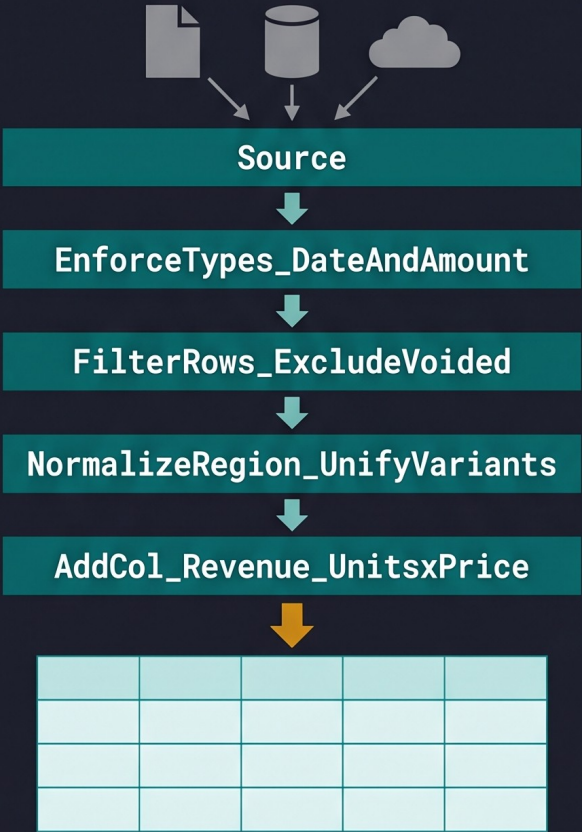


DATA TRANSFORMATION WITH POWER QUERY

*Self-Service ETL, Data Preparation Pipelines,
and the M Language for Analysts*



 *Refreshes automatically – every cycle, from the same pipeline*

FRU KINGSLY

GURUTECH CONSULTING INTERNATIONAL



HALF-TITLE PAGE

Data Transformation with Power Query

TITLE PAGE

DATA TRANSFORMATION WITH POWER QUERY

Self-Service ETL, Data Preparation Pipelines, and the M Language
for Analysts

Excel for Data Analytics & Business Intelligence Pathway

Volume 3 of 6

FRU Kingsly

GuruTech Consulting

2025

COPYRIGHT PAGE

© 2025 FRU Kingsly

All rights reserved.

No part of this book may be reproduced, stored in a retrieval system, or transmitted in any form or by any means – electronic, mechanical, photocopying, recording, or otherwise – without prior written permission of the author, except for brief quotations used in reviews or scholarly analysis.

Published by GuruTech Consulting International

ISBN: [To be assigned – Print]

ISBN: [To be assigned – eBook]

Series: Excel for Data Analytics & Business Intelligence Pathway

Microsoft Excel, Power Query, Power Pivot, Power BI, and related product names are trademarks or registered trademarks of Microsoft Corporation. This book is an independent publication and is not affiliated with, authorised by, endorsed by, or in any way officially connected with Microsoft Corporation.

This book was developed with the assistance of artificial intelligence tools in research and structural development. The analytical frameworks, pedagogical architecture, editorial judgments, and all final content decisions are the author's own.

All case illustrations, working examples, datasets, and scenario descriptions in this book are fictional composites constructed for pedagogical purposes. No case represents a specific named organisation, system, or event. Any resemblance to real organisations, persons, or events is coincidental and unintentional.

DEDICATION

For every analyst who has lost an afternoon to a spreadsheet that should have been a pipeline.

EPIGRAPH

"The most common source of error in analysis is not bad logic — it is data that arrived incorrectly prepared and was never questioned."

ACKNOWLEDGMENTS

A book about data preparation is, at its core, about the quality of the decisions made before analysis begins — and no analyst ever develops that quality alone.

I am indebted to the data analysts, business intelligence developers, reporting managers, and operations practitioners who shared their most persistent data preparation problems with me over more than a decade of analytical work. The situations this book describes — extract jobs run manually every Monday morning because nobody built the pipeline, date columns coerced through formula chains that break when the locale changes, join failures traced to trailing spaces that no one thought to remove, wide-format reports that resisted every aggregation attempt until someone finally unpivoted them — are drawn from real professional experience, even where the names and details have been changed. You trusted me with the problems that resisted your best tools. I hope this book gives others the discipline to reach the solution before the deadline.

Thank you to the Power Query engineering community, the M language practitioners who have built and published their knowledge in the open, and the broader data engineering scholarship whose foundational concepts this book's pipeline architecture is built upon — particularly those who insisted that professional-grade data preparation is not the exclusive domain of database engineers and that the discipline is available to any analyst willing to apply it. The intellectual debts are detailed in the Notes and References at the close of this volume.

Thank you to the reviewers who engaged draft material and challenged the formulations where the staging-layer logic was under-specified, the schema alignment standards where edge cases had

been insufficiently addressed, and the M language chapters where the expressions required more defensive handling than the narrative had supplied. This book is more deployable because you held it to the standard of someone who would actually need to use it under production conditions.

This book was developed with the assistance of artificial intelligence tools in research and structural development. The analytical frameworks, the pedagogical architecture, the editorial judgments, and all final content decisions are the author's own.

To every analyst who has stared at a source file too messy for a formula and too urgent for a consultant: the pipeline you needed was closer than you thought. This book is how to build it.

ABOUT THE SERIES

The *Excel for Data Analytics & Business Intelligence Pathway* is a six-volume Core Series designed to take an analyst from the behavioural foundations of clean Excel work through the full BI stack — calculation architecture, data transformation, relational modelling, and visual communication — to the integrative professional practice that closes the series. The Core Series is complete at six volumes; it is complemented by a separate two-book Technology Extensions Collection for readers who wish to pursue advanced mastery in specific technologies after completing the Core Series.

The series follows a single directional logic. Each volume owns one layer of the analytical stack and hands its outputs cleanly to the volume that follows. No capability is taught in more than one book as a primary subject. The series may be entered at any volume by a learner with the appropriate prerequisite skills, but it is designed to be read in order.

The Six Volumes

Volume 1 — *Excel for Analysts: The Analytical Foundation*

Subtitle: Data Structuring, Workbook Architecture, and the MO-210 Foundation for Analytical Work

Establishes Excel not as a spreadsheet tool but as a structured analytical environment. Every feature is taught in the analytical context where it actually matters. Learners who complete this volume do not merely pass the MO-210 benchmark — they think in structured data, named ranges, and layered workbooks. That mental model is the prerequisite for every subsequent volume in the series.

Certification alignment: Microsoft Office Specialist Excel Associate (MO-210) — primary.

Volume 2 — Analytical Formulas & Calculation Architecture

Subtitle: Advanced Logic, Lookup Mastery, Dynamic Arrays, and the MO-211 Calculation Engine

Builds the calculation engine. Formulas are treated as a professional craft discipline — not as a list of function syntax to be memorised. The volume covers lookup architecture, logical design, error-proofing, LET and LAMBDA, and the dynamic array functions that have fundamentally changed how Excel calculates. This volume delivers the most analytically powerful formula toolkit the platform offers.

Certification alignment: Microsoft Office Specialist Excel Expert (MO-211) — primary.

Volume 3 — Data Transformation with Power Query (This Volume)**

Subtitle: Self-Service ETL, Data Preparation Pipelines, and the M Language for Analysts

Automates the data preparation layer. Power Query, the M language, and the data preparation pipeline are the exclusive subjects. The volume reframes data transformation from a chore performed before the “real” analysis begins into a professional engineering discipline with its own architecture, audit standards, and production practices.

Certification alignment: None direct. Foundational to PL-300 and DP-600.

Volume 4 — Data Modeling with Power Pivot & DAX

Subtitle: Star Schema Design, Relational Thinking, and Measure-Driven Analysis

Introduces the relational modelling layer. The Data Model, Power Pivot, star schema design, and DAX are the exclusive subjects of this volume. This is the analytical engine room of the series —

where the prepared data arriving from Volume 3 is structured into relationships, and where measures replace formulas as the primary analytical instrument.

Certification alignment: Foundational to PL-300.

Volume 5 — Data Visualization & Dashboard Design

Subtitle: Chart Architecture, Dashboard Discipline, and Stakeholder-Facing Visual Communication

Delivers analytical outputs to decision-makers through structured visual communication. This is the only volume in the series whose central question is *how the work is read*, not how it is computed. Visualisation is treated as an analytical responsibility — the discipline of selection, structure, and signal clarity that determines whether insight reaches its audience.

Certification alignment: MO-210 and MO-211 charting and Pivot-Table domains — addressed at a level that significantly exceeds exam requirements.

Capstone Volume — The Excel Analyst at Work

Subtitle: End-to-End Analytical Practice, Professional Integration, and the Analyst's Ecosystem

Integrates every capability built across Volumes 1–5 into the kind of end-to-end analytical work the reader will encounter as a professional Excel analyst. Organised around realistic project scenarios — each requiring the full stack in sequence under realistic professional constraints. This is also the volume that looks outward: learners are asked for the first time to reason about Excel's position within the broader analytics ecosystem, and to recognise when a workload has outgrown it.

Certification alignment: None. Integrative practice and professional simulation.

The series is complete. No domain identified in the scope is left uncovered. No tool or capability is taught in more than one volume as a primary subject.

The Technology Extensions Collection

The six-volume Core Series is complete and closed. For readers who finish it and wish to pursue advanced mastery in specific technologies, a separate two-book **Technology Extensions Collection** continues the work. The extensions are a companion collection, not Volumes 7 and 8: the Core count remains six. The two extensions are independent of each other and may be read in either order after the Core Series.

Technology Extension TE-1 — *Mastering DAX for Analytical Intelligence*

Subtitle: Advanced Calculation Patterns, Context Mastery, and Analytical Modeling at Scale

A vertical extension of Volume 4. It takes a competent DAX writer — one who can author measures, use CALCULATE, and apply basic time intelligence — and builds an advanced analytical modeller who understands evaluation context, context transition, iterators, virtual tables, and performance optimisation, and who can implement sophisticated business calculation patterns at scale. Prerequisite: Volume 4 (mandatory); Volume 5 recommended.

Technology Extension TE-2 — *Excel Automation for Analysts*

Subtitle: VBA, Office Scripts, and Power Automate for Systematic Analytical Delivery

A horizontal automation layer over all six Core volumes. It treats VBA, Office Scripts, and Power Automate as a unified analytical automation toolkit, teaching analysts to systematise the production and delivery of the solutions they have already learned to build — report automation, scheduled refresh, and cross-platform workflows — with a principled framework for choosing among the tools. Prerequisite: all six Core Series volumes.

FOREWORD

Foreword by Dr. Amara Osei-Mensah

Associate Professor of Business Analytics and Data Systems, Whitmore Graduate School of Management, University of the Great Lakes, Toronto, Canada.

In fifteen years of teaching business analytics, I have developed a reliable diagnostic for the state of data preparation practice in any organisation: I ask the analyst responsible for the most important recurring report how long the preparation takes each cycle. Not the analysis — the preparation. The connecting, the cleaning, the reshaping, the coercing of dates into the format the formula expects, the removal of the blank rows that the ERP appended above the column headers.

The answers cluster between forty minutes and three hours. The median is around ninety minutes. And when I ask whether the process would survive the analyst's unplanned absence — whether a colleague could step in and produce the same output from the same source without guidance — the answer is, almost universally, no.

This is the problem that *Data Transformation with Power Query* is written to solve. It is not a small problem dressed in the language of operational efficiency. It is a structural failure in how analytical organisations treat the layer that everything they build depends on. The preparation layer is the foundation of every analytical output — every chart, every KPI, every dashboard, every model. When it is manual, undocumented, and person-dependent, every analytical output built on top of it inherits those properties. The correct KPI can be calculated from incorrectly prepared data. The well-

designed dashboard can be populated with incorrectly joined tables. The insight can be analytically sound and numerically wrong.

I have observed this failure mode across industries and across organisation sizes, and I have come to believe that the most consequential professional development an analyst can pursue is not the next visualisation tool, the next modelling technique, or the next certification — it is the discipline to build the preparation layer to a professional standard. That discipline is what this book teaches.

What distinguishes this book from the technical documentation that surrounds the Power Query platform is its insistence on the pipeline as a professional artefact. The ribbon experience — clicking through the interface, applying transformations by selecting options in dialogs — produces output. The professional discipline this book establishes produces output that is documented, auditable, and transferable. The distinction is the difference between an analyst who has solved the problem this month and an analyst who has built an infrastructure that solves the problem every month, for every analyst who inherits the workbook, for every auditor who asks how the number was produced.

The architecture this book builds — the staging layer, the shape layer, the combine layer, the M language layer, the production readiness standard — is an analytical engineering discipline. It applies the same professional standards to the transformation layer that software engineering has long applied to application code: modularity, documentation, testability, and handoff-readiness. These are not new principles. They are new in the context of the spreadsheet analyst who has been preparing data manually, without a framework that names what they are doing or specifies what doing it well looks like.

I have taught this material, or material adjacent to it, at the graduate level for a decade. What this book adds that most curricula lack is the integration of the technical instruction with the professional discipline. The step naming standard, the profiling log, the schema contract, the refresh runbook, the production readi-

ness audit – these are not supplementary to the technical content. They are what make the technical content professionally deployable. An analyst who knows every Power Query connector but documents nothing will build pipelines that cannot be inherited, audited, or trusted by anyone who did not build them. An analyst who has absorbed the discipline this book teaches will build pipelines that can be inherited, audited, and trusted by anyone who receives them.

The series this volume belongs to is an ambitious commitment to building analysts who are not merely technically capable but professionally complete. Volume 3 makes that commitment concrete at the layer where most analytical work silently fails. I commend it to practitioners who have sensed that there is a more professional way to prepare data than the one they have been using, and who are ready to build it.

Amara Osei-Mensah

Whitmore Graduate School of Management

University of the Great Lakes, Toronto

PREFACE

Data preparation is the most underestimated discipline in analytical work.

It is the work done before the charts, before the formulas, before the KPIs — the invisible foundation that either supports or undermines everything built on top of it. And yet, in most organisations, it is done manually: rows copied between sheets, columns cleaned with `TRIM(CLEAN())` chains, dates coerced with `TEXT()` workarounds, and the whole fragile construction re-executed from scratch the next time the source file arrives.

This book exists because that approach is neither necessary nor professional.

Power Query is a full ETL engine embedded inside Microsoft Excel. It connects to virtually any data source, transforms data through a structured, auditable pipeline of applied steps, and loads a clean, typed, reproducible output to wherever downstream analysis needs it. Every transformation is authored once. Every refresh replays the full pipeline automatically. Every step is visible, editable, and documented in the query itself.

These are not features. They are the behavioural foundations of professional data preparation — and learning them changes not just how analysts prepare data, but how they think about the relationship between raw inputs and analytical outputs.

This book teaches Power Query as a discipline, not as a ribbon experience. It does not cover every connector, every M function, or every edge case in the platform. What it covers, it covers completely: the structural thinking behind pipeline design, the professional habits that make queries readable and maintainable, the M language as an analytical reading and writing skill, and the produc-

tion practices that determine whether a pipeline survives contact with the Monday morning refresh.

The reader who finishes this book will not simply know how to use Power Query. They will know how to design a data preparation pipeline that is auditable, reproducible, parameterised, and ready to hand off — which is a different capability entirely, and a considerably more valuable one.

This volume is the third in the *Excel for Data Analytics & Business Intelligence Pathway*. It receives structured, clean workbooks from Book 1 and calculation-capable analysts from Book 2. It delivers prepared, typed, refresh-ready data to the Data Model that Book 4 will teach. The ETL layer it establishes is the foundation of the relational BI capability the series builds from here.

PREFACE TO THE SERIES

The *Excel for Data Analytics & Business Intelligence Pathway* was designed to answer a single question: what would it take to build a professional Excel analyst from the ground up, with no shortcuts and no gaps?

Not a power user. Not someone who knows a lot of functions. A professional analyst — someone who structures data with discipline, calculates with architectural intent, prepares pipelines that survive the Monday morning refresh, models data relationally, and communicates findings with the clarity that real decision-making requires. That is a complete and demanding capability set. No single book covers it. No existing series covers it without redundancy, inconsistency, or omission.

The answer was six books, each owning exactly one layer of the analytical stack, built in a strict dependency order from mindset to calculation to transformation to modelling to communication to integration. Every volume hands something clean and specific to the one that follows. No volume retreats to cover ground that a prior volume owns.

This series is the result of that design logic. Each volume is written to the standard that the next one demands of it.

HOW TO USE THIS BOOK

This book is designed to be read from front to back on a first pass. The six Parts are cumulative: Part I establishes the ETL discipline and the editor environment; Part II covers connection to data sources; Part III covers profiling and cleaning; Part IV covers shaping and combining; Part V introduces the M language as an analytical reading and writing skill; Part VI covers performance, refresh, and production readiness. Each Part assumes the concepts and vocabulary of every Part before it. Readers building this skill from the ground up should work through the book in order.

Chapter Structure

Each chapter is built around a consistent internal architecture. Understanding that architecture before beginning will help you use each chapter with maximum efficiency.

Chapter Introduction situates the chapter within the progression of the book — what is about to be taught, why it matters at this point in the pipeline, and how it connects to what came before.

Chapter Objectives state precisely what the reader should be able to do upon completing the chapter. Return to these objectives after completing the chapter as a self-assessment instrument.

Conceptual Sections (numbered X.1, X.2, X.3 within each chapter) form the intellectual backbone. They move from concept to professional application in sequence. Read them in order. Each conceptual section advances the argument of the chapter; none is merely supplementary.

Analyst's Corner sections appear in selected chapters and present practitioner-facing commentary — the texture of real-world application that supplements structured instruction. These are not

anecdotes; they are professional framings of the concepts the chapter has established.

Decision Lab sections are scenario-based applied challenges. A realistic analytical problem is presented in full professional context. The reader is walked through the reasoning required to solve it, and a complete author's analysis closes each lab. Decision Labs are not exercises — they are worked analytical cases. They should be read as fully as the conceptual sections.

Analytical Toolkit sections provide reference instruments that the analyst carries forward into professional practice: checklists, naming convention guides, decision matrices, connection registries, production readiness audits, and style guides. These are designed to be used on real work, not merely read in the context of instruction.

Practice Quiz sections test conceptual retention at the end of each chapter. Every question is followed by an answer key with full rationale, not merely a correct letter. Read the rationale for every question — including those answered correctly. The rationale is often where the nuanced understanding lives.

Exercises / Applied Practice sections provide hands-on tasks at three levels of cognitive demand:

- *Comprehension* exercises confirm that the reader can execute the mechanics just taught.
- *Application* exercises require using those mechanics on real analytical problems where the expected outcome is not pre-stated.
- *Judgment* exercises require the reader to recommend, defend, prioritise, or evaluate — which is what professional analytical work actually demands.

Capstone exercises appear in final chapters and require integrated application across multiple skills.

Reflection / Discussion Questions are designed for cohort-based learners and for individual readers who want to extend their thinking beyond the mechanics of execution into the reasoning behind professional practice.

What's Next closes each chapter with a transition that positions the reader clearly for what follows.

Figures

Every figure in this book serves an analytical purpose — either showing how a specific feature or behaviour appears inside the Power Query editor, or representing a structural or architectural concept that prose alone does not communicate effectively. Each figure is referenced in the text at the point where it supports the argument. Figures are not supplementary; several carry information that is not duplicated in the surrounding text.

Working Files

The exercises throughout this book are designed for use with provided working files containing synthetic datasets constructed for analytical realism. Working files are available at: <https://gurutech.consulting/resources/excel-series/book3>

Reading This Book as Part of the Series

This is Volume 3 of the *Excel for Data Analytics & Business Intelligence Pathway*. It assumes familiarity with the foundational capabilities established in Volumes 1 and 2. Readers who have not completed those volumes should review the Prerequisites section and confirm they carry the assumed capabilities before proceeding.

SERIES ARCHITECTURE

The *Excel for Data Analytics & Business Intelligence Pathway* is built on a single structural logic: each volume establishes one layer of the Excel analytical stack and passes its outputs cleanly to the volume that follows. The series does not spiral back to re-cover ground. It builds forward. A capability introduced as a primary subject in one volume is assumed, extended, and consumed in all volumes that follow — never re-taught from the beginning.

The progression is as follows.

Volume 1 — *Excel for Analysts: The Analytical Foundation*

Establishes the analytical mindset and workbook discipline that every subsequent volume assumes without re-teaching. The behavioural and structural foundation: structured data, layered workbooks, Table objects, named ranges, and the clean separation between raw data and analytical output. The output of this volume is an analyst who has already stopped thinking in unstructured ranges and formatting-first habits.



Volume 2 — *Analytical Formulas & Calculation Architecture*

Builds the calculation engine on the structured foundation Volume 1 established. Every formula is treated as a small argument about how data relates, designed for correctness, auditability, and resilience to change. The output of this volume is an analyst with a formula architecture capability — lookup hierarchies, logical structures, dynamic arrays, LET and LAMBDA abstractions, and the discipline of making calculation logic readable by others.



Volume 3 – Data Transformation with Power Query (This Volume)**

Automates the data preparation layer. The output of this volume is an analyst who can design and deploy a layered, parameterised, refresh-ready ETL pipeline – connecting to any supported source type, transforming with M, enforcing types, shaping structure, combining from multiple origins, and handing off clean, correctly structured data to the Data Model.



Volume 4 – Data Modeling with Power Pivot & DAX

Receives the prepared data from Volume 3 and builds the relational analytical model on top of it. Star schema design, Power Pivot relationships, DAX filter context, and explicit measures replace flat-table thinking and formula-driven aggregation. The output of this volume is an analyst with a relational analytical capability that is portable across the entire Microsoft BI ecosystem.



Volume 5 – Data Visualization & Dashboard Design

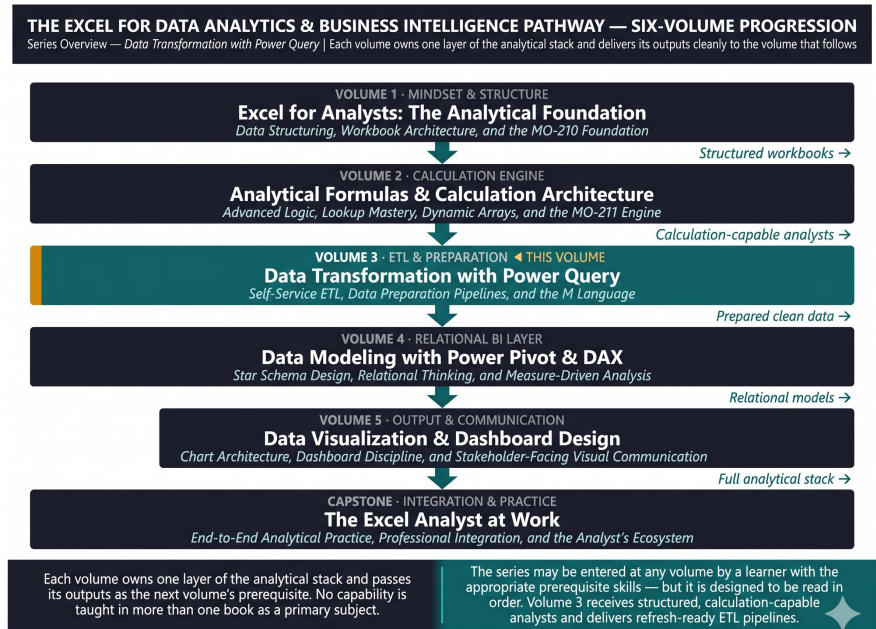
Receives the modelled data from Volume 4 and delivers it to decision-makers through structured visual communication. The output of this volume is an analyst who treats chart selection, layout, formatting, and interaction design as analytical judgments – and who can produce dashboards that communicate insight without requiring interpretation by their author.



