the data is clean, the next step is to ensure that it is in a suitable scale for analysis and modeling.

Learners should be familiar with numerical data, basic statistical concepts such as mean and standard deviation, and Python libraries like pandas and NumPy. This knowledge will help in understanding how scaling techniques transform data.

This unit focuses on improving data representation through normalization and standardization, which are essential for many machine learning and analytical tasks.

Keywords

Normalization, Standardization, Feature Scaling, Min-Max Scaling, Z-score, Data Transformation, Scaling Techniques, Mean, Standard Deviation



Discussion

2.2.1 Fundamentals of Data Normalization and Standardization

In real-world datasets, different features often have different scales. For example, consider a dataset containing:

- ◆ Age (0–100)
- ◆ Income (in thousands or millions)
- ◆ Marks (0–100)

These differences in scale can create problems during analysis, especially in machine learning models where certain features may dominate others due to their larger values.

To address this issue, feature scaling techniques such as normalization and standardization are used. These techniques transform data into a common scale without distorting the relationships between variables.

As illustrated in Figure 2.2.1, raw data with varying ranges is transformed into scaled data with uniform distribution.

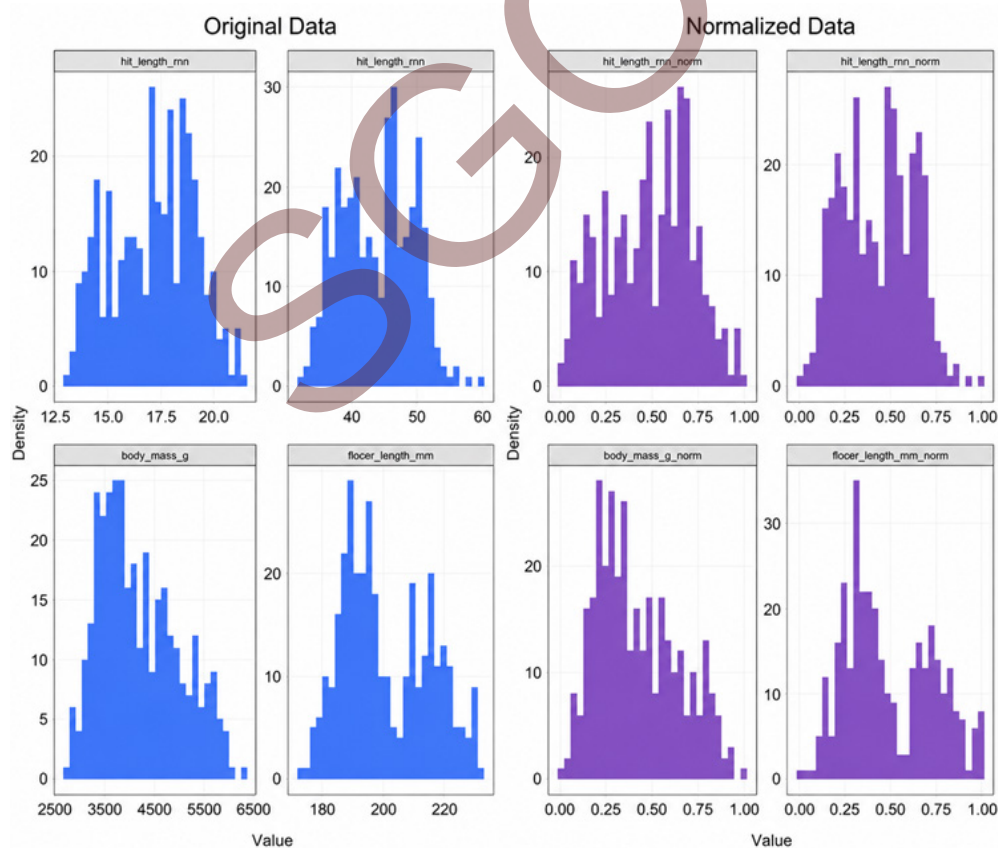


Fig. 2.2.1 Data before and after scaling

Understanding the Need for Scaling

As seen in Figure 2.2.1, unscaled data may have large variations, which can:

- ◆ Affect distance-based algorithms
- ◆ Slow down convergence in models
- ◆ Bias results toward larger values

Scaling ensures that all features contribute equally to analysis.

Normalization vs Standardization

Although both techniques are used for scaling, they differ in approach:

- ◆ Normalization: Rescales data to a fixed range (usually 0 to 1)
- ◆ Standardization: Transforms data to have mean = 0 and standard deviation = 1

Key Characteristics

Table 2.2.1 Normalization vs Standardization

Feature	Normalization	Standardization
Range	Fixed (0–1)	No fixed range
Method	Min-Max scaling	Z-score
Use Case	Bounded data	Normally distributed data
Sensitivity	Affected by outliers	Less affected

Real-World Context

Consider a dataset used for predicting house prices:

- ◆ Area (in square feet)
- ◆ Number of rooms
- ◆ Price

Without scaling, features like area may dominate the model due to larger values.

2.2.2 Feature Scaling Techniques (Min-Max Scaling)

In the previous section, we discussed the importance of feature scaling and introduced normalization and standardization. One of the most commonly used normalization techniques is Min-Max Scaling, which transforms data into a fixed range, typically between 0 and 1.

Min-Max Scaling is particularly useful when the goal is to bring all features to a common scale while preserving the relationships between data values.



As illustrated in Figure 2.2.1, scaling transforms data with varying ranges into a uniform scale, making it easier to compare and analyze.

2.2.2.1 Concept of Min-Max Scaling

Min-Max Scaling works by adjusting the values of a dataset so that they fall within a predefined range, usually $[0, 1]$.

The transformation is performed using the formula :

$$X_{scaled} = \frac{X - X_{min}}{X_{max} - X_{min}}$$

Where:

- X = Original value
- X_{min} = Minimum value in the dataset
- X_{max} = Maximum value in the dataset

This formula ensures that:

- ♦ The smallest value becomes 0
- ♦ The largest value becomes 1
- ♦ All other values lie between 0 and 1

2.2.2.2 Illustration of Min-Max Scaling

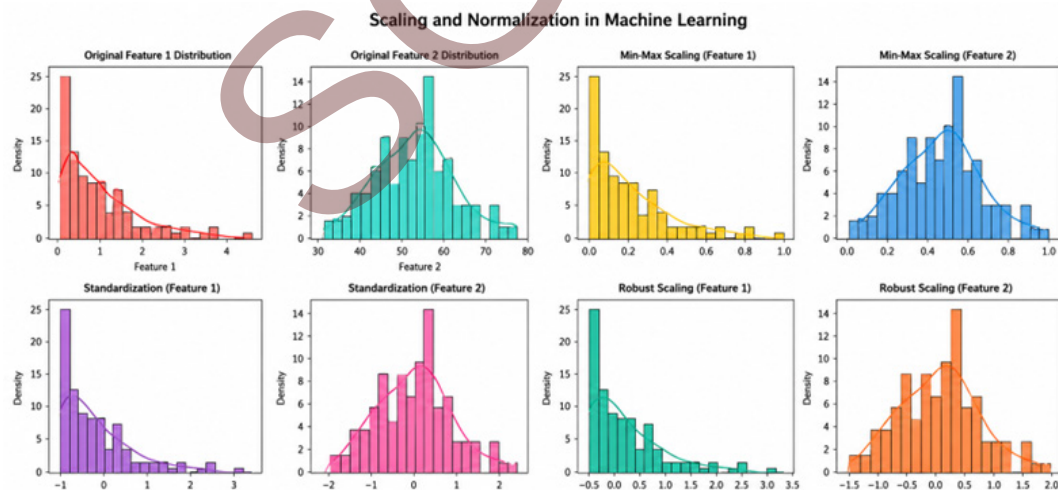


Fig.2.2.2 Min-max scaling transformation

As shown in Figure 2.2.2, original values are transformed proportionally into a smaller range without changing their relative distribution.

