



Claude



Claude for Enterprise Teams

**A Practical Playbook for Developers,
QA, PMs, and Business Leaders**

Copyright & Disclaimer Page

© 2026 Paramount Software Solutions, Inc. All rights reserved.

No part of this publication 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 publisher, except in the case of brief quotations in reviews or articles.

This guide is intended for informational purposes only. While every effort has been made to ensure the accuracy of the information at the time of publication, technology and AI platforms evolve quickly. Product features, pricing, and capabilities described here may change.

Organisations should consult their own legal, security, and compliance teams before implementing any AI solutions or policies. Nothing in this guide constitutes legal, financial, or professional advice.

Claude is a product of Anthropic. All product names, logos, and brands are property of their respective owners

Acknowledgements

This guide is the result of collaboration between Paramount's AI Centre of Excellence, engineering leaders, and clients who are experimenting with Claude in real projects. We are grateful to the developers, QA engineers, architects, PMs, SAP and functional consultants, and sales teams who shared their wins, failures, and questions about AI adoption. Their experiences shaped the practical patterns, checklists, and warnings you'll find in these pages.

How to Use This Guide

Who this guide is for

You will get the most value from this guide if you are a developer, QA engineer, architect, PM, SAP/functional consultant, sales leader, or team lead responsible for helping your organisation adopt AI in a safe and practical way.

How this guide is structured

- Chapters 1–3 set the foundation: why Claude matters, what it can and cannot do, and how to keep usage safe.
- Chapters 4–10 go role by role with concrete examples and prompts.
- Chapters 11–14 help you compare workflows, roll out Claude in your org, and govern it with clear policies and checklists.

How to read it

- If you're a **team lead or executive**, start with Chapters 1–3, then jump to Chapters 11–14.
- If you're a **practitioner** (Dev, QA, PM, SAP, Sales), skim Chapters 1–3, then go straight to the chapter for your role and the prompt library at the back.

You don't have to read this cover to cover. Treat it like a toolbox—pick the patterns and checklists that match your current stage of AI adoption.

Table of Contents

Part I – Foundations: Claude in the Enterprise8

1. Chapter 1 – The New AI Baseline: Why Claude, Why Now10

- 1.1 AI has quietly become infrastructure
- 1.2 Why Claude stands out for teams
- 1.3 The risk of ungoverned AI and “shadow AI”
- 1.4 The mindset shift: AI as a teammate, not a replacement
- 1.5 What this guide will—and will not—do

2. Chapter 2 – Claude Fundamentals for Teams13

- 2.1 Claude models, interfaces, and pricing tiers (Free, Pro, Team, Enterprise)
- 2.2 Key capabilities: reasoning, long context, multi-step tasks
- 2.3 Claude chat vs Claude Projects vs Claude Code: an overview
- 2.4 What Claude is great at—and where it struggles
- 2.5 Example end-to-end use case: from prompt to outcome

3. Chapter 3 – Working Safely: Accuracy, Hallucinations, and Guardrails.....19

- 3.1 What are hallucinations and why they matter
- 3.2 Prompt-level guardrails: asking for reasoning, citations, and “I don’t know”
- 3.3 Process-level guardrails: human review and cross-validation
- 3.4 Grounding Claude in your own data (Projects, RAG, MCP)
- 3.5 Quick checklist: safe use of Claude across roles

Part II – Claude in Daily Work: Role-Based Playbooks25

4. Chapter 4 – Claude for Software Developers & Programmer Analysts28

- 4.1 Daily workflows for developers (reading, planning, coding, debugging)
- 4.2 Designing your Claude Code environment and CLAUDE.md
- 4.3 Patterns for refactoring, tests, and code review with Claude
- 4.4 Adoption roadmap for dev teams (weeks 1–4)
- 4.5 Security & governance notes for developers
- 4.6 Prompt templates and examples for developers

5. Chapter 5 – Claude for QA Engineers & SDETs.....35

- 5.1 Test case and scenario generation from specs and user stories
- 5.2 Using Claude for test automation scripts and frameworks
- 5.3 Bug triage, clustering, and reporting with AI assistance
- 5.4 Adoption roadmap for QA and SDET teams
- 5.5 Security & governance notes for QA (logs, PII, environments)
- 5.6 Prompt templates and examples for QA & SDETs

