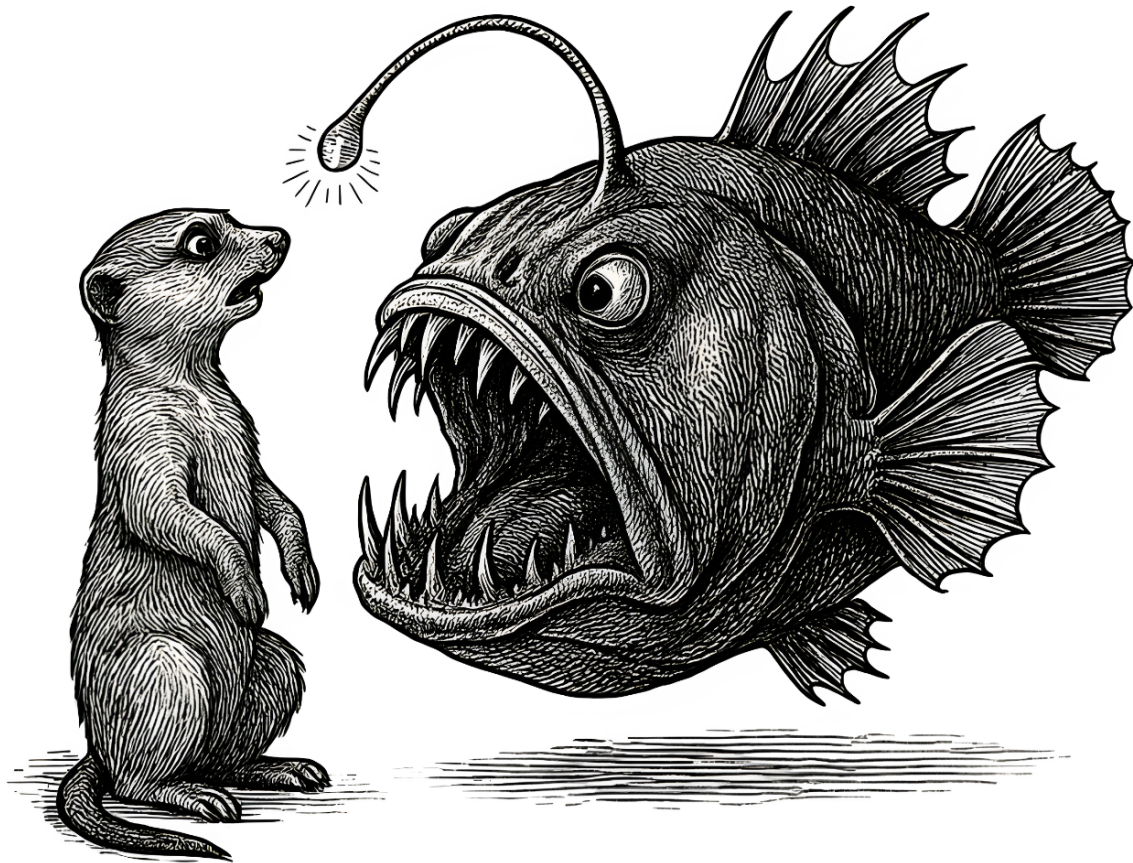


2nd Edition

On-Call In Action

Site Reliability Engineering Best Practices
for Building Resilient Systems



Quan Huynh

On-Call In Action

Site Reliability Engineering Best Practices for Building Resilient Systems

Quan Huynh

Table of contents

Table of contents.....	3
Forewords.....	7
Foundations: Why On-Call Matters & SRE Principles.....	9
1.1 The "Always On" World.....	9
1.2 Learning from Site Reliability Engineering (SRE).....	9
1.3 What This Book Will Cover.....	12
Anatomy of an Incident: The Management Lifecycle.....	14
2.1 Incident or Service Request.....	14
2.2 The Journey of an Incident.....	15
2.3 Not Always a Straight Path.....	19
Effective Alerting: Strategy and Routing.....	20
3.1 What Are Alerts and Why Do They Matter?.....	20
3.2 Core Alerting Concepts.....	21
3.2.1 Purpose of Alerting.....	21
3.2.2 The Danger: Alert Fatigue.....	21
3.2.3 Key Components of an Effective Strategy.....	23
3.3 Explain Centralized Alert Routing.....	33
3.4 Introduce Versus Incident as an Example.....	34
3.5 Practical Example: Basic Alert Routing with Versus Incident.....	35
3.5.1 Prerequisites: Getting Credentials.....	35
3.5.2 Versus Incident Configuration.....	36
3.6 From Strategy to Implementation.....	39
Integrating Monitoring Sources and Escalation Policies.....	41
4.1 The Necessity of Unified Alerting in Complex Systems.....	41
4.2 Case Study Deep Dive: The Book Info System on AWS.....	42
4.3 Defining Service Level Objectives for EKS Microservices.....	45
4.4 Practical Alerting Strategies for EKS Microservices.....	48
4.4.1 Metrics-Based Alerting (Prometheus/Alertmanager).....	48
4.4.2 Log-Based Alerting (Fluent Bit/Kibana/CloudWatch Logs).....	49
4.5 Defining Service Level Objectives for the Data Pipeline.....	50
4.6 Practical Alerting Strategies for the Data Pipeline.....	53
4.6.1 CloudWatch Metrics & Alarms.....	54
4.6.2 Log-Based Alerting (Optional but valuable).....	55
4.7 Crafting Actionable Templates.....	55

4.8 Designing Effective Escalation Policies.....	57