2.2.2.3 Step-by-Step Example

Consider a dataset of marks:

Table 2.2.2 dataset of Marks

Student	Marks
A	50
B	70
C	90

Here:

- ◆ Minimum = 50
- ◆ Maximum = 90

Applying Min-Max Scaling:

- ◆ $A \rightarrow (50-50)/(90-50) = 0$
- ◆ $B \rightarrow (70-50)/(90-50) = 0.5$
- ◆ $C \rightarrow (90-50)/(90-50) = 1$

Table 2.2.3 Example of Min-Max Scaling

Student	Original Marks	Scaled Value
A	50	0.0
B	70	0.5
C	90	1.0

2.2.2.4 Advantages of Min-Max Scaling

- ◆ Simple and easy to implement
- ◆ Preserves relationships between values
- ◆ Ensures all values lie within a fixed range
- ◆ Useful for algorithms that require bounded input

2.2.2.5 Limitations of Min-Max Scaling

- ◆ Sensitive to outliers
- ◆ Extreme values can distort scaling
- ◆ Requires knowledge of minimum and maximum values



2.2.2.6 Applications of Min-Max Scaling

Min-Max Scaling is widely used in:

- ◆ Machine learning algorithms (e.g., neural networks)
- ◆ Image processing (pixel value normalization)
- ◆ Data visualization
- ◆ Distance-based models

2.2.3 Feature Scaling Techniques (Z-score Normalization)

In addition to Min-Max Scaling, another widely used feature scaling technique is Z-score normalization, also known as standardization. Unlike normalization, which rescales data to a fixed range, standardization transforms data so that it has a mean of 0 and a standard deviation of 1.

This technique is particularly useful when the data follows a normal (Gaussian) distribution or when algorithms assume standardized input.

As illustrated in Figure 2.2.1, standardization reshapes the data distribution while preserving its relative structure.

2.2.3.1 Concept of Z-score Normalization

Z-score normalization transforms each data point based on how far it is from the mean of the dataset.

The formula used is:

$$Z = \frac{X - \mu}{\sigma}$$

Where:

- X = Original value
- μ = Mean of the dataset
- σ = Standard deviation

This transformation ensures that:

The dataset has a mean of 0

The spread of data is measured in terms of standard deviation

2.2.3.2 Interpretation of Z-scores

A Z-score indicates how far a data point is from the mean:

- ◆ $Z = 0 \rightarrow$ Value equals the mean
- ◆ $Z > 0 \rightarrow$ Value is above the mean
- ◆ $Z < 0 \rightarrow$ Value is below the mean

For example:

- ◆ $Z = +2 \rightarrow$ 2 standard deviations above the mean
- ◆ $Z = -1 \rightarrow$ 1 standard deviation below the mean

2.2.3.3 Illustration of Standardization

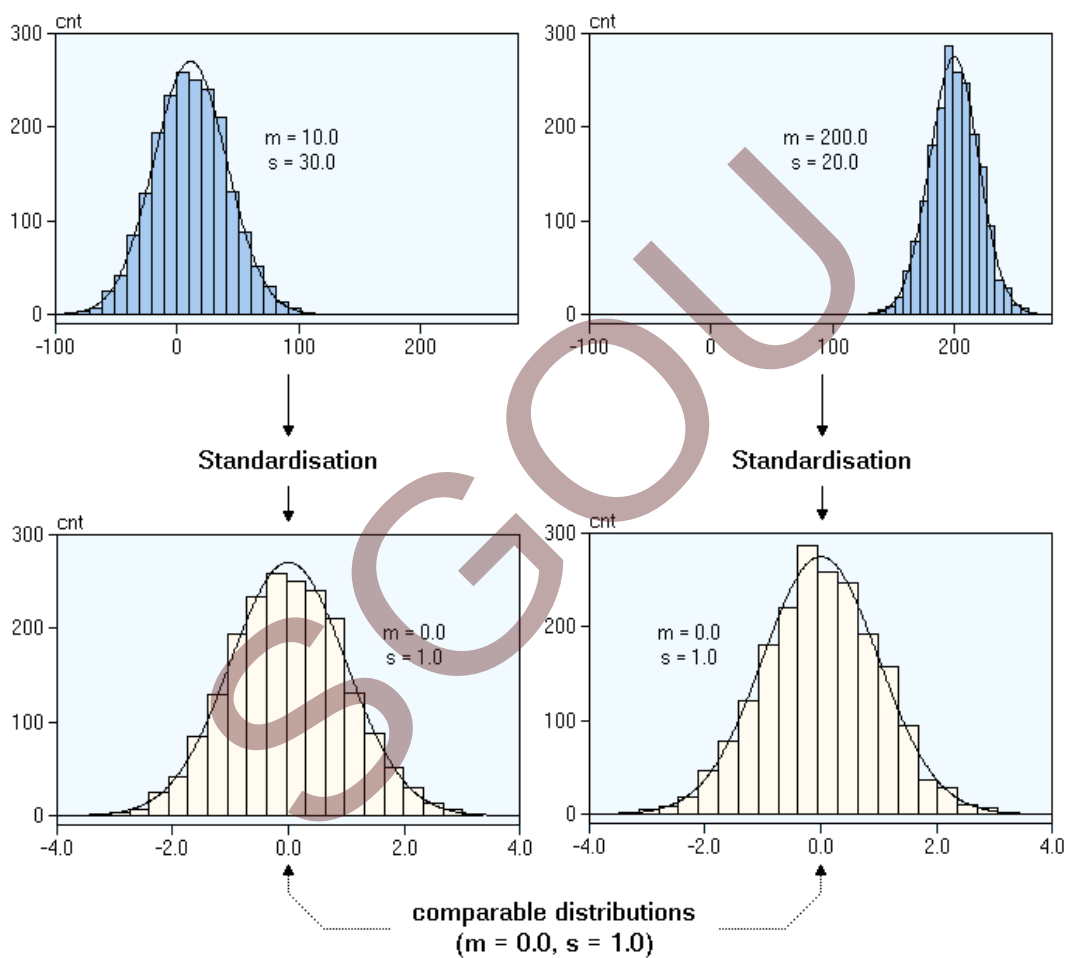


Fig. 2.2.3 Z-score Normalization (Standardization)

As shown in Figure 2.2.3, the original data distribution is transformed so that it is centered around zero with uniform variance.