Capstone – The Excel Analyst at Work

Integrates all five layers in realistic project scenarios under professional constraints. The output of this volume is an analyst who can scope, execute, document, and hand off an end-to-end analytical workflow – and who understands where Excel sits in the broader analytical ecosystem, and when a workload has exceeded what Excel should carry.

The six volumes above are the complete Core Series. Beyond them sits a separate companion collection — the two-book Technology Extensions Collection — for readers pursuing advanced mastery after the Core Series. The extensions are not new steps in the core sequence and not Volumes 7 and 8: the Core count stays at six. TE-1 (*Mastering DAX for Analytical Intelligence*) deepens the DAX layer introduced in Volume 4; TE-2 (*Excel Automation for Analysts*) adds a systematic automation layer over all six volumes. The two are independent and may be read in either order once the Core Series is complete.



Volume	Receives From Prior Volume	Delivers to the Next
Book 1	No prerequisite	Structured work-books, layered architecture, analytical discipline
Book 2	Structured work-books	Formula architecture, calculation-capable analysts
Book 3	Calculation-capable analysts	Refresh-ready ETL pipelines, prepared clean data
Book 4	Prepared clean data	Relational models, DAX measures
Book 5	Relational models	Insight-ready dashboards and reports
Capstone	All of the above	Integrated professional practice, ecosystem judgment

Certification Alignment Summary

The series addresses two Microsoft Office Specialist certifications as the quality benchmark and scope anchor for the foundational volumes – not as the ceiling of instruction. Both certifications are the floor of each volume’s ambition, not its destination.

Certification	Primary Volume	Also Addressed In
MO-210 Excel Associate	Book 1	Book 5 (charting and PivotTable applied domains)

Certification	Primary Volume	Also Addressed In
MO-211 Excel Expert	Book 2	Books 3-5 (extended and applied throughout)

Volumes 3, 4, and the Capstone are not certification-oriented. They target the professional capabilities that make an analyst effective in practice – capabilities that currently map most directly to the foundational preparation tier for PL-300 and DP-600, though those credentials are outside the scope of this series.

PREREQUISITES — What Books 1 and 2 Assumed and This Book Assumes

This volume is the third in a cumulative series. It does not re-teach the foundations established in Volumes 1 and 2. Before proceeding, readers should confirm they can perform the following without significant difficulty.

Carried Forward from Volume 1 — *Excel for Analysts: The Analytical Foundation*

- Work within the layered workbook architecture: maintaining the clean separation of the raw data layer, transformation layer, analysis layer, and presentation layer within a single workbook
- Convert ranges to Excel Tables and use structured reference syntax — `Table[Column]` notation — with confidence
- Apply named ranges to improve workbook readability and navigability
- Assign and interpret data types correctly — distinguishing numbers stored as text, dates stored as serial numbers, and the downstream consequences each produces
- Use data validation to enforce clean inputs at the data-entry stage
- Navigate multi-sheet workbooks with architectural intent, not merely by scrolling

If any of the above capabilities feels uncertain, Volume 1 should be completed before proceeding with this volume. The Power Query content here is not inherently difficult, but it is written for analysts

who already think structurally about data — and that structural thinking is the primary product Volume 1 delivers.

Carried Forward from Volume 2 — *Analytical Formulas & Calculation Architecture*

- Write and audit lookup formulas at the XLOOKUP level — and understand the structural reasons for choosing one approach over another
- Apply conditional aggregation using SUMIFS, COUNTIFS, and AVERAGEIFS across multiple simultaneous criteria
- Handle formula errors defensively using IFERROR, IFNA, and the LET function for intermediate logic isolation
- Read an unfamiliar formula of moderate complexity and describe what it calculates without executing it
- Understand the difference between a formula that produces a correct answer and a formula that is maintainable under changing requirements

Full fluency in dynamic arrays (FILTER, SORT, UNIQUE, SEQUENCE) or LAMBDA is not a prerequisite for this volume, though analysts who have completed Volume 2 in full will recognise useful parallels.

What This Book Does Not Assume

This book does not assume any prior exposure to Power Query, the M language, or ETL concepts. All Power Query content is introduced from first principles. Readers arriving from a strong Volume 1 and 2 foundation are fully prepared regardless of prior Power Query experience.

This book does not assume SQL knowledge, although analysts with SQL experience will recognise the structural parallels between relational query logic and Power Query transformation logic — and those parallels are noted explicitly throughout the text wherever they clarify rather than distract.

If You Are Entering the Series at This Volume

Readers entering the series at Volume 3 should confirm they are comfortable with Excel Tables and structured references (the core output of Volume 1) and can read a basic conditional aggregation formula without confusion (the core output of Volume 2). If either capability is uncertain, the relevant prior volume should be completed before proceeding. The assumptions this book carries are not arbitrary — they represent the structural thinking that makes the Power Query content learnable rather than merely executable.

Contents

HALF-TITLE PAGE	ii
TITLE PAGE	iii
COPYRIGHT PAGE	iv
DEDICATION	v
EPIGRAPH	vi
ACKNOWLEDGMENTS	vii
ABOUT THE SERIES	ix
FOREWORD	xiii
PREFACE	xvi
PREFACE TO THE SERIES	xviii
HOW TO USE THIS BOOK	xix
SERIES ARCHITECTURE	xxii
PREREQUISITES — What Books 1 and 2 Assumed and This Book Assumes	xxvii
Contents	xxx
List of Figures	lxii

INTRODUCTION — The Data Preparation Imperative	lxvi
I.1 Why Analysts Spend the Majority of Their Time Preparing Data — and Why That Is Not Inevitable	lxvi
I.2 The Limits of Formula-Based Cleaning — Why It Does Not Scale	lxvii
I.3 Power Query as an ETL Engine Inside Excel — and What ETL Actually Means	lxix
I.4 The Concept of the Reproducible, Auditable Data Pipeline	lxxi
I.5 The Five Stages of an ETL Workflow — Connect, Profile, Shape, Combine, Author	lxxiii
I.6 Pipeline Architecture as a Professional Discipline . . .	lxxiv
I.7 What This Book Deliberately Leaves to Book 4	lxxvi
I.8 The Path Forward — From Pipeline to Relational Model	lxxvii
 I The ETL Discipline and the Power Query Environment	1
 1 The Data Preparation Reality and the Case for Power Query	3
Chapter Introduction	3
Chapter Objectives	4
1.1 Why Analysts Spend the Majority of Their Time on Data Preparation	5
1.1.1 The Hidden 80% of Analytical Work — Time Spent, Errors Made, Reproducibility Lost	5
1.1.2 The Copy-Paste-Formula Workflow and Its Failure Modes	6
1.1.3 What Manual Preparation Costs — Time, Trust, and Institutional Memory	8
1.2 What Power Query Is and What It Is Not	10

1.2.1	The ETL Acronym — Extract, Transform, Load as a Professional Discipline	10
1.2.2	Power Query vs. Worksheet Formulas vs. Macros — Choosing the Right Tool	11
1.2.3	Power Query as a Layer in the Layered Workbook Architecture from Book 1	12
1.3	The ETL Paradigm — From One-Off Tasks to Reproducible Pipelines	14
1.3.1	The Concept of a Pipeline — Authored, Auditable, and Refreshable	16
1.3.2	Reproducibility, Auditability, and Transformation Traceability as Professional Standards . .	17
1.3.3	Transitioning from Spreadsheet Thinking to Pipeline Thinking	18
1.4	The Discipline Map for This Book	20
1.4.1	Connect, Profile, Shape, Combine, Author, Deploy — Six Stages of an ETL Workflow	20
1.4.2	Prerequisites Carried Forward from Books 1 and 2	23
1.4.3	What This Book Hands Off to Book 4	24
	Analyst's Corner 1 — A Day in the Life of an Analyst Without Power Query	24
	Analytical Toolkit — The ETL Workflow Self-Assessment Checklist	26
	Chapter Summary	29
	Key Takeaways	30
	Practice Quiz	31
	Multiple Choice Questions (12)	31
	Answer Key with Rationale	36

Exercises / Applied Practice	39
Exercise 1.1 — Auditing One of Your Current Workbooks for Manual Preparation Cost (<i>Comprehension</i>)	39
Exercise 1.2 — Mapping the Preparation Steps of an Existing Report Onto the Six ETL Stages (<i>Application</i>)	41
Exercise 1.3 — Defending the Case for Power Query to a Skeptical Colleague (<i>Judgment</i>)	42
Reflection / Discussion Questions	44
What's Next	45
2 The Power Query Editor as an Analytical Environment	46
Chapter Introduction	46
Chapter Objectives	47
2.1 Launching and Navigating the Editor	48
2.1.1 Get & Transform Data — Ribbon Entry Points and the Get Data Menu	48
2.1.2 The Queries Pane, Preview Pane, and the Editor Layout	49
2.1.3 The Ribbon — Home, Transform, Add Column, and View Tabs	51
2.2 The Anatomy of a Query	53
2.2.1 Source, Applied Steps, and the Final Output — What Each Component Does	54
2.2.2 The Three Load Destinations — Worksheet, Connection-Only, and Data Model	55
2.2.3 Why the Connection-Only Query Is the Plumbing of Multi-Stage Pipelines	58

2.2.4	The Formula Bar and the Advanced Editor — Reading Generated M from the Start	59
2.3	Configuring the Editor for Professional Work	61
2.3.1	Query Options, Privacy Settings, and Locale Configuration	61
2.3.2	Query Properties, Renaming, and Grouping Queries for Navigation	63
2.3.3	Building a Consistent Editor Setup Across a Team	65
	Analyst's Corner 2 — Why the Preview Pane Lies — and What to Do Instead	66
	Analytical Toolkit — The Editor Configuration Reference .	67
	Chapter Summary	70
	Key Takeaways	71
	Practice Quiz	72
	Multiple Choice Questions (12)	72
	Answer Key with Rationale	77
	Exercises / Applied Practice	81
	Exercise 2.1 — Loading a CSV Three Ways — to Work- sheet, Connection-Only, and Data Model (<i>Com- prehension</i>)	81
	Exercise 2.2 — Configuring Locale and Privacy Set- tings for a Mixed-Source Workbook (<i>Applica- tion</i>)	82
	Exercise 2.3 — Reading the Anatomy of an Inherited Query Without Modifying It (<i>Application</i>) . . .	83
	Exercise 2.4 — Documenting All Connections in a Pro- vided Workbook Using the Source Audit Tem- plate (<i>Judgment</i>)	85

Reflection / Discussion Questions	86
What's Next	87
3 The Applied Steps Pane as a Transformation Audit Trail	89
Chapter Introduction	89
Chapter Objectives	90
3.1 The Applied Steps Pane as a Professional Document	91
3.1.1 Default Step Names vs. Professional Step Names — The Audit Gap	91
3.1.2 The Five-Second Readability Test — What a Well-Named Pipeline Looks Like	92
3.1.3 Step Naming Conventions — Rules, Patterns, and Examples	94
3.2 Managing the Step Sequence — Dependencies, Re- ordering, and Deletion	97
3.2.1 How Steps Reference Their Predecessors — The Dependency Chain	97
3.2.2 Inserting, Moving, and Deleting Steps With- out Breaking the Pipeline	99
3.2.3 The Gear Icon and Step-Level Editing — Mod- ifying a Step Without Rebuilding It	100
3.3 The Applied Steps Pane as a Communication Instru- ment	101
3.3.1 Writing the Pipeline for the Analyst Who In- herits It	102
3.3.2 The Step Comment — Inline Documentation in the M Language	103
3.3.3 The Applied Steps Pane as a Review and Ap- proval Surface	105

Analyst's Corner 3 — The Audit That Was and Wasn't . . .	106
Analytical Toolkit — The Step Naming Style Guide and Pipeline Readability Checklist	107
Chapter Summary	112
Key Takeaways	113
Practice Quiz	114
Multiple Choice Questions (12)	114
Answer Key with Rationale	119
Exercises / Applied Practice	123
Exercise 3.1 — Renaming All Steps in a Provided Query to Professional Standard (<i>Comprehension</i>) . .	123
Exercise 3.2 — Inserting a Step in the Middle of a Pipeline Without Breaking Downstream Steps (<i>Application</i>)	125
Exercise 3.3 — Writing an Applied Steps Pane That Passes the Five-Second Readability Test (<i>Judgment</i>)	126
Reflection / Discussion Questions	128
What's Next	129
II Connecting to Data: The Source Layer	133
4 File and Folder Connections — Excel, CSV, JSON, and Folder Sources	136
Chapter Introduction	136
Chapter Objectives	137
4.1 The Source Layer — What Belongs Here and Why . .	138
4.1.1 Staging Queries as the Connection Architecture	138

4.1.2	What the Source Step Contains and What It Should Not Contain	140
4.1.3	Path Portability — The Case Against Hard-Coded File Paths	141
4.2	Connecting to Excel Workbooks and CSV Files	142
4.2.1	The Excel Connector — Sheets, Tables, Named Ranges, and the Navigation Step	142
4.2.2	The CSV Connector — Delimiter Detection, Encoding, and Locale at Import	143
4.2.3	Promoting Headers and the First Applied Step Pattern	146
4.3	Connecting to JSON and XML Sources	147
4.3.1	JSON Structure in Power Query — Records, Lists, and Tables	147
4.3.2	Expanding Records and Lists — The Expand Column Operation	147
4.3.3	XML Sources and the Document Structure Navigation Pattern	149
4.4	The Folder Connector — Combining Many Files Into One Pipeline	150
4.4.1	How the Folder Connector Works — The Sample File and Transformation Function Pattern	150
4.4.2	Building a Robust Folder Pipeline — File Filtering, Column Selection, and Source Column Management	151
4.4.3	When Folder Sources Break and How to Prevent It	153
4.5	Path Parameterization — Building Portable, Shareable Connections	154

4.5.1	Parameter Queries — Definition, Creation, and Reference in Source Steps	154
4.5.2	The Parameter Query Pattern for File Paths and Folder Paths	155
4.5.3	Updating Parameters Without Opening the Editor	157
	Analyst's Corner 4 — The Folder That Grew: When a Two-File Source Becomes Twenty	157
	Analytical Toolkit — The Connection Architecture Checklist and Source Audit Template	159
	Chapter Summary	161
	Key Takeaways	162
	Practice Quiz	164
	Multiple Choice Questions (12)	164
	Answer Key with Rationale	169
	Exercises / Applied Practice	173
	Exercise 4.1 — Building a Parameterized CSV Staging Query (<i>Comprehension</i>)	173
	Exercise 4.2 — Building a Folder Pipeline for Monthly Sales Files (<i>Application</i>)	174
	Exercise 4.3 — Migrating a Hard-Coded Workbook to Parameterized Connections (<i>Judgment</i>)	176
	Reflection / Discussion Questions	177
	What's Next	179
5	Database, Cloud, and Web Connections	180
	Chapter Introduction	180
	Chapter Objectives	181
5.1	Connecting to Relational Databases	182

5.1.1	The SQL Server and PostgreSQL Connectors — Server, Database, and Credential Configuration	182
5.1.2	Query Folding — What It Is, Why It Matters, and How to Verify It	183
5.1.3	Native Database Queries — When SQL Belongs in Power Query	185
5.1.4	IT Governance and Database Credentials — Working Within Organizational Constraints .	186
5.2	Connecting to SharePoint and OneDrive	187
5.2.1	SharePoint Lists — Connection, Navigation, and the Column Expansion Pattern	188
5.2.2	SharePoint Folder and OneDrive — Cloud-Based Folder Connector Patterns	190
5.2.3	Authentication and Tenant Permissions for Mi- crosoft 365 Sources	191
5.3	Connecting to Web Pages and REST APIs	192
5.3.1	The Web Connector — HTML Tables, Static Pages, and Reliability Limits	192
5.3.2	The REST API Connector — HTTP Requests, Authentication Headers, and JSON Responses	193
5.3.3	Handling Paginated API Responses — The List.Generate Pattern	196
5.4	Privacy Levels in Mixed-Source Pipelines	198
5.4.1	The Formula Firewall — Why It Exists and What Triggers It	198
5.4.2	Assigning Privacy Levels Correctly Across Source Types	199
5.4.3	Mixed-Source Pipeline Configuration — The Professional Standard	199

Analyst's Corner 5 — The API That Changed Its Mind . . .	202
Analytical Toolkit — The Mixed-Source Privacy Level Assignment Matrix	203
Chapter Summary	205
Key Takeaways	206
Practice Quiz	208
Multiple Choice Questions (12)	208
Answer Key with Rationale	213
Exercises / Applied Practice	217
Exercise 5.1 — Building a SQL Server Staging Query with Query Folding Verification (<i>Application</i>)	217
Exercise 5.2 — Connecting to a SharePoint List and Building the Expansion Pipeline (<i>Application</i>)	219
Exercise 5.3 — Building a Paginated REST API Pipeline (<i>Judgment</i>)	220
Reflection / Discussion Questions	222
What's Next	223
III Profiling and Shaping Data	226
6 Data Profiling — Diagnosing the Source Before Transforming It	229
Chapter Introduction	229
Chapter Objectives	230
6.1 The Profiling Discipline — Why Diagnosis Precedes Transformation	231
6.1.1 Profiling as an Engineering Practice, Not a Convenience Feature	231
6.1.2 The Cost of Transforming Without Profiling .	232

6.1.3	The Profiling Workflow — What to Look for Before Writing a Step	233
6.2	Column Quality — Identifying Validity, Errors, and Empty Values	234
6.2.1	What Column Quality Measures — Valid, Error, and Empty Percentages	234
6.2.2	Interpreting Column Quality Results — What Each Pattern Means	235
6.2.3	Extending the Column Quality Sample to the Full Dataset	237
6.3	Column Distribution — Understanding Value Spread and Cardinality	238
6.3.1	The Distribution Histogram — Reading Frequency and Range	238
6.3.2	Cardinality and Distinct Value Count — What These Tell the Analyst	239
6.3.3	Spotting Anomalies in the Distribution — Outliers, Spikes, and Unexpected Gaps	240
6.4	Column Profile — Deep Inspection of a Single Column	241
6.4.1	The Value Distribution Table — Every Distinct Value and Its Count	241
6.4.2	Statistics Panel — Min, Max, Average, Count, and Standard Deviation	242
6.4.3	Using Column Profile to Document Transformation Decisions	244
6.5	The Profiling Workflow — From Findings to Transformation Plan	245
6.5.1	The Profiling Log — Documenting What Was Found Before Anything Is Changed	245

6.5.2 Translating Profiling Findings Into Transformation Steps	249
6.5.3 Re-Profiling After Transformation — Verifying That the Pipeline Cleaned What Was Found	250
Analyst's Corner 6 — The Column That Looked Clean . . .	251
Analytical Toolkit — The Pre-Transformation Profiling Checklist and Profiling Log Template	252
Chapter Summary	254
Key Takeaways	255
Practice Quiz	257
Multiple Choice Questions (12)	257
Answer Key with Rationale	262
Exercises / Applied Practice	266
Exercise 6.1 — Full Profiling of a Provided Dataset Before Any Transformation (<i>Comprehension</i>) . .	266
Exercise 6.2 — Building a Transformation Plan from Profiling Findings and Implementing It (<i>Application</i>)	267
Exercise 6.3 — Profiling an Inherited Pipeline — Verifying That the Cleaning Was Complete (<i>Judgment</i>)	268
Reflection / Discussion Questions	270
What's Next	271
7 Shaping Data — Row and Column Operations	273
Chapter Introduction	273
Chapter Objectives	274
7.1 Row Operations — Filtering, Deduplication, and Row Management	275

7.1.1	Filter Rows — Condition Types, Multi-Condition Logic, and the Profiling Connection	275
7.1.2	Keep Rows and Remove Rows — Top N, Bottom N, Range, and Alternating Patterns	277
7.1.3	Remove Duplicates — On One Column, Multiple Columns, and the Grain Question	278
7.1.4	Sort Rows — When It Belongs in the Pipeline and When It Does Not	279
7.2	Type Enforcement — The Professional Standard . . .	280
7.2.1	Why Type Enforcement Is Not Optional . . .	280
7.2.2	Enforcing Types Correctly — Column by Column vs. Batch Assignment	281
7.2.3	Using Locale for Date and Number Types — The Complete Pattern	282
7.2.4	Handling Type Errors Gracefully — Try and Otherwise	284
7.3	Value Normalization — Replace, Clean, Trim, and Format	285
7.3.1	Replace Values — Exact Match, Partial Match, and Case Sensitivity	285
7.3.2	Trim and Clean — Removing Whitespace and Non-Printable Characters	286
7.3.3	Text Case Normalization — Upper, Lower, Proper	287
7.3.4	Number and Date Formatting — What Belongs in Power Query and What Belongs in the Presentation Layer	288
7.4	Null and Error Handling	289
7.4.1	Null Values — Replace, Filter, Fill Down, Fill Up	289
7.4.2	Error Values — Replace Errors, Remove Error Rows, and the Try-Otherwise Pattern	290

7.4.3	Null vs. Empty String vs. Zero — Distinguishing and Handling Each	292
7.5	Column Management	294
7.5.1	Remove Columns vs. Keep Columns — Which to Use and When	294
7.5.2	Rename and Reorder Columns — Building the Analytical Schema	295
7.5.3	Split Column — By Delimiter, Character Count, and Position	295
7.5.4	Merge Columns — Concatenation and the Separator Pattern	296
	Analyst's Corner 7 — The Trim That Wasn't Enough	297
	Analytical Toolkit — The Row and Column Operations Quick Reference	299
	Chapter Summary	303
	Key Takeaways	304
	Practice Quiz	305
	Multiple Choice Questions (12)	305
	Answer Key with Rationale	311
	Exercises / Applied Practice	315
	Exercise 7.1 — Full Row and Column Cleaning of a Provided Dataset (<i>Comprehension</i>)	315
	Exercise 7.2 — Building a Type-Safe Pipeline with Locale and Error Handling (<i>Application</i>)	316
	Exercise 7.3 — Cleaning an Inherited Query Against Its Profiling Log (<i>Judgment</i>)	317
	Reflection / Discussion Questions	319
	What's Next	320