6. Chapter 6 – Claude for Architects & Senior Engineers	42
6.1 Using Claude as a first-pass architecture reviewer	
6.2 Documenting systems: maps, flows, and ADRs with Claude	
6.3 Supporting onboarding and knowledge transfer	
6.4 Adoption roadmap for architecture & platform teams	
6.5 Security & governance notes for architecture work	
6.6 Prompt templates and examples for architects	
7. Chapter 7 – Claude for Project Managers & Scrum Leads.....	49
7.1 Planning with Claude: charters, milestones, and task breakdowns	
7.2 Running meetings: agendas, notes, and action-item tracking	
7.3 Reporting: weekly updates, RAG, and stakeholder summaries	
7.4 Adoption roadmap for PMO and Scrum teams	
7.5 Security & governance notes for PMs (docs, stakeholders, data)	
7.6 Prompt templates and examples for PMs & Scrum leads	
8. Chapter 8 – Claude for SAP, Functional Consultants & Business Analysts	56
8.1 Drafting functional specs, config docs, and impact analyses	
8.2 Designing test cases and UAT scripts with Claude	
8.3 Client-facing communication and training materials	
8.4 Adoption roadmap for SAP and functional teams	
8.5 Security & governance notes for SAP/BA work (config & data)	
8.6 Prompt templates and examples for SAP & BA roles	
9. Chapter 9 – Claude for Sales & Business Development.....	63
9.1 Account research and meeting preparation	
9.2 Messaging: emails, LinkedIn outreach, and proposals	
9.3 Customer success: QBRs, renewal plans, and risk reviews	
9.4 Adoption roadmap for sales and GTM teams	
9.5 Security & governance notes for sales (client, pricing, claims)	
9.6 Prompt templates and examples for sales & BD	
10. Chapter 10 – Claude for Your Career Growth	71
10.1 Using Claude as a personal tutor and practice environment	
10.2 Building your portfolio, CV, and LinkedIn with AI support	
10.3 Interview preparation and promotion cases	
10.4 Personal adoption roadmap: 30–60–90 days with Claude	
10.5 Security & privacy notes for individuals	
10.6 Prompt templates for learning, branding, and interviews	

Part III – Workflows, Rollout, and Governance78

11. Chapter 11 – Comparing AI Workflows: Standard Prompting vs Projects vs Claude Code	79
11.1 Standard prompting: strengths and limits	
11.2 Claude Projects: when to use workspaces with persistent context	
11.3 Claude Code: when to bring AI into your IDE and repos	
11.4 Choosing the right mode for each role and task	
11.5 Example “day in the life” workflows combining all three modes	
12. Chapter 12 – Building a Team Adoption Roadmap	86
12.1 Start small: selecting the first use cases and teams	
12.2 Designing a 4-week pilot: goals, metrics, and feedback loops	
12.3 Scaling from pilot to multi-team rollout	
12.4 Change management and training: upskilling your people	
12.5 Adoption roadmaps by role: Dev, QA, PM, SAP, Sales	
12.6 Common pitfalls in AI adoption—and how to avoid them	
13. Chapter 13 – Security, Data Privacy, and Governance for Claude in the Enterprise	92
13.1 Claude Enterprise security and compliance overview	
13.2 Data classification: allowed, restricted, and prohibited data types	
13.3 Access control across chat, Projects, and Claude Code	
13.4 Monitoring, logging, and cost governance for Claude usage	
13.5 Managing hallucinations as a safety and risk issue	
13.6 Policy checklist: AI usage, data privacy, and security controls	
14. Chapter 14 – Implementation Checklists and Templates	99
14.1 Claude Security & Data Privacy Policy Checklist	
14.2 Role-based AI usage policies (Dev, QA, PM, SAP, Sales, Individuals)	
14.3 Risk assessment worksheet for AI use cases	
14.4 Data classification template for AI tools	
14.5 Sample adoption plan for a 90-day Claude rollout	
14.6 Internal communication templates (emails, FAQ, training invites)	
15. About Paramount Software Solutions	105

PART I

Foundations: Working With Claude, Not Under It



Claude is not here to replace your judgment; it is here to amplify it. Part I gives you the mental models and basic techniques you need before you trust it with real work.

You will learn how Claude actually reasons, how to give it the right kind of instructions, and how to get outputs you can rely on instead of “AI-shaped noise.” By the end of this Part, you’ll know when to ask, what to share, and how to stay in control of the answers you get.

Chapters in this Part

- Chapter 1 – Why Claude, Why Now
- Chapter 2 – The Core Claude Playbook (Prompts, Projects, Code)
- Chapter 3 – Getting Reliable Output (Grounding, Checking, Iterating)

Chapter 1: The New AI Baseline: Why Claude. Why Now?

1.1 AI has quietly become infrastructure

In 2023 and 2024, most teams experienced AI through demos and side projects. A few power users experimented with chatbots for code snippets, email drafts, or research. But day-to-day work remained largely unchanged.

By 2026, that has shifted. AI is no longer a toy sitting in a separate browser tab; it is becoming part of how code is written, how requirements are documented, how test suites are generated, and how status updates reach leadership. Teams that treat AI as infrastructure—something embedded into workflows, tools, and processes—are starting to move faster and build more resilient systems than those that don't.

