

Observability -Driven Testing

Testing in Production

THE MODERN QA ENGINEER SKILLSET

Observability-Driven Testing: Testing in Production

Logs, Metrics, Traces, and Quality in Live Systems

Yuri Syuganov

Modern QA Course

<https://modernQAcourse.com>

Observability-Driven Testing: Testing in Production

Logs, Metrics, Traces, and Quality in Live Systems

Book 6 of 26 in the series *The Modern QA Engineer Skillset*

Copyright © 2026 Yuri Syuganov. All rights reserved.

No part of this publication may be reproduced, distributed, or transmitted in any form or by any means without the prior written permission of the publisher, except for brief quotations in reviews and certain noncommercial uses permitted by copyright law.

All product names, logos, and brands mentioned in this book are the property of their respective owners. Use of these names does not imply endorsement.

First Edition · Version 1.0.0 · 2026-07-11

Set in Charter, IBM Plex Sans, and IBM Plex Mono. Typeset with Pandoc and XeLaTeX.

Modern QA Course · <https://modernQAcourse.com> · support@modernQAcourse.com

Contents

Introduction	1
PART I THE PARADIGM SHIFT	7
1. Why Testing Before Deployment Is Not Enough	9
About This Sample	23
About the Author	25
Also in This Series	27
Stuck on a Real Problem?	29

Introduction

A Comprehensive Self-Guided Learning Textbook for QA Engineers

First Edition, 2026

“The traditional testing model – write tests, run them before deployment, gate the release – is necessary but insufficient. Modern systems are too complex, too distributed, and too dynamic for pre-production testing alone to guarantee quality.”

About This Book

This book represents a fundamental shift in how quality assurance professionals think about testing. For decades, the QA discipline has been built on a simple premise: test everything before it reaches production, and if you test well enough, production will be safe. That premise is no longer sufficient.

Modern software systems are distributed across dozens of microservices, depend on third-party APIs that change without notice, serve users across global regions with different network conditions, and operate at scales where rare edge cases become daily occurrences. No test suite, no matter how comprehensive, can simulate the full complexity of production.

This book teaches you to extend the quality feedback loop beyond the deployment gate and into production itself. You will learn to use feature flags, canary deployments, A/B testing, structured logging, distributed tracing, metrics, synthetic monitoring, and AI-powered analysis as quality tools – not just operational tools.

This is not “testing in production” in the reckless sense. It is a disciplined practice of controlled exposure with automated safety nets. Observability-driven testing supplements, not replaces, pre-production testing.

By the end of this book, you will have built real, working systems that monitor production quality continuously, detect regressions before users notice them, and create a feedback loop that makes your test suite smarter over time.

Who This Book Is For

This book is written for **QA engineers, SDETs, test architects, and quality leads** who want to evolve their practice beyond traditional test-then-deploy workflows. You will get the most from this book if you:

- Have at least 1-2 years of experience in software testing or quality assurance
- Are comfortable reading and writing code (Python and TypeScript are used throughout)
- Understand basic concepts of web applications, APIs, and databases
- Have some familiarity with CI/CD pipelines (GitHub Actions, Jenkins, or similar)

- Are curious about how modern tech companies achieve high reliability at scale

You do not need:

- Prior experience with observability tools (we start from scratch)
- DevOps or SRE experience (though it helps)
- A production environment to practice with (we build local setups for all exercises)
- Deep statistical knowledge (we explain the math as we go)

Prerequisites

Required software (install before starting Chapter 1):

Tool	Version	Purpose
Python	3.10+	Primary language for examples
Node.js	18+	Playwright and TypeScript examples
Docker	24+	Running local observability stacks
Docker Compose	2.20+	Multi-container setups
Git	2.40+	Version control
A code editor	Any	VS Code recommended

Accounts you will need (free tiers are sufficient):

Service	Purpose	When Needed
GitHub	CI/CD workflows	Part I
LaunchDarkly or Flagsmith	Feature flags	Part II
Grafana Cloud (free tier)	Dashboards and alerting	Part III
OpenAI API	AI analysis exercises	Part V

How to Use This Book

Structure

The book is organized into **six parts**, each building on the previous:

- **Part I: The Paradigm Shift** – Why traditional testing is insufficient and what observability-driven testing looks like
- **Part II: Production Testing Strategies** – Feature flags, canary deployments, and A/B testing as quality gates
- **Part III: The Three Pillars of Observability** – Structured logging, distributed tracing, and metrics
- **Part IV: Monitoring and Alert Design** – Synthetic monitoring and alert engineering
- **Part V: AI-Powered Observability** – Using LLMs for anomaly detection and incident correlation
- **Part VI: Integration and Mastery** – The correlation framework, capstone project, and career guidance

Learning Features

Throughout the book, you will encounter these elements:

PRO TIP: Practical advice from real-world experience that will save you time and prevent mistakes.