8	Structural Transformations and Derived Columns	322
	Chapter Introduction	322
	Chapter Objectives	323
8.1	Unpivot — Converting Wide Tables to Long Format .	324
8.1.1	Why Wide Tables Are Analytically Problematic	324
8.1.2	Unpivot Columns — Selected Columns, Other Columns, and All Columns	325
8.1.3	Naming the Attribute and Value Columns . . .	327
8.1.4	Handling Nulls After Unpivot	328
8.2	Pivot — Converting Long Tables to Wide Format . . .	328
8.2.1	When Pivot Is and Is Not Appropriate	328
8.2.2	The Pivot Column Operation — Values Col- umn and Aggregation Function	329
8.2.3	Pivot Output Schema and the Downstream Schema Dependency Risk	330
8.3	Structural Transformations — Transpose, Group By, and Fill	331
8.3.1	Transpose — When Row-Column Inversion Is Required	331
8.3.2	Group By — Aggregation at the Query Layer .	332
8.3.3	Fill Down and Fill Up Revisited — Structural Context	334
8.4	Derived Columns — Custom Column and Column From Examples	334
8.4.1	Custom Column — Writing M Expressions for Column Derivation	334
8.4.2	Column From Examples — Inferring Transfor- mation Logic from Sample Values	336

8.4.3	Date Part Extraction — Year, Month, Quarter, Day of Week	336
8.4.4	Conditional Columns — If-Then-Else Logic at the Row Level	337
8.4.5	Text Derivations — Concatenation, Extraction, and Pattern Operations	339
8.5	The Shape Stage in the Six-Stage Workflow — What Leaves Part III	340
8.5.1	The Clean, Typed, Correctly Structured Table as the Shape Stage Output	340
8.5.2	The Schema Contract Between Staging/Shape and Combine/Model Layers	341
8.5.3	Pre-Combination Checklist	342
	Analyst's Corner 8 — Unpivoting the Report That Was Never a Dataset	343
	Analytical Toolkit — The Structural Transformation Deci- sion Guide	345
	Chapter Summary	348
	Key Takeaways	349
	Practice Quiz	350
	Multiple Choice Questions (12)	350
	Answer Key with Rationale	355
	Exercises / Applied Practice	358
	Exercise 8.1 — Unpivoting a Wide Monthly Sales Re- port (<i>Comprehension</i>)	358
	Exercise 8.2 — Building a Derived Column Set from Date and Categorical Inputs (<i>Application</i>)	360
	Exercise 8.3 — Designing the Full Shape Stage for a Multi-Problem Source (<i>Judgment</i>)	361

Reflection / Discussion Questions	363
What's Next	364
IV Combining Data: The Combine Layer	367
9 Append — Combining Same-Structure Sources	370
Chapter Introduction	370
Chapter Objectives	370
9.1 The Append Pattern — What It Does and When It Applies	371
9.1.1 Append vs. Union — The Same Operation	371
9.1.2 When Append Is the Correct Combination Operation	372
9.1.3 Schema Alignment — The Prerequisite for Reliable Append	372
9.2 Append Queries — Two Tables and Multiple Tables	373
9.2.1 Append Two Queries	373
9.2.2 Append Multiple Queries — The Append As New Operation	375
9.2.3 The Staging Query Pattern for Multi-Source Append	376
9.3 Schema Alignment Discipline	377
9.3.1 Column Name Matching — Case Sensitivity and Exact Match Requirements	377
9.3.2 Data Type Alignment — What Happens When Types Differ	377
9.3.3 Column Order — Why It Doesn't Determine Matching but Does Affect Output	378

9.4 Append in Production — Building a Resilient Multi-Source Pipeline	379
9.4.1 The Reference Query Pattern — Referencing vs. Duplicating	379
9.4.2 Source Provenance — Adding and Managing the Source Identifier Column	380
9.4.3 Append with Folder Connector vs. Manual Append — Choosing the Right Pattern	380
Analyst's Corner 9 — The Append That Lost 10,000 Rows	381
Analytical Toolkit — The Append Schema Alignment Checklist	382
Chapter Summary	384
Key Takeaways	385
Practice Quiz	386
Multiple Choice Questions (12)	386
Answer Key with Rationale	390
Exercises / Applied Practice	393
Exercise 9.1 — Appending Three Regional Sales Queries (Comprehension)	393
Exercise 9.2 — Building a Resilient Multi-Source Append Pipeline (Application)	395
Exercise 9.3 — Diagnosing and Repairing a Broken Append Pipeline (Judgment)	396
Reflection / Discussion Questions	397
What's Next	399
10 Merge — Joining Related Tables	400
Chapter Introduction	400
Chapter Objectives	400

10.1	The Merge Pattern — What It Does and When It Applies	401
10.1.1	Merge vs. VLOOKUP — What Power Query Merge Adds	401
10.1.2	The Six Join Types — Left Outer, Right Outer, Full Outer, Inner, Left Anti, Right Anti	402
10.1.3	Single-Key and Multi-Key Joins	405
10.2	Performing a Merge — From Dialog to Applied Step .	405
10.2.1	The Merge Queries Dialog — Table Selection, Key Columns, and Join Type	405
10.2.2	The Expand Step — Selecting Columns from the Merged Table	406
10.2.3	Naming the Merge and Expand Steps	407
10.3	Join Type Selection — Analytical Logic for Each Join	408
10.3.1	Left Outer Join — The Most Common Production Join	408
10.3.2	Inner Join — Strict Match, Silent Exclusion Risk	408
10.3.3	Left Anti Join — Finding Non-Matching Records	409
10.3.4	Full Outer, Right Outer, Right Anti — When Each Applies	409
10.4	Join Key Discipline — The Properties That Determine Join Correctness	410
10.4.1	Type Alignment on Join Keys — The Most Common Merge Failure	410
10.4.2	Null Values in Join Keys — Silent Record Loss	411
10.4.3	Case Sensitivity in Text Join Keys	411
10.4.4	Leading/Trailing Whitespace in Join Keys . .	412
10.5	Merge in Production — Multi-Key Joins, Non-Matches, and Performance	412

10.5.1 Multi-Key Joins — When a Single Column Is Not a Sufficient Key	412
10.5.2 Detecting and Handling Non-Matches After a Left Outer Join	413
10.5.3 Merge Performance and Query Folding Implications	415
Analyst's Corner 10 — The Join That Matched Everything Wrong	415
Analytical Toolkit — The Join Key Discipline Checklist . .	416
Chapter Summary	418
Key Takeaways	419
Practice Quiz	420
Multiple Choice Questions (12)	420
Answer Key with Rationale	425
Exercises / Applied Practice	429
Exercise 10.1 — Performing a Left Outer Join with Expand and Non-Match Detection (<i>Comprehension</i>)	429
Exercise 10.2 — Multi-Key Join with Non-Match Detection and Row Multiplication Check (<i>Application</i>)	430
Exercise 10.3 — Diagnosing a Merge That Produces Unexpected Results (<i>Judgment</i>)	431
Reflection / Discussion Questions	433
What's Next	435
V The M Language: Reading, Writing, and Authoring	437
11 The M Language: Structure, Syntax, and the let-in Expression	440

Chapter Introduction	440
Chapter Objectives	440
11.1 M as a Functional Language — The Mental Model . .	441
11.1.1 What Makes M Different from Procedural Languages and Formulas	441
11.1.2 The Evaluation Model — Lazy Evaluation and Expression Trees	442
11.1.3 M Inside Power Query — Where It Lives and How It Runs	443
11.2 The let-in Expression — Building a Query as a Sequence of Named Values	443
11.2.1 Let, In, and the Final Expression	444
11.2.2 Named Steps as Named Variables — The Connection to the Applied Steps Pane	446
11.2.3 Expression Chaining — Each Step References Its Predecessor	447
11.3 M Syntax Fundamentals	448
11.3.1 Literals — Text, Number, Date, Logical, Null .	448
11.3.2 Identifiers and Quoting — When to Use #”...” .	448
11.3.3 Comments — Single-Line and Block	449
11.3.4 Operators — Arithmetic, Comparison, Logical, Text Concatenation	450
11.4 Reading Generated M — The Formula Bar and Advanced Editor	451
11.4.1 Reading a Five-Step Query from Start to Finish	451
11.4.2 Identifying the Function, Its Arguments, and Its Return Value	454
11.4.3 Tracing a Value Through the Pipeline	455

Analyst's Corner 11 — From Reader to Writer: The Natural Progression	455
Analytical Toolkit — The M Reading Checklist	456
Chapter Summary	457
Key Takeaways	458
Practice Quiz	458
Multiple Choice Questions (12)	459
Answer Key with Rationale	462
Exercises / Applied Practice	464
Exercise 11.1 — Reading a Generated Query in Full (Comprehension)	464
Exercise 11.2 — Annotating an Inherited Query's M Code (Application)	464
Exercise 11.3 — Modifying a Query by Editing M Di- rectly (Judgment)	465
Reflection / Discussion Questions	465
What's Next	467
12 Working with M Types: Records, Lists, Tables, and Functions	468
Chapter Introduction	468
Chapter Objectives	468
12.1 The Four Container Types	469
12.1.1 Records — Named Field Collections	469
12.1.2 Lists — Ordered Value Collections	470
12.1.3 Tables — Structured Row-Column Collections	470
12.1.4 Functions — Parameterized Expressions . . .	470
12.2 Operating on Records	472

12.2.1	Record Field Access — [FieldName] and Record[FieldName]	472
12.2.2	Record Construction and Transformation	472
12.2.3	Records in Power Query Context — Row-Level Expressions	473
12.3	Operating on Lists	473
12.3.1	List Construction — {value1, value2, ...}	473
12.3.2	List Functions — List.Select, List.Transform, List.Contains	473
12.3.3	List.Accumulate and Generating Dynamic Value Sets	474
12.4	Operating on Tables	475
12.4.1	Table Functions for Row Operations	475
12.4.2	Table Functions for Column Operations	475
12.4.3	Table.NestedJoin and Table.Combine — The M Underpinnings of Merge and Append	476
Analyst's Corner 12 — Understanding the Container You're Working With		477
Analytical Toolkit — The M Container Type Quick Reference		478
Chapter Summary		479
Key Takeaways		480
Practice Quiz		480
Multiple Choice Questions (10)		480
Answer Key with Rationale		483
Exercises / Applied Practice		484
Exercise 12.1 — Working with Records and Lists in Custom Column (<i>Comprehension</i>)		484
Exercise 12.2 — Table Operations in M (<i>Application</i>)		485
Reflection / Discussion Questions		485

What's Next	486
13 Writing M for Column Derivation and Custom Transformation	487
Chapter Introduction	487
Chapter Objectives	487
13.1 Custom Column Expressions — Beyond the Ribbon .	488
13.1.1 When the Ribbon Cannot Produce the Required Result	488
13.1.2 Text Functions — Text.Replace, Text.Split, Text.Format	489
13.1.3 Number Functions — Number.Round, Number.Abs, Number.Mod	490
13.1.4 Date Functions — Date.AddDays, Date.IsInCurrentMonth, Date.ToText	490
13.2 Conditional Logic in M	491
13.2.1 The if-then-else Expression — Single and Multi-Branch	491
13.2.2 Nested Conditionals and the Readability Problem	491
13.2.3 List.Contains for Multi-Value Category Testing	492
13.3 The each Shorthand — Writing Row-Level Functions	493
13.3.1 What each Does — The Implicit Parameter . .	493
13.3.2 each in Table.SelectRows, Table.AddColumn, and List.Select	493
13.3.3 When each Is Insufficient and Explicit Parameters Are Required	494
13.4 Error Handling in Custom Expressions	495
13.4.1 Try-Otherwise Revisited — Full Syntax and Nesting	495

13.4.2 Error.Record — Capturing and Logging Error Detail	495
13.4.3 Defensive Expression Design — Writing M That Fails Predictably	496
Analyst's Corner 13 — The Expression That Broke on Row 50,001	497
Analytical Toolkit — The Custom M Expression Style Guide	498
Chapter Summary	499
Key Takeaways	500
Practice Quiz	501
Multiple Choice Questions (10)	501
Answer Key with Rationale	504
Exercises / Applied Practice	505
Exercise 13.1 — Writing Custom Text and Date Ex- pressions (<i>Comprehension</i>)	506
Exercise 13.2 — Building a Multi-Branch Classifica- tion Expression (<i>Application</i>)	506
Exercise 13.3 — Writing a Defensive Column Deriva- tion for a Problematic Source (<i>Judgment</i>)	507
Reflection / Discussion Questions	508
What's Next	509
14 Parameterization, Functions, and Query Architecture	510
Chapter Introduction	510
Chapter Objectives	510
14.1 Parameter Queries Revisited — M-Level Parameter- ization	511
14.1.1 Parameter Values in M Expressions — Refer- encing Parameters Beyond File Paths	512

14.1.2	Dynamic Date Parameters — Building Refresh-Relative Date Filters	513
14.1.3	Environment Parameters — Switching Between Dev, Test, and Production Sources	514
14.2	Custom Function Queries	514
14.2.1	What a Function Query Is and When to Build One	514
14.2.2	Writing a Simple Function Query — (Parameter) => Expression	515
14.2.3	Invoking a Function Query — Across a Table Column	516
14.3	Advanced Patterns	518
14.3.1	List.Generate Revisited — Full Syntax Explanation	518
14.3.2	Table.ExpandTableColumn and the All Rows Group Pattern	519
14.3.3	Dynamic Column Selection with Table.SelectColumns	521
14.4	Query Architecture in M	521
14.4.1	Layered Query Design — What Belongs in Each Layer	522
14.4.2	Naming, Documentation, and the Architecture Standard	522
14.4.3	Query Dependencies and the Refresh Order .	523
	Analyst's Corner 14 — The Function That Eliminated Forty Staging Queries	523
	Analytical Toolkit — The M Architecture Standard Reference	525
	Chapter Summary	527
	Key Takeaways	527

Practice Quiz	528
Multiple Choice Questions (10)	528
Answer Key with Rationale	531
Exercises / Applied Practice	533
Exercise 14.1 — Building a Dynamic Date Filter Pa- rameter (<i>Comprehension</i>)	533
Exercise 14.2 — Building a Custom Function Query for File Transformation (<i>Application</i>)	534
Exercise 14.3 — Refactoring a Multi-Query Pipeline Using Function Queries (<i>Judgment</i>)	534
Reflection / Discussion Questions	536
What's Next	537
 VI Performance and Production Readiness	 540
 15 Query Performance and Optimization	 543
Chapter Introduction	543
Chapter Objectives	543
15.1 Understanding Query Evaluation — Folding, Caching, and the Evaluation Engine	544
15.1.1 Query Folding Revisited — The Full Picture . .	544
15.1.2 The Evaluation Cache — How Power Query Avoids Re-Evaluation	545
15.1.3 Background Refresh vs. Manual Refresh — Op- erational Differences	546
15.2 Diagnosing Slow Queries — The Diagnostic Tools . .	547
15.2.1 Query Diagnostics — Enabling and Reading the Trace	547

15.2.2	Reading Query Plan Output — What the Fold Indicator Tells You	548
15.2.3	Step-by-Step Timing — Identifying the Slow Step	549
15.3	Performance Optimization Patterns	550
15.3.1	Push Filters Early — The Fold-Before-Transform Discipline	550
15.3.2	Remove Unnecessary Columns Early	550
15.3.3	Avoid Privacy Level Workarounds — Their Performance Cost	551
15.3.4	Connection-Only Queries and Evaluation Sharing	551
15.4	Data Quality Monitoring at the Pipeline Layer	552
15.4.1	Building Row Count Checks Into the Pipeline	552
15.4.2	Building Null Rate Checks Into the Pipeline	553
15.4.3	The Data Quality Check Query — A Connection-Only Monitoring Layer	556
	Analyst's Corner 15 — The Refresh That Took Four Hours	556
	Analytical Toolkit — The Performance Optimization Checklist	557
	Chapter Summary	559
	Key Takeaways	559
	Practice Quiz	560
	Multiple Choice Questions (10)	560
	Answer Key with Rationale	565
	Exercises / Applied Practice	567
	Exercise 15.1 — Diagnosing Query Folding in a Database Pipeline (<i>Comprehension</i>)	567

Exercise 15.2 — Optimizing a Slow Multi-Source Pipeline (<i>Application</i>)	567
Exercise 15.3 — Building a Data Quality Monitoring Layer (<i>Judgment</i>)	568
Reflection / Discussion Questions	569
What's Next	570
16 Production Readiness, Handoff, and the Refresh Run- book	571
Chapter Introduction	571
Chapter Objectives	571
16.1 The Production Readiness Standard — What “Done” Means for a Pipeline	572
16.1.1 The Full Production Readiness Audit — Assem- bling the Checklists	572
16.1.2 Pipeline Documentation Standard — The Com- plete Deliverable Set	573
16.1.3 The Schema Contract as a Production Artifact	575
16.2 The Refresh Runbook — Documenting the Opera- tional Procedure	576
16.2.1 What a Refresh Runbook Contains	576
16.2.2 Documenting the Normal Refresh Procedure	577
16.2.3 Documenting the Failure Diagnostic Procedure	577
16.3 Error Monitoring and Alerting	578
16.3.1 Types of Refresh Failure — Connection, Type, Business Rule, and Silent Failures	578
16.3.2 The Manual Monitoring Standard — What to Check After Every Refresh	579

16.3.3 When to Escalate — The Data Quality Escalation Decision Tree	580
16.4 Handoff Protocols	581
16.4.1 The Handoff Package — What the Recipient Receives	581
16.4.2 The Handoff Session — What Is Demonstrated vs. Documented	581
16.4.3 Post-Handoff Support — The First Thirty Days	583
Analyst's Corner 16 — The Workbook That Could Not Be Handed Off	584
Analytical Toolkit — The Production Readiness Audit (Full Consolidated Checklist)	585
Chapter Summary	590
Key Takeaways	591
Practice Quiz	592
Multiple Choice Questions (10)	592
Answer Key with Rationale	596
Exercises / Applied Practice	598
Exercise 16.1 — Completing a Production Readiness Audit on a Provided Pipeline (<i>Comprehension</i>)	598
Exercise 16.2 — Writing the Refresh Runbook for a Production Pipeline (<i>Application</i>)	598
Exercise 16.3 — Designing the Complete Handoff Package for a Multi-Source Pipeline (<i>Judgment</i>) . .	599
Reflection / Discussion Questions	600
What's Next	601
CONCLUSION — The Data Preparation Layer, Completed	604

C.1 What the Reader Can Now Build — and What That Changes About Their Analytical Work	604
C.2 The 80% of Preparation Time That Is Now Automated, Reproducible, and Auditable	605
C.3 The Habits That Carry Forward Into Books 4 and 5 — What the Data Model Expects to Receive	607
C.4 Power Query in the Broader Microsoft BI Ecosystem — From Excel to Power BI and Beyond	609
Appendices and Back Matter	614
APPENDIX A — The Step Naming Style Guide (Complete Reference)	615
APPENDIX B — The M Function Quick Reference	619
APPENDIX C — The Power Query Architecture Standard	625
APPENDIX D — The Production Readiness Audit Checklist	629
GLOSSARY	634
ABOUT THE AUTHOR	638
ABOUT THE SERIES	640

List of Figures

SA.1 Series Progression Map	xxiv
I.1 The Preparation Time Paradox: Manual vs. Pipeline-Based Workflows	lxviii
I.2 Power Query as an ETL Engine Inside Excel: The Three-Stage Architecture	lxxi
I.3 The Five Stages of a Professional ETL Workflow . . .	lxxv
I.4 The Pipeline-to-Model Handoff: Where This Book Ends and Book 4 Begins	lxxviii
1.1 The Three Costs of Manual Data Preparation	7
1.2 Tool Selection Matrix: Power Query vs. Worksheet Formulas vs. VBA Macros	13
1.3 Power Query in the Layered Workbook Architecture	15
1.4 The Six-Stage ETL Discipline Map and Its Correspondence to This Book's Parts	22
2.1 The Power Query Editor Layout: Annotated Interface Map	50
2.2 The Four Power Query Ribbon Tabs and Their Primary Operation Domains	52
2.3 The Three Power Query Load Destinations and Their Architectural Roles	56
2.4 The M Language Structure in the Formula Bar and Advanced Editor	61

3.1	Default Step Names vs. Professional Step Names: The Audit Gap Side by Side	93
3.2	A Well-Named Pipeline Passing the Five-Second Readability Test	95
3.3	The Applied Steps Dependency Chain: How Each Step References Its Predecessor	98
3.4	M Step Comments: Inline Pipeline Documentation in the Formula Bar and Advanced Editor	105
4.1	The Source Layer Architecture: Staging Queries as the Connection Boundary	139
4.2	The CSV Connection Configuration Dialog: Delimiter, Encoding, and Locale Decisions	145
4.3	JSON Structure in Power Query: Records, Lists, and Tables	148
4.4	The Folder Connector Architecture: Sample File, Transformation Function, and Combined Output	151
4.5	The Parameter Query Pattern: From Hard-Coded Path to Portable Connection	156
5.1	Query Folding Indicators in the Applied Steps Pane .	184
5.2	SharePoint List Column Expansion Pattern: From Metadata-Heavy Default to Analytical Schema	189
5.3	REST API Connection Architecture: Request, Authentication, and JSON Response Flow	195
5.4	Privacy Level Assignment Matrix for Common Source Types	200
6.1	Column Quality Indicator: Reading the Valid/Error/Empty Bar	236

6.2	Column Distribution Histograms: Reading Four Common Patterns	239
6.3	Column Profile: Value Distribution Table and Statistics Panel	243
6.4	From Profiling Finding to Named Transformation Step	248
7.1	Filter Rows Multi-Condition Logic: Dialog, Generated M, and Named Step	277
7.2	Null Handling Decision Tree: Choosing the Correct Null Response	291
7.3	Split Column by Delimiter: Source, Dialog, and Output	297
8.1	Wide vs. Long Table Format: The Unpivot Transformation	326
8.2	Pivot Column: Input Configuration and Output Schema	330
8.3	Group By: Configuration, Input, and Output	333
8.4	Date Part Extraction: From Date Column to Multiple Analytical Dimensions	338
9.1	Append Schema Alignment: Matched, Partially Misaligned, and Fully Misaligned Sources	374
10.1	The Six Join Types: Visual Representation and Row Inclusion Rules	404
10.2	The Merge Queries Dialog: Annotated Wireframe . .	406
10.3	Non-Match Detection After a Left Outer Join	414
11.1	The let-in Expression Structure: Annotated Diagram	445
11.2	Reading a Complete M Query: Pipeline Chain and Function Anatomy	453

12.1 The Four M Container Types: Structure and Power Query Context	471
13.1 The each Shorthand: Expansion and Row-Level Evaluation	494
14.1 Custom Function Query Syntax: Function Definition and Invocation	517
14.2 List.Generate: Four Arguments and Iteration Logic .	519
15.1 Query Diagnostics Output: Reading Fold Status and Step Timing	548
15.2 Data Quality Monitoring Layer Architecture	554
16.1 The Production Readiness Audit: Seven Categories and Their Source Chapters	574
16.2 The Handoff Package Components and Their Recipients	582
C.1 Power Query Across the Microsoft BI Ecosystem . .	610

INTRODUCTION — The Data Preparation Imperative

Every analyst who has worked on a real dataset has lived a version of the same experience. A file arrives — exported from an ERP, extracted from a CRM, downloaded from a government database, emailed by a colleague who extracted it from something else. It is messy in ways that are both predictable and tedious: dates formatted as text, currency values carrying symbols that prevent arithmetic, category fields with three spellings of the same value, columns that represent data and columns that represent layout, rows that are subtotals mixed with rows that are transactions. The analyst opens it, sighs quietly, and gets to work.