The question is no longer “Should we use AI?” It is “How do we use AI in a way that is practical, secure, and aligned with how our teams actually work?”

1.2 Why Claude stands out for teams

There are several capable AI platforms today. This guide focuses on Claude because of three qualities that matter deeply in enterprise environments:

- 1. Strong reasoning with long context**Claude can work across long documents, multi-file codebases, and complex instructions thanks to its large context windows and model design. That makes it particularly suited for tasks like architecture reviews, PRD analysis, and multi-step coding workflows.
- 2. Enterprise-grade security and governance options**Claude's enterprise offerings are built with security controls, compliance, SSO, audit logging, and data-handling options that matter for regulated or risk-conscious organisations. Combined with your internal policies, this makes it possible to move beyond “playground” usage to governed, production-adjacent workflows.
- 3. Specialised tools: Claude Projects and Claude Code**Whereas simple chat interfaces work well for ad-hoc questions, most real work happens in **projects** and **repos**. Claude Projects let you ground the model in your documents and instructions, and Claude Code lets AI work inside your actual development environment. These two pieces turn Claude from a generic chatbot into a real teammate for developers, QA, PMs, and others.

1.3 The risk of ungoverned AI

If your people are already using AI tools on their own laptops, you already have AI adoption—just not necessarily the kind you want. This “shadow AI” shows up as:

- Sensitive information pasted into free tools
- Multiple vendors and models in use without security review
- Outputs that look convincing but have never been validated
- No clear record of how AI influenced decisions or code

Industry guidance is clear: without governance, AI adoption creates fragmented risk and inconsistent results rather than sustainable productivity. At the same time, banning AI outright simply pushes usage underground.

Claude offers you a way to bring AI into the *official* toolkit—on enterprise plans, through governed Projects and Claude Code—while wrapping it in your own policies, training, and oversight.

1.4 The mindset shift: AI as a teammate, not a replacement

The most successful teams we’ve observed do not treat Claude as a magical replacement for expertise. Instead, they adopt a simple mental model:

- **Humans own direction, decisions, and accountability.**
- **Claude accelerates analysis, drafting, and exploration.**

Developers still own the code that ships to production, but Claude accelerates reading, drafting, and refactoring. QA engineers still own test strategy, but Claude helps generate test ideas and structure suites. PMs still own scope and stakeholder alignment, but Claude turns messy notes into clear narratives.

Throughout this book, you’ll see this pattern repeated: let Claude handle the heavy reading, the first draft, and the structured thinking—then let humans critique, refine, and decide.

1.5 What this guide will (and will not) do

This guide **will**:

- Show role-specific, concrete ways to use Claude in daily work
- Give you language for policies, checklists, and governance
- Help you design a realistic, phased rollout instead of a big-bang launch

This guide will **not**:

- Promise that AI is flawless or always correct
- Replace your own security, legal, or compliance advice
- Tell you that one model/vendor can solve everything

Claude is powerful, but it is still a tool. The value comes from pairing it with good process, thoughtful constraints, and teams that understand where to trust it and where to be sceptical.

The chapters that follow will help you do exactly that—starting with the fundamentals of how Claude works for teams, and then diving into what it means for each role in your organisation.

Chapter 2 – Claude Fundamentals for Teams

2.1 What Claude actually is (beyond “a chatbot”)

If your only exposure to Claude has been in a browser tab, it is easy to think of it as “just another chatbot.” In reality, Claude is a **family of large language models** (LLMs) with different performance and speed tiers, accessible through chat, desktop apps, and APIs. As of mid-2026, Anthropic offers higher-end Opus models for the most complex reasoning tasks, Sonnet as the primary workhorse model on `claude.ai`, and Haiku as a lighter, faster option for high-volume or latency-sensitive scenarios.

What matters for teams is not memorising model names, but understanding that you can **trade off power vs. speed vs. cost**: use the strongest models for complex architecture reviews or deep analysis, and use faster ones for iterative drafting, summarisation, or batch operations.

On top of the raw models, Anthropic ships **experience layers** that change how your people interact with Claude:

- Chat in the browser or desktop app
- **Claude Projects**, which act like long-term workspaces
- **Claude Code**, which brings Claude into your editor and working directory

Understanding these layers is the key to using Claude as more than a one-off Q&A tool.

2.2 Claude’s main interfaces and plans

Most organisations today interact with Claude in three ways:

1. Claude on the web / desktop

- A chat-like interface where users can talk to Claude, upload files, and, if enabled, access Projects.
- Good for early pilots, research, writing, and some role-based workflows (PMs, SAP consultants, sales, etc.).

2. Claude Projects

- Dedicated workspaces that hold files, instructions, and conversations tied to a specific project, client, or product area.
- Essential once you move from experimentation to ongoing work.