COMMON MISTAKE: Patterns that look reasonable but lead to problems. Learn from others' errors.

Each chapter ends with:

- **Key Takeaways** – The most important points to remember
- **Self-Assessment Quiz** – Test your understanding (answers in Appendix A)
- **Hands-On Exercises** – Graded by difficulty: Beginner, Intermediate, and Advanced
- **Further Reading** – Curated resources to go deeper

Conventions

Code examples use real, working code (not pseudocode). File names are shown in comments at the top of each code block. When a code example builds on a previous one, the file name will match.

Terminal commands are prefixed with \$:

```
$ docker compose up -d  
$ python scripts/quality_gate.py --window 15m
```

Configuration files show the complete file, not just fragments, so you can copy them directly into your projects.

Table of Contents

Part I: The Paradigm Shift

- Chapter 1: Why Testing Before Deployment Is Not Enough
- Chapter 2: The Observability-Driven Testing Model

Part II: Production Testing Strategies

- Chapter 3: Feature Flags – Decoupling Deployment from Release
- Chapter 4: Canary Deployments – Statistical Validation with Live Traffic
- Chapter 5: A/B Testing as Quality Gates

Part III: The Three Pillars of Observability

- Chapter 6: Structured Logging – The Foundation
- Chapter 7: Distributed Tracing with OpenTelemetry
- Chapter 8: Metrics and Alerting with Prometheus

Part IV: Monitoring and Alert Design

- Chapter 9: Synthetic Monitoring – Building a Production Watchdog
- Chapter 10: Alert Design – Detecting Real Problems Without Fatigue

Part V: AI-Powered Observability

- Chapter 11: AI-Powered Log Analysis and Anomaly Detection
- Chapter 12: The Correlation Framework – Closing the Feedback Loop

Part VI: Integration and Mastery

- Chapter 13: Capstone Project – Building a Complete Observability-Driven Testing Pipeline
- Chapter 14: Career Guide – Applying These Skills in the Real World

Appendices

- Appendix A: Quiz Answers
- Appendix B: Glossary
- Appendix C: Tool Comparison Reference
- Appendix D: Further Reading and Resources

PART I

THE PARADIGM SHIFT

Why Testing Before Deployment Is Not Enough

1.1 The Traditional Testing Model

For most of software engineering history, the testing model has looked like this:

```
[Write Code] -> [Write Tests] -> [Run Tests in CI] -> [Deploy to Production] -> [Hope]
```

This model assumes that if you test thoroughly enough in a pre-production environment, production will be safe. It is a reasonable assumption for simple, monolithic applications running on a single server. It breaks down in the modern world for reasons that are worth understanding deeply.

The Monolith Era

In the monolith era (roughly 1990-2010 for most organizations), the testing model worked reasonably well because:

1. **The system was self-contained.** A single application, a single database, a single server. If your tests covered the application logic, they covered most of what could go wrong.
2. **Dependencies were few and stable.** You might depend on an email service and a payment gateway. These changed rarely and had well-documented behaviors.

- 3. **Scale was predictable.** You knew approximately how many users you had, and traffic patterns were relatively stable.
- 4. **Deployment was infrequent.** Deploying once a month or once a quarter meant each deployment was a major event with extensive manual testing.

The Microservices Era

The transition to microservices (accelerating from 2010 onward) broke every one of those assumptions:

- 1. **The system is distributed.** A single user request might traverse 5, 10, or 20 services. Testing each service in isolation does not guarantee the system works as a whole.
- 2. **Dependencies are numerous and volatile.** Your application depends on dozens of internal services, third-party APIs, CDNs, DNS providers, certificate authorities, and cloud infrastructure components. Any of them can change behavior, degrade, or fail at any time.
- 3. **Scale is unpredictable.** Viral moments, seasonal spikes, and bot traffic create load patterns that no test environment can fully simulate.
- 4. **Deployment is continuous.** Deploying dozens of times per day means each deployment is a small change, but the aggregate risk across all services is substantial.