An hour later — sometimes two, sometimes more — there is a cleaned version. Formulas have been applied. Values have been trimmed. Types have been coerced. The analyst can now begin the work they were actually asked to do. The preparation is invisible in the output. No one asks about it. No one knows how it was done. No one could replicate it if they needed to.

And next month, when the same file arrives again with slightly different problems, the whole process begins from the beginning.

This book is a systematic argument against that pattern — not on efficiency grounds alone, though efficiency matters, but on professional grounds. A data preparation process that is undocumented, unrepeatable, and inseparable from its author is not a professional analytical asset. It is a liability: fragile, unauditable, and impossible to hand off. Power Query exists to replace it with something that is none of those things. This Introduction makes that case in full before the technical instruction begins.

I.1 Why Analysts Spend the Majority of Their Time Preparing Data — and Why That Is Not Inevitable

The most frequently cited finding in applied data analytics is that analysts spend somewhere between sixty and eighty percent of their time preparing data rather than analyzing it. The specific figure varies with the study, the industry, and the definition of “preparation,” but the direction of the finding does not. In virtually every professional context where the question has been measured, preparation consumes the majority of analytical time — and the actual analytical work that justifies the analyst’s existence consumes the minority.

This finding is commonly presented as a lamentable fact of analytical life, something to be acknowledged and endured. It should not be. The ratio is not a natural property of the work. It is a property of how the work is organized — specifically, of whether data preparation is treated as an engineering discipline with reproducible tooling and documented workflows, or as an ad hoc manual activity performed anew each time data arrives.

Organizations that have adopted systematic data preparation tooling — whether through ETL platforms, data engineering pipelines, or the kind of Power Query infrastructure this book teaches — consistently report material reductions in the proportion of analytical time spent on preparation. The preparation does not disappear. The source data does not become cleaner. What changes is that the preparation is done once, documented in the tool, and replayed automatically thereafter. The analyst who built the pipeline spends one hour on it rather than one hour every reporting cycle. The analyst who inherits it does not have to rebuild it from scratch because they cannot decipher what their predecessor did.

This is not an aspirational scenario. It is what Power Query makes possible today, inside Excel, with no additional infrastructure and no additional tooling license. The sixty-to-eighty percent is not inevitable. It is the cost of not yet having built the pipeline.

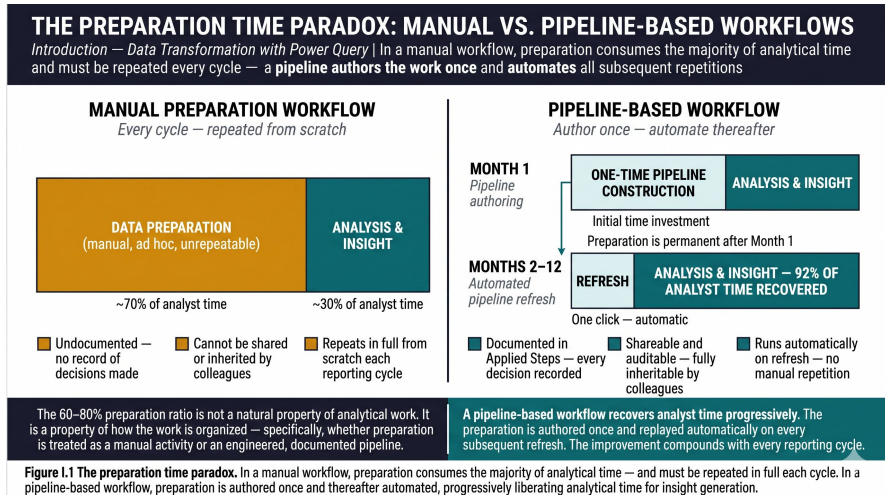


Figure I.1. The preparation time paradox. In a manual workflow, preparation consumes the majority of analytical time — and must be repeated in full each cycle. In a pipeline-based workflow, preparation is authored once and thereafter automated, progressively liberating analytical time for insight generation.

I.2 The Limits of Formula-Based Cleaning — Why It Does Not Scale

Before Power Query, Excel analysts who wanted systematic data cleaning had two options: do it manually, cell by cell, or build formula chains to automate the most repetitive transformations. The formula-based approach is better than fully manual work, but it introduces a set of limitations that become increasingly costly as the scale, complexity, or sharing requirements of the analytical work increase.

The first limitation is fragility. A formula-based cleaning workflow is built on top of the data it cleans. When the source data changes structure — a column added, a column renamed, a date format from a new regional locale — the formulas built on that structure break in ways that may produce wrong answers rather than visible errors. The analyst must inspect the output and notice the problem, which requires knowing what the output should look like. In

production environments where the pipeline runs automatically and the output feeds downstream reports, silent formula failures are among the most serious quality risks in the analytical workflow.

The second limitation is opacity. A sequence of TRIM, SUBSTITUTE, TEXT, IFERROR, and VALUE formulas spread across several helper columns conveys no structural intent to anyone reading it later. It is a series of operations without architecture — impossible to audit at a glance and difficult to modify without understanding every dependency. The formula author understands it. No one else does, and the author may not either in six months.

The third limitation is irreproducibility. A formula-based cleaning workflow does not document the reasoning behind its transformations. It encodes the result of that reasoning in the formulas themselves, but the reasoning — why certain rows were filtered, why that date column was treated as a text value before conversion, why that substitution was applied — lives only in the analyst's memory. When the analyst leaves or the workbook is inherited, the reasoning leaves with them. The formulas remain, but their justification does not.

Power Query addresses all three limitations structurally. Transformations are applied through an explicit, named, sequenced pipeline of steps that is documented in the query itself. The steps are visible, editable, and labelled by the analyst. When the source changes, the analyst can inspect exactly which step was affected and why. The transformation logic is not hidden behind formula syntax; it is exposed as authored decisions, arranged in the order they were made, accessible to anyone who opens the editor.

This is not merely a usability difference. It is an architectural difference — and it determines whether a data preparation process is a professional asset or a personal workaround.

I.3 Power Query as an ETL Engine Inside Excel — and What ETL Actually Means

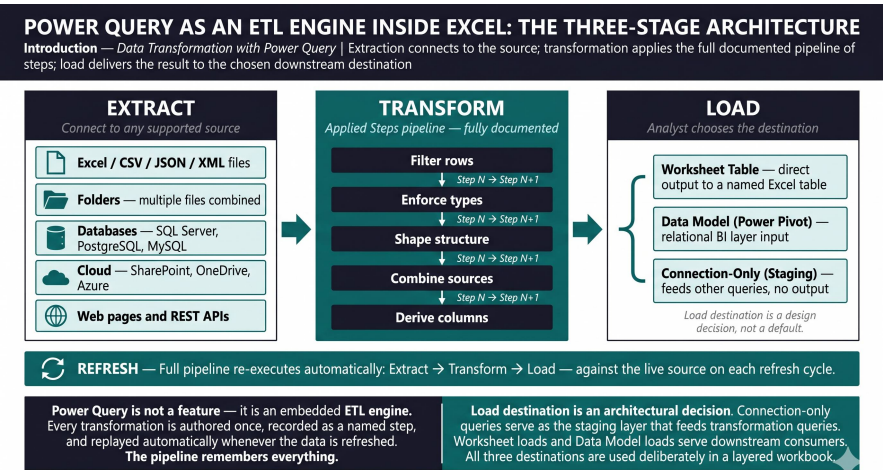
Power Query is described throughout this book as an ETL engine. The acronym — Extract, Transform, Load — comes from enterprise data engineering, where it describes the architecture of industrial-scale data pipeline systems. Understanding what it means, and why it applies accurately to what Power Query does inside Excel, is the conceptual foundation for everything that follows.

Extract is the process of connecting to a data source and reading data from it. This is the Connect stage of the pipeline. Power Query supports extraction from an exceptionally broad range of source types: local files in Excel, CSV, JSON, XML, and fixed-width formats; folders containing many files of the same or similar structure; relational databases including SQL Server, PostgreSQL, MySQL, and Oracle; cloud platforms including SharePoint, OneDrive, and Azure; web pages and HTML tables; and REST APIs using HTTP connections with configurable authentication and pagination. In every case, the Extract stage produces a representation of the source data inside the editor — a structured table that the transformation steps will operate on.

Transform is the process of shaping that extracted data into an analytically useful form. This is the majority of what Power Query does and the majority of what this book teaches. Transformation encompasses an enormous range of operations: filtering rows, removing duplicates, enforcing data types, splitting and merging columns, replacing values, handling nulls and errors, pivoting and unpivoting structures, grouping and aggregating, combining from multiple queries, and constructing derived columns through M expressions. Every transformation is recorded as a named step in the Applied Steps pane — which means every transformation is documented, auditable, and replayable.

Load is the process of delivering the transformed data to its downstream destination. In Power Query, a query can be loaded to a worksheet table, to the Excel Data Model for use with Power Pivot and DAX, or as a connection-only query that exists inside the workbook but produces no direct output — instead serving as an intermediate staging layer that feeds other queries. Load destination is an architectural decision, not a default to be accepted without thought.

What makes Power Query an ETL engine rather than simply a transformation tool is the combination of these three stages into a single, coherent, refresh-replayable pipeline. The pipeline remembers where it connects, what it does to the data it finds there, and where it delivers the result. When the source data is updated and the analyst clicks Refresh, the entire pipeline re-executes from Extract through Load automatically. The analyst does not intervene. The cleaned, transformed, correctly typed output is reproduced without manual effort.



I.4 The Concept of the Reproducible, Auditable Data Pipeline

The most important conceptual shift this book asks the reader to make is not technical. It is architectural. It concerns the difference between a data preparation workflow and a data preparation pipeline.

A workflow is a sequence of actions an analyst performs on data. It may be consistent across repetitions — the same analyst doing the same things each month — but the consistency is behavioral, not structural. It lives in the analyst's memory, their habits, and their notes. If they are absent, the workflow is absent. If they forget a step, the step is skipped. If someone else needs to execute the workflow, they must be trained — and the training is an approximation of what the original analyst actually does. There is no authoritative record.

A pipeline is a documented, executable program that transforms data from its source form to its target form through a specified set of operations, in a specified order, with a specified set of decisions made explicit at each step. The pipeline is not a description of what an analyst does — it is the thing itself. It can be re-run. It can be inspected. It can be modified at a specific step without touching the others. It can be inherited by a colleague who has never seen it before and understood within minutes. It can be audited when something goes wrong and corrected at the exact point of failure.

Power Query builds pipelines of the second kind. Every transformation the analyst applies is recorded as an Applied Step. The step has a name — which the analyst should choose deliberately — and the full M expression behind it is visible in the formula bar. The sequence of steps is the pipeline. It is the deliverable, not merely the means to a deliverable.

This framing has practical consequences. It means that naming steps well is not cosmetic — it is documentation. It means that the order of steps is a design decision, not an accident of the se-

quence in which transformations were applied. It means that inserting a new step in the middle of an existing pipeline, or reorganizing steps for performance reasons, or replacing one transformation with a more robust equivalent, are all legitimate and traceable professional activities. The pipeline can be maintained. That is the property that distinguishes professional data preparation from sophisticated manual work.

1.5 The Five Stages of an ETL Workflow — Connect, Profile, Shape, Combine, Author

Power Query supports a broad range of transformations, and the breadth can feel disorienting to a learner approaching the tool for the first time. This book organizes that breadth into five stages that reflect the natural sequence of a professional ETL workflow. Understanding the stages before the technical instruction begins orients each chapter within the larger architecture.

Connect is the first stage. Before any transformation can be applied, the analyst must establish a connection to the data source, navigate its structure, and configure the connection for stability, portability, and future refresh reliability. Authentication, privacy levels, locale settings, file path parameterization, and connector-specific behavior are all concerns of the Connect stage. Getting the connection right is not a mechanical step before the “real” work begins — it is an architectural decision whose consequences propagate across the entire pipeline lifetime.

Profile is the second stage. Before transforming an unfamiliar source, the professional analyst profiles it: examining column quality, value distributions, null rates, error distributions, and the grain of the data — the level at which each row represents a unique fact. Profiling reveals problems that transformation planning must address. Cleaning without profiling is improvisation. Profiling before writing a single transformation step is engineering. The profiling tools inside Power Query — Column Quality, Column Distribution, and

Column Profile — make this stage systematic rather than approximate.

Shape is the third stage, encompassing the majority of cleaning and transformation work. Shaping includes row-level operations (filtering, deduplication, error handling), column-level operations (renaming, reordering, splitting, merging), structural operations (pivot, unpivot, transpose, group by), type enforcement, value replacement, and derived column construction. The Shape stage is where the raw source becomes an analytically usable table.

Combine is the fourth stage. Real analytical work rarely involves a single source. The Combine stage encompasses the operations that bring multiple queries together: Append for stacking same-structure sources into a single longer table, and Merge for joining related tables on matching key columns — the relational join operations available in the Power Query layer before the Data Model is introduced.

Author is the fifth stage, encompassing the M language capabilities that go beyond what the ribbon can produce: reading and understanding generated M code, writing custom column expressions, parameterizing queries for dynamic behavior, and authoring small reusable functions for transformations the interface cannot produce by clicking. Authoring is the stage that distinguishes a Power Query practitioner from a Power Query user.

1.6 Pipeline Architecture as a Professional Discipline

Professional Power Query work is not a matter of applying transformations until the data looks right. It is a matter of designing a pipeline with a structure that is maintainable, performant, and comprehensible to someone other than its author. That design discipline has a name: layered architecture.

A well-structured Power Query workbook organizes its queries into logical layers, each with a distinct and bounded responsibility.

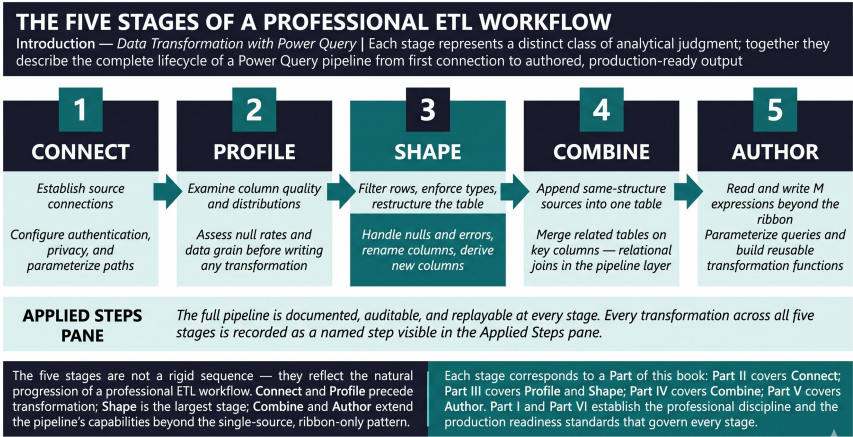


Figure I.3 The five stages of a professional ETL workflow as organized in this book. Each stage represents a distinct class of analytical judgment. Together they describe the complete lifecycle of a Power Query pipeline from first connection to authored, production-ready output.

Figure I.3. The five stages of a professional ETL workflow as organized in this book. Each stage represents a distinct class of analytical judgment. Together they describe the complete lifecycle of a Power Query pipeline from first connection to authored, production-ready output.

The staging layer holds connection-only queries that reflect the source data with minimal transformation — renaming for readability, establishing correct types, and nothing more. The transformation layer applies the business logic: the filtering, shaping, combining, and derivation that turn staged data into analytical inputs. The load layer holds the final queries that deliver typed, structured data to their downstream destination — whether a worksheet, the Data Model, or another consumer query.

The separation of concerns this architecture enforces has immediate practical consequences. When the source changes its schema — a column renamed, a new field added — the fix is isolated to the staging layer. The transformation and load layers are unaffected. When a business rule changes and a filtering condition must be updated, the change is isolated to the transformation layer and traced in a single named step. When a new downstream report needs a slightly different projection of the same data, a new load query can consume the existing transformation layer without duplicating the transformation logic.

This architecture is not imposed by Power Query — the tool does not require it. It is a professional discipline that must be chosen and applied deliberately. Most Power Query workbooks in practice are not layered. They are collections of queries in which connection, transformation, and output concern are interleaved indiscriminately, producing pipelines that are difficult to debug, expensive to maintain, and impossible to audit systematically. The difference between those workbooks and the pipelines this book teaches is not a technical difference. It is an architectural one — and it is entirely within the analyst's control.

Pipeline architecture as a discipline encompasses naming conventions, layer structure, step documentation, query grouping, parameter design, and the production readiness standards that determine whether a pipeline is fit for handoff. Each of those elements is covered in the chapters that follow.

I.7 What This Book Deliberately Leaves to Book 4

Several capabilities that are technically within Power Query's reach are deliberately excluded from this volume's primary instruction and left to Volume 4.

The most significant exclusion is the Data Model as a relational analytical environment. This book introduces the Data Model as a load destination — the analyst can load a query's output to the model rather than to a worksheet — but it does not teach Power Pivot table relationships, DAX calculations, or the dimensional modeling concepts that make the Data Model analytically powerful. Those are the exclusive subject of Volume 4, and they require their own dedicated treatment. Teaching them here would undercut both books.

Calendar tables are a related exclusion. A complete ETL pipeline for time-series analysis often includes a generated or imported calendar table to enable proper time intelligence. This book covers date-related transformations in Power Query — date part ex-

traction, date arithmetic, fiscal period derivation — but stops short of designing the full calendar table structure that Volume 4's DAX time intelligence functions require. That design is a Data Model concept, and it belongs there.

Query folding — the mechanism by which Power Query pushes transformation logic back to the source database engine rather than processing it in memory — is introduced as a concept in this volume's connection and performance chapters. But the advanced query folding analysis that becomes necessary when designing large-scale Power BI data models is a concern of the PL-300 and DP-600 pathways, not of this series.

Knowing what this book does not teach is as important as knowing what it does. The boundaries are intentional, not accidental, and they reflect a series architecture in which each volume's scope is exactly what the next volume needs to build on.

I.8 The Path Forward — From Pipeline to Relational Model

When a reader completes this book, they will have a data preparation pipeline capable of connecting to virtually any source Excel supports, transforming the data it finds there through a documented and reproducible sequence of operations, and delivering typed, clean, correctly structured output to the Excel Data Model.

That output is the precise input that Volume 4 requires.

Volume 4 — *Data Modeling with Power Pivot & DAX* — begins where this book ends. It picks up the prepared data sitting in the Data Model's load destination, establishes relationships between tables, designs the dimensional structure of a star schema, and builds the DAX measures that turn a relational data model into a dynamic analytical environment. Every technique in Volume 4 depends on the pipeline infrastructure this book teaches being sound. A Data Model built on unprepared, incorrectly typed, or structurally inconsistent Power Query output will produce errors, performance

problems, and analytical failures that are expensive to diagnose and impossible to prevent after the fact.

The work done in this book is not preparatory to the “real” analytical work. It is the foundation of it. The cleanliness of the pipeline determines the integrity of the model. The integrity of the model determines the reliability of the measures. The reliability of the measures determines the trustworthiness of the dashboard. That chain of dependency runs from this book all the way to the Capstone — and it begins with the first connection the analyst makes in Chapter 1.

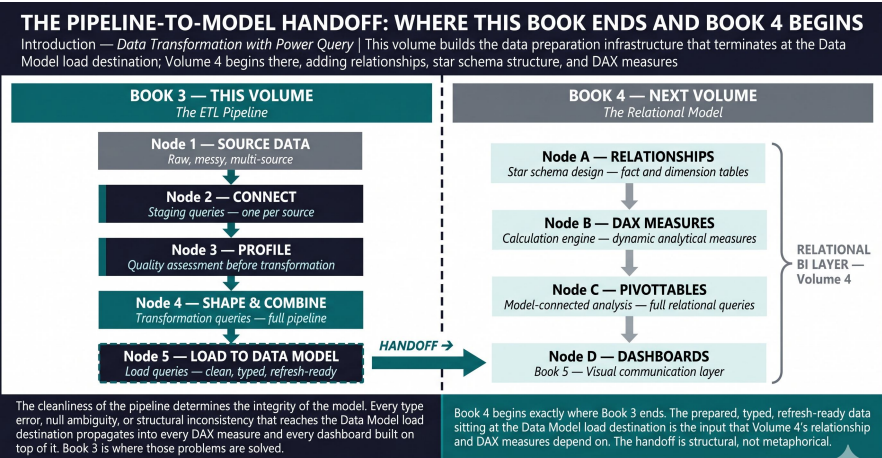


Figure I.4 The handoff from pipeline to model. This volume builds the data preparation infrastructure that terminates at the Data Model load destination. Volume 4 begins there, adding relationships, star schema structure, and DAX measures to the clean, typed data this volume delivers.

Figure I.4. The handoff from pipeline to model. This volume builds the data preparation infrastructure that terminates at the Data Model load destination. Volume 4 begins there, adding relationships, star schema structure, and DAX measures to the clean, typed data this volume delivers.

PART I

The ETL Discipline and the Power Query Environment

Part I Introduction

Most analysts who open Power Query for the first time do so under mild duress. A file is too large to clean by hand, or a colleague has recommended it, or a tutorial has surfaced in a search for how to combine several spreadsheets into one. The analyst tries it, makes something work, and returns to their usual workflow. They leave with a new tool they do not fully understand, sitting inside an activity they have not yet reframed.

Part I addresses that reframing before the technical instruction begins. It makes the argument — not assumed, not asserted, but argued — that data preparation is a professional engineering discipline with its own standards, architecture, and craft requirements, and that Power Query is the instrument through which that discipline is practiced inside Excel. By the end of Part I, the reader will not merely understand what the Power Query editor looks like and how to navigate it. They will understand why the way they have been preparing data imposes costs that are invisible until they are named, and why the alternative this book teaches is not simply faster but structurally different in ways that matter for professional analytical work.

Chapter 1 makes the case for Power Query by diagnosing the manual preparation pattern it replaces. Chapter 2 orients the reader inside the editor as an analytical environment, covering its layout, its load destinations, and its configuration for professional use. Chapter 3 establishes the Applied Steps pane — not as a feature of the interface, but as the most important audit artifact in the analyst's toolkit and the mechanism through which a Power Query workflow becomes a documented, defensible pipeline.

Every subsequent chapter in this book assumes the orientations Part I builds. They are not re-established. They are applied.

CHAPTER 1

The Data Preparation Reality and the Case for Power Query

Chapter Introduction

This chapter does not begin with Power Query. It begins with the problem Power Query solves.

That sequencing is deliberate. An analyst who understands the full cost of manual data preparation — not merely in time, but in trust, reproducibility, and institutional memory — approaches Power Query as a professional answer to a professional problem, not as a ribbon feature to be explored when convenient. The re-framing that Chapter 1 accomplishes is the most important foundation this book lays, and it is worth the chapter required to lay it.

