

Table of contents

Student Edition	1
1 Data Mining in the Age of AI	3
1.1 Learning Outcomes	3
1.2 Introduction	3
1.3 History of Data Mining	5
1.4 Types of Data in the AI Era	9
1.5 AI-Enhanced Data Mining Process	11
1.6 Modern Tools for AI-Enhanced Data Mining	15
1.7 Ethics and Limitations in the Age of AI	18
1.8 Key Terms	20
1.9 Review Questions	21
1.10 Lab 1 Setup for Python	22
1.11 Conclusion	28
1.12 References	28
2 Descriptive Statistics	29
2.1 Learning Outcomes	29
2.2 Introduction	30
2.3 Introduction to Descriptive Statistics	31
2.4 Frequency Tables and Distributions	34
2.5 Measures of Central Tendency	38
2.6 Measures of Dispersion	41
2.7 Shape of Distributions	45
2.8 Summary Statistics with Tools	47
2.9 Key Terms	52
2.10 Review Questions	52
2.11 Lab 2: Descriptive Statistics	53
2.12 Conclusion	63
2.13 References	64
3 Obtaining and Exploring Data	65
3.1 Learning Outcomes	65

- 3.2 Introduction 66
 - 3.3 Data Sources in the AI Era 67
 - 3.4 File Types and AI-Enhanced Processing 69
 - 3.5 Problems with Imported Data and AI Solutions 72
 - 3.6 Data Exploration Strategies 75
 - 3.7 Advanced Data Integration Techniques 78
 - 3.8 Data Quality Assessment and Monitoring 79
 - 3.9 Ethical Considerations in Data Collection and Use 80
 - 3.10 Key Terms 81
 - 3.11 Review Questions 83
 - 3.12 Lab 3: Obtaining and Exploring Data 84
 - 3.13 Conclusion 97
- 4 Data Preprocessing 99
 - 4.1 Learning Outcomes 99
 - 4.2 Introduction 99
 - 4.3 Data Cleaning 100
 - 4.4 Data Modification 102
 - 4.5 Data Restructuring 103
 - 4.6 Normalizing Data 105
 - 4.7 Reducing Dimensionality 107
 - 4.8 Principal Components Analysis 108
 - 4.9 Key Terms 116
 - 4.10 Review Questions 117
 - 4.11 Lab 4: Data Preprocessing 117
 - 4.12 Conclusion 134
- 5 Cluster Analysis 137
 - 5.1 Learning Outcomes 137
 - 5.2 Introduction 137
 - 5.3 Understanding Unsupervised Learning 138
 - 5.4 Mathematical Foundations: Distance and Similarity 139
 - 5.5 Data Preparation: Normalization and Scaling 143
 - 5.6 Clustering Methodology Overview 145
 - 5.7 Hierarchical Clustering 146
 - 5.8 K-Means Clustering: Partitioning Data into Distinct Groups . . . 151
 - 5.9 DBSCAN: Density-Based Clustering for Complex Patterns 156
 - 5.10 Comparative Analysis of Clustering Methods 160
 - 5.11 Key Terms 162
 - 5.12 Review Questions 163
 - 5.13 Lab 5: Cluster Analysis 164
 - 5.14 Conclusion 185
- 6 Classification Analysis 187

6.1	Learning Outcomes	187
6.2	Introduction	187
6.3	Supervised Learning	188
6.4	The Classification Process	190
6.5	Logistic Regression	192
6.6	Decision Trees	194
6.7	K-Nearest Neighbors (KNN)	197
6.8	Advanced Methods Preview	200
6.9	Model Evaluation and Performance Metrics	201
6.10	Comparing Classification Methods	205
6.11	Key Terms	206
6.12	Review Questions	207
6.13	Lab 6: Classification Analysis	207
6.14	Conclusion	238
7	Association Analysis	241
7.1	Learning Outcomes	241
7.2	Introduction	241
7.3	Introduction to Association Analysis	242
7.4	Itemsets and Frequent Patterns	244
7.5	Support, Confidence, and Lift	246
7.6	The Apriori Algorithm	247
7.7	Key Terms	249
7.8	Review Questions	250
7.9	Lab 7: Association Analysis	251
7.10	Conclusion	272
8	AI-Enhanced Data Mining	275
8.1	Learning Outcomes	275
8.2	Introduction	276
8.3	The AI Revolution in Data Mining	276
8.4	Automated Feature Engineering	278
8.5	Natural Language Processing: Unlocking Text Data	280
8.6	Computer Vision: Mining Visual Data	283
8.7	Ethical AI and Responsible Data Mining	285
8.8	Emerging Trends and Future Possibilities	287
8.9	Key Terms	290
8.10	Review Questions	291
8.11	Lab 8: AI Enhanced Data Mining	291
8.12	Conclusion	311