3. Claude Code

- An AI coding agent that runs in your development environment (IDE/CLI), sees your working directory, and can write, edit, and run code.
- Especially powerful for developers, QA engineers, and architects.

On top of this, many organisations use **Claude API** or enterprise deployments, giving them more control over networking, data retention, and governance.

From a planning perspective, think in terms of:

- **Free / Pro** – good for early exploration, individuals, and low-risk workflows.
- **Team / Enterprise** – required for any serious organisational rollout, giving you SSO, admin controls, security and compliance features, and usage visibility.

The rest of this book assumes you are either on Team/Enterprise or planning to move there as adoption grows.

2.3 Claude chat vs Projects vs Claude Code

To use Claude effectively as a team, you need to understand the differences between **standard chats**, **Projects**, and **Claude Code**. Think of them as three modes of working with the same underlying intelligence.

2.3.1 Standard chats

Standard chats are:

- **Best for:** quick questions, one-off analysis, small pieces of code, brainstorming, simple content drafts.
- **Strengths:** zero setup, fast, flexible; great for getting started or ad-hoc needs.
- **Limitations:** context is short-lived and fragmented; you constantly re-explain; hard to build repeatable workflows or team-shared processes.

This mode is ideal when the task is small, low-risk, and self-contained. It should not be your primary mode for complex projects.

2.3.2 Claude Projects

Projects turn Claude into a **persistent collaborator** that remembers your documents and instructions for that project. A Project typically includes:

- A name and description (e.g., “Payments Platform – Architecture”, “Client X SAP Rollout”, “Claude Rollout PMO”)
- Custom instructions describing what Claude should do and how it should behave for this context
- Uploaded files: PRDs, specs, architecture diagrams, process maps, test plans, transcripts, templates, etc.
- Chats that all share that same context

Projects are:

- **Best for:** multi-week initiatives, teams that need shared context, and workflows where you always refer back to the same set of documents.
- **Strengths:** huge effective context window, file-grounded reasoning, consistent instructions, reusability across chats and collaborators.
- **Limitations:** still chat-based; for actual code execution and repo traversal, you need Claude Code.

For example, a PM might maintain a “Product X – Roadmap & Reporting” Project; a QA lead might have a “Regression & UAT – Release Y” Project; and SAP consultants might keep a “Client Z – FI/CO Rollout” Project with all relevant artifacts.

2.3.3 Claude Code

Claude Code is an **AI coding agent** that integrates into your development environment (IDE, terminal) and works on your real files and commands, not just pasted snippets. It can:

- See your working directory and repository structure
- Open, read, and edit files
- Run commands (e.g., tests, build scripts) and see outputs
- Follow multi-step plans to perform complex coding tasks

Claude Code is:

- **Best for:** engineers and SDETs working on real codebases—refactors, feature work, test generation, debugging, and scaffolding.
- **Strengths:** deep repo awareness, command execution, reusable routines and skills, ability to behave like a “coding coworker.”
- **Limitations:** requires setup and permissions; primarily for technical users; needs strong guardrails to avoid risky commands.

For practical purposes: **Projects manage knowledge and context; Claude Code manages code and commands.** Together, they support most of what engineering-heavy teams need.

2.4 Key capabilities teams should actually care about

Rather than listing every feature, it helps to focus on capabilities that change how your people work.

2.4.1 Deep reasoning on long, messy inputs

Claude can ingest long PRDs, design docs, logs, and code files and then:

- Summarise them for different audiences (developer, manager, executive)
- Compare versions, spot inconsistencies, and highlight risks

- Answer targeted questions like “Where do we mention non-functional requirements?” or “Which modules touch payment retries?”

This **reasoning over large context** is at the core of its value in complex organisations.

2.4.2 Multi-step planning and execution

Claude is designed to handle multi-step tasks when prompted correctly. You can ask it to:

- Break a request into steps
- Propose a plan
- Execute each step with explicit checkpoints

In Claude Code, this planning can extend to actions like “analyse tests”, “apply diff”, “run tests again”, which makes it feel more like an autonomous junior engineer—especially when combined with routines.

2.4.3 Working with structured and semi-structured data

Through Projects and file uploads, Claude can:

- Read CSVs, logs, JSON exports, and spreadsheets
- Extract and transform structured information (e.g., turning meeting notes into a risk register)
- Propose data views and summaries that feed into dashboards and reports

This is particularly useful for PMs, SAP/BA teams, and sales leaders who live in tools like Jira, SAP, CRM, and spreadsheets.

2.4.4 Agentic extensions: routines, skills, and scheduled tasks

A newer layer on top of Claude Code and Projects is **agentic workflows**—things like routines, scheduled tasks, and reusable skills.