4.9 Bridging Practice and Principle.....	61
Chapter 5 - Measuring Reliability: SLIs, SLOs, and Error Budgets.....	63
5.1 Service Level Indicators (SLIs): Quantifying User Happiness.....	63
5.1.1 Formal Definition.....	63
5.1.2 Purpose - Measuring What Matters.....	64
5.1.3 Common SLI Types.....	64
5.1.4 SLI Specification vs. Implementation.....	65
5.1.5 Choosing Good SLIs - Practical Advice.....	65
5.1.6 Aggregation Methods.....	66
5.1.7 SLI Examples.....	67
5.1.8 Good vs. Bad SLI Examples.....	68
5.2 Service Level Objectives (SLOs): Setting Reliability Targets.....	69
5.2.1 Formal Definition.....	69
5.2.2 Purpose.....	70
5.2.3 Setting SLO Targets - Key Considerations.....	70
5.2.4 The "Good Enough" Principle (Why Not 100%?).....	71
5.2.5 SLO Specification Details.....	73
5.2.6 Table: SLO Examples for Book Info System.....	74
5.4 SLOs vs. SLAs: Internal Goals vs. External Promises.....	75
5.4.1 Defining Service Level Agreements (SLAs).....	75
5.4.2 Why the Distinction Matters for On-Call.....	77
5.5 Error Budgets: The Currency of Reliability.....	77
5.5.1 Formal Definition.....	77
5.5.2 Purpose - Balancing Reliability and Velocity.....	77
5.5.3 How Error Budgets Work.....	78
5.5.4 Error Budget Policies - Driving Decisions.....	78
5.5.5 Error Budgets and On-Call.....	79
5.5.6 Table: Error Budget Calculation Examples.....	79
5.5.7 Table: Budget Policy based on Error Budget Status.....	80
5.6 Managing Error Budget Consumption: Budget Burn Rate.....	81
5.6.1 Defining Budget Burn Rate.....	81
5.6.2 The Importance of Monitoring Burn Rate.....	82
5.6.3 Burn Rate Calculation and Interpretation.....	83
5.6.4 Connecting Burn Rate to Alerting Strategies.....	84
5.6.5 Burn Rate Alerting Example: The "14x over 1 Hour" Rule.....	86
5.7 Practical Implementation: Choosing SLIs and Setting SLOs.....	88
5.7.1 Methodologies Revisited.....	88
5.7.2 SLI Implementation Examples (Connecting to Tools).....	89
5.7 SLOs and Error Budgets as Active Control Mechanisms.....	90
5.8 Bridging Reliability to Alerting.....	92