Data Mining and Exploration

Student Edition

By Kimberly Seefeld, EdD

© 2026 Data Analytics Curriculum LLC

Published under **DataJoyAI** imprint

All rights reserved. No part of this publication may be reproduced, distributed, or transmitted in any form or by any means, including photocopying, recording, or other electronic or mechanical methods, without prior written permission from the publisher, except in the case of brief quotations for scholarly review or educational use.

Companion resources for this book are available at:

The logo for DataJoyAI, featuring the word "Data" in dark blue, "Joy" in teal, and "AI" in orange.

<https://datajoyai.com>

Edition: First Edition

ISBN: 978-1-969233-31-9

Chapter 1

Data Mining in the Age of AI

1.1 Learning Outcomes

1-1. Explain the concept of data mining and its evolution from manual analysis to AI-enhanced techniques.

1-2. Identify and differentiate types of data (numeric, binary, text, image, audio, video, categorical, time series, and spatial) and describe how AI improves analysis for each.

1-3. Describe the AI-enhanced data mining process, including data exploration, preprocessing, mining, postprocessing, and application of knowledge.

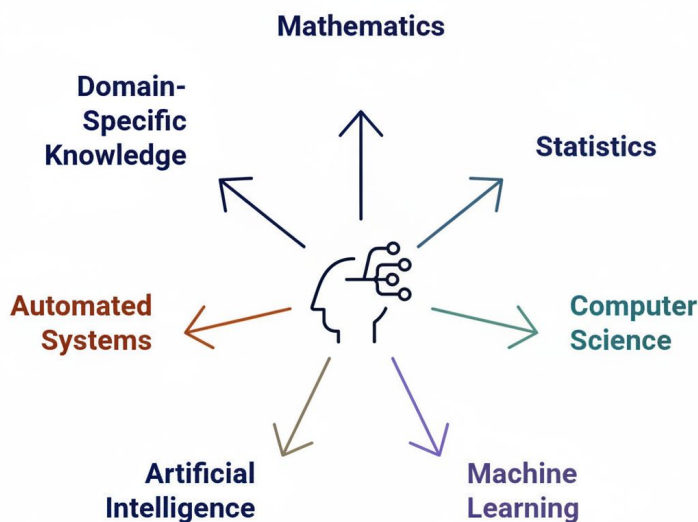
1-4. Recognize modern tools and platforms for AI-powered data mining, including traditional software, programming languages, AutoML, and no-code/low-code solutions.

1-5. Discuss ethical considerations and limitations in AI-enhanced data mining, including bias, privacy, transparency, and responsible AI practices.

1.2 Introduction

In this chapter we'll delve into the essence of data mining and cover some fundamental aspects that have been transformed by artificial intelligence. Our primary goal is to articulate what data mining entails in today's AI-enhanced landscape and set the stage for further study.

Data mining can be defined in a few ways depending on the perspective, but in 2025, it has become deeply intertwined with artificial intelligence and machine learning. Firstly, data mining can be viewed as a process. From this angle, data mining is akin to a search for patterns within a dataset, but now we have AI algorithms helping us find patterns we never could have discovered manually. Whether consumer data revealing purchasing behaviors through recommendation systems or medical data unveiling disease-symptom correlations via machine learning, the essence remains the same: identifying patterns within data. This perspective emphasizes the core aspect of data mining—uncovering meaningful patterns. In essence, data mining still mirrors the process of mining for gold, where the objective is to extract valuable insights from a vast expanse of data, but now we have much more sophisticated tools to help us dig.