Examples:

- A routine that, every Monday, pulls the latest log files, summarises key errors, and drafts a QA focus list
- A skill that reads your README, CLAUDE.md, and architecture docs and answers “what’s the best place to add this new feature?”
- Scheduled tasks in a Project that generate weekly status updates or stakeholder reports.

You don’t have to start here, but it is important to know that Claude can grow from “chat” into **automation** once you have stable workflows.

2.5 Where Claude shines—and where it struggles

2.5.1 Strengths

Claude tends to perform particularly well in:

- **Complex reasoning tasks** that require understanding relationships across multiple documents or code modules.
- **Code understanding, refactoring, and debugging** at the system level rather than just single functions.
- **Structured writing**: docs, specs, reports, and communication that follow a clear template or pattern.
- **Multi-role workflows** where the same core context (e.g., a PRD or client account) must be seen differently by Dev, QA, PM, and Sales.

This makes it ideal as the “AI backbone” for engineering-heavy organisations.

2.5.2 Limitations

Claude, like all LLMs, has important limitations:

- It **does not have a real-time system state** unless you connect it via tools and integrations. It won't know your latest deployment status or live revenue metrics by default.
- It can **hallucinate** when asked for facts without sources or when pushed beyond its training data, even if it sounds confident.
- It is **not a decision-maker**—it cannot understand business risk, politics, or long-term strategy in the way humans can.

This is why we emphasise throughout this guide: *AI accelerates; humans decide.*

2.6 A simple end-to-end example

To make all of this concrete, imagine a single feature: adding “soft delete” to a core service.

With Claude:

1. A **PM** uploads the PRD and existing user stories into a Project and asks Claude to highlight unclear requirements and edge cases.
2. An **architect** uses the same Project to ask Claude for potential design options, risks, and impacts on related services.
3. A **developer** opens Claude Code in the repo, asks it to locate all modules that handle deletes, and drafts a plan to introduce soft delete without breaking invariants.

-
4. A **QA engineer** uses the Project to generate regression and UAT test cases, focusing on restore flows and data retention.
 5. A **PM** asks Claude to summarise the change and testing approach into a stakeholder-friendly update.

Behind the scenes, your **security and governance** settings ensure the work happens in a controlled environment with appropriate data classification and access.

The rest of this book unpacks exactly how to design these flows for each role, and how to roll them out in an intentional way—not as scattered experiments, but as part of your operating model.

Chapter 3 – Working Safely: Accuracy, Hallucinations, and Guardrails

Claude can make your teams dramatically faster. It can also be confidently wrong. This chapter is about building the “seatbelts and guardrails” that let you keep the speed without losing control.

We’ll look at what hallucinations really are, why they happen, and how to reduce them with a combination of better prompts, better grounding, and better processes.

3.1 What hallucinations are (and why they matter in real work)

In everyday terms, a **hallucination** is when Claude produces an answer that sounds plausible and confident—but is wrong, fabricated, or not supported by any reliable source. This might be a made-up API, a non-existent error code, an incorrect regulation, or a “case study” that never happened.

Hallucinations occur because LLMs generate the most likely sequence of tokens based on patterns in their training data, not because they “look up” facts the way a database does. When the model doesn’t have enough signal, or when the prompt is vague, it will still try to produce something that fits the pattern—even if reality doesn’t match.

In a consumer chat, a hallucination might be amusing. In an enterprise context, hallucinations can:

- Introduce bugs or security issues into production code
- Mislead architects during design reviews
- Distort test coverage and risk assessments
- Misstate facts to customers or regulators

Preventing hallucinations is not about achieving perfection. It is about **reducing risk to an acceptable level** and making sure your teams know when to trust, when to verify, and when to say “we need more information.”

3.2 Why hallucinations happen: core causes

Most hallucinations show up when one or more of these conditions are true:

1. **Lack of grounding data**Claude is forced to rely only on general training data instead of your actual systems, documents, or code. Without relevant, up-to-date context, it fills gaps with guesses.
2. **Vague or overloaded prompts**Prompts that are broad (“Tell me everything about our payments architecture”) or multi-layered (“Write the code and also test and deploy it”) push the model to improvise rather than reason carefully.
3. **No rules for uncertainty**If you don’t give Claude permission to say “I don’t know” or “I can’t tell from this information,” it will often default to answering anyway.
4. **No verification step**When outputs are used directly—pasted code, unreviewed tests, unverified facts—there is no safety net to catch errors.

The good news: all of these failure modes can be mitigated with **guardrails** at the prompt level, system level, and process level.

3.3 Prompt-level guardrails: how to talk to Claude

Well-designed prompts are the first line of defence against hallucinations. Anthropic’s own guidance emphasises clear instructions, examples, and letting Claude “think out loud” before answering.

3.3.1 Be specific about task, scope, and sources

Instead of:

“Explain our system and suggest improvements.”