Chapter 6 - Practical Examples of Unified Alerting & Templating.....	93
6.1 Integrating Diverse Monitoring Sources.....	93
6.1.1 Prometheus via Alertmanager.....	93
6.1.2 AWS CloudWatch via SNS.....	95
6.1.3 Grafana.....	98
6.1.4 Fluent Bit (Example: Forwarding JSON Application Logs).....	99
6.1.5 Sentry.....	101
6.2 Crafting a Unified Notification with Go Templates.....	103
Step 1: Identifying the Alert Source.....	104
Step 2: Extracting and Mapping Key Information.....	105
Step 3: The Complete Go Template Example.....	107
Step 4: The Unified Notification Output.....	107
6.3 Bridging to Blameless Postmortems.....	108
Chapter 7 - Learning from Failure: Blameless Postmortem.....	110
7.1 The Philosophy of Blamelessness: Shifting Focus from People to Systems.....	110
7.2 The Value Proposition: Benefits of a Blameless Approach.....	111
7.3 Implementing Blameless Postmortems: A Practical Guide.....	112
7.4 Root Cause Analysis (RCA) Techniques.....	115
7.5 Sustaining the Blameless Culture: Overcoming Challenges.....	116
7.6 Example Postmortem Report: Orders Service Outage.....	117
7.7 Learning from the Field: Case Studies.....	121
7.8 Transitioning from Blamelessness to Sustainable Operations.....	123
Chapter 8 - Sustainable On-Call: Scheduling and Managing Burnout.....	124
8.1 The Toll of Unsustainable On-Call.....	124
8.2 Sustainable Schedules: Fairness, Flexibility, and Recovery.....	125
8.3 Beyond Rotations: Safety & Support.....	127
8.4 Leadership: Championing Well-being.....	128
8.5 Measuring On-Call Health.....	129
8.6 Continuous On-Call Improvement.....	130
8.7 The Virtuous Cycle: On-Call & Innovation.....	131
8.8 A Case Study.....	132
8.9 From Sustainable Practices to Effective Incident Response.....	133
Chapter 9 - Effective Incident Communication.....	135
9.1 Understanding the Basics: Key Parts of Communication.....	135
9.1.1 Knowing Your Audiences: Who Needs Updates?.....	135
9.1.2 Defining Roles: Who Does What for Communication?.....	136
9.1.2 Choosing Channels: How Do You Send Updates?.....	137
9.2 Strategy and Planning: Getting Ready Before Problems Hit.....	138
9.3 Communicating Well When Things Go Wrong.....	140
9.4 Real-World Example: Book Info System Database Down.....	141
9.5 Communication is Your On-Call Superpower.....	143

Chapter 10 - The On-Call Ecosystem.....	145
10.1 Understanding the On-Call Ecosystem.....	145
10.2 The Tooling Landscape.....	147
10.2.1 Commercial Platforms.....	147
10.2.2 Open-Source Alternatives.....	149
10.2.3 Key Considerations for Tool Selection.....	150
10.3 Future Trends Shaping the On-Call Ecosystem.....	151
10.3.1 AIOps and Event Intelligence Solutions.....	151
10.3.2 Advanced Observability Trends.....	152
10.3.3 Anticipated Innovations in Tooling and Practices.....	152
10.4 Conclusion.....	153
Chapter 11 - On-Call in Action: Digital Customer Onboarding in Banking.....	154
11.1 The Digital Customer Onboarding Journey.....	154
11.2 Our Example: VKBank's System for DCO.....	155
11.3 Defining Meaningful Reliability: SLIs, SLOs, and Error Budgets.....	157
11.4 Alerting Strategy.....	162
11.5 Responding to Problems: Using AWS Incident Manager for On-Call.....	163
11.6 Integrating Versus Incident with AWS Incident Manager.....	164
11.6.1 Setting up AWS Incident Manager for the DCO SRE Team.....	164
11.6.2 Configuring Versus Incident.....	167
11.6.3 The Incident Lifecycle: From Alert to Resolution.....	169
11.6 Sustainable On-Call Practices for the DCO Team.....	171
11.7 Effective Incident Communication.....	172
11.8 The Journey to Operational Maturity.....	172
Acknowledgments.....	174

Forewords

These days, it feels like we're always connected in an "always-on" world. From the moment we wake up to the moment we go to sleep (and sometimes even while we're asleep!), digital services are running in the background, shaping our lives.