2.2.3.4 Step-by-Step Example

Consider a dataset of marks:



Student	Marks
A	50
B	70
C	90

Step 1: Calculate mean

$$\mu = \frac{50 + 70 + 90}{3} = 70$$

Step 2: Calculate standard deviation (approx.) = 16.33

Step 3: Apply formula:

- ◆ $A \rightarrow (50 - 70) / 16.33 \approx -1.22$
- ◆ $B \rightarrow (70 - 70) / 16.33 = 0$
- ◆ $C \rightarrow (90 - 70) / 16.33 \approx +1.22$

Table 2.2.4 Example of Z-score Normalization

Student	Original Marks	Z-score
A	50	-1.22
B	70	0
C	90	+1.22

2.2.3.5 Advantages of Z-score Normalization

- ◆ Less sensitive to outliers compared to Min-Max scaling
- ◆ Centers data around zero
- ◆ Suitable for normally distributed data
- ◆ Improves performance of many machine learning algorithms

2.2.3.6 Limitations of Z-score Normalization

- ◆ Does not bound values within a fixed range
- ◆ Assumes data distribution is approximately normal
- ◆ Requires calculation of mean and standard deviation

2.2.3.7 Applications of Z-score Normalization

Z-score normalization is widely used in:

- ◆ Machine learning algorithms (e.g., regression, clustering)
- ◆ Statistical analysis
- ◆ Data preprocessing for predictive models
- ◆ Financial and scientific data analysis

Recap

- ◆ Data normalization and standardization are important preprocessing techniques used to scale data
- ◆ Feature scaling ensures that all variables contribute equally during analysis
- ◆ Normalization rescales data to a fixed range, typically between 0 and 1
- ◆ Min-Max Scaling preserves relationships between values while adjusting scale
- ◆ Standardization transforms data to have mean = 0 and standard deviation = 1
- ◆ Z-score normalization measures how far a value is from the mean
- ◆ Min-Max Scaling is sensitive to outliers, whereas Z-score is more robust
- ◆ Proper scaling improves performance of machine learning models and analytical methods

Objective Questions

1. What is feature scaling?
2. What is normalization?
3. What is standardization?
4. Which technique scales data between 0 and 1?
5. What does Z-score represent?
6. Which method uses mean and standard deviation?
7. What is the range of Min-Max scaling?

8. Which scaling technique is less sensitive to outliers?
9. What does a Z-score of 0 indicate?
10. Which technique is used for normally distributed data?
11. What is the purpose of scaling data?
12. Which method is used in neural networks commonly?
13. What is the main limitation of Min-Max scaling?
14. What does standard deviation measure?
15. Which technique centers data around zero?

Answers

1. Scaling data to a common range
2. Rescaling data to a fixed range
3. Transforming data to mean 0 and std 1
4. Min-Max Scaling
5. Distance from mean
6. Z-score normalization
7. 0 to 1
8. Z-score normalization
9. Value equals mean
10. Standardization
11. Improve model performance
12. Min-Max Scaling
13. Sensitive to outliers
14. Spread of data
15. Standardization

Assignments

1. Explain the concept of feature scaling and its importance.
2. Compare normalization and standardization with examples.
3. Describe the Min-Max scaling technique with advantages and limitations.
4. Explain Z-score normalization and its applications.
5. Discuss the impact of scaling on machine learning models.

Reference

1. Python Software Foundation. pandas Documentation
2. Python Software Foundation. NumPy Documentation

Suggested Reading

1. Géron, A. (2019). *Hands-On Machine Learning with Scikit-Learn & TensorFlow*. O'Reilly Media.
2. McKinney, W. (2018). *Python for Data Analysis*. O'Reilly Media.
3. VanderPlas, J. (2016). *Python Data Science Handbook*. O'Reilly Media

UNIT 3

Introduction to ETL Processes

Learning Outcomes

At the end of this unit, the learner will be able to :

- ◆ explain the concept of ETL (Extract, Transform, Load) processes
- ◆ describe the stages involved in ETL workflows
- ◆ understand the importance of ETL in data engineering
- ◆ identify different data sources used in ETL processes
- ◆ explain how data is extracted from files and databases

Prerequisite

In previous units, learners explored how to clean, preprocess, and scale data using Python. Once data is prepared, the next step is to move it across systems for storage, analysis, and decision-making. This is achieved through ETL processes.

Learners should be familiar with concepts such as data cleaning, transformation, file handling, and Python libraries like pandas. Understanding how data is structured and manipulated will help in grasping ETL workflows.

This unit introduces ETL as a systematic approach to handling data movement and transformation in real-world data engineering environments

Keywords

ETL, Extract, Transform, Load, Data Pipeline, Data Sources, Data Integration, Data Warehouse, Data Processing, Data Flow

Discussion

2.3.1 Introduction to ETL Processes

In modern data-driven organizations, data is generated from multiple sources such as applications, databases, sensors, and external systems. However, this data is often stored in different formats and locations, making it difficult to analyze directly. To make data usable, it must be collected, processed, and stored systematically. This is achieved through ETL processes.

ETL stands for Extract, Transform, and Load, which represents a sequence of steps used to move data from source systems to a target system, such as a data warehouse or database.

As illustrated in Figure 2.3.1, ETL acts as a pipeline that connects data sources to data storage systems.



Fig. 2.3.1 ETL Process Workflow

Understanding ETL

The ETL process consists of three main stages:

- ◆ **Extract:** Collecting Data from Various Sources In this initial phase, raw data is retrieved from multiple, often incompatible sources such as relational databases (SQL), NoSQL stores, flat files (CSV, XML), and external APIs. This stage requires establishing secure connections and identifying the specific data points needed for analysis without disrupting the performance of the source systems.
- ◆ **Transform:** Transformation is the most critical and complex stage. Once extracted, the data is placed in a “staging area” where it undergoes several operations:
 - **Cleaning:** Removing duplicates, handling missing values, and correcting errors.
 - **Standardization:** Ensuring all dates, currencies, and units follow a consistent format.
 - **Logic Application:** Performing calculations (e.g., total sales) or filtering data based on specific business rules.

- **Mapping:** Aligning source data fields with the structure of the target system.
- ◆ **Load:** The final stage involves moving the newly structured and cleaned data into its permanent destination, typically a Data Warehouse or Data Lake. This enables business intelligence tools and data scientists to access a “single source of truth” for reporting and visualization.

These stages work together to ensure that raw data is converted into meaningful and structured information.

Importance of ETL Processes

As shown in Figure 2.3.1, ETL plays a critical role in data engineering by:

- ◆ Integrating data from multiple sources
- ◆ Ensuring data consistency and quality
- ◆ Supporting analytics and decision-making
- ◆ Enabling data warehousing

Without ETL, data would remain scattered and difficult to analyze.

Real-World Context

Consider an e-commerce company:

- ◆ Customer data from a website
- ◆ Sales data from transactions
- ◆ Inventory data from warehouses

ETL processes combine and transform this data into a unified system for analysis.

2.3.2 Extracting Data from Various Sources

The first stage of the ETL process is data extraction, where data is collected from different sources for further processing. In real-world scenarios, data is rarely stored in a single location. Instead, it is distributed across multiple systems such as files, databases, and online platforms.

Extraction is a critical step because the quality and completeness of extracted data directly influence the effectiveness of subsequent transformation and analysis.

As illustrated in Figure 2.3.1, the extraction stage gathers data from various sources and feeds it into the ETL pipeline.

2.3.2.1 Types of Data Sources

Data can be extracted from a variety of sources depending on the application and system architecture.

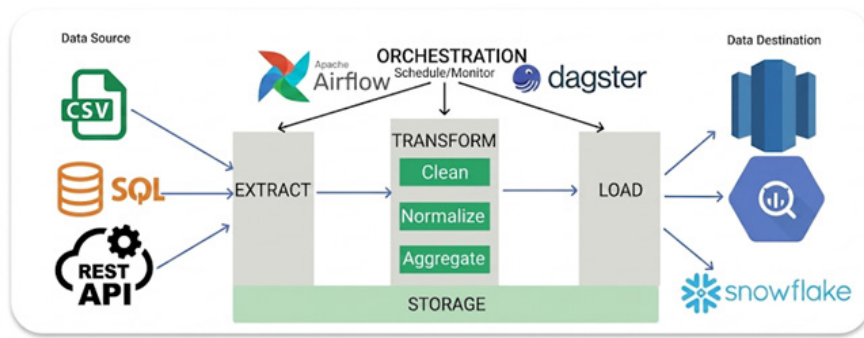


Fig.2.3.2 Data extraction from multiple sources

As shown in Figure 2.3.2, common data sources include:

1. Files

Files are one of the most common data sources.