The chapter moves through four arguments in sequence. First, it establishes the factual reality of how analysts spend their time and why that distribution is not inevitable. Second, it examines the two dominant approaches to data preparation that predate Power Query — the manual copy-paste workflow and the formula-based cleaning workflow — and diagnoses their structural failure modes. Third, it introduces the concept of a reproducible, auditable pipeline as the professional alternative. Fourth, it maps the disciplines of that alternative across the six stages that organize this book, so that every subsequent chapter can be understood in relation to the whole.

By the end of this chapter, the reader should be able to name what is structurally wrong with their current data preparation practice,

articulate why Power Query is an architectural answer rather than a convenience feature, and orient themselves within the discipline map that the rest of this book will teach.

Chapter Objectives

Upon completing this chapter, the reader will be able to:

- Describe the hidden costs of manual data preparation — in analyst time, analytical trust, and institutional memory — and explain why those costs are not inherent to the work but to the method
- Distinguish between the manual workflow, the formula-based workflow, and the pipeline-based workflow, and identify the structural failure mode of each non-pipeline approach
- Explain what ETL means — Extract, Transform, Load — as a professional discipline and how Power Query implements each stage inside Excel
- Position Power Query correctly in relation to worksheet formulas and VBA macros, using the appropriate tool for the appropriate analytical task
- Locate Power Query within the layered workbook architecture introduced in Volume 1 of this series
- Describe the concept of a reproducible, auditable data pipeline and explain what distinguishes it from a data preparation workflow
- Map the six stages of the ETL discipline introduced in this chapter — Connect, Profile, Shape, Combine, Author, Deploy — to the Parts of this book

1.1 Why Analysts Spend the Majority of Their Time on Data Preparation

1.1.1 The Hidden 80% of Analytical Work — Time Spent, Errors Made, Reproducibility Lost

The figure appears in nearly every survey of applied analytical work, and it is worth examining rather than merely citing. Studies of data analyst and data scientist time allocation have consistently found that a substantial majority of working time — figures ranging from 60 to 80 percent across different industries and job levels — is consumed by activities that precede the analysis itself: locating data, extracting it from source systems, cleaning it, reshaping it into a usable form, and validating that the result reflects the source accurately. The analysis — the charts, the models, the recommendations, the insight — consumes the minority.

These findings are often presented with a tone of resigned fatalism, as though the ratio describes a natural property of analytical work that practitioners must learn to endure more efficiently. It does not. The ratio is not a property of the work. It is a property of how the work is organized — and specifically of whether data preparation is treated as an improvised manual activity performed anew each time data arrives, or as an engineered, documented, automated discipline that is built once and replayed thereafter.

The distinction matters for reasons beyond time. Manual preparation does not merely consume analyst hours — it produces analytical risk at each repetition. Every time a human being performs a sequence of filtering, trimming, formatting, and copy-pasting operations from memory, the risk of inconsistency accumulates. A filter criterion applied last month may not match the criterion applied this month if the analyst's memory of last month's decision is imperfect. A SUBSTITUTE formula applied to fix a recoding inconsistency in the source may be applied to the wrong column if the source has added a new field since the last cycle. A subtotal row removed manually may be missed in a particular refresh if the

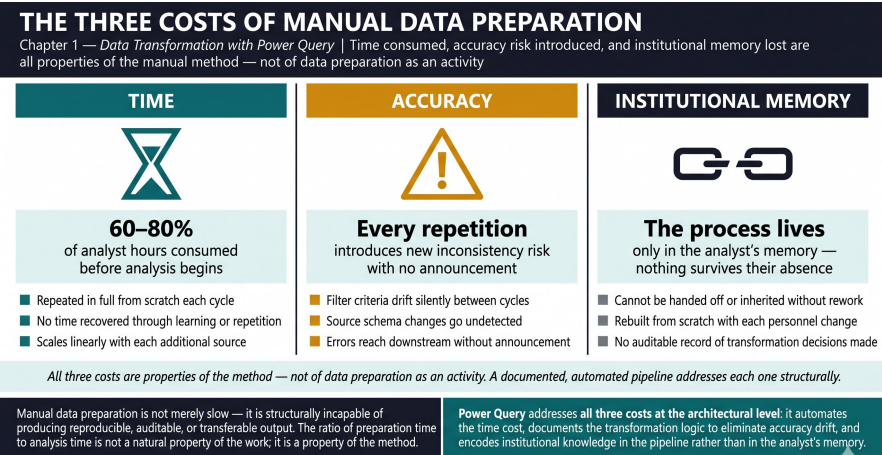
analyst is working under time pressure. None of these errors announces itself. The output looks clean. The numbers proceed to downstream analysis and, eventually, to a report that is read by someone who has no way of knowing the preparatory decisions that produced it.

The third cost — institutional memory — is the least visible and often the most expensive. When data preparation is manual and undocumented, the analyst who built the process is the process. When that analyst leaves, takes a different role, or is simply unavailable on a deadline, the preparation cannot be handed off, because there is nothing to hand. A new analyst inherits a workbook, looks at the output, and has no reliable way of determining what was done to the source to produce it. They begin again, making their own decisions, introducing their own inconsistencies, at the cost of their own hours. Institutional analytical knowledge is lost and reconstructed repeatedly, in each cycle and with each personnel change, because it was never encoded anywhere that outlasted its author.

Power Query does not reduce the amount of analytical work that needs to be done. It does not make source data cleaner. What it does is make the preparation work — in full, with every decision recorded — into a replayable, auditable, transferable artifact. The hidden 80% does not disappear; it is automated, documented, and rendered permanent.

1.1.2 The Copy-Paste-Formula Workflow and Its Failure Modes

The dominant data preparation approach in professional Excel practice takes a recognizable form. A source file arrives — from an export, an email, a system extract. The analyst opens it, opens the workbook they are building, and begins moving data between them: copying ranges, pasting values, applying a column of TRIM formulas to strip leading spaces, a SUBSTITUTE formula to remove stray characters, a TEXT-to-date conversion to fix the date fields that arrived as text strings, and an IFERROR wrapper to suppress



put is used, which requires knowing what the output should look like, which is exactly the kind of institutional memory that only the original analyst possesses.

The second failure mode is the paste-over-delete pattern. The act of pasting formula outputs over themselves to remove formula references, and then deleting the helper columns, is a one-way operation. Once it is done, the transformation history is gone. There is no way to look at the cleaned data and determine what was done to it. The audit trail — whatever it was — has been destroyed in the act of producing the output. If the output is later questioned, the analyst must either reconstruct the transformation from memory or re-execute the entire cleaning process from the source to verify that the output is correct. In either case, time is consumed for a purpose that a documented pipeline would eliminate entirely.

The third failure mode is the repeatability gap. Each cycle of the copy-paste-formula workflow requires the analyst to execute the same sequence of steps from memory. Small decisions accumulate across cycles — which column to TRIM, which values to substitute, which error condition to suppress — and the memory of those decisions degrades over time. By the sixth or twelfth monthly cycle, the workflow the analyst executes is not identical to the workflow they executed in the first cycle. It is a memory-based approximation. The outputs of the two cycles are not guaranteed to be produced by the same logic. Comparing results across cycles thus carries a latent inconsistency that the analysis does not account for and cannot detect.

1.1.3 What Manual Preparation Costs — Time, Trust, and Institutional Memory

The analysis above identifies individual failure modes. But the cumulative cost of the manual preparation pattern is best understood not per-failure but across the full lifecycle of an analytical product — a monthly report, a quarterly dashboard, a recurring regulatory return.

Consider a report that is produced monthly. The preparation cycle takes three hours each month. Over twelve months, that is thirty-six hours of analyst time consumed by a sequence of operations that will be performed again in substantially the same form next month. Thirty-six hours is not the end of the cost. When the analyst who owns the report takes annual leave in month eight, another analyst must produce the report in their absence. They spend an additional four hours — two trying to understand what the original analyst did, two actually doing it based on their best inference. The report is produced on time. Whether it is produced correctly is unknown, because there is no documented standard to verify it against.

At month fourteen, the analyst who built the original workflow leaves the organization. Their replacement is given the workbook and told that the report runs monthly. The workbook contains cleaned data but no record of how it was cleaned. The replacement analyst rebuilds the preparation process over the course of a week, making their own decisions at each ambiguous step, some of which align with their predecessor's decisions and some of which do not. The trend data in the report now carries a discontinuity at month fourteen that the report's consumers cannot see and the report's author does not know to disclose.

This is not a failure of individual analysts. It is the predictable output of a method in which institutional knowledge about data transformation is stored in human memory rather than in a documented, executable artifact. The cost is not a one-time penalty for building the wrong tool; it is a compounding cost that accumulates with every cycle, every personnel transition, and every question about whether last month's number was produced by the same logic as this month's. Trust in an analytical product, once undermined, is expensive to rebuild and never fully recovered.

A pipeline built in Power Query addresses this cost at the structural level: the preparation is documented, automated, and independent of the individual who built it. The knowledge is in the tool, not in a person's memory.

1.2 What Power Query Is and What It Is Not

1.2.1 The ETL Acronym — Extract, Transform, Load as a Professional Discipline

Power Query is described throughout this book as an ETL tool. The term comes from enterprise data engineering, where it describes the architecture of industrial pipeline systems that move data from operational source systems into analytical data warehouses. Understanding the acronym in its original context clarifies what Power Query does — and why applying that framing to a tool embedded in Excel is accurate rather than aspirational.

Extract describes the process of connecting to a source system and retrieving data from it. In enterprise ETL platforms, this means writing connectors to databases, flat files, APIs, and legacy systems, handling authentication, managing the read load on source systems, and retrieving data in a form the transformation layer can process. In Power Query, extraction is accomplished through connectors — dozens of them, covering a range that spans local Excel files and CSV exports through relational databases, SharePoint lists, OneDrive folders, REST APIs, and web pages. The analyst selects a connector, configures the connection, navigates the source structure, and retrieves data to the editor. Extract is complete when the data appears in the preview pane as a structured table.

Transform describes every operation applied to the extracted data before it reaches its destination. In enterprise ETL, transformation encompasses business rule application, data quality enforcement, referential integrity checks, aggregation, derived field construction, and structural reshaping — all implemented in a pipeline that is versioned, tested, and deployed. In Power Query, transformation is accomplished through the Applied Steps pane: every operation the analyst applies — filtering, type enforcement, column splitting, value replacement, pivoting, merging — is recorded as a named, sequenced step in the pipeline. The full transformation is

documented in the query as an executable M expression, visible and editable at any time.

Load describes the delivery of the transformed data to its destination. In enterprise ETL, this means writing to a target database table, updating a data mart, or inserting records into a warehouse. In Power Query, the destination choices are three: a worksheet table for direct use in formulas and charts, the Excel Data Model for use with Power Pivot and DAX, or a connection-only query that exists in memory as an intermediate stage feeding other downstream queries without producing any direct output. Load destination is an explicit design decision, not a default to be accepted without consideration.

What makes Power Query an ETL tool rather than simply a set of transformation features is the integration of all three stages into a single, coordinated, refresh-replayable pipeline. When the source is updated and the analyst triggers a refresh, all three stages re-execute in sequence — the connection re-reads the source, the full transformation pipeline re-applies every step, and the output is redelivered to the destination. The analyst does not intervene. The pipeline is the process.

1.2.2 Power Query vs. Worksheet Formulas vs. Macros — Choosing the Right Tool

Microsoft Excel provides three distinct automation paradigms for analytical work: worksheet formulas, VBA macros, and Power Query. Each has its appropriate domain, and the professional discipline of using Excel for analytical work includes knowing which tool belongs to which task.

Worksheet formulas operate on individual cells or ranges within the Excel calculation engine. They are the correct instrument for deriving new values from existing data within the analysis layer of the workbook: calculating running totals, applying lookup logic, constructing conditional flags, building dynamic summaries. For-

mulas are live: they recalculate automatically whenever their inputs change. They are wrong for data preparation precisely because of this live quality — a formula applied to raw source data is dependent on the structure and content of that source, and changes to the source break the formula silently rather than through a documented step that the analyst can inspect and repair.

VBA macros are recorded or authored programs that automate sequences of Excel operations, including operations that formulas cannot accomplish: navigating between files, running external queries, manipulating workbook structure, and controlling the Excel application itself. Macros are appropriate for automating repetitive operational workflows that do not fit the formula or pipeline paradigm — generating formatted reports, managing file exports, automating print sequences. They are wrong for data preparation because their transformation logic is embedded in procedural code that is difficult to audit without VBA proficiency, impossible to inspect at the level of an individual transformation step, and not natively refreshable against a live data source.

Power Query is the correct instrument for data preparation: acquiring data from a source, transforming it from its source form into an analytically usable form, and loading the result to a destination that formulas and models can consume. It is not a replacement for formulas in the analysis layer. It is not a replacement for macros in operational workflows. It is the appropriate tool for the layer between raw source data and clean analytical input — a layer that formulas were never designed to occupy and macros address with unnecessary complexity.

1.2.3 Power Query as a Layer in the Layered Workbook Architecture from Book 1

Readers who completed Volume 1 of this series — *Excel for Analysts: The Analytical Foundation* — will recognize the concept of the layered workbook architecture: the discipline of separating the raw data layer, the transformation layer, the analysis layer, and

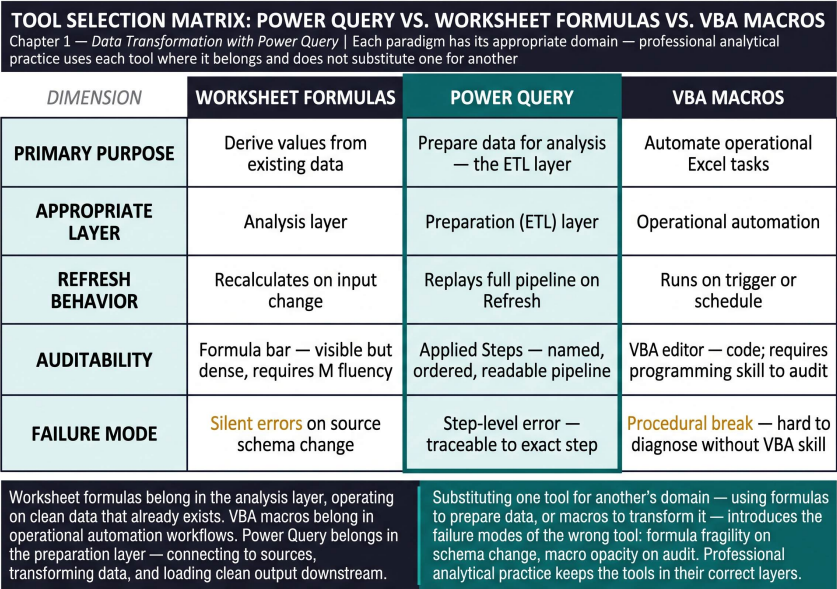


Figure 1.2 The tool selection matrix for Excel analytical automation. Each paradigm has its appropriate domain: formulas for the analysis layer, Power Query for the data preparation layer, and macros for operational automation. Professional analytical practice uses each tool where it belongs and does not substitute one for another.

Figure 1.2. The tool selection matrix for Excel analytical automation. Each paradigm has its appropriate domain: formulas for the analysis layer, Power Query for the data preparation layer, and macros for operational automation. Professional analytical practice uses each tool where it belongs and does not substitute one for another.

the presentation layer within a single workbook, so that each concerns itself exclusively with its own responsibility and does not contaminate the responsibilities of adjacent layers.

Power Query occupies a specific and bounded position in that architecture. It is the mechanism of the transformation layer.

In the Volume 1 workbook architecture, the transformation layer sat between the raw data layer and the analysis layer. Raw source data arrived in the raw data layer — unchanged, as it came from the source. The transformation layer applied cleaning and reshaping logic — manually, or through formula chains — and produced a clean structured table for the analysis layer to consume. The analysis layer then applied formulas, PivotTables, and aggregation logic to that clean table. The presentation layer took the output of the analysis layer and formatted it for its audience.

Power Query replaces the manual or formula-based transformation layer with a documented, automated pipeline. The raw data layer now exists as a source external to or minimally present in the workbook — a file, a database, a cloud connection — rather than a pasted range. Power Query connects to it, transforms it through the pipeline, and loads the result to the workbook at the point where the analysis layer expects to find it. The analysis layer does not change: it still consumes a clean, structured table. What changes is how that table arrives and how reliably it reflects the transformation decisions that were made.

The architecture introduced in Volume 1 remains intact. Power Query does not displace it; it fulfills one layer of it with an engineering-grade tool.

1.3 The ETL Paradigm — From One-Off Tasks to Reproducible Pipelines

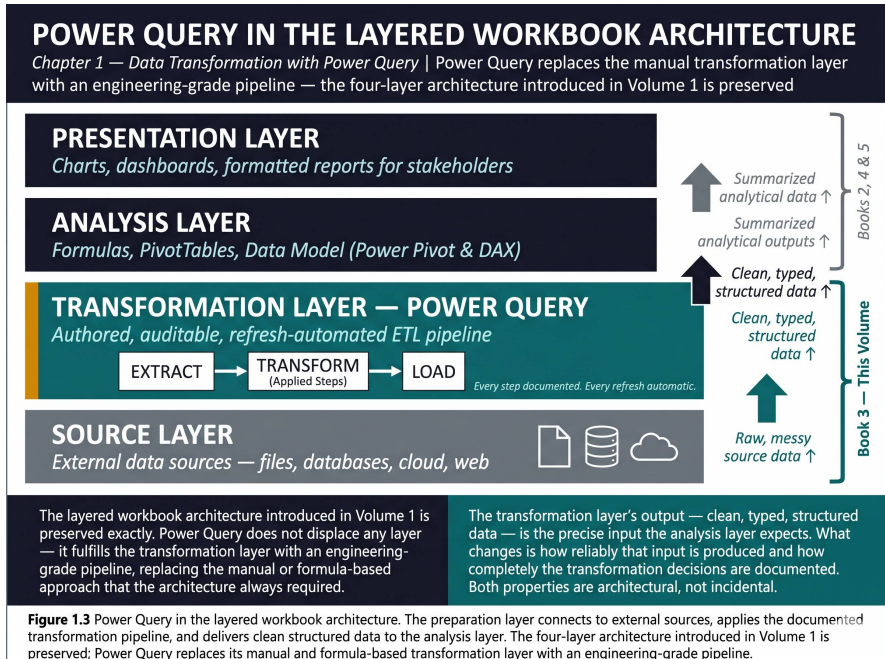


Figure 1.3. Power Query in the layered workbook architecture. The preparation layer connects to external sources, applies the documented transformation pipeline, and delivers clean structured data to the analysis layer. The four-layer architecture introduced in Volume 1 is preserved; Power Query replaces its manual and formula-based transformation layer with an engineering-grade pipeline.

1.3.1 The Concept of a Pipeline — Authored, Auditable, and Refreshable

The word “pipeline” is used frequently in data engineering contexts, often loosely. In the specific context of this book, it has a precise meaning, and that precision matters for understanding what Power Query builds and why.

A pipeline, as used here, is a sequence of explicitly authored, named, and ordered transformation operations that converts input data from its source form to its target form, can be re-executed against updated source data without modification, and preserves a complete and inspectable record of every transformation decision it applies.

Each element of that definition is doing work.

Explicitly authored means that every transformation in the pipeline was placed there by deliberate decision. It was not applied automatically by the tool, not inferred from the data, and not left over from a previous version. The analyst made a decision — “remove rows where this column is null,” “split this field on the semi-colon delimiter,” “rename this column to match the downstream schema” — and that decision is now encoded in the pipeline as a step.

Named and ordered means that the sequence of steps is visible, labelled with names the analyst has chosen, and arranged in a specific order that the analyst controls. The order is not arbitrary — it reflects the logic of the transformation, the performance implications of early versus late filtering, and the dependency relationships between steps. A well-ordered, well-named pipeline can be read from top to bottom as a plain-language description of what was done to the data and why.

Re-executable against updated source data is the operational definition of a reproducible pipeline. The source data changes — a new month’s records arrive, a database table is updated, a folder receives a new file. The analyst clicks Refresh, or the workbook

is opened with background refresh enabled, and the full pipeline re-executes against the updated source, applying every transformation in exactly the same sequence, delivering the same structural output. The analyst does not intervene. There is nothing to remember and no opportunity to make a different decision than last month's decision.

Preserves a complete and inspectable record is the operational definition of an auditable pipeline. Every step is visible in the Applied Steps pane. Every step's underlying M expression is accessible in the formula bar or the Advanced Editor. When the output is questioned — “did this month's figure use the same category mapping as last month's?” — the answer is available in the pipeline itself, without requiring the analyst's memory or a separate documentation artifact.

These three properties — reproducibility, auditability, and refresh automation — are what distinguish a data preparation pipeline from a data preparation workflow. The workflow produces the same output, or approximately the same output. The pipeline produces the same output, documented, automatically, every time.

1.3.2 Reproducibility, Auditability, and Transformation Traceability as Professional Standards

In professional analytical environments, reproducibility and auditability are not merely operational conveniences — they are the conditions under which analytical outputs can be trusted.

Reproducibility means that the same pipeline applied to the same source data produces the same output. This is a necessary condition for comparing outputs across time periods: if the transformation logic changes between cycles — even slightly, even inadvertently — then what appears to be a trend in the data may be a trend in the analyst's decisions. A pipeline that is documented and re-executed mechanically is not subject to this risk. A workflow that is executed from memory is.

Auditability means that the logic of a transformation can be inspected and verified by someone other than its author. In regulated industries — finance, healthcare, public administration — this is often a compliance requirement rather than a professional preference. In all industries, it is the property that allows analytical outputs to be questioned, challenged, and ultimately trusted. An output whose preparation logic cannot be inspected cannot be fully verified. An output whose preparation is documented step by step in a readable pipeline can be reviewed, approved, and signed off on as a professional deliverable.

Transformation traceability is the specific capability to follow a value in the output back through the transformation steps that produced it and forward through the downstream calculations that consume it. If a figure in a report appears anomalous, traceability means the analyst can determine whether the anomaly originates in the source data, in a transformation step, or in a downstream formula — and can pinpoint the exact location. Without traceability, debugging an anomaly requires a full re-examination of everything that could have affected the figure, which is expensive in time and unreliable in result.

Power Query's Applied Steps pane is the mechanism for all three properties. Each step is a traceable, inspectable, reproducible transformation decision. Together they constitute the transformation audit trail — the professional record that justifies the output to anyone who has the right and the need to question it.

1.3.3 Transitioning from Spreadsheet Thinking to Pipeline Thinking

The shift from manual data preparation to pipeline-based data preparation requires a corresponding shift in how the analyst reasons about their work. This shift is as important as any technical skill the book teaches, and it is worth articulating explicitly before the technical instruction begins.

Spreadsheet thinking approaches data preparation as a series of individual actions performed on a dataset: open the file, sort the column, filter out the blanks, trim the text, fix the dates, delete the helper columns. Each action is complete in itself. The result accumulates cell by cell, column by column, until the dataset looks right. The analyst's attention is on the current state of the data — what it looks like now — and on the next action needed to improve it.