Try:

“You are a senior backend engineer. Based on the attached payments-architecture.md and ADR-014-soft-delete.docx, summarise how the payments service currently handles deletes. Then list three potential risks of introducing soft delete without changing downstream consumers. Do not speculate beyond what is in these documents.”

Good prompts:

- Specify **role** (“You are a senior QA in a fintech company...”)
- Define **scope** (“Focus only on the Checkout and Payments services”)
- Constrain **sources** (“Use only the attached files; if they are insufficient, say so”)

This reduces the space where Claude can improvise.

3.3.2 Explicitly allow “I don’t know”

Anthropic’s own docs recommend giving Claude explicit permission to admit uncertainty, which can “drastically reduce false information.”

You can bake this into your standard instructions:

“If you are not sure, or if the information is not available in the provided documents, say ‘I don’t know from these documents’ and stop. Do not guess.”

This is especially important in Projects and Claude Code where users might otherwise assume that any answer is grounded in real data.

3.3.3 Ask for reasoning and evidence

Several best-practice guides recommend “show your work” prompting—chain-of-thought or step-by-step reasoning—so that humans can inspect the logic, not just the conclusion.

Examples:

- “First list the assumptions you are making. Then reason step by step. Only then give a final answer.”
- “For each claim, provide a direct quote from the attached documents that supports it, with section or file name.”

In Claude’s own docs, Anthropic suggests:

- Extracting **word-for-word quotes** first, then reasoning based on them
- Requiring **citations** for each factual claim
- Using **chain-of-thought verification** to reveal faulty logic

This makes responses auditable and easier to review.

3.4 Grounding Claude in your own data

Prompt-level improvements help, but the single most powerful way to reduce hallucinations is to **ground Claude’s answers in your actual data and documents**, rather than letting it rely solely on general training.

You have three main tools for this in a Claude-centric setup:

3.4.1 Use Projects for document grounding

By putting PRDs, specs, architecture docs, test plans, and policies into a Claude Project, you continuously give the model a relevant “source of truth” to reason from.

Best practices include:

- Grouping documents by domain (e.g., one Project per product, client, or platform area)

- Adding clear Project-level instructions (“Use these files as your primary source; if they conflict, prefer the latest version”)
- Asking Claude to **quote the exact text** from these files before summarising or recommending actions

This is a lightweight form of **retrieval-augmented generation (RAG)** and significantly reduces hallucinations compared to open-ended Q&A.

3.4.2 Integrate structured retrieval (RAG) where needed

For more advanced setups—especially through your AI CoE or platform team—you can connect Claude to enterprise data sources via RAG patterns and services.

Common patterns include:

- Indexing internal documents and knowledge bases and retrieving relevant passages for each query
- Filtering data by metadata (date, source, owner) to prioritise the most trustworthy information
- Instructing Claude to **only** use retrieved snippets when forming its answer (“external knowledge restriction”).

Studies and practitioner reports consistently show that properly configured RAG significantly reduces hallucinations and improves groundedness.

3.4.3 Restricting external knowledge when appropriate

Anthropic explicitly recommends an “external knowledge restriction” pattern for high-risk tasks: telling Claude to ignore its general training and use only the provided documents.

For example:

“Answer using only the content of the attached policy PDFs. If you need information that is not present in these files, say ‘Not specified in the policy documents provided.’ Do not rely on your general knowledge.”

This is particularly useful for compliance, internal process, and architecture questions.

3.5 Process-level guardrails: putting humans back in the loop

Even with strong prompts and good grounding, you still need **process guardrails**. These are practices and policies that ensure humans stay in control of high-impact decisions.

3.5.1 Human review based on risk level

Most enterprise AI guides recommend **tiered review** based on risk:

- **Low-risk** (brainstorming, drafts, internal summaries): light review by the user is usually enough.
- **Medium-risk** (internal test cases, internal documentation, non-prod code): require a peer or lead to review AI-generated outputs before adoption.
- **High-risk** (production code, security controls, legal or regulatory content, external comms): require a senior owner and documented sign-off.

The key is to be explicit: define which categories of work are allowed with Claude, and what level of human review they require.

3.5.2 Second-pass verification

Several sources recommend using a **second pass** to catch hallucinations and logic errors:

- Ask Claude to **critique its own answer**: “Now review your previous response. List three possible mistakes or assumptions, and correct them if needed.”
- Use **best-of-N**: run the same prompt multiple times and compare outputs to identify inconsistencies.
- Use **cross-model validation**: compare Claude’s answer with another model or a specialised tool (e.g., static analysis for code, linters, domain-specific validators).

These techniques are especially effective for engineering tasks where you can run tests or linters as an objective check.