Examples:

- ◆ CSV files
- ◆ Text files
- ◆ Excel files

These files are widely used because they are simple, portable, and easy to process using Python.

2. Databases

Databases store structured data in tables and are widely used in organizations.

Examples:

- ◆ MySQL
- ◆ PostgreSQL
- ◆ SQLite

Data can be extracted using queries, allowing selective retrieval of required information.

3. APIs (Application Programming Interfaces)

APIs allow data to be accessed from external systems over the internet.

Examples:

- ◆ Weather APIs

- ◆ Social media APIs
- ◆ Financial data APIs

APIs provide real-time data and are essential for dynamic applications.

4. Logs and System Data

Systems generate logs that contain valuable information about operations.

Examples:

- ◆ Server logs
- ◆ Application logs

These logs are useful for monitoring and analysis.

2.3.2.2 Methods of Data Extraction

Data extraction strategies are typically chosen based on the urgency of the information and the volume of the data being processed. Below are the two primary methods described in detail:

1. Batch Extraction

Batch extraction is the process of collecting data in large groups at specific, predefined times.

- ◆ **Scheduled intervals:** Data is gathered and moved during specific windows, such as nightly or weekly, to minimize the impact on system performance during peak hours.
- ◆ **Suitability for large datasets:** This method is highly efficient for processing massive volumes of historical data where immediate availability is not a priority.
- ◆ **Periodic data processing:** It is most commonly used in traditional business reporting and data warehousing, where data from the previous day or month is analyzed in bulk.

2. Real-Time Extraction

Real-time extraction, also known as streaming extraction, involves the continuous capture of data as soon as it is created or updated.

- ◆ **Continuous generation:** Data is extracted and moved almost instantaneously, ensuring the destination system is always synchronized with the source.
- ◆ **Suitability for real-time analytics:** This method is essential for high-stakes environments like fraud detection, stock market monitoring, or live system health tracking.

- ◆ **Streaming applications:** It is heavily utilized in modern applications that rely on tools like Apache Kafka or AWS Kinesis to handle live data feeds.

2.3.2.3 Extraction Using Python

Python serves as a versatile language for data management due to its robust ecosystem of libraries and modules. Here are the primary ways Python facilitates data extraction:

- ◆ **Reading CSV and text files using pandas:** The pandas library offers high-level functions like `read_csv()` and `read_excel()`, which allow for the swift ingestion of flat files into structured DataFrames for immediate analysis.
- ◆ **Connecting to databases using libraries:** Python utilizes powerful connectors such as SQLAlchemy, psycopg2, or mysql-connector-python to establish secure links with relational databases, enabling the execution of complex queries to retrieve specific datasets.
- ◆ **Fetching data from APIs using HTTP requests:** By using the requests library, developers can easily interact with RESTful APIs to extract data in JSON or XML formats, facilitating the collection of real-time information from web services.
- ◆ **Scraping data from web pages:** For sources without official APIs, Python provides tools like BeautifulSoup and Scrapy to parse HTML and extract relevant content directly from the web.

2.3.2.4 Challenges in Data Extraction

- ◆ **Data inconsistency across sources:** Discrepancies in how the same information is recorded—such as varying naming conventions or conflicting values for the same attribute—can lead to confusion and inaccurate analysis.
- ◆ **Different data formats:** Extraction often involves dealing with a mix of structured data (like SQL tables), semi-structured data (like JSON or XML), and unstructured data (like PDFs), making it difficult to standardize the input.
- ◆ **Missing or incomplete data:** The presence of null values or partially filled records can compromise the quality of the dataset, requiring sophisticated cleaning techniques to fill gaps or filter out unusable entries.
- ◆ **Network or access issues:** Reliable extraction is often hindered by connectivity problems, restricted permissions, or API rate limits, which can delay the data pipeline and lead to incomplete transfers.
- ◆ **Scalability concerns:** As data volumes grow, extraction processes may struggle to maintain performance, potentially slowing down the entire ETL pipeline and affecting real-time decision-making.

Properly managing these challenges is essential for building a robust foundation for any data-driven system.

2.3.2.5 Importance of Data Extraction

Data extraction is a critical phase in data management for several reasons:

- ◆ **Gather data from multiple sources:** It allows organizations to collect information from various origins, such as databases, flat files, and web services, ensuring a comprehensive data pool.
- ◆ **Enables integration of diverse datasets:** By extracting raw information, the system can combine disparate data types into a unified format, which is essential for cross-functional analysis.
- ◆ **Forms the foundation of ETL processes:** It serves as the initial and most vital step of the **Extract, Transform, Load (ETL)** cycle, providing the raw material necessary for subsequent cleaning and processing.
- ◆ **Supports data-driven decision-making:** By ensuring that relevant and accurate data is available for analysis, extraction empowers leadership to make strategic choices based on factual insights rather than intuition.

Recap

- ◆ ETL stands for Extract, Transform, and Load
- ◆ ETL processes are used to collect, process, and store data systematically
- ◆ The Extract stage involves gathering data from multiple source
- ◆ Data sources include files, databases, APIs, and system logs
- ◆ Batch extraction processes data at scheduled intervals
- ◆ Real-time extraction processes data continuously
- ◆ Python supports data extraction through libraries like pandas and API tools
- ◆ ETL enables data integration, consistency, and supports analytics

Objective Questions

1. What does ETL stand for?
2. Which stage involves collecting data?
3. What is the purpose of ETL?
4. Name one source of data extraction.

5. What type of file is commonly used for data storage?
6. What is batch extraction?
7. What is real-time extraction?
8. Which language is commonly used for ETL processes?
9. What are APIs used for?
10. Which stage comes after extraction?
11. What type of data do databases store?
12. What is a data pipeline?
13. Which type of extraction is continuous?
14. Name one challenge in data extraction.
15. What is the role of ETL in data engineering?

Answers

1. Extract, Transform, Load
2. Extract
3. Data processing and integration
4. CSV file
5. CSV file
6. Periodic data extraction
7. Continuous data extraction
8. Python
9. Accessing external data
10. Transform
11. Structured data
12. Data flow system
13. Real-time extraction
14. Data inconsistency
15. Managing data flow

Assignments

1. Explain the concept of ETL processes with examples.
2. Describe the stages of ETL in detail.
3. Discuss different data sources used in ETL processes.
4. Compare batch and real-time data extraction.
5. Explain how Python is used in data extraction

Reference

1. Python Software Foundation. pandas Documentation
2. Apache Software Foundation. ETL and Data Pipeline Documentation

Suggested Reading

1. Kleppmann, M. (2017). *Designing Data-Intensive Applications*. O'Reilly Media.
2. McKinney, W. (2018). *Python for Data Analysis*. O'Reilly Media.
3. Kimball, R. & Ross, M. (2013). *The Data Warehouse Toolkit*. Wiley.