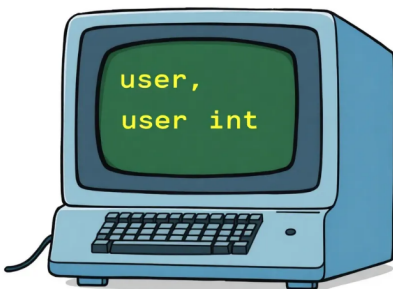
Academically, data mining is an interdisciplinary field that employs techniques from mathematics, statistics, computer science (including database management and programming), machine learning, artificial intelligence, and automated systems, along with domain-specific knowledge from fields like business and healthcare. Contrary to popular belief, data mining isn't solely a realm of mathematics or computer science, and it certainly isn't just about AI either. To set the stage for this book, it's crucial to understand that much like data analytics itself, data mining uses techniques from many academic disciplines, with AI serving as a powerful enhancement rather than a replacement

for fundamental analytical thinking.

1.3 History of Data Mining

To understand the history of data mining, it's essential to recognize that it's not entirely a novel field and has evolved over decades, with the most recent evolution being the integration of artificial intelligence. Data mining has roots that stretch back quite a while, though not in the AI-enhanced context we typically associate with it today. Let's review some key highlights to provide perspective on its evolution.

1.3.1 Early Computing (1960s–1980s)



The origins of data mining can be traced back to the 1960s. Early computing machines, though primitive by today's standards, hinted at the potential for automating analytical tasks. Before this period, data management primarily relied on manual processes, with information stored in physical documents within office settings.

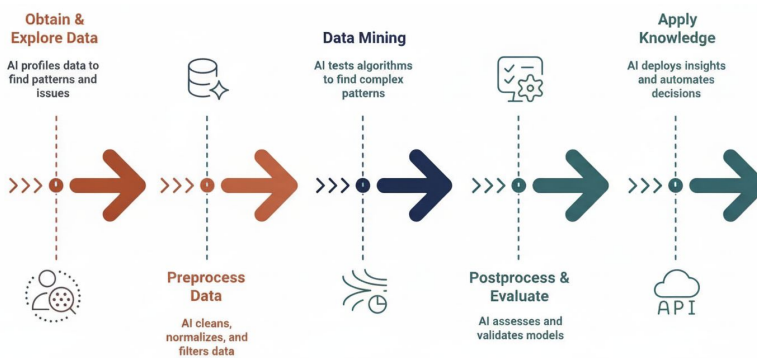
However, the landscape began to shift in the 1970s, marking the initial steps towards what we now recognize as data mining. As we moved into the 1970s, there was a gradual shift towards computerized data management. While still in its infancy, the concept of automating analytical processes became increasingly feasible with the emergence of rudimentary computing technology.

tal studies. AI has enhanced our ability to analyze location-based patterns and optimize routes and logistics.

Each type of data presents its own set of challenges and opportunities for analysis. Understanding these distinctions is essential for effectively leveraging data analytics techniques to derive actionable insights. The key difference in the AI era is that techniques of data mining can now be applied to all kinds of data more effectively, with AI serving as a powerful tool to uncover patterns that would be impossible to find manually.

1.5 AI-Enhanced Data Mining Process

The data mining process we're following is inspired by the CRISP-DM (Cross-Industry Standard Process for Data Mining) methodology, but it has evolved significantly with the integration of artificial intelligence and automation. While CRISP-DM provides a specific algorithm for data mining, our approach is adapted and modified to incorporate modern AI capabilities. The general pattern remains the same, though the specifics have been enhanced by intelligent automation.