What Pre-Production Testing Cannot Catch

Here is a non-exhaustive list of production issues that no pre-production test suite can reliably detect:

Issue	Why Pre-Production Misses It
Third-party API degradation	You test against mocks or sandboxes, not the live API
Regional connectivity problems	Your CI runs in one region; users are global
Certificate expiration	Certificates are environment-specific

continued on next page

Issue	Why Pre-Production Misses It
CDN misconfiguration	Your staging CDN config differs from production
Database connection pool exhaustion under real load	Test environments have different pool sizes and load
Memory leaks that manifest after hours of operation	CI tests run for minutes, not hours
Data-dependent bugs (specific user data triggers a code path)	Test data is synthetic and limited
Race conditions under real concurrency	Test environments have lower concurrency
Cascading failures across services	Test environments rarely have all services at production scale
Configuration drift between staging and production	Inevitable in any environment with manual configuration

This is not an argument against pre-production testing. Pre-production testing is essential. It catches the vast majority of bugs. But it is an argument that pre-production testing alone is insufficient for modern systems.

1.2 The Cost of Production Failures

Before we discuss the solution, let us understand the cost of the problem. When a production issue occurs that pre-production testing did not catch:

Direct costs:

- Revenue loss during the outage (for e-commerce, this can be thousands of dollars per minute)
- Engineering time to diagnose and fix the issue
- Customer support costs handling user complaints
- Potential SLA violation penalties

Indirect costs:

- User trust erosion (users who experience failures are less likely to return)
- Brand reputation damage (especially if the outage makes social media or news)
- Engineer morale impact (being woken at 3 AM to fix an issue that “should have been caught”)
- Opportunity cost (engineers fixing incidents are not building features)

Compounding costs:

- Mean Time to Detection (MTTD): How long before anyone notices the problem?
- Mean Time to Resolution (MTTR): How long before the problem is fixed?

Without observability, MTTD can be hours – you find out about the problem when a customer complains. With observability, MTTD can be seconds – automated monitoring detects the anomaly immediately.

1.3 The Real World: A Day Without Observability

Imagine this scenario. Your team deploys a new version of the checkout service at 2:00 PM on a Tuesday. The deployment passes all CI tests. The staging environment looks fine. The deployment goes out to all production instances simultaneously.

At 2:15 PM, a subtle bug begins manifesting: the new version has a race condition in the payment processing logic that occurs in approximately 1% of transactions. For 99% of users, everything works perfectly. For 1%, their payment is charged but the order is not created.

Without observability:

- The first customer complaint arrives at 2:45 PM via the support chat
- Support escalates to engineering at 3:15 PM
- Engineering begins investigating at 3:30 PM
- The bug is identified at 4:30 PM

- A rollback is initiated at 4:45 PM
- The rollback completes at 5:00 PM

Total impact: 2 hours and 45 minutes of exposure. With 100 transactions per minute and a 1% failure rate, approximately 165 users were affected.

With observability-driven testing:

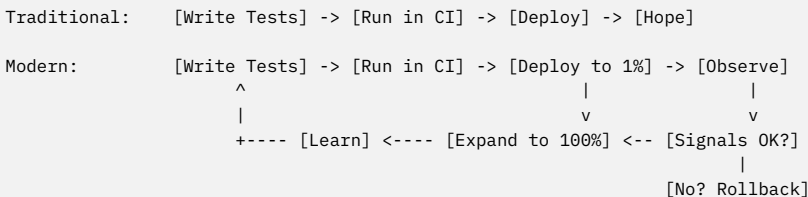
- At 2:00 PM, the deployment goes to 5% of traffic (canary)
- At 2:05 PM, the automated canary analysis detects that the error rate on the canary is 1% vs. 0.01% on the stable version
- At 2:06 PM, the canary is automatically rolled back
- At 2:10 PM, the team receives an alert with full context: error rate spike, affected trace IDs, and the specific error message
- By 2:30 PM, the bug is identified and a fix is in progress

Total impact: 6 minutes of exposure to 5% of traffic. Approximately 0.3 users were affected.

That is the difference observability-driven testing makes. Not “we test better before production” but “we detect and contain problems in production before they affect most users.”

1.4 The Paradigm Shift

The shift from traditional testing to observability-driven testing is not about abandoning pre-production testing. It is about extending the testing feedback loop into production:



Key principles of the new model:

1. **Deployment is not release.** You can deploy code to production without exposing it to users. Feature flags and canary deployments make this possible.