UNIT 4

Transforming Data using Python, Loading Data, and Security Measures

Learning Outcomes

At the end of this unit, the learner will be able to :

- ◆ explain the concept of data transformation in ETL processes
- ◆ apply Python libraries and functions for transforming data
- ◆ describe methods for loading transformed data into target systems
- ◆ understand the importance of data loading in ETL workflows
- ◆ identify basic data security measures in Python

Prerequisite

In the previous unit, learners explored how data is extracted from various sources in ETL processes. However, extracted data is often raw and unsuitable for direct use. It must be transformed into a structured and meaningful format before being stored or analyzed.

Learners should be familiar with Python libraries such as pandas and concepts like data cleaning, preprocessing, and scaling. Understanding ETL workflows will help in grasping how transformation and loading fit into the overall process.

This unit focuses on the final stages of ETL—transforming data, loading it into target systems, and ensuring data security

Keywords

Data Transformation, ETL, Data Loading, Data Pipeline, pandas, Data Processing, Data Security, Encryption, Access Control



Discussion

2.4.1 Basics of Data Transformation

After data is extracted from various sources, it is often not in a usable format. It may contain inconsistencies, missing values, or incompatible structures. The process of converting this raw data into a clean and structured format is known as data transformation.

Data transformation is a crucial stage in ETL because it ensures that data is ready for analysis, reporting, or storage.

As illustrated in Figure 2.4.1, transformation acts as an intermediate step between extraction and loading, refining raw data into meaningful information.

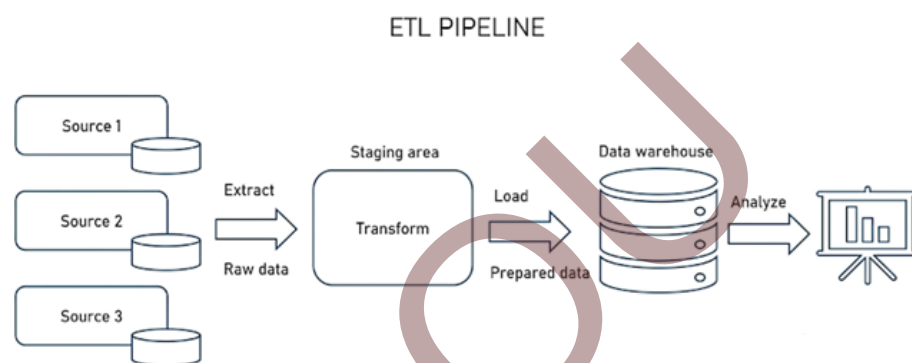


Fig. 2.4.1 Data transformation in ETL process

Understanding Data Transformation

Data transformation is the pivotal stage where raw, often messy data is refined into a high-quality asset. This process ensures that the information is consistent, accurate, and ready for analytical modeling. The core operations involved include:

- ◆ **Cleaning data:** This step involves identifying and rectifying errors in the dataset, such as removing duplicate entries, handling null or missing values, and correcting typos. It acts as a quality control gate to prevent “garbage in, garbage out” scenarios.
- ◆ **Formatting data:** To ensure compatibility across different systems, data must be converted into a standardized layout. This includes synchronizing date formats (e.g., changing “DD/MM/YYYY” to “YYYY-MM-DD”), normalizing currency units, and ensuring consistent character encoding.
- ◆ **Aggregating data:** This operation involves summarizing raw data to provide high-level insights. Examples include calculating the average monthly temperature, totaling daily sales, or counting the number of active users in a specific region.
- ◆ **Filtering and sorting data:** Filtering removes irrelevant records that do not meet specific criteria (e.g., excluding transactions below a certain value), while sorting organizes the remaining data into a logical sequence (e.g., chronological order or alphabetical order) to enhance readability and search efficiency.

Purpose of Data Transformation

The main objectives of transformation are:

- ◆ **Improve data quality:** Transformation serves as a critical cleaning phase where errors, duplicates, and outliers are removed. By addressing missing values and correcting inaccuracies, the process ensures that the resulting information is trustworthy for high-stakes decision-making.
- ◆ **Ensure consistency across datasets:** When gathering data from multiple sources, the same information might be represented differently. Transformation reconciles these differences—such as unifying different naming conventions or measurement units—so that the data can be accurately compared and integrated.
- ◆ **Convert data into a standard format:** This objective focuses on structural uniformity. By standardizing date formats, character encodings, and currency symbols, transformation ensures that the data is compatible with various software tools and programming environments.
- ◆ **Prepare data for analysis and storage:** Beyond cleaning, transformation reshapes data to fit the specific schema of a data warehouse or the requirements of a machine learning model. This optimization allows for faster querying, efficient storage, and more effective data visualization.

Real-World Context

Consider a sales dataset:

- ◆ Dates stored in different formats
- ◆ Missing values in transaction records
- ◆ Duplicate entries

Transformation ensures that such issues are resolved before loading the data into a system.

2.4.2 Transforming Data using Python Libraries and Functions

In the previous section, we discussed the importance of data transformation in ETL processes. In practical data engineering workflows, transformation is carried out using powerful programming tools, among which **Python** plays a central role.

Python provides a wide range of libraries and built-in functions that enable efficient data transformation. These tools allow data engineers to clean, modify, and restructure datasets with minimal effort.

As illustrated in Figure 2.4.1, transformation involves multiple operations applied sequentially to convert raw data into a structured format.

2.4.2.1 Key Python Libraries for Data Transformation

Python offers several libraries that simplify data transformation tasks:



1. pandas

pandas is the most widely used library for data manipulation and transformation.

It provides:

- ◆ DataFrame structure for tabular data
- ◆ Functions for cleaning, filtering, and aggregating data
- ◆ Easy handling of missing values

2. NumPy

NumPy supports numerical transformations such as:

- ◆ Mathematical operations
- ◆ Array manipulation
- ◆ Statistical calculations

3. Built-in Python Functions

Python's built-in functions are also useful for transformation:

- ◆ `map()` → Apply functions to data
- ◆ `filter()` → Select data based on conditions
- ◆ `lambda` functions → Perform quick transformations

2.4.2.2 Common Transformation Operations

Using Python libraries, several transformation tasks can be performed efficiently.

1. Data Cleaning

- ◆ Removing duplicates
- ◆ Handling missing values
- ◆ Correcting inconsistencies

2. Data Filtering

Selecting specific rows or columns based on conditions.

Example:

- ◆ Selecting records where `sales > 1000`

3. Data Aggregation

Combining data to produce summary statistics:

- ◆ Sum
- ◆ Mean
- ◆ Count

4. Data Formatting

- ◆ Changing the structure or format of data:
- ◆ Converting data types
- ◆ Standardizing date formats

5. Feature Engineering

- ◆ Creating new variables from existing data:
- ◆ Calculating total sales
- ◆ Extracting year from date

2.4.2.3 Transformation Workflow Using Python

Python provides a streamlined environment for data transformation, allowing developers to build reproducible pipelines. The following steps outline a standard transformation workflow using libraries such as Pandas and NumPy:

- ◆ **Load raw data:** The process begins by ingesting data from the extraction phase. Python can handle diverse inputs, such as using `pd.read_csv()` for flat files, `pd.read_json()` for web data, or `pd.read_sql()` to pull directly from relational databases into a DataFrame.
- ◆ **Clean and preprocess:** Once loaded, the data is audited for quality issues. This involves using functions like `.drop_duplicates()` to remove redundant records and `.fillna()` or `.dropna()` to manage missing information, ensuring the dataset is “clean” before any logic is applied.
- ◆ **Apply transformations:** This is the core stage where business logic is implemented. Python allows for complex operations such as:
 - **Formatting:** Using `.astype()` to ensure correct data types (e.g., converting strings to integers).
 - **Aggregation:** Using `.groupby()` and `.agg()` to calculate summaries like totals or averages.
 - **Feature Engineering:** Creating new columns based on existing data to provide deeper insights.