1.5.1 Step 1: Obtain and Explore Data with AI Assistance

The initial phase involves gathering and examining the data, but now we have AI tools to help us do this more efficiently and thoroughly. Typically, data still comes from a database or data warehouse system maintained by IT professionals. An analyst will either receive the data from IT or retrieve it directly using SQL or related technology. Once the data is obtained, the analyst uploads it

Chapter 6

Classification Analysis

6.1 Learning Outcomes

- 6-1. Explain the concept of classification analysis as a supervised learning method and distinguish it from clustering and regression.
- 6-2. Describe the classification process, including data preparation, feature selection, training/testing split, model training, and validation.
- 6-3. Compare and contrast logistic regression, decision trees, and K-nearest neighbors (KNN) in terms of interpretability, accuracy, and applicability.
- 6-4. Evaluate classification models using performance metrics such as confusion matrices, accuracy, precision, recall, F1-score, ROC curves, and cross-validation.
- 6-5. Identify practical considerations for algorithm selection and understand when to explore advanced classification methods like neural networks or ensemble approaches.

6.2 Introduction

Classification analysis represents one of the most fundamental and widely used techniques in data mining. Unlike clustering, which works with unlabeled data, classification is a supervised learning method that uses labeled training data to predict categories or classes for new, unseen data. This chapter focuses on three core classification methods that form the foundation of predictive analytics, with

hands-on examples using modern AI-enhanced tools.

6.3 Supervised Learning

Classification falls under the umbrella of supervised learning, a branch of machine learning where algorithms learn from labeled training data to make predictions about new data. The fundamental difference between supervised and unsupervised learning lies in the presence of target variables or labels that guide the learning process.

6.3.1 Labeled Data and Target Variables



In supervised learning, we work with datasets that contain both input features (independent variables) and known outcomes (dependent variables or labels). For example, if we're building a model to predict whether an email is spam or not spam, our training data would include email features like sender information, subject line characteristics, and word frequencies, along with the known classification of each email as "spam" or "not spam."

The target variable, also called the response variable or dependent variable, represents what we want to predict. In classification problems, this target variable is categorical, meaning it represents discrete classes or categories rather than continuous numerical values. Common examples include predicting customer churn (yes/no), medical diagnosis (disease/no disease), or product categories

(electronics/clothing/books).

Modern AI tools can now automatically detect the type of target variable in your dataset and suggest appropriate classification algorithms. Platforms like Orange, WEKA, or even business intelligence tools like Tableau can analyze your data structure and recommend whether classification or regression techniques are most suitable.

6.3.2 Classification vs. Regression



While both classification and regression are supervised learning techniques, they differ in their target variables. Classification predicts discrete categories or classes, while regression predicts continuous numerical values. For instance, predicting house prices would be a regression problem (continuous values), while predicting house price categories (low/medium/high) would be a classification problem (discrete categories).

Understanding this distinction is crucial because it determines which algorithms and evaluation metrics are appropriate for your analysis. Classification algorithms are specifically designed to handle categorical outcomes and use different mathematical approaches than regression algorithms.

ing or decreasing the likelihood of a particular outcome. This makes the model especially useful when understanding the relationship between variables is just as important as making predictions. Logistic regression also works best when the classes can be reasonably separated using a linear boundary, meaning a straight-line (or flat surface in higher dimensions) can divide the groups.

6.13.1.6.1 Implementing Logistic Regression

In this section, the logistic regression model is created and trained using the training dataset. The model learns patterns by adjusting its internal coefficients so that the predicted probabilities match the actual class labels as closely as possible. The `fit()` function performs this learning process.

After training, the model is used to make predictions on new, unseen data stored in the test set. The `predict()` function returns the final predicted class for each observation, while `predict_proba()` provides the probability that each observation belongs to each possible class. This distinction is important because it allows us to see not just what the model predicts, but how confident it is in those predictions.

Finally, the model's performance is evaluated using accuracy scores on both the training and testing datasets. The training accuracy shows how well the model learned from the data it was given, while the testing accuracy indicates how well it generalizes to new data.

```
from sklearn.linear_model import LogisticRegression

log_reg = LogisticRegression(random_state=42, max_iter=1000)

# Create and train logistic regression model
print("LOGISTIC REGRESSION ANALYSIS")
print("=" * 50)

# Initialize the model
log_reg = LogisticRegression(random_state=42, max_iter=1000)

# Train the model
```

```
log_reg.fit(X_train, y_train)

# Make predictions
y_pred_log = log_reg.predict(X_test)
y_pred_proba_log = log_reg.predict_proba(X_test)

print("Model Training Complete!")
print(f"Training Accuracy: {log_reg.score(X_train, y_train):.3f}")
print(f"Testing Accuracy: {log_reg.score(X_test, y_test):.3f}")
```