Pipeline thinking approaches data preparation as the design of a sequence of documented operations that transforms a class of inputs into a class of outputs, reliably and automatically, regardless of which specific instance of that input arrives. The analyst's attention is not on the current state of the data but on the design of the transformation architecture — what the pipeline must do, in what order, to any version of this source that conforms to the expected schema. The individual transformation step matters not because it fixes the current problem, but because it is the correctly placed, correctly named decision that will be replayed against next month's data with no intervention.

This is a genuine cognitive shift, and it does not happen automatically when the analyst opens the Power Query editor for the first time. It is possible to use Power Query in the spreadsheet-thinking mode — clicking through the interface, accepting default step names, building queries that work once but are fragile on refresh — and to miss the pipeline-thinking discipline entirely. This book is designed to prevent that outcome. The naming conventions, the layer architecture, the Applied Steps discipline, the production readiness standards — all of these are expressions of the pipeline-thinking orientation, and all of them are introduced and reinforced throughout the text.

The transition is worth making. Analysts who have completed it describe the experience of inheriting an undocumented manual preparation process as qualitatively different from before — not merely slower, but structurally wrong in a way they can now articulate and remedy.

1.4 The Discipline Map for This Book

1.4.1 Connect, Profile, Shape, Combine, Author, Deploy — Six Stages of an ETL Workflow

This book is organized around six stages that represent the complete professional lifecycle of a Power Query pipeline. Each stage is a distinct class of analytical judgment, and the chapters that teach each stage form a coherent discipline rather than a collection of features.

Stage 1: Connect. Before any transformation can occur, the analyst must establish a connection to the data source. Connection is not a mechanical preliminary — it is a set of architectural decisions whose consequences propagate across the pipeline’s entire lifetime. Which connector to use, how to configure authentication, what privacy level to assign, how to parameterize the source path for portability and refresh resilience: these are all design decisions made at the Connect stage, and they determine the fragility or robustness of everything that follows. Part II of this book (Chapters 4 and 5) covers connection to file, folder, database, cloud, and web sources in full professional depth.

Stage 2: Profile. Once connected, the analyst’s first act on an unfamiliar source should not be transformation — it should be diagnosis. The profiling tools in Power Query — Column Quality, Column Distribution, and Column Profile — reveal the actual state of the data: null rates by column, error rates, value distributions, cardinality, and the granularity of the dataset. Transformation decisions made without profiling are guesses. Transformation decisions made after profiling are responses to documented findings. Part III (Chapter 6) covers data profiling as a professional analytical discipline.

Stage 3: Shape. The Shape stage encompasses the majority of what most analysts think of as “data cleaning”: filtering rows, removing duplicates, enforcing data types, renaming and reorder-

ing columns, splitting and merging fields, replacing values, handling null and error conditions, and deriving new columns through expressions. It also includes structural transformations — pivot, unpivot, transpose, group by — that reshape data from its source form into an analytically appropriate form. Part III (Chapters 7 and 8) covers shaping in full.

Stage 4: Combine. Real analytical work rarely operates on a single source. The Combine stage encompasses the operations that bring multiple queries together: Append for stacking same-structure sources into a unified longer table, and Merge for joining related tables on matching key columns — the relational join operations available in the Power Query layer, before the full relational model is introduced in Volume 4. Part IV (Chapters 9 and 10) covers both combination approaches.

Stage 5: Author. The Author stage encompasses the M language capabilities that go beyond what the ribbon interface can produce: reading and understanding the M expressions that Power Query generates from ribbon operations, writing custom column expressions that ribbon operations cannot produce, parameterizing queries for dynamic behavior, building reusable functions for transformations applied across multiple queries, and structuring queries into the layered staging-transformation-load architecture. Part V (Chapters 11 through 14) covers M as an analytical reading and writing skill.

Stage 6: Deploy. A pipeline that works correctly in a development context is not the same as a pipeline that is fit for production. The Deploy stage covers the capabilities that determine whether a pipeline survives contact with a live operational environment: query folding and performance optimization, error handling at refresh time, data quality checks within the pipeline, and the production readiness standards — documentation, naming, handoff protocols — that allow a pipeline to be shared, inherited, and maintained. Part VI (Chapters 15 and 16) covers production deployment in full.

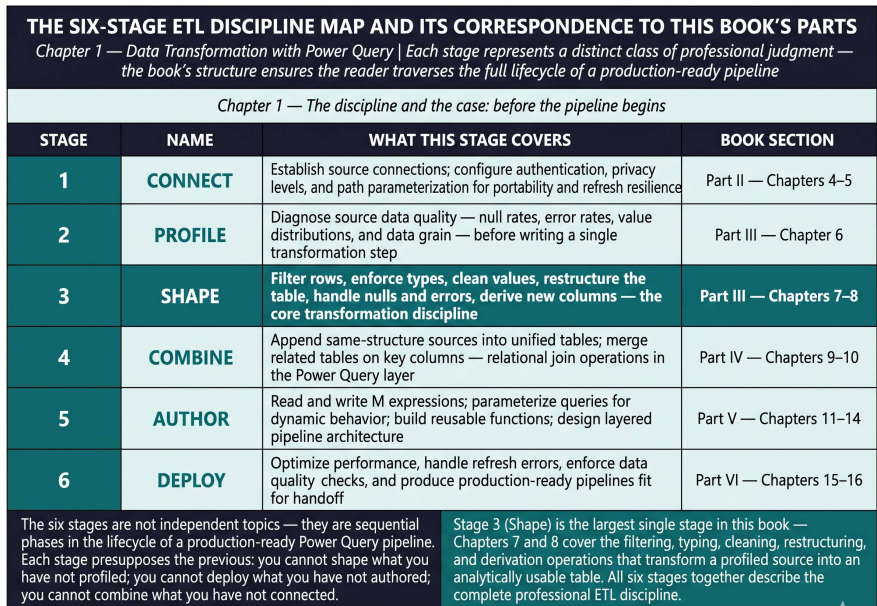


Figure 1.4 The six-stage ETL discipline map and its correspondence to the Parts and chapters of this book. Each stage represents a distinct class of professional judgment. The book's structure ensures that by the final chapter, the reader has traversed the full lifecycle of a production-ready Power Query pipeline.

Figure 1.4. The six-stage ETL discipline map and its correspondence to the Parts and chapters of this book. Each stage represents a distinct class of professional judgment. The book's structure ensures that by the final chapter, the reader has traversed the full lifecycle of a production-ready Power Query pipeline.

1.4.2 Prerequisites Carried Forward from Books 1 and 2

This book assumes specific capabilities that were the primary subjects of Volumes 1 and 2. They are not re-taught here. They are assumed and applied.

From Volume 1, the following are assumed: the analyst works in the layered workbook architecture, understands the role of Excel Tables and structured references, assigns data types with awareness of their analytical consequences, and distinguishes between raw data, transformation logic, and analytical outputs as conceptually and structurally separate concerns. The Power Query pipeline constructed in this book is the fulfillment of the transformation layer that Volume 1 identified as necessary. Readers who did not complete Volume 1 should review the Prerequisites section of this book's front matter and confirm they carry the relevant capabilities before proceeding.

From Volume 2, the following are applied in the context of this book without re-introduction: the ability to read a formula of moderate complexity and describe what it calculates, the understanding of how data types interact with formula operations, and the discipline of writing expressions that are correct, readable, and maintainable. These capabilities are applied when reading M expressions in Part V, when designing custom column formulas, and when reasoning about the interaction between Power Query outputs and downstream formula architecture.

The Power Query content in this book is not formally dependent on any specific formula function from Volume 2. But the habit of analytical reasoning that Volume 2 builds — treating each expression as a documented decision with consequences, designing logic that is readable by others, handling errors as design considerations rather than afterthoughts — is directly applied in the pipeline-thinking orientation that this book teaches.

1.4.3 What This Book Hands Off to Book 4

Precision about scope boundaries is a professional discipline, not a limitation. This section identifies what this book explicitly does not cover, and confirms that what it does not cover is addressed completely in the volume for which that coverage is appropriate.

The Data Model as a relational analytical environment — table relationships, star schema design, cardinality, DAX calculations — is outside the scope of this book. The Data Model appears here only as a load destination: Power Query can deliver its output to the model, and the analyst should understand the implications of that destination choice. What happens inside the model after that delivery is Volume 4's subject.

Calendar table design, time intelligence, and date dimension architecture are excluded for the same reason. Power Query in this book performs date-level transformations — extracting year, quarter, month, and week values from date columns, calculating date differences, deriving fiscal period fields — but it stops short of designing the full calendar table that Volume 4's DAX time intelligence functions will require.

The relational join semantics that Power Query's Merge operation performs — inner join, left outer, right outer, full outer, left anti, right anti — are covered in Part IV at the level of practical usage and analytical judgment. The full dimensional modeling implications of those join patterns, and the specific join logic required to build correct star schema relationships, belong to Volume 4.

What this book delivers to Volume 4 is a fully prepared, correctly typed, refresh-ready output in the Data Model — tables that are clean, granular, consistently structured, and free of the schema drift and refresh failures that would corrupt the relational model built on top of them. Volume 4 begins at that point and builds upward.

Analyst's Corner 1 — A Day in the Life of an Analyst Without Power Query

Picture the Monday morning of an analyst who manages a weekly sales performance report for a regional sales team. The report is read by five managers and two directors on Monday afternoon, and it must reflect Friday's closing numbers.

At 8:00 AM the analyst opens their email and finds that the weekly extract from the CRM has arrived. It is a CSV file with 4,200 rows and 22 columns. The analyst opens it in Excel alongside the report workbook and begins their standard preparation.

By 8:45 AM they have removed the seven columns that are not used in the report, filtered out the test accounts, trimmed the customer name column, and fixed the date format in the order date field — the CRM exports dates as text strings in DD/MM/YYYY format, and the report needs them as Excel date serials. This requires a column of `DATEVALUE(TEXT(MID(A2,4,2)&"/"&LEFT(A2,2)&"/"&RIGHT(A2,4),"MM/DD/YYYY"))` formulas, which the analyst has memorized from the first time they built it eight months ago. They paste the formula column over itself as values and delete the helper column.

At 9:10 AM the analyst discovers that one of the regional codes in the territory field appears as "SW" this week instead of the standard "South-West" that the report's SUMIFS formulas expect. They do not know if this is a one-time CRM entry error or a permanent code change. They use Find & Replace to normalize it to "South-West," make a note to follow up with the CRM administrator, and continue.

At 9:35 AM the analysis is complete. The report looks correct. The analyst sends it to the distribution list.

At 2:15 PM one of the directors replies: "The South-West figures look low compared to last week. Can you verify?" The analyst opens the previous week's extract — which they no longer have,

because it was not archived — and cannot directly compare the preparation logic applied to each week's data. They spend forty minutes reconstructing the preparation process from memory, applying it to a backup of the report file they happen to have, and confirming that the figures are correct. The apparent discrepancy was in the director's expectation, not the data. The forty minutes was consumed producing evidence for a conclusion that a documented, replayable pipeline would have provided in two minutes.

On Thursday, the CRM administrator responds to confirm that "SW" is now the permanent regional code. The analyst updates their SUBSTITUTE formula memorization for next Monday.

This is not an unusual analytical Monday. It is a very ordinary one. Power Query does not change the territory code. It does not prevent the director's question. What it does is replace every manually executed step in that sequence — the column deletion, the date conversion, the SUBSTITUTE, the paste-over, the reconstruction in the afternoon — with a documented, named, automatically refreshable pipeline that runs in its entirety when the analyst opens the workbook, produces the same output it produced last week from the same logic, and shows the director in thirty seconds exactly which transformation normalized the regional code and exactly what the raw value in the source was before it was normalized.

The Monday does not get easier because the work gets simpler. It gets easier because the work gets documented.

Analytical Toolkit — The ETL Workflow Self-Assessment Checklist

Use this checklist to evaluate the data preparation practices you currently apply to a workbook or report you manage regularly. Answer each question honestly and use the results to identify which stages of this book's discipline are most immediately applicable to your work.

Stage: Preparation Method

My current data preparation for this report involves manual copy-paste operations performed at each reporting cycle.

My current preparation involves formula chains (TRIM, SUBSTITUTE, TEXT conversions) applied to source data within the workbook.

My current preparation has no documented record of what was done — it is in my memory or a private set of notes.

My current preparation could not be re-executed identically by a colleague without my active guidance.

If you checked any of the above, you are operating in the manual or formula-based preparation paradigm. Every checked item represents a cost that a Power Query pipeline eliminates.

Stage: Connection and Source Management

My source file paths are hard-coded (absolute paths that break when the file moves or the workbook is opened on a different machine).

I do not know the authentication requirements for all data sources I use regularly.

I have experienced a “refresh failure” and had to spend time diagnosing its cause rather than proceeding immediately to the fix.

My workbook consumes data from more than one source, and I have no documented record of what those sources are and what each contributes.

Stage: Data Profiling

I begin transforming a new source dataset without first examining its column quality, null rates, or value distributions.

I have encountered errors in a downstream report caused by problems in source data that I did not detect before processing.

I do not have a standard first-pass checklist for evaluating an unfamiliar source before writing any transformation logic.

Stage: Transformation Documentation

My transformation steps are not named — they carry Power Query's default names (Changed Type, Removed Columns, etc.).

A colleague could not determine what my queries do by reading the Applied Steps pane without asking me.

My queries mix staging logic, transformation logic, and output delivery in a single unstructured query.

Stage: Production Readiness

My workbook does not have a documented refresh procedure.

I do not have an explicit error-handling strategy for refresh failures — I discover failures when the output is wrong.

I have no data quality checks built into my pipeline that would catch silent data loss between refresh cycles.

Interpreting Your Results

Every checked item in this list corresponds to a professional discipline addressed in a specific chapter of this book. Readers who check nothing face a different challenge: confirm that each unchecked item reflects genuine practice, not aspiration, and that a colleague would describe the same workbook in the same terms.

The purpose of this checklist is not to produce a score. It is to give the learning in this book a concrete address — a specific inadequacy in a specific real workflow — so that the techniques taught in

each chapter are not abstract but immediately applicable to work that exists and matters.

Chapter Summary

This chapter made the case for Power Query as a professional answer to a structural problem. The argument moved through four stages.

First, the hidden costs of manual data preparation were identified and named: the time consumed at each reporting cycle, the accuracy risk introduced by every manual repetition, and the institutional memory lost when preparation logic is stored in the analyst's memory rather than in a documented, executable artifact. The sixty-to-eighty percent figure was contextualized not as an inherent cost of analytical work but as the cost of a method — a method that is replaceable.

Second, the dominant pre-Power Query approaches to data preparation — the copy-paste workflow and the formula-based cleaning workflow — were examined for their structural failure modes: schema sensitivity, irreproducibility through the paste-over-delete pattern, and the repeatability gap that accumulates across cycles. Both approaches produce output, and sometimes correct output. Neither produces a pipeline.

Third, Power Query was introduced as an ETL engine implementing the Extract, Transform, Load paradigm inside Excel. The ETL framing was extended to the concept of a pipeline — a sequence of authored, named, ordered transformation steps that is reproducible, auditable, and refresh-automated. The distinction between spreadsheet thinking and pipeline thinking was drawn and identified as the cognitive shift this book facilitates.

Fourth, the six-stage discipline map for this book was introduced — Connect, Profile, Shape, Combine, Author, Deploy — and mapped to the Parts and chapters that teach each stage. The prerequi-

sites from Volumes 1 and 2 were confirmed, and the explicit scope boundary with Volume 4 was established.

The reader now has the conceptual framework within which every subsequent chapter's technical instruction will be placed. Chapter 2 enters the Power Query editor itself.

Key Takeaways

- Data preparation consuming sixty to eighty percent of analyst time is not an inherent property of analytical work. It is the cost of the manual and formula-based preparation method. A documented, automated pipeline addresses this cost structurally, not incrementally.
- The copy-paste-formula workflow fails in three systematic ways: schema sensitivity (breaking silently when the source changes structure), the paste-over-delete pattern (destroying the transformation audit trail), and the repeatability gap (producing slightly different logic across cycles because the logic lives in memory rather than in a document).
- Power Query implements the ETL paradigm — Extract, Transform, Load — as a three-stage pipeline inside Excel. The pipeline is authored once, replays automatically on refresh, and preserves a complete, inspectable record of every transformation decision.
- Power Query is the correct tool for the data preparation layer — the layer between raw source data and clean analytical input. Worksheet formulas belong in the analysis layer. VBA macros belong in operational automation. Each tool has an appropriate domain; professional practice uses each where it belongs.
- Power Query occupies the transformation layer of the layered workbook architecture introduced in Volume 1. It does

not displace that architecture; it fulfills one layer of it with an engineering-grade tool.

- A data preparation pipeline is distinguished from a data preparation workflow by three properties: it is reproducible (the same logic applies identically each refresh), it is auditable (every transformation decision is visible and inspectable), and it is refresh-automated (the full pipeline re-executes against updated source data without analyst intervention).
- The six-stage discipline map — Connect, Profile, Shape, Combine, Author, Deploy — organizes the full professional life-cycle of a Power Query pipeline. This book is structured to traverse all six stages completely.

Practice Quiz

Multiple Choice Questions (12)

Question 1. Research studies consistently find that professional data analysts spend the majority of their working time on data preparation rather than analysis. Which of the following best explains why this ratio is not an inevitable property of analytical work?

- A) Analysts have insufficient training in statistical methods, which slows the analysis phase.
- B) The ratio reflects the manual and undocumented methods used for preparation, not the inherent requirements of the work itself.
- C) Data sources are increasingly complex, and preparation time will always dominate regardless of the tools used.
- D) The analysis phase is inherently faster than preparation because formulas calculate automatically.

Question 2. An analyst prepares a monthly sales report by copying data from a CRM export, applying TRIM and SUBSTITUTE formulas, pasting the results over themselves as values, and deleting

the helper columns. Which of the following is the most significant structural failure of this approach?

- A) The use of TRIM formulas is computationally inefficient for large datasets.
- B) Copying data between workbooks introduces file size constraints that limit scalability.
- C) The paste-over-delete pattern destroys the transformation audit trail, leaving no inspectable record of what was done to the source data.
- D) Formulas cannot handle date conversion reliably across different Excel versions.

Question 3. In the context of Power Query and data engineering, what does the “Extract” stage of the ETL process accomplish?

- A) It applies transformation rules to clean and reshape data in preparation for analysis.
- B) It connects to the data source, navigates its structure, and retrieves data to the Power Query editor.
- C) It delivers the transformed data to its destination — a worksheet, the Data Model, or a connection-only query.
- D) It calculates derived fields and aggregations from the source data.

Question 4. An analyst is deciding whether to use Power Query, a worksheet formula, or a VBA macro to combine data from thirty monthly files stored in a shared folder into a single analytical table. Which choice is most appropriate, and why?

- A) A VBA macro, because combining files requires automating file navigation, which formulas cannot accomplish.
- B) A worksheet formula using INDIRECT and VLOOKUP, because these functions can reference external workbooks.

- C) Power Query, because combining same-structure files from a folder is precisely the data preparation task for which Power Query's ETL pipeline is the appropriate instrument.
- D) Power Query or a macro are equally appropriate; the choice is a matter of analyst preference.

Question 5. Which of the following correctly describes the position of Power Query within the layered workbook architecture introduced in Volume 1 of this series?

- A) Power Query replaces the analysis layer, performing calculations that would otherwise require worksheet formulas.
- B) Power Query occupies the transformation layer, connecting to external sources, applying the preparation pipeline, and delivering clean data to the analysis layer.
- C) Power Query occupies the presentation layer, formatting analytical outputs for stakeholder consumption.
- D) Power Query is an alternative to the layered workbook architecture, replacing it with a single-query approach.

Question 6. A data preparation pipeline is described as “auditable.” What does auditability mean in this context?

- A) The pipeline has been reviewed and approved by a compliance officer.
- B) The transformation logic of the pipeline can be inspected and verified by someone other than the analyst who built it.
- C) The pipeline was built using a version-controlled development environment.
- D) The pipeline produces output that matches the source data exactly, with no transformation applied.

Question 7. An analyst discovers that a value in a published report looks anomalous compared to the same value in the previous month's report. The analyst built both reports using a Power

Query pipeline. What capability of the pipeline most directly supports the investigation of this anomaly?

- A) Refresh automation, which allows the pipeline to be re-run against the previous month's data instantly.
- B) Transformation traceability through the Applied Steps pane, allowing the analyst to follow the value through the sequence of transformation steps that produced it.
- C) Connection-only query support, which isolates the anomaly to a specific source.
- D) The Data Model load destination, which retains historical versions of the output.

Question 8. In the six-stage ETL discipline map introduced in this chapter, which stage encompasses the operations of filtering rows, enforcing data types, splitting columns, handling null values, and constructing derived columns?

- A) Connect
- B) Profile
- C) Shape
- D) Author

Question 9. Which of the following best describes the difference between “spreadsheet thinking” and “pipeline thinking” as introduced in this chapter?

- A) Spreadsheet thinking uses formulas; pipeline thinking uses macros.
- B) Spreadsheet thinking focuses on the current state of the data and the next individual action; pipeline thinking focuses on designing a documented sequence of operations that will transform any conforming instance of the source correctly.
- C) Spreadsheet thinking is appropriate for large datasets; pipeline thinking is appropriate for small datasets.

D) Pipeline thinking requires SQL knowledge; spreadsheet thinking does not.

Question 10. Which of the following is explicitly identified in this chapter as outside the scope of this book and reserved for Volume 4 of this series?

A) Connecting Power Query to a REST API with authentication headers.

B) Reading and understanding the M expressions generated by Power Query ribbon operations.

C) The Excel Data Model as a relational analytical environment — table relationships, DAX calculations, and star schema design.

D) The Folder connector for combining many same-structure source files into one query.

Question 11. An analyst inherits a workbook that produces a weekly report. The preparation queries have default Power Query step names — “Changed Type,” “Removed Columns,” “Filtered Rows” — and the queries are organized in a single flat list with no grouping or naming convention. Which professional standard identified in this chapter is most clearly absent?

A) Refresh automation: the pipeline cannot be refreshed with default step names.

B) The auditability standard: the Applied Steps pane cannot be read as a plain-language description of the transformation decisions because the steps are not meaningfully named.

C) The reproducibility standard: default step names cause the pipeline to apply different logic on each refresh.

D) The Extract stage: connection-only queries require named steps to connect to sources correctly.

Question 12. A Power Query pipeline is described in this chapter as having three properties that distinguish it from a manual preparation workflow: reproducibility, auditability, and refresh automation. Which combination correctly defines all three?