- ◆ **Generate structured output:** The final step involves converting the refined DataFrame into a format suitable for the **Load** phase. The data can be exported using `to_csv()`, `to_parquet()`, or written back to a production database, ensuring it adheres to the target schema.

2.4.2.4 Advantages of Using Python for Transformation

Python is widely considered the industry-standard language for data engineering and ETL processes due to the following characteristics:

- ◆ **Simple and readable syntax:** Python's design philosophy emphasizes code clarity and readability, which allows developers to write complex transformation logic using fewer lines of code compared to languages like Java or C++.
- ◆ **Extensive library support:** The ecosystem includes powerful, specialized libraries such as Pandas for data manipulation, NumPy for numerical computations, and SQLAlchemy for database abstraction, providing pre-built solutions for almost any data task.
- ◆ **Efficient handling of large datasets:** Through integration with high-performance tools like PySpark or Dask, Python can scale from processing small local files to managing massive, distributed datasets across cloud clusters.
- ◆ **Integration with ETL pipelines:** Python acts as the “glue” in modern data stacks, seamlessly connecting with orchestration tools like Apache Airflow or Prefect to automate the movement and scheduling of data tasks.
- ◆ **Flexibility for custom transformations:** Unlike rigid GUI-based ETL tools, Python provides the freedom to write bespoke functions and complex conditional logic, making it ideal for handling unique or non-standard data transformation requirements.

2.4.2.5 Applications in Real-World Systems

Python-based transformation is widely used in:

- ◆ Data analytics pipelines
- ◆ Business intelligence systems
- ◆ Machine learning workflows
- ◆ Data warehousing

2.4.3 Loading Transformed Data into Target Systems

After data has been extracted and transformed, the final stage in the ETL process is loading. This stage involves transferring the processed data into a target system where it can be stored, accessed, and used for analysis.

Loading is a critical step because it ensures that transformed data is made available for decision-making, reporting, and further processing.

As illustrated in Figure 2.4.1, loading represents the final stage of the ETL pipeline, where structured data is stored in a destination system.

2.4.3.1 Concept of Data Loading

Data loading is the final and crucial stage of the ETL process, where refined information is officially committed to its destination. This step ensures that the work done during extraction and transformation is preserved in a structured environment for long-term use.

- ◆ **Databases:** In many operational environments, data is loaded into relational databases (like PostgreSQL or MySQL) to support day-to-day application functions and transactional queries.
- ◆ **Data warehouses:** This involves moving highly structured data into centralized repositories like Amazon Redshift or Snowflake. These systems are specifically optimized for complex analytical queries and business intelligence reporting.
- ◆ **Data lakes:** For organizations handling massive volumes of raw or semi-structured data, loading involves storing files in systems like Hadoop or Amazon S3. This allows data scientists to access vast amounts of information for deep learning and exploratory research.
- ◆ **Cloud storage systems:** Modern workflows frequently load data into scalable cloud environments (such as Google Cloud Storage or Azure Blob Storage), providing high availability and seamless integration with various cloud-based analytics tools.

This stage ensures that data is organized and ready for use by analysts, applications, or machine learning models.

2.4.3.2 Types of Data Loading

Data loading can be performed in different ways depending on system requirements.

1. Full Load

- ◆ Entire dataset is loaded into the target system
- ◆ Used when initializing a system or replacing existing data

Example : Loading all customer records into a new database

2. Incremental Load

Only new or updated data is loaded

More efficient for large datasets

Example : Adding daily transaction records to an existing dataset

3. Batch Loading

- ◆ Data is loaded at scheduled intervals
- ◆ Suitable for large-scale processing

4. Real-Time Loading

- ◆ Data is loaded continuously as it is processed
- ◆ Used in real-time analytics systems

2.4.3.3 Loading Using Python

Python provides tools and libraries to load data into various systems:

- ◆ Writing data to CSV or text files
- ◆ Inserting data into databases using connectors
- ◆ Uploading data to cloud storage systems

These capabilities make Python an essential tool for implementing ETL workflows.

2.4.3.4 Data Loading Workflow

The loading stage involves:

1. Receiving transformed data
2. Validating data
3. Storing data into target systems
4. Making data available for use

2.4.3.5 Challenges in Data Loading

Some common challenges include:

- ◆ Data consistency issues
- ◆ Handling large volumes of data
- ◆ Maintaining performance during loading
- ◆ Ensuring data integrity

Proper planning and optimization are required to address these challenges.

2.4.3.6 Importance of Data Loading

Data loading is essential because:

- ◆ It completes the ETL process

- ◆ Makes data available for analysis
- ◆ Supports business intelligence systems
- ◆ Enables data-driven decision-making

2.4.4 Implementing Data Security Measures in Python

In data engineering workflows, handling data securely is as important as processing it efficiently. As data moves through ETL pipelines—from extraction to transformation and loading—it may contain sensitive information such as personal details, financial records, or confidential business data. Therefore, it is essential to implement data security measures to protect data from unauthorized access, misuse, or loss.

As illustrated in Figure 2.4.2, data security acts as a protective layer across all stages of the ETL process.