We chat with friends across the globe, manage our finances with the phone, binge-watch our favorite shows, and depend on critical systems that need to be online 24/7. We've come to expect all of this to just work, all the time. This isn't just about convenience; it's a real demand that businesses have to meet if they want to succeed and keep their users happy. It puts a lot of pressure on organizations to keep everything running smoothly, protect their income, and build trust with their customers.

To achieve this kind of reliability, you need robust systems, careful monitoring, and, most importantly, a team of people ready to jump in when things go wrong. That's where "on-call" comes in. It means having tech experts—like Site Reliability Engineers (SREs), DevOps engineers, or operations staff—available when needed to handle problems quickly. They're the first line of defense when systems falter or something unexpected happens, and they're the ones who figure out what's wrong, halt the damage, fix the issue, or get the right people involved.

But being on-call isn't just about reacting to alerts. It also involves doing routine maintenance during shifts and, more strategically, working to make the systems better over time. The things you learn while handling incidents can help you improve the way systems are designed, make monitoring more effective, automate responses, and, ultimately, prevent the same problems from happening again.

Unfortunately, on-call practices are often put together without much thought, which can lead to messy processes, slow response times, and a lot of stress for engineers, sometimes even causing burnout. To make on-call work well and be sustainable, you need a plan.

This book, "On-Call In Action," is meant to be your guide to navigating the tricky world

of modern on-call. We'll cover the basics, explore practical ways to implement strategies and tools, with the open-source project as a key example, and dive into important concepts like Service Level Objectives (SLOs) for measuring reliability. We'll also talk about the human side of things, like scheduling, preventing burnout, and creating a supportive culture, all of which are essential for building on-call practices that are effective and sustainable. Whether you're a developer, operator, or manager dealing with systems that need on-call support, this book will give you the knowledge and best practices to create or improve an on-call system that actually works.

Foundations: Why On-Call Matters & SRE Principles

Welcome to the world of building reliable systems! In this chapter, we'll explore why having technical experts "on-call" is so important today and how a set of ideas called Site Reliability Engineering, or SRE, can help make this challenging job more manageable and effective.

1.1 The "Always On" World

Think about the apps and websites you use every day. Whether it's for banking, shopping, talking to friends, or getting work done, we all expect these digital services to be working perfectly, any time we need them – day or night, weekday or weekend. If a popular online store suddenly goes offline during a big sale, or if your company's internal communication tool stops working, it can cause a lot of frustration and even cost businesses money.

Because everyone expects services to be "always on," companies need a way to fix problems quickly when they happen. This is where being on-call comes in. Being on-call means that certain technical people in a company are ready to respond to issues, even if it's outside of their normal working hours. These are often engineers, perhaps with titles like Site Reliability Engineers (SREs) or DevOps Engineers. They are the first line of defense when something unexpected goes wrong with a live service. Their job is to jump in, figure out the problem, and get things back to normal as fast as possible.

1.2 Learning from Site Reliability Engineering (SRE)

Dealing with systems that need to be available all the time can be stressful.

Fortunately, there's a way of thinking about and managing technology that can make this much better. It's called **Site Reliability Engineering (SRE)**. SRE started at Google, and it's a set of practices and a mindset that treats the work of keeping systems running (often called "operations") like a software engineering challenge. This means using data, automation, and careful planning to build and maintain reliable services.

SRE offers some really useful ideas that can make on-call work more effective and less painful. Let's look at a few key ones:

It's Okay Not to Be Perfect

It might sound strange, but aiming for 100% perfect uptime for every service, all the time, is usually not a realistic or even a good goal. Why? Because trying to achieve absolute perfection is incredibly expensive and can slow down the process of adding new features or making improvements.

SRE recognizes this. Instead of chasing perfection, we define how much *unreliability* is acceptable. This "acceptable amount of downtime or errors" is called an error budget. For example, if a service aims to be available 99.9% of the time, its error budget is the remaining 0.1%. This budget gives teams a clear understanding of how much risk they can take. If the service is performing well and has plenty of error budget left, they might feel more confident releasing new features. If the error budget is used up, it's a signal to focus on stability and fix problems before making more changes. It's all about finding a healthy balance.