- A) Reproducibility: produces approximately similar output on each run. Auditability: was built by a qualified analyst. Refresh automation: runs without the analyst opening Excel.
- B) Reproducibility: the same transformation logic applies identically to each instance of the source. Auditability: every transformation decision is visible and inspectable in the Applied Steps pane. Refresh automation: the full pipeline re-executes against updated source data without analyst intervention.
- C) Reproducibility: the pipeline has been tested on multiple datasets. Auditability: the pipeline was reviewed by a second analyst. Refresh automation: the pipeline triggers automatically when the source file changes.
- D) Reproducibility: the pipeline produces no errors. Auditability: the pipeline has been documented in a separate file. Refresh automation: the workbook opens at a scheduled time.

Answer Key with Rationale

Question 1 — Correct answer: B

The sixty-to-eighty percent figure describes the cost of the manual and undocumented preparation method, not a necessary feature of analytical work. Organizations that adopt systematic pipeline tooling consistently reduce the preparation-to-analysis time ratio. Option A confuses the preparation problem with the analysis problem. Option C is incorrect because the ratio is method-dependent, not source-complexity-dependent. Option D confuses formula recalculation (a feature of the analysis layer) with preparation efficiency.

Question 2 — Correct answer: C

The paste-over-delete pattern is the most significant structural failure because it is irreversible and destroys information. Once formula results are pasted as values and helper columns are deleted, there is no record of what was done to the data or why. Options A,

B, and D describe real concerns but none is as structurally critical as the loss of the audit trail, which affects every subsequent use of the output regardless of dataset size or Excel version.

Question 3 – Correct answer: B

The Extract stage is the connection and retrieval stage. It connects to the source, navigates its structure (selecting the relevant table, sheet, or endpoint), and makes the data available inside the Power Query editor. Option A describes the Transform stage. Option C describes the Load stage. Option D describes derived column construction, which is a Transform-stage activity.

Question 4 – Correct answer: C

Combining many same-structure files from a folder is a data preparation task – acquiring and combining source data into a unified analytical input – which is precisely the domain of the Power Query ETL pipeline. Option A is incorrect because this is not an operational automation task; it is a data preparation task. Option B is technically possible for small numbers of files but does not scale, is fragile on schema change, and destroys the audit trail. Option D is incorrect: this is not a matter of preference – tool selection should be based on the domain of the task.

Question 5 – Correct answer: B

Power Query occupies the transformation layer: it connects to the source (extraction), applies the pipeline (transformation), and delivers clean data to the analysis layer (load). It does not replace the analysis layer, the presentation layer, or the architecture itself. The layered workbook architecture from Volume 1 is preserved; Power Query fulfills one layer of it.

Question 6 – Correct answer: B

Auditability, as used in this context, means that the transformation logic can be inspected and verified by any analyst with access to the workbook, not just the one who built it. This is achieved through the Applied Steps pane and the M expressions it exposes. Options A and C describe process controls external to the pipeline.

Option D describes a pipeline that performs no transformation, which is not the analytical goal.

Question 7 — Correct answer: B

Transformation traceability through the Applied Steps pane is the capability that directly supports anomaly investigation. The analyst can step through each transformation applied to the data, identify where the value changed from the source to the output, and determine whether the change is intentional or erroneous. Option A is useful but supports comparison rather than diagnosis. Options C and D do not directly address transformation traceability.

Question 8 — Correct answer: C

The Shape stage encompasses row-level operations (filtering, deduplication, error handling), column-level operations (renaming, splitting, merging), type enforcement, null and error handling, and derived column construction. These are the operations that convert source data from its raw form into an analytically usable form. Connect is Stage 1. Profile is Stage 2. Author covers M language and custom expressions.

Question 9 — Correct answer: B

The chapter explicitly frames the distinction as a difference in where the analyst's attention is directed: spreadsheet thinking is focused on the current state and the next action; pipeline thinking is focused on designing a documented, replayable architecture. Options A and C introduce incorrect tool and dataset-size associations. Option D incorrectly requires SQL knowledge for pipeline thinking; this book's pipeline discipline uses M, not SQL, and introduces M from first principles.

Question 10 — Correct answer: C

The chapter explicitly reserves the Data Model as a relational analytical environment — table relationships, DAX, and star schema design — for Volume 4. REST API connections (option A) are cov-

ered in Part II. M expression reading (option B) is covered in Part V. The Folder connector (option D) is covered in Part II.

Question 11 — Correct answer: B

Default step names (“Changed Type,” “Removed Columns”) render the Applied Steps pane unreadable as an audit trail. Someone inheriting the pipeline cannot determine from the step names what transformation decisions were made, in what order, or why. This is a failure of the auditability standard. Default step names do not affect refresh execution (option A) or transformation logic (option C). Option D incorrectly links step naming to connection behavior.

Question 12 — Correct answer: B

Option B correctly defines all three properties as introduced in the chapter: reproducibility as identical logic applied each time, auditability as inspectability of transformation decisions in the Applied Steps pane, and refresh automation as the full pipeline re-executing automatically against updated source data. The other options contain inaccuracies in one or more definitions — approximate output (A), external review or automatic file-change triggering (C), and separate documentation files (D) are not the definitions this chapter uses.

Exercises / Applied Practice

Exercise 1.1 — Auditing One of Your Current Workbooks for Manual Preparation Cost (Comprehension)

Objective: Identify and quantify the manual data preparation cost in a workbook you currently manage or have access to, using the concepts and vocabulary introduced in this chapter.

Instructions:

Select a workbook you work with regularly that involves data preparation — ideally one that is refreshed on a recurring schedule (weekly, monthly, quarterly). If you do not have access to such a workbook,

use any Excel file that contains data that arrived from an external source and was cleaned or reshaped before use.

1. Estimate the time spent on data preparation for this workbook at each refresh cycle. Include all manual activities: downloading or locating source files, copying and pasting data, applying formula-based cleaning, fixing encoding or date format issues, deleting helper columns, and any other preparatory steps performed before the analytical work begins. Record this estimate in minutes.
2. Using the three cost categories introduced in Section 1.1.1 — Time, Accuracy, and Institutional Memory — evaluate your workbook against each:
 - Time: Multiply your per-cycle estimate by the number of cycles per year. Record the annual manual preparation time in hours.
 - Accuracy: Identify at least two specific points in your preparation workflow where a different decision in this cycle versus last cycle would produce a different output without any indication that the outputs are inconsistent.
 - Institutional Memory: Determine whether a colleague could reproduce your preparation process, correctly and completely, with no guidance from you. Write one paragraph describing what they would need to know that is not documented anywhere in the workbook.
1. Complete the ETL Workflow Self-Assessment Checklist from the Analytical Toolkit section of this chapter, applied to the specific workbook you have selected.

Deliverable: A written summary (approximately one page) covering your per-cycle and annual time estimates, your accuracy risk assessment at two specific workflow points, your institutional memory evaluation, and your checklist results. This document

is your baseline audit — it establishes what Chapter 16’s production readiness standards will measure your work against when this book is complete.

Exercise 1.2 — Mapping the Preparation Steps of an Existing Report Onto the Six ETL Stages (*Application*)

Objective: Translate the vocabulary of a current data preparation workflow into the six-stage discipline map introduced in Section 1.4.1, identifying which stages are present, which are absent, and which are performed without discipline.

Instructions:

Using the same workbook selected in Exercise 1.1, or a different workbook if that one involved no meaningful preparation:

1. Write out every distinct step of your current preparation process as a numbered list, in the order you perform them. Be specific: “Open CRM export,” “Delete columns B, D, F through H,” “Apply TRIM formula to column C,” “Paste column C as values and delete helper column,” and so on. Do not summarize or group steps unless they are genuinely a single indivisible action.
2. Assign each step in your list to one of the six ETL stages: Connect, Profile, Shape, Combine, Author, Deploy. For steps that do not map cleanly to any stage — because they are workarounds, ad hoc fixes, or activities that a pipeline would not require — label them as “Manual overhead.”
3. For each of the six stages, write one sentence describing the degree to which that stage is currently performed with discipline in your workflow. Use the following scale:
 - *Absent:* This stage is not performed at all.
 - *Informal:* The stage is performed but without documentation, consistency standards, or repeatability.

- *Partial*: Some elements of the stage are performed with discipline; others are not.
 - *Disciplined*: The stage is performed consistently, documentably, and reproducibly.
1. Identify the one stage whose current absence or informality poses the greatest risk to the quality and reliability of the output. Write two to three sentences justifying your selection with specific reference to the failure modes described in this chapter.

Deliverable: The numbered step list with stage assignments, the six stage-level discipline assessments, and the risk identification paragraph. This map will serve as a reference point throughout the book as each stage's professional standard is introduced.

Exercise 1.3 — Defending the Case for Power Query to a Skeptical Colleague (*Judgment*)

Objective: Construct a professional argument for adopting Power Query as the data preparation standard for a specific analytical workflow, addressed to an audience with valid reasons to be skeptical.

Context: A colleague who manages the quarterly financial variance report for your organization has been doing so for three years using a manual copy-paste and formula-based workflow. The report is accurate, it is delivered on time, and the colleague is proficient and experienced. They are skeptical of the argument that they should change their method. Their objections include: “My process works fine. Learning a new tool takes time I don’t have. The report is only quarterly — it’s not worth automating.” You have been asked to respond to these objections substantively.

Instructions:

Write a professional communication — a short internal memo or structured talking points, approximately 400 to 600 words — that makes the case for Power Query adoption in the specific context of this quarterly report. Your argument must:

1. Acknowledge the validity of the colleague's current approach: the report is accurate and delivered on time. Do not characterize the current method as failing — it is not failing. Identify what it costs.
2. Address each of the three stated objections directly, using the arguments from this chapter:
 - “My process works fine.” — Address this with the distinction between a process that produces output and a process that is documented, auditable, and resilient to personnel change.
 - “Learning a new tool takes time I don't have.” — Address this with the concept of one-time pipeline authoring cost versus recurring manual preparation cost, using the annual time calculation concept from Exercise 1.1.
 - “The report is only quarterly — it's not worth automating.” — Address this with the institutional memory and accuracy risk arguments, scaled specifically to a quarterly high-stakes deliverable such as a financial variance report.
1. Close with one concrete, specific proposal: a single first step your colleague could take with Power Query that would produce a visible improvement in one specific aspect of their current workflow — not a full pipeline, but a meaningful beginning.

Deliverable: The professional communication in full, written as if it would actually be sent. This is a judgment exercise: there is no single correct argument, but there are better and worse arguments by the standards introduced in this chapter. A strong submission applies the chapter's conceptual vocabulary accurately, responds to the actual objections rather than generic versions of

them, and proposes a first step that is realistic and appealing to a skeptical but reasonable professional.

Reflection / Discussion Questions

1. This chapter argues that the sixty-to-eighty percent preparation time ratio is a property of the method, not the work. In your professional context, does this claim hold? Can you identify specific preparation activities in your current work that would genuinely be eliminated — not merely accelerated — by a well-built Power Query pipeline? Where does the argument break down, if anywhere?
2. The Analyst's Corner described a scenario in which a four-hundred-minute re-investigation was triggered by a director's question that a documented pipeline would have answered in two minutes. In your experience, how much of the time analysts spend on report defense, output verification, and anomaly investigation is a consequence of undocumented preparation — that is, time spent answering “what did we do?” rather than “what does this mean?”
3. The chapter distinguishes between spreadsheet thinking and pipeline thinking as different orientations toward the same work. Where on the spectrum between those orientations would you honestly place your current practice? What would have to change — about how you approach a new dataset, how you name what you do, how you think about the recipient of your work — for the pipeline-thinking orientation to become your default?
4. The case for Power Query in this chapter is presented almost entirely in terms of professional standards: auditability, reproducibility, institutional memory, handoff quality. Are these arguments persuasive to decision-makers in your organization, or does the argument that will actually change practice need to be made differently — in terms of time, er-

ror rates, or risk to the organization? What version of the case would you make to your own manager?

5. The scope boundary with Volume 4 was explicit in Section 1.4.3: the Data Model as a relational environment, DAX, and calendar table design are left entirely to the next volume. Given that, what do you expect Power Query to be able to accomplish that you currently handle with formula arrays or manual aggregation? What specific analytical capabilities are you expecting to gain, and from which stage of the six-stage map?

What's Next

Chapter 1 made the case. Chapter 2 enters the environment where the pipeline is built.

The Power Query editor is not a simplified interface bolted onto Excel's ribbon. It is a full analytical environment with its own layout, its own navigation logic, and its own set of professional configuration choices that determine how reliably the pipeline behaves in a shared or production setting. Chapter 2 orients the reader inside that environment with the deliberateness that a professional workstation deserves: understanding the editor layout as an analytical instrument, distinguishing the three load destinations and their respective architectural roles, reading the formula bar and Advanced Editor without fear, and configuring the editor's options, privacy settings, and locale parameters for the work that follows.

The editor is where the pipeline lives. Chapter 2 is the reader's first walk through it — with the lights on and everything labeled.

CHAPTER 2

The Power Query Editor as an Analytical Environment

Chapter Introduction

Chapter 1 established why Power Query is the correct instrument for the data preparation layer and what kind of thinking it requires. This chapter establishes where that thinking takes place: the Power Query editor itself.

The editor is not a dialog box or a wizard. It is a full-featured analytical environment with its own interface regions, its own ribbon, its own configuration options, and its own logic for how the work of data preparation is organized and represented. An analyst who approaches the editor without understanding its architecture will use it reactively — clicking through operations, accepting defaults, building queries that work once but are fragile — and will miss the professional discipline the environment is designed to support.

This chapter orients the reader inside the editor before any transformation is applied. It proceeds in three movements. First, it establishes how the editor is accessed and how its major regions relate to each other as a coherent analytical workspace. Second, it dissects the anatomy of a single query — the components that every query possesses, the three destinations to which a query can deliver its output, and the M language surface that the editor exposes from the first moment of use. Third, it covers the professional configuration of the editor: the options, privacy settings, locale parameters, query naming conventions, and team-

facing setup decisions that separate a workbook built for production from one built in a hurry.

By the end of this chapter, the reader will be able to navigate the editor with intention, locate any query component without searching, choose a load destination with architectural awareness, and configure a new workbook's editor environment to the professional standard this book maintains throughout.

Chapter Objectives

Upon completing this chapter, the reader will be able to:

- Access the Power Query editor through the correct ribbon entry points and identify the appropriate path for each source type
- Identify and describe the purpose of each major region of the editor interface: the Queries pane, the Preview pane, the Applied Steps pane, the formula bar, and the ribbon
- Distinguish the four ribbon tabs — Home, Transform, Add Column, View — by their organizational logic and primary function domain
- Describe the three components of every query — Source, Applied Steps, and Final Output — and explain the role each plays in the pipeline
- Choose the correct load destination — worksheet table, Data Model, or connection-only — based on the query's role in the broader workbook architecture
- Explain what a connection-only query is, why it exists, and how it functions as the intermediate staging mechanism in a multi-stage pipeline
- Open the formula bar and Advanced Editor, read a basic generated M expression, and identify its structural components without yet writing M code

- Configure Query Options settings — background refresh behavior, privacy levels, locale — appropriate to a professional working environment
- Rename, describe, and group queries according to the naming conventions this book maintains throughout

2.1 Launching and Navigating the Editor

2.1.1 Get & Transform Data — Ribbon Entry Points and the Get Data Menu

The Power Query editor does not open from a standalone application or a separate menu. It opens from within Excel through the Data tab on the main Excel ribbon, in the group labeled Get & Transform Data. Understanding the entry points in this group is the correct starting point because the path taken into the editor determines the initial structure of the query that opens — and different paths are appropriate for different source types.

The primary entry point for most analytical work is the Get Data button, which opens a hierarchical menu organized by source category. The top-level categories include: From File (covering Excel workbooks, CSV and text files, XML, JSON, PDF, and folders of files), From Database (covering SQL Server, Access, Oracle, MySQL, PostgreSQL, and other relational databases), From Azure and From Online Services (covering SharePoint, OneDrive, and cloud platforms), From Other Sources (covering web pages, OData feeds, blank queries, and the ODBC and OLE DB connectors for legacy systems), and Combine Queries (for append and merge operations on existing queries). Navigating Get Data by category — rather than searching for a connector by name — reinforces the habit of thinking about the source type before the connection, which is the first architectural decision of every pipeline.

Several secondary entry points are available for common operations. From Table/Range connects directly to an Excel Table or

named range that is already present in the active workbook. From Text/CSV is a single-click entry point for the most common flat-file source type. Recent Sources lists the last several data sources accessed, useful during iterative pipeline development when the analyst returns to the same source multiple times. Existing Connections opens a panel showing all current query connections in the workbook, which is the correct starting point when the goal is to inspect or modify an existing pipeline rather than create a new one.

Each of these entry points opens the Power Query editor with a query already partially built — one that reflects the source type selected and the navigation choices made during the connection dialog. The analyst does not begin in a blank environment; they begin with a Source step already present in the Applied Steps pane, representing the connection just made.

2.1.2 The Queries Pane, Preview Pane, and the Editor Layout

The Power Query editor occupies its own window, separate from the main Excel application. It has a specific layout that is consistent across source types and operations, and understanding that layout as a coherent workspace — rather than as a collection of panels to navigate by trial and error — is the foundation of working in the editor with professional efficiency.

The **Queries pane** occupies the left side of the editor window. It lists every query in the current workbook, organized into groups if the analyst has established a group structure, or as a flat list if not. The active query — the one currently displayed in the Preview pane and the Applied Steps pane — is highlighted. Clicking any query in the Queries pane switches the editor's focus to that query: its preview, its Applied Steps, and its formula bar context all update to reflect the selected query. Right-clicking any query in the pane provides the full query management menu: rename, delete, duplicate, reference, move to group, enable or disable load, and access

THE POWER QUERY EDITOR LAYOUT: ANNOTATED INTERFACE MAP

Chapter 2 — Data Transformation with Power Query | Each region has a specific analytical function — understanding their relationship before applying any transformation is the foundation of working in the editor with professional intention

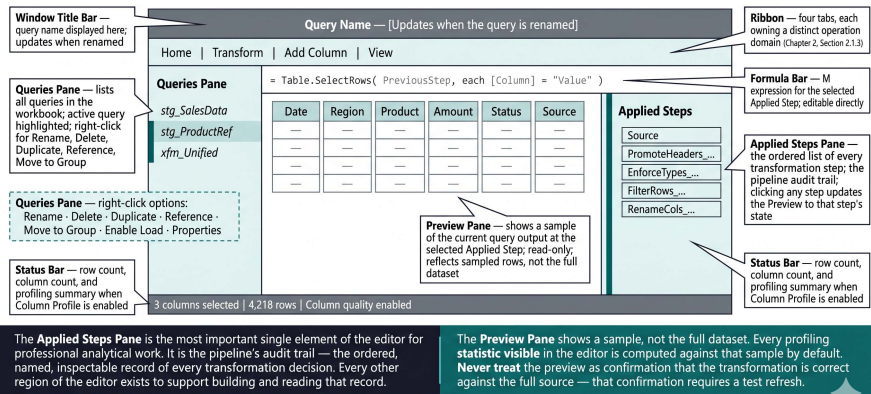


Figure 2.1 The Power Query editor interface, annotated by region. Each region has a specific analytical function; understanding their relationship before applying any transformation is the foundation of working in the editor with professional intention.

Figure 2.1. The Power Query editor interface, annotated by region. Each region has a specific analytical function; understanding their relationship before applying any transformation is the foundation of working in the editor with professional intention.

to query properties. The Queries pane is the workbook-level navigation surface; the rest of the editor provides query-level detail.

The **Preview pane** occupies the center and majority of the editor window. It displays a tabular sample of the current query's output at the currently selected Applied Step — which means it shows what the data looks like after all steps up to and including the selected step have been applied. Columns are displayed with their names in the header row and their data types indicated by an icon to the left of each column name. A number icon indicates a numeric type; a calendar icon indicates a date type; an “ABC” icon indicates a text type; a true/false icon indicates a logical type. Clicking a column header selects the column and makes column-level operations available in the ribbon. Clicking a cell in the preview does not select the cell for editing — the preview is read-only; clicking a cell selects the row for row-level operations, or may reveal additional information in the status bar. The important limitation of the preview — that it shows a sample rather than the full dataset — is addressed in full in Analyst's Corner 2.

The **Applied Steps pane** occupies the right side of the editor window. It is, as Chapter 3 will establish in detail, the most important single element of the editor for professional analytical work: the ordered, named, auditable list of every transformation step in the pipeline. At this stage, it is enough to observe its position and function: clicking any step in the Applied Steps pane updates the Preview pane to show the data as it appears after that step is applied, allowing the analyst to inspect the state of the data at any point in the transformation sequence. Adding a new transformation appends a new step to the bottom of the list. The gear icon to the right of certain step names opens the configuration dialog for that step — allowing its parameters to be changed without deleting and rebuilding the step. The × icon deletes the step. Both the gear and × icons appear on hover.

The **formula bar** runs below the ribbon, spanning the full width of the editor above the Preview pane. It displays the M expression that corresponds to the currently selected Applied Step. For steps applied through ribbon operations — which generate M automatically — the formula bar shows the generated M for that operation. For steps that the analyst has written directly in M, the formula bar shows exactly what was written. The formula bar is where the analyst can read, edit, or override the M expression for any step without opening the Advanced Editor. Its use is introduced in Section 2.2.4.

2.1.3 The Ribbon — Home, Transform, Add Column, and View Tabs

The Power Query ribbon contains four tabs, each organized around a distinct domain of operations. Understanding the organizational logic of each tab — not memorizing its contents, but grasping the principle by which it is populated — allows the analyst to find the operation they need without scanning every button.

The **Home tab** is the operational hub of the editor. It contains the operations the analyst uses most frequently across all pipeline

THE FOUR POWER QUERY RIBBON TABS AND THEIR PRIMARY OPERATION DOMAINS			
Chapter 2 — Data Transformation with Power Query Each tab is organized around a distinct class of analytical activity — understanding the organizational logic allows navigation by intent rather than by memorized button location			
HOME	TRANSFORM	ADD COLUMN	VIEW
■ Output — Close & Load — load destination choices ■ Queries — New source, recent sources, data source settings ■ Manage Columns — Choose Columns, Remove Columns ■ Reduce Rows — Keep Rows, Remove Rows, Filter Rows, Sort ■ Transform — Data type assignment and type-change shortcuts ■ Combine — Append Queries, Merge Queries Workbook-level operations, the most-used cleaning operations, and pipeline output configuration	■ Table — Transpose, Reverse Rows, Count Rows ■ Any Column — Data type, Replace Values, Fill Up/Down, Detect Type ■ Text Column — Split, Format, Trim, Clean, Extract ■ Number Column — Statistics, Standard, Rounding, Scientific ■ Date & Time — Date, Time, Duration part extraction ■ Structured Column — Expand, Aggregate, Extract Values (Lists/Records/Tables) Column-type-aware transformations on the selected column or table — modifies in place	■ General — Custom Column (M expression) — Column From Examples (inferred logic) — Invoke Custom Function ■ Text / Number / Date — Same operations as Transform tab — but creates a new column instead of modifying △ Add Column preserves the source column — Transform does not. Use Add Column when the original is needed downstream Non-destructive column derivation — all operations add a new column; the source column is preserved	■ Data Preview — Formula Bar toggle — Column Quality, Column Distribution, Column Profile — Always Allow (extends profiling beyond 1,000 rows) ■ Query — Advanced Editor (full M code) — Query Dependencies (visual pipeline map) ■ Parameters — Manage Parameters — link to Part V ■ Layout — Query Settings visibility Editor display configuration, profiling tools, and the Advanced Editor access point
The critical distinction between Transform and Add Column is not cosmetic. Transform replaces the selected column with the transformed result. Add Column creates a new column and leaves the source unchanged. When the original column is needed downstream — for verification, for further derivation, or as a check — Add Column is the correct choice.		The View tab's Column Quality, Column Distribution, and Column Profile options are the profiling tools covered in full in Part III. Enabling them on a new source before applying any transformation is the first professional discipline of data preparation. The Advanced Editor is the full M authoring surface developed in Part V.	