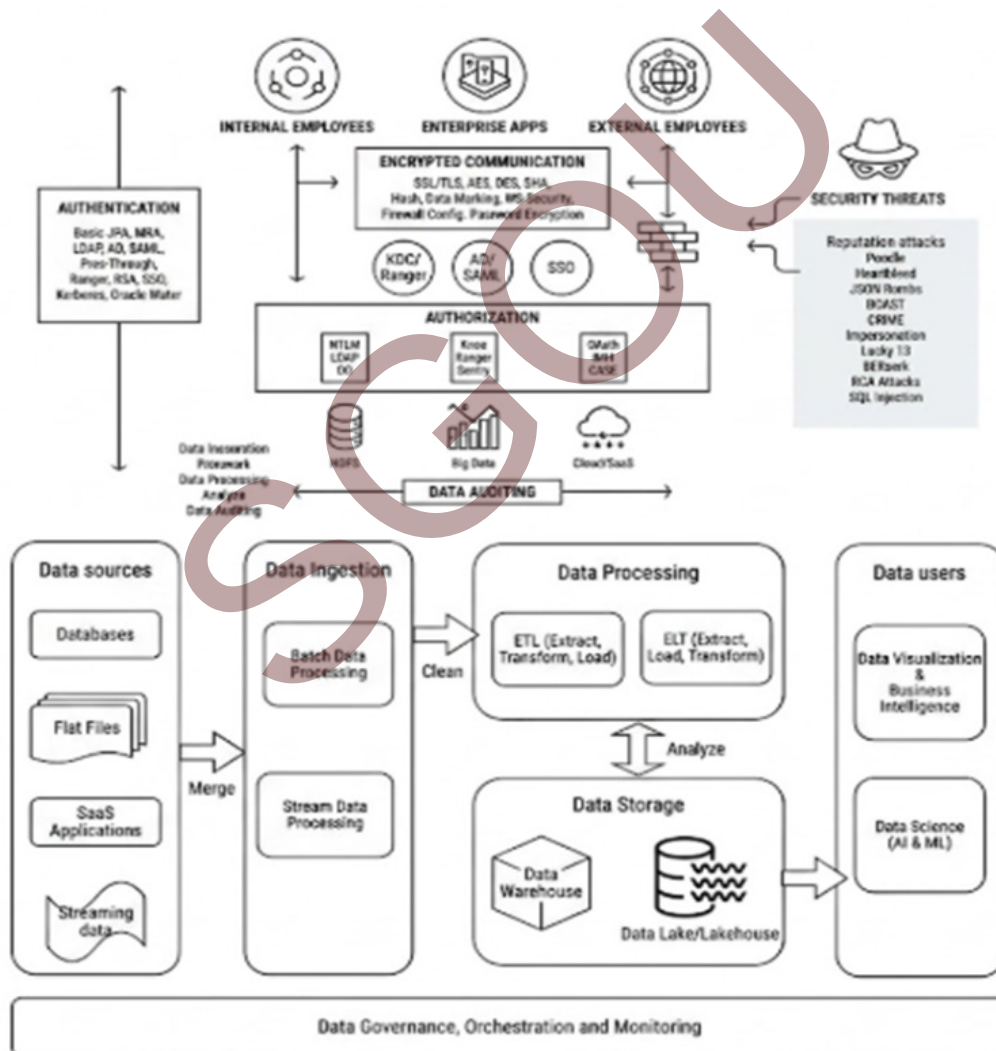


Fig. 2.4.2 Data security in ETL processes

This diagram, labelled as Fig 2.4.2 Data Security in ETL Processes, outlines a secure data architecture that integrates protection mechanisms with a standard data lifecycle.

The top portion focuses on security layers, showing how internal employees, apps, and external users must pass through authentication (using protocols like MFA and SAML), authorization (via tools like Ranger or OAuth), and encrypted communication while being monitored for security threats like SQL injection or malware. The lower portion details the technical pipeline, beginning with data sources like databases and streaming data that are merged during ingestion. The data then moves through processing—using ETL (Extract, Transform, Load) or ELT—where it is cleaned and analysed before being held in storage solutions like a Data Warehouse or Data Lake. Finally, the secured and processed data is delivered to data users for business intelligence or AI applications, all underpinned by a foundation of data governance, orchestration, and monitoring.

2.4.4.1 Importance of Data Security

Data security ensures that:

- ◆ Sensitive data is protected from unauthorized access
- ◆ Data integrity is maintained
- ◆ Data privacy regulations are followed
- ◆ Systems are protected from cyber threats

Without proper security measures, data breaches can lead to serious consequences such as financial loss and reputational damage.

2.4.4.2 Common Data Security Threats

During ETL processes, data may be exposed to various risks:

- ◆ Unauthorized access
- ◆ Data leakage
- ◆ Data corruption
- ◆ Malware and cyber attacks

Understanding these threats helps in implementing effective security measures.

2.4.4.3 Key Security Measures in Python

Python provides several methods and practices to ensure data security.

1. Data Encryption

Encryption converts data into a secure format that cannot be easily read without a key.

- ◆ Protects sensitive information
- ◆ Ensures confidentiality

Example applications:

- ◆ Encrypting files before storage
- ◆ Securing data during transmission

2. Access Control

Access control ensures that only authorized users can access data.

- ◆ Role-based access
- ◆ Authentication mechanisms

This prevents unauthorized usage of data.

3. Secure File Handling

While reading and writing files:

- ◆ Avoid exposing sensitive data
- ◆ Use secure file paths
- ◆ Restrict permissions

4. Data Masking

Sensitive data can be masked to protect privacy.

Example:

- ◆ Displaying only last four digits of a phone number

5. Secure API Communication

When using APIs:

- ◆ Use secure protocols (HTTPS)
- ◆ Authenticate requests
- ◆ Avoid exposing API keys

2.4.4.4 Best Practices for Data Security

To ensure effective data protection:

- ◆ Use strong passwords and authentication
- ◆ Encrypt sensitive data
- ◆ Regularly update systems and libraries

- ◆ Monitor data access and usage
- ◆ Backup data securely

2.4.4.5 Role of Security in ETL Workflows

As shown in Figure 2.4.4, security must be integrated at every stage:

- ◆ During extraction → secure data sources
- ◆ During transformation → protect data processing
- ◆ During loading → ensure safe storage

This ensures end-to-end data protection.

Recap

- ◆ Data transformation is the process of converting raw data into a structured and usable format
- ◆ Python libraries such as pandas and NumPy are widely used for data transformation
- ◆ Common transformation operations include cleaning, filtering, aggregation, and formatting
- ◆ Data loading is the process of storing transformed data into target systems such as databases and data warehouses
- ◆ Different loading types include full load, incremental load, batch loading, and real-time loading
- ◆ Python enables loading data into files, databases, and cloud systems
- ◆ Data security is essential in ETL processes to protect sensitive information
- ◆ Security measures include encryption, access control, data masking, and secure API communication
- ◆ Implementing security ensures data integrity, confidentiality, and reliability

Objective Questions

1. What is data transformation?
2. Which stage comes after transformation in ETL?
3. Name one Python library used for data transformation.
4. What is data loading?
5. Which type of loading replaces the entire dataset?
6. What is incremental loading?
7. Which loading type processes data continuously?
8. What is the purpose of data security?
9. What is encryption?
10. What is access control?
11. Which protocol is used for secure API communication?
12. What is data masking?
13. Which library is commonly used for data manipulation in Python?
14. What is batch loading?
15. Why is data security important in ETL?

Answers

1. Converting data into usable format
2. Load
3. pandas
4. Storing processed data
5. Full load
6. Loading new or updated data
7. Real-time loading
8. Protect data
9. Converting data into secure form
10. Restricting access



11. HTTPS
12. Hiding sensitive data
13. pandas
14. Scheduled data loading
15. Protects data integrity and privacy

Assignments

1. Explain the concept of data transformation in ETL processes.
2. Describe how Python libraries are used for data transformation.
3. Discuss different types of data loading with examples.
4. Explain the importance of data security in ETL workflows.
5. Describe various security measures used in Python for data protection

Reference

1. Python Software Foundation. pandas Documentation
2. Apache Software Foundation. ETL and Data Pipeline Documentation

Suggested Reading

1. Kleppmann, M. (2017). *Designing Data-Intensive Applications*. O'Reilly Media.
2. McKinney, W. (2018). *Python for Data Analysis*. O'Reilly Media.
3. Kimball, R. & Ross, M. (2013). *The Data Warehouse Toolkit*. Wiley.



SREENARAYANAGURU OPEN UNIVERSITY

QP CODE:

Reg. No. :

Name:

END SEMESTER DEGREE EXAMINATION
SKILL ENHANCEMENT COURSE - II
B24DS02SE - DATA ENGINEERING WITH PYTHON (CBCS - UG)
2024-25 - Admission Onwards
MODEL QUESTION PAPER- SET A

Time: 3 Hour

Max Marks: 70

SECTION A

*Answer any **ten** questions of the following. Each question carries **one** mark.*

(10 × 1 = 10 Marks)

1. What is the process of file Input/Output (I/O).
2. Which Python library is mainly used for numerical computing?
3. Which tool is used for distributed data processing in Python?
4. What is the role of Data Engineering in data-driven organizations?
5. Write any two advantages of using Python for data transformation?
6. Which stage of ETL involves collecting data from different sources?
7. Write any two advantages of Min-Max Scaling.
8. What is data preprocessing?
9. What does “a” denote in Python file modes?
10. What are the primary data structures in Pandas ?
11. Name one Python library commonly used for data analysis.
12. What is batch loading?
13. Name the three stages of ETL
14. List any two applications of Z-score Normalization.
15. What is noise in data?