Output:

```
Model Training Complete!
Training Accuracy: 0.931
Testing Accuracy: 0.933
```

6.13.1.6.2 Interpreting Logistic Regression Coefficients

In this section, the model's coefficients are examined to understand how each feature influences the predictions. When the model was trained, it assigned a weight to every feature, and these weights are stored in `log_reg.coef_`. By organizing these values into a table, we can clearly see how each variable contributes to predicting each class.

```
# Analyze feature importance through coefficients
feature_names = X.columns
classes = log_reg.classes_

print("\nLogistic Regression Coefficients Analysis")
print("=" * 45)

# Create coefficient interpretation table

coef_df = pd.DataFrame(
    log_reg.coef_.T,
    index=feature_names,
    columns=[f'Coef_{c}' for c in classes]
```

1 Decision Tree 0.933

2 K-Nearest Neighbors 0.933

```
# Detailed classification reports
for model_name, (predictions, _) in models.items():
    print(f"\n{model_name} - Detailed Performance:")
    print("-" * 45)
    print(classification_report(y_test, predictions))
```

Output:

Logistic Regression - Detailed Performance:

	precision	recall	f1-score	support
Excellent	0.93	1.00	0.97	140
Needs Improvement	0.00	0.00	0.00	1
Satisfactory	0.00	0.00	0.00	9
accuracy			0.93	150
macro avg	0.31	0.33	0.32	150
weighted avg	0.87	0.93	0.90	150

Decision Tree - Detailed Performance:

	precision	recall	f1-score	support
Excellent	0.93	1.00	0.97	140
Needs Improvement	0.00	0.00	0.00	1
Satisfactory	0.00	0.00	0.00	9
accuracy			0.93	150
macro avg	0.31	0.33	0.32	150
weighted avg	0.87	0.93	0.90	150

K-Nearest Neighbors - Detailed Performance:

	precision	recall	f1-score	support
Excellent	0.93	1.00	0.97	140
Needs Improvement	0.00	0.00	0.00	1
Satisfactory	0.00	0.00	0.00	9
accuracy			0.93	150
macro avg	0.31	0.33	0.32	150
weighted avg	0.87	0.93	0.90	150

6.13.1.9.1 Confusion Matrix Analysis

The final part focuses on confusion matrix analysis, which visualizes where each model is making correct versus incorrect predictions. A confusion matrix shows actual classes along one axis and predicted classes along the other. The diagonal values represent correct predictions, while off-diagonal values indicate misclassifications. The code plots all three confusion matrices side by side using heatmaps, making it easier to compare model performance visually.

Overall, this workflow helps not only to quantify how well each model performs but also to interpret the types of errors each model makes. By combining accuracy scores, classification reports, and confusion matrices, you get a comprehensive picture of model strengths and weaknesses, enabling informed decisions about which model to choose for deployment or further tuning.

```
# Create confusion matrices for all models
print("\nCONFUSION MATRICES (TABLE FORMAT)")
print("=" * 40)

for model_name, (predictions, _) in models.items():
    cm = confusion_matrix(y_test, predictions)

    print(f"\n{model_name}")
    print("-" * len(model_name))
```

```
print(pd.DataFrame(
    cm,
    index=[f"Actual {cls}" for cls in log_reg.classes_],
    columns=[f"Pred {cls}" for cls in log_reg.classes_]
))
```

Output:

Logistic Regression

	Pred Excellent	Pred Needs Improvement	\
Actual Excellent	140	0	
Actual Needs Improvement	1	0	
Actual Satisfactory	9	0	

	Pred Satisfactory
Actual Excellent	0
Actual Needs Improvement	0
Actual Satisfactory	0

Decision Tree

	Pred Excellent	Pred Needs Improvement	\
Actual Excellent	140	0	
Actual Needs Improvement	1	0	
Actual Satisfactory	9	0	