Setting Clear Goals

To know if a service is reliable enough, we need clear goals. In SRE, these goals are called Service Level Objectives (SLOs). An SLO is a specific, measurable target for how well a service should perform from a user's perspective. For instance, an SLO might be "99.9% of login attempts should be successful" or "95% of search results should appear in under half a second."

SLOs are important because they:

- Tell everyone what "good enough" looks like.
- Help teams prioritize work (if an SLO is at risk, fixing it becomes a high priority).
- Provide a shared understanding between different teams (like development and

operations) about reliability expectations.

Good SLOs are based on what users actually care about.

Cutting Down on Repetitive Work

Imagine an on-call engineer having to manually restart a server every few hours because of a recurring minor issue, or spending a lot of time copying and pasting data for reports. This kind of manual, repetitive, and often boring work that doesn't add much long-term value is what SRE calls toil.

A big focus in SRE is to reduce or eliminate toil, usually by automating these tasks. If a task can be automated, it probably should be. Reducing toil is great because it:

- Frees up engineers from tedious work.
- Allows them to spend more time on projects that make the systems better and more reliable in the long run (like fixing the root cause of that recurring server issue!).
- Reduces the chance of human error in repetitive tasks.
- Makes on-call work less frustrating.

Learning Without Blame

When something goes wrong – and in complex systems, things inevitably will – it's natural to want to know why. However, it's very important *how* we approach this. SRE promotes a blameless culture. This means that when an incident happens, the focus is on understanding *what* happened in the system and *how* the processes or tools might have contributed to the problem, not on blaming individuals.

A blameless approach assumes that everyone was doing their best with the information and tools they had at the time. Why is this so important?

- It creates psychological safety: People feel safe to report problems, admit mistakes, and share information openly without fear of punishment.
- It leads to better learning: When people are honest, the team can get to the real root causes of issues and make meaningful improvements.
- It helps prevent the same problems from happening again.

Instead of asking "Who messed up?", we ask "What can we learn from this to make our systems and processes better?"

Mixing Project Work with Daily Tasks

To truly improve reliability and reduce future on-call pain, SRE teams aim to spend a significant portion of their time on engineering projects that have a long-term impact. This could be automating manual tasks, improving monitoring, redesigning parts of the system for better resilience, or developing new tools.

A common guideline in SRE (like at Google) is that engineers should spend at least 50% of their time on this kind of project work, with the other 50% (or less) on operational duties, including handling incidents. If a team is constantly firefighting and has no time for improvement projects, they get stuck in a reactive cycle, and the on-call burden rarely gets lighter. This balance is key to making on-call sustainable.

Fair Pay for On-Call Work

Being on-call is a significant responsibility. It can disrupt personal time, affect sleep, and add stress. SRE principles generally suggest that engineers who take on this extra duty should be compensated fairly for it. This acknowledges the effort and the importance of the role in maintaining service reliability.

1.3 What This Book Will Cover

Now that you have a basic understanding of why on-call is essential and how SRE principles can guide us, you might be wondering how to put all of this into practice. That's what the rest of this book is about!

We'll take you on a journey, chapter by chapter, through the key aspects of building and managing an effective and sustainable on-call system.

We will cover:

- **The Incident Lifecycle:** Understanding the stages from detection to resolution and learning (Chapter 2).
- **Effective Alerting:** Strategies for creating meaningful, actionable alerts while minimizing noise, including practical implementation using tools like Versus Incident (Chapter 3).
- **Integration and Escalation:** Connecting diverse monitoring sources and defining robust escalation policies (Chapter 4).
- **Measuring Reliability:** Defining and using SLIs, SLOs, and Error Budgets to drive decisions (Chapter 5).

- **Putting It All Together:** Practical Examples of Unified Alerting & Templating (Chapter 6).
- **Learning from Failure:** Implementing blameless postmortems for continuous improvement (Chapter 7).
- **The Human Element:** Designing fair schedules, managing burnout, and ensuring effective communication (Chapters 8 & 9).
- **The Ecosystem:** Navigating the tooling landscape and looking at future trends (Chapter 10).
- **Digital Customer Onboarding:** A case study illustrating the journey of building an on-call practice (Chapter 11).

Our goal is to give you the knowledge and confidence to build on-call practices that not only keep your services running smoothly but also support the well-being and growth of your engineering teams. Let's get started!