2. **Production is a testing environment.** Not in the reckless sense, but in the sense that production generates the most realistic test data possible: real users, real traffic, real dependencies.
3. **Observation is automated.** You do not watch dashboards manually. Automated quality gates compare the new version's metrics against the baseline and make promote/rollback decisions.
4. **The feedback loop is closed.** Production incidents inform test strategy. Test results predict production behavior. The system gets smarter over time.
5. **Rollback is a success, not a failure.** Catching a problem before it reaches all users is the system working as designed. Celebrate rollbacks; mourn undetected incidents.

1.5 What You Will Build in This Book

By the end of this book, you will have built:

1. **A feature flag system** with quality-gated progressive rollouts (Chapter 3)
2. **A canary deployment pipeline** with automated statistical analysis (Chapter 4)
3. **A/B testing quality gates** with statistical significance calculations (Chapter 5)
4. **A structured logging pipeline** with PII detection and log-based testing (Chapter 6)
5. **An OpenTelemetry instrumentation setup** with trace-based test assertions (Chapter 7)
6. **A Prometheus metrics and alerting system** with multi-burn-rate SLO alerts (Chapter 8)
7. **A Playwright synthetic monitor** running 24/7 against production (Chapter 9)
8. **An alert design system** with runbooks and fatigue prevention (Chapter 10)
9. **An AI-powered anomaly detection pipeline** using LLMs (Chapter 11)
10. **A test-production correlation framework** that makes your test suite smarter (Chapter 12)

11. **A complete capstone project** integrating all of the above (Chapter 13)

Key Takeaways – Chapter 1

1. Pre-production testing is necessary but insufficient for modern distributed systems.
2. Production environments generate failure modes that no test environment can simulate.
3. The cost of undetected production failures includes revenue loss, trust erosion, and engineer burnout.
4. Observability-driven testing extends the quality feedback loop into production with controlled exposure and automated safety nets.
5. The paradigm shift is from “test then deploy” to “deploy to a small percentage, observe, then expand.”

Self-Assessment Quiz – Chapter 1

Q1.1: Name three categories of production issues that pre-production testing cannot reliably detect.

Q1.2: What is the difference between Mean Time to Detection (MTTD) and Mean Time to Resolution (MTTR)?

Q1.3: In the observability-driven testing model, what is the role of a “quality gate”?

Q1.4: True or False: Observability-driven testing replaces the need for unit tests and integration tests.

Q1.5: Why is “rollback is a success, not a failure” an important cultural principle?

Hands-On Exercises – Chapter 1

Exercise 1.1 (Beginner): Production Failure Inventory

Think about the last 3-5 production incidents at your organization (or a previous one). For each incident, answer: (a) Could pre-production testing have caught it? (b) If not, what kind of production observability would have detected it sooner? Create a simple table documenting your findings.

Exercise 1.2 (Intermediate): Calculate the Cost of an Incident

Choose one real or hypothetical production incident. Estimate the cost using these categories: (a) revenue lost during the incident, (b) engineering hours spent diagnosing and fixing, (c) customer support hours, (d) estimated user trust impact. Compare this cost against the cost of implementing canary deployments.

Exercise 1.3 (Advanced): Map Your Testing Gaps

Create a diagram of your current system architecture. For each service and integration point, categorize your test coverage as: (a) well-covered by pre-production tests, (b) partially covered, (c) not covered. Identify the three highest-risk gaps and propose an observability-driven approach for each.

Career Translation**Resume phrasing**

- Championed the shift from gate-only QA to observability-driven testing, reducing mean time to detection (MTTD) from hours to under 5 minutes across production services.
- Designed and implemented production feedback loops combining canary analysis, synthetic monitoring, and structured logging to extend quality coverage beyond pre-deployment pipelines.
- Quantified the cost of production incidents and built the business case for controlled-exposure deployment strategies, decreasing user-facing outage impact by over 90%.

Cover letter framing

I believe the future of quality assurance lies in extending the testing feedback loop into production – not recklessly, but with disciplined controlled exposure and automated safety nets. In my work, I have seen firsthand how organizations that treat deployment as a binary gamble suffer longer incidents and higher costs than those that adopt progressive rollouts with observability. I bring both the technical foundation and the strategic mindset to help teams make that transition.

Interview framing

“I approach production quality as a continuous signal, not a one-time gate. My starting point is always the question: what can production tell us that staging cannot? From there, I design systems that deploy to small cohorts first, observe real-user metrics, and promote or roll back automatically. The trade-offs include additional infrastructure complexity and the cultural shift of treating rollback as a success rather than a failure.”