SECTION B

*Answer any **five** questions of the following. Each question carries **two** marks.*

(5×2 =10 Marks)

16. Write the Key Components of Data Engineering.
17. What is the purpose of interacting with databases in data engineering?
18. Write about the concept of reshaping data with one example.
19. Explain how data is read from files in Python.
20. Define data quality.
21. Define data normalization and standardization
22. Explain about the ETL Processes.
23. Explain the significance of data security measures in Python-based data processing.
24. Describe the major challenges faced in data engineering environments.
25. Explain min-max scaling.

SECTION C

*Write a short note on any **five** questions of the following.
Each question carries **four** marks.*

(5×4 = 20 Marks)

26. Explain the steps involved in the data cleaning and transformation process.
27. Explain how data is read from files in Python with suitable examples.
28. Describe different techniques used for handling missing values.
29. Write about Min-Max Scaling.
30. How does PySpark help in handling big data applications?
31. Explain the types of Data Sources.
32. Explain the key components of Data Engineering.
33. Explain the types of data loading used in ETL.
34. Explain grouping and aggregation operations in pandas.
35. Write about data quality and outliers.

SECTION D

*Answer any **two** questions of the following. Each question carries **fifteen** marks.*

(2×15 =30 Marks)

36. Explain the concept of data transformation and discuss its importance in data preprocessing. Also explain how data can be transformed using Python libraries and functions with suitable examples.
37. Explain File I/O operations in detail. Also, write about reading from a file and writing to a file.
38. Discuss different file I/O operations and explain their role in storing and retrieving data efficiently.
39. Explain the concept of data transformation in ETL processes.



SREENARAYANAGURU OPEN UNIVERSITY

QP CODE:

Reg. No. :

Name:

END SEMESTER DEGREE EXAMINATION
SKILL ENHANCEMENT COURSE - II
B24DS02SE - DATA ENGINEERING WITH PYTHON (CBCS - UG)
2024-25 - Admission Onwards
MODEL QUESTION PAPER- SET B

Time: 3 Hour

Max Marks: 70

SECTION A

*Answer any **ten** questions of the following. Each question carries **one** mark.*

(10 × 1 = 10 Marks)

1. Which process is commonly used for data extraction, transformation, and loading?
2. What are the stages of the Data Engineering Lifecycle?
3. What does the term “Extract” mean in ETL?
4. Which file mode is used for reading a file in Python?
5. List any four key Concepts in data cleaning and preprocessing.
6. Which method is used in neural networks commonly?
7. Which type of extraction is continuous?
8. Write any common data security threats.
9. What is Data Engineering?
10. Write any 2 sources used for data generation.
11. Which library is used for tabular data manipulation in Python?
12. Which technique is used to handle missing values in datasets?
13. Which scaling technique is less sensitive to outliers?
14. Write about data extraction using Python
15. What is the purpose of data transformation?

SECTION B

Answer any **five** questions of the following. Each question carries **two** marks.

(5×2 =10 Marks)

16. What are big data processing frameworks? Give examples.
17. What is the role of ETL in the data engineering lifecycle?
18. Write any two differences between Python List and NumPy Array.
19. What are the different file modes in Python?
20. Define Outliers.
21. Why is feature scaling important in machine learning algorithms?
22. What is the importance of data extraction?
23. Differentiate between data transformation and data loading in ETL processes.
24. Write about the purpose of the **Extract (E)** and **Transform (T)** stages in the ETL process.
25. What are the challenges in data extraction

SECTION C

Write a short note on any **five** questions of the following.
Each question carries **four** marks.

(5×4 = 20 Marks)

26. Describe the structure and functionality of pandas DataFrame.
27. Explain how data is written to files in Python with suitable examples.
28. What are the techniques for handling outliers?
29. Explain the concept of Min-Max Scaling.
30. Describe the key components of a Data Engineering system.
31. What are the types of Data Loading.
32. Explain the stages of the data engineering lifecycle in detail.
33. Compare batch and real-time data extraction.
34. Explain about data security in ETL process.
35. Discuss the challenges caused by missing values and outliers and explain how preprocessing techniques address them.

SECTION D

*Answer any **two** questions of the following. Each question carries **fifteen** marks.*

(2×15 =30 Marks)

36. Write a detailed note on Data Cleaning and Transformation.
37. Explain the different techniques used for handling missing values and outliers in data preprocessing. Discuss their importance, methods, advantages, and suitable examples.
38. Explain how data is read from files in Python and discuss the importance of file writing in data workflows.
39. Explain Min-Max Scaling in detail with formula, step-by-step example, advantages, limitations, and applications.

സർവ്വകലാശാലാഗീതം

വിദ്യാൽ സ്വതന്ത്രരാകണം
വിശ്വപൗരരായി മാറണം
ഗ്രഹപ്രസാദമായ് വിളങ്ങണം
ഗുരുപ്രകാശമേ നയിക്കണം

കുരിശുട്ടിൽ നിന്നു ഞങ്ങളെ
സൂര്യവീഥിയിൽ തെളിക്കണം
സ്നേഹദീപ്തിയായ് വിളങ്ങണം
നീതിവൈജയന്തി പാറണം

ശാസ്ത്രവ്യാപ്തിയെന്നുമേകണം
ജാതിഭേദമാകെ മാറണം
ബോധരശ്മിയിൽ തിളങ്ങുവാൻ
ജ്ഞാനകേന്ദ്രമേ ജ്വലിക്കണം

കുരിപ്പുഴ ശ്രീകുമാർ

SREENARAYANAGURU OPEN UNIVERSITY

Regional Centres

Kozhikode

Govt. Arts and Science College
Meenchantha, Kozhikode,
Kerala, Pin : 673002
Ph: 04952920228
email:rckdirector@sgou.ac.in

Thalassery

Govt. Brennen College
Dharmadam, Thalassery
Kannur, Pin : 670106
Ph: 04952990494
email:rctdirector@sgou.ac.in

Tripunithura

Govt. College
Tripunithura, Ernakulam
Kerala, Pin : 682301
Ph: 04842927436
email:rcdirector@sgou.ac.in

Pattambi

Sree Neelakanta Govt. Sanskrit College
Pattambi, Palakkad
Kerala, Pin : 679303
Ph: 04662912009
email:rcddirector@sgou.ac.in

BEFORE YOU EVEN THINK OF RAGGING



Download

**ANTI
RAGGING**

App



Humiliation

Blacklisting

Ruined Career

Don't just stand and watch. Stop Ragging! Show Character

Visit UGC Website i.e. www.ugc.ac.in & www.antiragging.in to UGC Anti Ragging regulations.



सत्यमेव जयते

Ministry of Education
Government of India



ज्ञान-विज्ञान विमुक्तये

Are You Being Ragged ?

विश्वविद्यालय अनुदान आयोग
University Grants Commission
quality higher education for all

**DON'T LET IT
BE TOO LATE**

**SAY
NO
TO
DRUGS**

**LOVE YOURSELF
AND ALWAYS BE
HEALTHY**



SREENARAYANAGURU OPEN UNIVERSITY

The State University for Education, Training and Research in Blended Format, Kerala

Data Engineering with Python

COURSE CODE : B24DS02SE

SGOU



YouTube



Sreenarayanaguru Open University

Kollam, Kerala Pin- 691601, email: info@sgou.ac.in, www.sgou.ac.in Ph: +91 474 2966841

ISBN 978-81-997556-7-3



9 788199 755673