3.5.3 Confidence and escalation

In RAG and enterprise patterns, teams often add a light confidence mechanism:

- If Claude can’t find supporting quotes or relevant retrieved passages, it must **lower its confidence and defer**.
- Low-confidence outputs are either discarded or **escalated** to human experts for manual handling.

You can approximate this even in chat by instructing Claude:

“If your confidence is below 70% because the documents don’t fully cover the question, say ‘I’m not confident enough to answer this reliably’ and suggest follow-up questions instead of guessing.”

3.6 Embedding guardrails into Claude Code and Projects

Because this book focuses heavily on Claude Code and Projects, it is worth calling out a few patterns specific to those environments.

3.6.1 Guardrails for Claude Code

For Claude Code, combine AI guardrails with your existing SDLC controls.

- **No direct pushes to main:** treat AI-generated code like any other change—require code review, tests, and CI.
- **Command review:** require developers to **approve every command** before it runs, especially anything that touches production or infra.
- **Limit scope:** start with repositories and environments that are safe for experimentation (non-prod, sandboxes).

Claude Code's own security guidance emphasises reviewing commands, avoiding secrets in prompts, and isolating high-risk operations. Those same patterns help reduce the impact of hallucinations as well.

3.6.2 Guardrails for Projects

For Projects:

- **Keep Projects scoped:** avoid dumping entire organisations' worth of documents into one Project. Use one per product, client, or domain, so context is relevant and responses stay grounded.
- **Write Project-level instructions** that include safety rules (e.g., "If requirements conflict, ask for clarification. Do not assume.").
- **Keep documents current:** outdated specs can cause "hallucinations by omission," where Claude is accurate to old information but wrong for the current system.

Over time, you can treat each Project as a living, governed knowledge space, with owners responsible for data quality and scope.

3.7 A practical checklist your teams can follow

To make this chapter actionable, here is a short checklist you can embed into training, prompts, and policy docs.

For every Claude session, ask:

- 1. Have I given Claude enough relevant context?**
 - Did I upload or reference the key documents, code, or data?
 - Did I specify which sources it should use or avoid?
- 2. Have I told it what to do when it doesn't know?**
 - Did I explicitly allow "I don't know" and forbid guessing?
- 3. Have I asked for reasoning and evidence?**
 - Step-by-step explanation?
 - Quotes and citations from our docs where needed?
- 4. Is there a human review step appropriate to the risk?**
 - Who is reviewing this before it reaches production, customers, or regulators?
- 5. Am I using the right mode?**
 - For complex, ongoing work: Projects instead of one-off chats.
 - For real code: Claude Code with proper SDLC controls.

If your teams learn to answer these five questions consistently, you will have addressed most of the real-world causes of hallucinations and unreliable AI usage in day-to-day work.

In the next chapters, we will see how these guardrails show up **inside** each role's workflow—so developers, QA engineers, PMs, SAP consultants, and sales teams all have their own, concrete patterns for using Claude safely.

PART II

Claude for Roles and Workflows: From Product Teams to Your Career



Most value from Claude comes not from heroic one-off prompts, but from everyday use in real workflows. Part II shows how developers, QA, product managers, designers, architects, PMs, Scrum leads, SAP/BA teams, sales, and individuals turn Claude into a teammate across their day-to-day work.

You'll see concrete patterns for turning specs into code, tickets into tests, workshops into specs, flows into UAT scripts, CRM noise into account briefs, and scattered updates into sharp status reports and sales messages. This Part also helps you use Claude for your own career growth and to choose the right workflow—chat, Projects, or Claude Code—for the work in front of you.

Chapters in this Part

- **Chapter 4** – Claude for Developers & QA
- **Chapter 5** – Claude for Product Managers & Designers
- **Chapter 6** – Claude for Architects & Senior Engineers
- **Chapter 7** – Claude for Project Managers & Scrum Leads
- **Chapter 8** – Claude for SAP, Functional Consultants & Business Analysts
- **Chapter 9** – Claude for Sales & Business Development
- **Chapter 10** – Claude for Your Career Growth
- **Chapter 11** – Comparing Workflows: Chat, Projects, and Claude Code

Chapter 4 – Claude for Software Developers & Programmer Analysts

For developers and programmer analysts, Claude is not just a smarter autocomplete. When used well, it becomes a **junior-to-mid engineer with superhuman reading speed**: great at absorbing large codebases, drafting implementations, and explaining complex flows—while you remain in charge of design, correctness, and production decisions.

This chapter shows how to integrate Claude into your daily development workflow, from understanding legacy systems to refactoring, testing, and code review. It also explains how to set up CLAUDE.md and other context files so that Claude understands your codebase and your standards.

4.1 How Claude changes the developer workflow

Traditional development involves a lot of “hidden” work: reading existing code, chasing dependencies, understanding architecture, and translating requirements into concrete tasks. Claude excels at exactly these activities.