What not to say

- “We just need more unit tests.” – Ignores the entire class of failures that only manifest in production under real traffic and real dependencies.
- “Testing in production is reckless.” – Conflates undisciplined ad-hoc releases with controlled-exposure strategies that reduce risk.
- “Our staging environment is identical to production.” – It never is; configuration drift, traffic patterns, and third-party dependencies always differ.
- “If CI passes, we are good.” – Dismisses the reality that CI cannot simulate real concurrency, third-party failures, or regional connectivity issues.
- “Observability is an ops concern, not a QA concern.” – Modern quality requires owning the signal from code commit through production behavior.

Interview Depth Check

Question 1

Prompt: Your team deploys a checkout service update that passes all CI tests. Thirty minutes later, 2% of transactions are silently failing – payments are charged but orders are not created. Walk me through how you would have prevented or contained this with observability-driven testing.

What a strong answer should cover:

- The deployment should have gone to a canary (1-5% of traffic) first, not 100%.
- Automated quality gates should compare error rate and business metrics (order creation rate) between canary and stable.

- Structured logging with `trace_id` would make the silent failure visible.
- Synthetic monitoring would detect the order-not-created condition between deployments.
- Automated rollback on quality gate failure would limit blast radius.

Example answer:

- “First, I would never deploy to 100% simultaneously. I would configure a canary at 5% traffic with a quality gate checking error rate, latency, and – critically – a business metric like order-confirmation-rate relative to payment-success-rate.”
- “The quality gate would detect the divergence within the 10-minute observation window and trigger an automatic rollback. Total impact: 5% of traffic for 10 minutes instead of 100% for 2.5 hours.”
- “Post-incident, I would add a synthetic monitor that places a test order every 5 minutes and asserts that both the payment charge and order creation succeed, closing the coverage gap permanently.”

Question 2

Prompt: A colleague argues that investing in observability-driven testing is overkill because your staging environment already has good test coverage. How do you make the case for extending quality into production?

What a strong answer should cover:

- Concrete categories of failures staging misses (third-party APIs, real concurrency, regional issues, data-dependent bugs).
- Cost analysis: MTTD with vs. without observability.
- The feedback loop argument: production signals make the test suite smarter over time.
- That observability-driven testing supplements, not replaces, pre-production testing.

Example answer:

- “I would start with data. I would pull the last 10 production incidents and categorize how many could have been caught by staging tests. In most organizations, at least half involve third-party API degradation, configuration drift, or real-traffic race conditions that staging cannot simulate.”
- “Then I would calculate the cost: average incident takes 3 hours to detect and resolve, costs \$X in revenue and engineering time. Canary deployment with automated quality gates would have contained 80% of those to under 10 minutes of partial-traffic exposure.”
- “The clincher is the feedback loop: every incident that slips through teaches the system what to watch for next time. Without production observability, you only learn from failures after they fully impact users.”

Question 3

Prompt: You are asked to design the rollout strategy for a new AI-powered recommendation engine replacing the existing rule-based system. The new system has fundamentally different latency and accuracy characteristics. What is your approach?

What a strong answer should cover:

- Feature flags for application-level control, not just infrastructure canary.
- Progressive rollout stages: internal dogfood, beta, gradual percentage increase.
- Quality gates must include both technical metrics (latency, error rate) and business metrics (click-through rate, conversion).
- A/B testing with statistical significance for business metrics since canary alone misses UX impact.
- Fallback mechanism within the code, not just rollback of the deployment.

Example answer:

- “I would use feature flags rather than a canary deployment because we need to measure user-level engagement metrics, not just infrastructure health. Stage 1: internal employees for 3 days, checking error rate and latency. Stage 2: 5% of beta users for a week, adding click-through rate and conversion as quality gates with statistical significance testing.”
- “Stage 3: gradual rollout from 10% to 100% over two weeks, with 24-hour holds at each step. The code itself has a fallback – if the new engine’s quality score falls below 0.7, it silently returns the rule-based result and tracks the degradation.”
- “The trade-off is speed versus safety. This rollout takes three weeks, but for a recommendation engine that directly impacts revenue, the slower cadence is justified.”

Question 4

Prompt: Explain the difference between “deployment is not release” and traditional deployment. Why does this distinction matter for quality assurance?