Figure 2.2 The Power Query ribbon tabs and their primary operation domains. Each tab is organized around a distinct class of analytical activity. Understanding the organizational logic of each tab allows the analyst to navigate by intent rather than by memorized button location.

Figure 2.2. The four Power Query ribbon tabs and their primary operation domains. Each tab is organized around a distinct class of analytical activity. Understanding the organizational logic of each tab allows the analyst to navigate by intent rather than by memorized button location.

types: closing the editor and delivering the query output (Close & Load and its variants), adding new queries or modifying existing source connections, the most common column management operations (Choose Columns, Remove Columns), the most common row operations (Keep Rows, Remove Rows, Filter Rows), and the data type assignment controls. The Home tab also provides short-cut access to Append and Merge — the Combine stage operations covered in Part IV. As a general rule, if an operation applies to the entire table or involves the query's relationship to the workbook, it is on the Home tab.

The **Transform tab** applies operations to the selected column or to the table as a whole, modifying data that already exists. The tab is context-sensitive: the operations available depend on the data type of the selected column. With a text column selected, text-specific operations (Split Column, Format, Trim, Clean, Extract) are enabled. With a numeric column selected, arithmetic and statistical operations are enabled. With a date column selected, date-part extraction operations are enabled. The critical

distinction between the Transform tab and the Add Column tab — one that new Power Query users frequently miss — is that Transform operations modify the selected column in place. The original column is replaced by the transformed version. When that is not the desired behavior, the Add Column tab is the correct choice.

The **Add Column tab** mirrors many of the Transform tab's operations but produces a new column rather than modifying the existing one. The result of an Add Column operation is a new column appended to the right of the table, containing the derived values, while the source column remains unchanged. This non-destructive characteristic makes the Add Column tab the correct choice whenever the original column may be needed downstream — either as a check against the derived value or as an input to a further derivation. The Add Column tab also contains Custom Column (for writing M expressions directly), Column From Examples (for inferring transformation logic from sample values), and Invoke Custom Function (for applying a parameterized function query to each row, introduced in Part V).

The **View tab** controls the editor's display configuration and provides access to the two most important diagnostic surfaces: the profiling tools and the Advanced Editor. The formula bar toggle enables or disables the formula bar. The Column Quality, Column Distribution, and Column Profile options toggle the profiling indicators that appear below each column header in the preview — the profiling tools that Part III covers as a professional discipline. The Always Allow option overrides the editor's default behavior of restricting profiling to the sampled preview rows, enabling profiling against the full dataset. The Advanced Editor button opens the full M expression for the active query as an editable code window — the surface used for authoring complex M logic, introduced in Section 2.2.4 and developed in depth in Part V.

2.2 The Anatomy of a Query

2.2.1 Source, Applied Steps, and the Final Output — What Each Component Does

Every Power Query query — regardless of its source type, its transformation complexity, or its load destination — is composed of the same three structural components: a Source, a sequence of Applied Steps, and a Final Output. Understanding these components as a coherent architecture, rather than as interface elements encountered during use, is the foundation of reading and working with any query, including ones built by others.

The **Source** is the first Applied Step in every query. It contains the M expression that establishes the connection to the data origin: the file path and format parameters for a file connection, the server name and database name for a database connection, the URL and query parameters for a web or API connection. The Source step is also where authentication credentials are referenced, where privacy level context is established, and where the connection is defined in a form that Power Query can re-execute on every refresh. Everything the pipeline does depends on the Source step executing without error. When a pipeline fails to refresh, the Source step is the first place to examine.

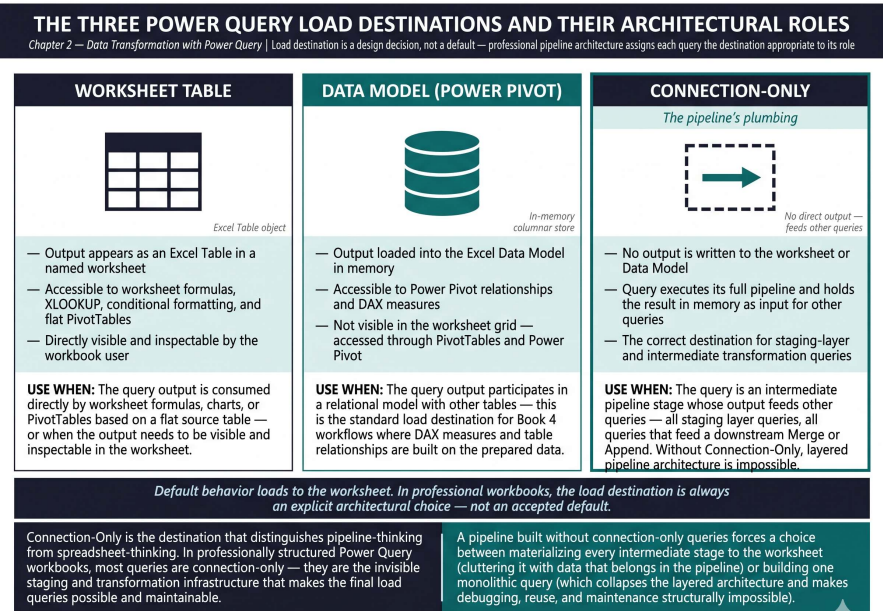
The **Applied Steps** are the sequence of transformation operations that follow the Source. Each step takes the output of the previous step as its input, applies one transformation, and passes the result to the next step. The steps are ordered: their sequence is the pipeline's logic. Removing a step from the middle of the sequence causes every subsequent step to recalculate based on what the data looked like before the removed step was applied — which may break downstream steps that depended on a column or structure that the removed step introduced. Reordering steps is possible but carries the same dependency risk: a step that references a column must come after the step that created that column. Managing these dependencies is part of the professional discipline of step management, addressed in Chapter 3.

The **Final Output** is the result of applying all Applied Steps in sequence to the data retrieved by the Source. It is what the query delivers to its load destination. The Final Output is not stored inside the query; it is computed each time the query is evaluated — either when the editor is open and the preview is rendering, or when a refresh is triggered from the closed-editor workbook context. The Final Output is a structured table with a fixed schema: a specified set of columns, each with a specified data type, in a specified order. The stability and correctness of this schema is what allows downstream consumers — worksheet formulas, PivotTables, the Data Model — to depend on the query's output without breaking when the source data changes in content but not in structure.

This three-component architecture is the conceptual model for reading any query. When the analyst opens an inherited workbook and examines an unfamiliar query, the investigation follows the same three steps: what is the Source (where does the data come from and how is it connected?), what are the Applied Steps (what transformations were applied, in what order, and with what decision logic?), and what is the Final Output (what schema does the query produce, and is it the schema that downstream analysis expects?). These three questions answer, between them, almost everything an analyst needs to know about a query before modifying it.

2.2.2 The Three Load Destinations — Worksheet, Connection-Only, and Data Model

A Power Query query is not complete until its load destination has been chosen. The load destination is not a cosmetic setting — it is an architectural decision that determines what role the query plays in the workbook, what downstream capabilities it supports, and whether its output is visible in the worksheet or operates as invisible infrastructure. There are three choices, and professional analytical work has a clear logic for choosing between them.



Load to worksheet delivers the query's Final Output as an Excel Table in a specified worksheet. The table is formatted as a standard Excel Table object — the same object introduced in Volume 1 — with a header row, data rows, and the structured reference syntax that worksheet formulas use. When the query is refreshed, the table updates in place: rows are added, removed, or changed to reflect the updated source. Downstream formulas that reference the table by structured reference update automatically. This destination is appropriate when the query's output is consumed by formulas in the same workbook, used as the data source for a chart or PivotTable based on a flat table, or needs to be visible and inspectable by the workbook user. It is not appropriate for queries that feed other queries in a multi-stage pipeline, because materializing intermediate stages to the worksheet adds unnecessary data to the workbook and breaks the layered architecture that separates staging, transformation, and output concerns.

Load to Data Model delivers the query's Final Output to the Excel Data Model — the in-memory columnar store that Power Pivot operates on and that DAX measures are written against. Queries loaded to the Data Model do not appear as visible tables in any worksheet; they are accessible only through Power Pivot's Diagram View, through PivotTables that use the Data Model as their source, and through DAX measure expressions. This destination is the standard choice for queries whose output participates in a relational model — which is the architecture Volume 4 introduces. For pipelines built in this book, loading to the Data Model means the prepared, clean output will be ready to receive the relationships and measures that Volume 4 will add.

Connection-only is the load destination that most distinguishes a pipeline-thinking analyst from a spreadsheet-thinking one. A connection-only query executes its full pipeline — Source through every Applied Step — but delivers its Final Output nowhere visible. The output exists in memory and is available as the input to other queries: specifically, any query that references a connection-only query as its source will receive that query's output as its starting

point. Connection-only is the correct destination for every staging layer query — every query whose purpose is to clean, type-enforce, and standardize a source so that transformation queries can consume it — and for every intermediate transformation that feeds a downstream merge or append operation. Without connection-only queries, multi-stage pipelines cannot be structured into layers; every stage must either materialize to the worksheet (cluttering it with intermediate data) or be merged into a single monolithic query (which defeats the purpose of layered architecture).

2.2.3 Why the Connection-Only Query Is the Plumbing of Multi-Stage Pipelines

The connection-only query deserves more sustained attention than its modest interface footprint suggests. It is invisible in the workbook — no table appears in any worksheet, no grid shows its output — and this invisibility makes it easy to overlook or underuse. In professionally structured Power Query workbooks, connection-only queries are not the exception; they are the majority of the queries in the pipeline.

Consider a workbook that combines sales transactions from three regional sources, joins them against a product reference table, applies a set of business rules, and delivers a unified analytical table to the Data Model. In a naively built workbook, this might be one large query that connects to each source, applies each transformation, and combines everything in a single long Applied Steps sequence. That query works. It is also impossible to debug efficiently — when it fails, the failure can be anywhere in the sequence. It is impossible to reuse — if another report needs the product reference table without the transaction data, the logic must be duplicated. It is impossible to maintain — a change to the business rules requires editing a query that also contains connection logic and source-specific cleaning, creating the risk of inadvertently altering something that was not supposed to change.

A pipeline built with connection-only queries structures the same work differently. Three staging queries — one per regional source — connect to their respective sources, apply source-specific cleaning and type enforcement, and load connection-only. A fourth staging query connects to the product reference table and does the same. A fifth query — a transformation query — appends the three regional staging queries into a single unified transactions table, connection-only. A sixth query merges the unified transactions against the product reference, applies the business rules, and loads to the Data Model. The final output is produced by a clean sequence of well-bounded, independently inspectable, individually refreshable queries. When something fails, the failure is isolatable to the query where it occurred. When a business rule changes, exactly one query is modified. When another report needs the product reference table, it references the existing product staging query rather than duplicating the connection logic.

This structure is the pipeline architecture this book teaches. Connection-only queries are its infrastructure — the plumbing that makes layered, maintainable, auditable pipelines possible.

2.2.4 The Formula Bar and the Advanced Editor — Reading Generated M from the Start

Every operation the analyst performs through the Power Query ribbon produces a corresponding M expression. That expression is the actual instruction that Power Query executes when the query is evaluated. The ribbon is a graphical interface for authoring M; the M expression is what is actually run. Understanding this relationship — that ribbon operations and M expressions are two representations of the same thing — is the conceptual foundation for Part V's instruction on the M language.

The **formula bar** is the analyst's first and most accessible window into the M language. When a step is selected in the Applied Steps pane, the formula bar displays the M expression for that step. For a step that filters rows to keep only records where the Region col-

umn equals “North,” the formula bar displays an expression that calls `Table.SelectRows` with the table and a condition function as its arguments. The analyst does not need to understand M syntax to read this expression at a useful level — the function name, the column reference, and the comparison value are all visible in plain form. Reading the formula bar expression for each step as transformations are applied builds M reading fluency before writing a single line of M code.

The **Advanced Editor** opens the complete M expression for the active query as a single, editable code block. It is accessible from the View tab on the ribbon. Where the formula bar shows the M expression for one step at a time, the Advanced Editor shows the entire query as a `let...in` expression: a sequence of named intermediate calculations (`let`) whose final value is returned as the query’s output (`in`). Each Applied Step corresponds to one named variable in the `let` block. The name of the variable is the name of the Applied Step — which is one reason why step names matter: they are variable names in the underlying M code, and a step named “Changed Type” produces the variable name `#"Changed Type"` in the M, while a step named “Enforce_Source_Types” produces the cleaner variable name `Enforce_Source_Types`.

Reading the Advanced Editor without yet writing M is a productive and achievable first step. The analyst can identify each step, trace the pipeline chain from Source through to the final `in` expression, and confirm that the step names visible in the Applied Steps pane correspond exactly to the variable names in the code. This reading practice is introduced early in this book — before any M writing instruction — because fluency with M reading is the prerequisite for M writing, and that reading fluency is most efficiently built through exposure to generated code rather than through abstract syntax instruction.

Part V returns to the Advanced Editor and the formula bar with writing instruction. The analyst who has been reading generated M from the beginning of their Power Query practice will find that

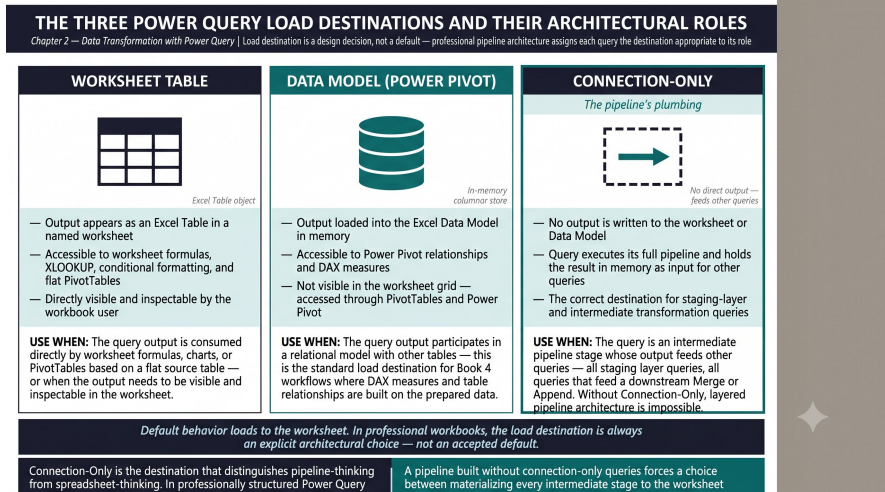


Figure 2.4. The M language as seen in the formula bar (single-step view) and the Advanced Editor (full query view). Every ribbon operation produces an M expression. Reading these expressions from the first use of the editor builds the M literacy that Part V develops into a writing skill.

Part V's instruction requires far less effort than it would if the Advanced Editor had been avoided until that point.

2.3 Configuring the Editor for Professional Work

2.3.1 Query Options, Privacy Settings, and Locale Configuration

The Power Query editor has a set of configuration options — accessed through File → Options and settings → Query Options in the editor, or through the Options dialog in the Excel Data tab — that determine how the editor behaves across all queries in the workbook. Two of these option categories have direct consequences for pipeline reliability in professional analytical environments: privacy levels and locale.

Privacy levels are a data governance mechanism built into Power Query's evaluation engine. When a workbook combines data from

sources of different trust categories — a public web source and a private internal database, for instance — Power Query's Formula Firewall prevents the evaluation of expressions that could inadvertently expose private data to a public source. The Formula Firewall error ("Formula.Firewall: Query references other queries or steps, so it may not directly access a data source") is one of the most common and confusing errors new Power Query users encounter. It is not a bug; it is the security mechanism working as designed. It appears when queries at different privacy levels are combined without explicit privacy level configuration.

Privacy levels are assigned per data source through the Data Source Settings dialog. The three levels are: Public (the source contains no sensitive data and its contents can be shared with any other query), Organizational (the source contains data appropriate for internal use but not for external exposure), and Private (the source contains sensitive data that must not be combined with any other source unless explicitly permitted). Assigning the correct privacy level to each source is the resolution to Formula Firewall errors — not the workaround of setting all sources to Public or disabling the privacy check entirely, which is the approach most tutorials recommend and which carries genuine data exposure risk in environments where sources of mixed sensitivity are used.

The correct approach is to assign the appropriate level to each source based on what it actually contains, and to understand that Power Query will then allow combination only where the levels are compatible. Chapter 5 addresses privacy level configuration in the context of specific connector types where Formula Firewall errors are most likely to arise.

Locale configuration determines how Power Query interprets values that are locale-sensitive: dates, numbers with decimal separators, and currency formats. A CSV file exported from a system configured for the French locale may represent dates as DD/M-M/YYYY and numbers with a comma as the decimal separator. An Excel workbook on a UK-configured machine will interpret those values differently than the same workbook on a US-configured

machine — which is the root cause of the “date imported as text” problem that affects most analysts who work with international data sources at some point in their careers.

The locale setting for a specific data source can be configured during the import step (the “Using locale” option in the data type assignment dialog) or globally for the workbook through Query Options. Setting the locale at the query level — at the point where types are enforced — is the more robust approach: it makes the type enforcement locale-explicit, which means the query will interpret dates and numbers from that source correctly regardless of the locale of the machine running the workbook. This is the correct configuration for any workbook that will be shared across machines with different regional settings.

2.3.2 Query Properties, Renaming, and Grouping Queries for Navigation

Every query in the Queries pane has a name and a set of properties. Query naming and organization are not cosmetic decisions — they are the workbook-level expression of the pipeline-thinking discipline introduced in Chapter 1. In a workbook with four queries, poor naming is an inconvenience. In a workbook with twenty-five queries across three layers of a complex pipeline, poor naming and no organizational structure make the workbook’s architecture invisible to anyone who opens it.

Renaming queries is accomplished by double-clicking the query name in the Queries pane or through the Query Properties dialog. Query names should follow a naming convention that signals the query’s layer and role in the pipeline. This book maintains the following convention throughout:

- `stg_` prefix for staging layer queries that connect to a source and apply minimal cleaning. Example: `stg_SalesTransactions`, `stg_ProductReference`.

- `xfm_` prefix for transformation layer queries that apply business logic, combine staged sources, or produce a shaped analytical output. Example: `xfm_UnifiedSales`, `xfm_SalesWithProduct`.
- `lod_` prefix for load layer queries that deliver the final output to a worksheet or the Data Model. Example: `lod_SalesFact`, `lod_ProductDim`.
- Parameter queries use a `prm_` prefix. Example: `prm_ReportingYear`, `prm_FilePath`.

This convention is not the only acceptable convention — teams should adopt and document their own standard — but it is consistent, readable, and used without deviation throughout this book. Appendix C provides the complete naming standard as a reference document.

Query properties — accessible by right-clicking a query and selecting Properties, or through the Query Settings panel visible when the editor is open — include the query name and a description field. The description field is where inline query-level documentation lives: a one-to-three-sentence description of what the query does, what source it connects to, and any special configuration it relies on. Well-populated description fields mean that the Queries pane communicates the pipeline’s architecture at a glance. Empty description fields mean that the next analyst to open the workbook must read every query’s Applied Steps to determine its role.

Grouping queries organizes the Queries pane into a navigable structure when the query count grows beyond what a flat list manages clearly. Groups are created by right-clicking in the Queries pane and selecting New Group, or by right-clicking a query and selecting Move to Group. This book groups queries by their layer: a group named “Staging” contains all `stg_` queries, a group named “Transformation” contains all `xfm_` queries, and a group named “Load” contains all `lod_` queries. Parameter queries occupy their own group named “Parameters.” Appendix C documents this structure as part of the full architecture standard.

2.3.3 Building a Consistent Editor Setup Across a Team

A shared analytical workbook is not used by one analyst on one machine. It is opened, refreshed, and potentially modified by multiple colleagues in different working environments. Building a consistent editor setup across a team is the configuration discipline that determines whether a shared workbook behaves predictably in all of those environments.

The primary sources of cross-machine inconsistency are locale settings, privacy level assignments, and file path parameterization. Locale settings were addressed in Section 2.3.1: making type enforcement locale-explicit in the query rather than relying on the machine's regional settings eliminates the most common source of date and number import inconsistency. Privacy level assignments are workbook-level metadata that travel with the workbook file — they do not need to be configured per machine — but they do need to be set intentionally rather than left at the default, and any team member who opens the workbook on a machine where the data source has not been previously accessed will be prompted to assign a privacy level for that source. A team-level document specifying the correct privacy level for each source the workbook accesses prevents this from becoming an ad hoc decision made under deadline pressure.

File path parameterization is addressed in depth in Part II when parameterized queries are introduced, but the principle is worth establishing here: any file path in a query's Source step that is an absolute path tied to one machine's directory structure will fail when the workbook is opened on a different machine where that path does not exist. The solution — replacing hard-coded paths with parameter queries that can be updated centrally — is the connection-layer standard this book applies throughout, and it is the first configuration decision that should be made before any shared workbook is handed off to a second team member.

Documenting these configuration decisions — the privacy level specification, the locale settings applied per source, the param-

ter query values and the file locations they point to — is part of the query documentation standard addressed in Chapter 16. A workbook whose configuration is undocumented cannot be supported by anyone other than its author.

Analyst’s Corner 2 — Why the Preview Pane Lies — and What to Do Instead

The Power Query editor’s Preview pane shows data. It does not show all the data. This is the single most consequential misunderstanding that new Power Query users bring to the editor, and it produces errors that are difficult to diagnose precisely because the error is not visible in the environment where it occurs.

By default, Power Query evaluates queries against the first one thousand rows of the source data. The Preview pane displays a sample from within those one thousand rows. Every profiling indicator — the column quality percentages, the value distribution histograms, the null counts — is computed against that sample. When the analyst applies a filter, removes rows, or counts distinct values, the result they see in the preview reflects the sampled portion of the dataset, not the full dataset.

This default is a reasonable engineering trade-off: evaluating the full dataset on every keystroke would make the editor unusably slow for large sources. The problem arises when the analyst treats the preview as if it were the full dataset — when they profile a column and observe zero null values in the preview, and conclude that the column has no nulls in the source, when in fact the nulls are distributed in rows beyond the first thousand. Or when they observe that a category field has six distinct values in the preview and build a transformation that handles exactly six values, when the full dataset contains a seventh value that the sampled rows did not happen to include.

There are two correct responses to this limitation, and both should be habitual for the professional analyst.