	Pred Satisfactory
Actual Excellent	0
Actual Needs Improvement	0
Actual Satisfactory	0

K-Nearest Neighbors

```

        'Support': f"{rule['support']:.3f}",
        'Confidence': f"{rule['confidence']:.3f}",
        'Lift': f"{rule['lift']:.3f}",
        'Count': rule['count']
    })

    return pd.DataFrame(summary_data)

# Run the Apriori implementation
apriori = AprioriAlgorithm(min_support=0.2, min_confidence=0.6)
apriori.fit(transactions)

# Display results
print(f"\nFINAL ASSOCIATION RULES")
print("=" * 30)
rules_summary = apriori.get_rules_summary()
print(rules_summary.to_string(index=False))

```

When the Apriori algorithm starts, it first prints out the settings and size of your dataset. Here, the minimum support is set at 0.2, meaning only itemsets that appear in at least 20% of transactions will be considered “frequent.” The minimum confidence is 0.6, so only rules that are correct at least 60% of the time will be kept. The algorithm sees that there are 10 transactions in total.

7.9.1.4.1 Phase 1: Finding Frequent Itemsets

The algorithm begins by examining individual items (1-itemsets) to see which ones appear frequently enough. It finds 7 frequent 1-itemsets, such as bread, milk, and butter. Then it combines these frequent items to form 2-itemsets (pairs) and checks which of these pairs meet the minimum support. For example, bread & milk occurs in 40% of transactions, so it is included. This process continues to 3-itemsets (triples), such as bread, butter, milk, which occurs in 30% of transactions. The algorithm stops when no larger frequent combinations can be found.

By the end of Phase 1, the algorithm has 18 frequent itemsets: 7 single items, 8

8.9 Key Terms

AI-Enhanced Data Mining – Using artificial intelligence to improve or automate data mining tasks.

Automated Feature Engineering – AI-driven creation of new variables to improve predictive performance.

Feature Selection – Identifying the most relevant variables for analysis to avoid overfitting and improve interpretability.

Natural Language Processing (NLP) – AI techniques for analyzing and extracting insights from text data.

Sentiment Analysis – Determining emotional tone or opinion from text.

Named Entity Recognition (NER) – Identifying and classifying key elements in text, such as names, dates, or products.

Topic Modeling – Discovering hidden themes in large text datasets.

Computer Vision – AI methods to interpret and analyze images and videos.

Image Classification / Object Detection / Image Segmentation / OCR – Core computer vision techniques for analyzing visual data.

Ethical AI – AI designed to be fair, transparent, accountable, and aligned with human values.

Explainable AI (XAI) – Methods that make AI decision-making understandable to humans.

Federated Learning – Collaborative machine learning without centralizing data.

Generative AI – AI that can create new content or synthetic data for training and analysis.

Human-AI Collaboration / Augmented Analytics – Combining AI capabilities with human expertise for better insights and decisions.

8.10 Review Questions

1. How does AI enhance traditional data mining techniques like clustering and classification?
2. What is automated feature engineering, and why is it important for predictive modeling?
3. Explain the difference between NLP and computer vision in AI-enhanced data mining.
4. What are some key techniques used in NLP for data mining applications?
5. How does computer vision transform analysis of image and video data?
6. Why is ethical AI important, and what types of bias can occur in AI systems?
7. Describe strategies for mitigating bias and ensuring transparency in AI-enhanced data mining.
8. How can generative AI be used to create synthetic data or generate insights?
9. What is federated learning, and why is it useful for privacy-sensitive domains?
10. Explain the concept of human-AI collaboration and how it supports augmented analytics.

8.11 Lab 8: AI Enhanced Data Mining

This lab demonstrates three core areas:

Automated Feature Engineering – Learn how AI-inspired methods can generate new variables from raw data, capture complex relationships, and improve the predictive potential of datasets. Students will practice creating mathematical transformations, interaction features, aggregation metrics, and categorical bins, followed by intelligent feature selection using statistical methods.

Natural Language Processing (NLP) for Text Mining – Examine how unstructured text data, such as customer reviews, can be transformed into structured