What a strong answer should cover:

- Deployment puts code on production servers; release exposes it to users.
- Feature flags and canary deployments enable this decoupling.
- QA benefits: test on real infrastructure without user exposure, controlled blast radius, instant rollback via flag toggle.
- Cultural shift: rollback is a success, not a failure.

Example answer:

- “In traditional deployment, putting code on production and exposing it to users are the same action. That makes every deployment a binary bet. Decoupling them means I can deploy code to production servers on Monday but only release it to 1% of users on Wednesday after verifying infrastructure health.”

- “For QA, this is transformative. I can run my full test suite against production infrastructure with the flag off, validate with internal users first, then progressively expose real users while monitoring quality gates. If anything degrades, I toggle the flag in milliseconds – no rollback deployment needed.”
- “The cultural shift matters too. When rollback is a flag toggle, teams are more willing to experiment. The cost of catching a problem drops dramatically, so the incentive to hide or delay releases disappears.”

About This Sample

You have just read **Chapter 1 of Observability-Driven Testing: Testing in Production** — the free sample. The complete book continues with every remaining chapter, the interview Q&A, and all of the code.

Get the full book

Sample readers get a discount at

<https://leanpub.com/modernqa-06-observability-driven-testing/c/SAMPLE10>

Read online first

Chapter one of **every** book in the series is free to read at

<https://modernqacourse.com/library>

About the Author

Yuri Syuganov is a software quality engineer with more than twenty years of hands-on testing experience, spanning the arc from manual test design through Selenium-era automation to today's AI-augmented, agent-driven frameworks. He is the author of *The Modern QA Engineer Skillset*, a 26-book series built on a blunt analysis of what companies actually expect from QA engineers, and the creator of modernQAcourse.com, where the course behind the series is free and readers can bring real testing problems to experienced practitioners. He revises the series as the industry changes — because it does.

Also in This Series

The Modern QA Engineer Skillset — a 26-book course in modern software quality assurance.

AI-Augmented QA (Books 1–10)

- **Book 1:** Agent Skills for Browser Automation: The Complete Guide
- **Book 2:** AI-Augmented Test Design: From Prompts to Test Suites
- **Book 3:** Agentic Testing Architectures: Building Autonomous QA Systems
- **Book 4:** API and Contract Testing with AI: Schema to Shield
- **Book 5:** Performance and Chaos Engineering: Breaking Things on Purpose
- **Book 6:** Observability-Driven Testing: Testing in Production
- **Book 7:** Security Testing for AI Applications: The New Attack Surface
- **Book 8:** Infrastructure as Code Testing: Terraform to Kubernetes
- **Book 9:** Mobile and Cross-Platform Testing: 24,000 Devices, One Strategy
- **Book 10:** Visual and Accessibility Testing: Pixels and People

Technical Foundations (Books 11–19)

- **Book 11:** Manual Testing Fundamentals: The Bedrock
- **Book 12:** Programming for QA: Python, JavaScript, and TypeScript from Zero to Fluent
- **Book 13:** Browser Test Automation: Playwright, Patterns, and the Death of Fragile Tests

- **Book 14:** API Testing Fundamentals: Postman to Pytest
- **Book 15:** SQL and Database Testing: Query Your Way to Quality
- **Book 16:** CI/CD Pipelines: From Commit to Confidence
- **Book 17:** Git and Version Control: Beyond Push and Pull
- **Book 18:** Test Management Tools: Making Quality Visible
- **Book 19:** Linux and Command Line: Your Production Survival Kit

Professional Skills (Books 20–26)

- **Book 20:** Agile and Scrum for QA: Quality in Every Sprint
- **Book 21:** Communication and Stakeholder Management for QA: Influence Without Authority
- **Book 22:** Test Strategy and Quality Metrics: Measuring What Matters
- **Book 23:** QA Leadership and Mentoring: Multiplying Your Impact
- **Book 24:** Technical Writing for QA: Documents That Drive Action
- **Book 25:** Interview Preparation and Career Strategy: Land the Job
- **Book 26:** Testing Like a Senior: Patterns & Anti-Patterns Across Software QA Career Levels

Visit <https://modernQAcourse.com> for the full course catalog, updates, and companion resources.

Stuck on a Real Problem?

The full course behind this series is free to read at <https://modernQAcourse.com> — and it stays free.

When a chapter meets the real world and something breaks, post it on the **Help Board** at <https://modernQAcourse.com>. Asking is free, and experienced practitioners answer right in the thread. Need hands-on help? Book a 1:1 screen-share session with a QA expert.