With Claude Code, developers can:

- **Navigate and understand codebases faster** – by asking questions across multiple files and services.
- **Plan work more thoroughly** – by having Claude break features into steps, highlight risks, and suggest tests.
- **Draft and refactor code** – using the plan as a guide, in small, reviewable increments.
- **Debug issues** – by correlating logs, stack traces, and code to propose likely root causes and fixes.

Used this way, Claude is not “writing your app for you.” It’s **amplifying** the parts of your job that involve reading, pattern-matching, and structured reasoning, so that you can spend more of your time on design, integration, and trade-offs.

4.2 Designing your Claude Code setup and CLAUDE.md

Before you throw Claude at a complex repo, you need to give it a map. That map typically lives in a file called CLAUDE.md at or near the root of your project.

4.2.1 What CLAUDE.md is for

A good CLAUDE.md does three things:

- **WHAT** – describes the tech stack, project structure, and key directories.
- **WHY** – explains the purpose of the project and the role of major components.
- **HOW** – tells Claude how to work on the project: commands, tools, tests, and constraints.

You can think of it as Claude’s onboarding document—exactly what you would give a new engineer joining the team.

4.2.2 Best practices for CLAUDE.md

Practitioner guides and Claude power users recommend keeping CLAUDE.md concise and focused:

- Aim for **under 200–300 lines**, focusing on information that is always relevant.
- Use **progressive disclosure**: instead of dumping all your knowledge, tell Claude where to look for more detail (e.g., “See /docs/architecture/ for detailed diagrams.”).
- Capture your **non-obvious conventions**, like using bun instead of node, custom scripts, codegen tooling, or domain-specific linters.
- Include **verification instructions**: how to run tests, what “done” looks like, and how to check changes (e.g., “Run npm test and npm run lint before considering a change complete.”).

You can optionally create additional CLAUDE.md files in subdirectories (e.g., /services/payments/CLAUDE.md) for large monorepos to give Claude local context for each service.

4.2.3 Example structure for CLAUDE.md

A typical outline might look like:

- Project purpose and domain
- Tech stack and main frameworks
- Repo structure and key directories
- How to run the app locally
- How to run tests and linters
- Coding guidelines and style rules
- Things Claude must never do (e.g., “Never touch infra scripts in /ops/ without explicit instruction.”)

Once you have this in place, Claude Code can orient itself in your project far more effectively.

4.3 Daily workflows: how developers actually use Claude

Let’s walk through the main parts of a developer’s day and how Claude fits into each.

4.3.1 Understanding existing code and architecture

When you work on an unfamiliar service or feature, the first problem is **orientation**. Claude Code can help by:

- Explaining how a particular endpoint, class, or function works, referencing all related files.
- Tracing data flow: “Show me how a checkout request moves from the API gateway to the payment provider.”
- Generating architecture summaries or sequence descriptions you can use as documentation or onboarding material.

Because Claude Code can see the full directory and follow imports, it can pull together context that would otherwise take you hours of manual searching.

4.3.2 Planning a feature or change

Before writing code, developers can ask Claude to collaborate on planning:

- Translating a PRD or ticket into a **step-by-step implementation plan**, including which files will likely be touched.
- Identifying **risks and edge cases** given existing patterns and dependencies.
- Suggesting **tests** and validation steps as part of the plan.

Many Claude Code power users recommend a specific pattern:

1. Start in Plan mode (or an equivalent prompt)
2. Ask Claude to write a **phase-by-phase plan** with clear checkpoints
3. Commit the plan somewhere (docs or repo)
4. Execute one phase at a time, with review in between

This makes Claude’s work more predictable and easier to supervise.

4.3.3 Writing and refactoring code

Claude is highly effective at implementing well-specified tasks in existing patterns.

Examples:

- Implementing a new function or endpoint given a detailed spec and existing examples
- Refactoring a module for clarity or performance while maintaining the same behaviour
- Applying the same change pattern across multiple files (e.g., adding feature flags, updating logging, changing an API signature)

The key is to:

- Keep tasks **small and focused** (“Implement phase one only,” “Refactor this component, but don’t change its public interface”).
- Always **review diffs** and explanations before accepting changes.
- Run tests and linters after each change to catch regressions.

This aligns with best-practice guidance to treat Claude Code like a mid-level engineer you still supervise closely.

4.3.4 Debugging issues

When debugging, Claude can:

- Read stack traces and logs and correlate them with relevant code paths
- Propose hypotheses for root causes and suggest logging or instrumentation to confirm them
- Draft patches and test cases to verify that the issue is fixed

Developers report that Claude is especially effective when you:

- Provide **all relevant context** (error messages, sample inputs, environment details)
- Ask it to **explain its reasoning** and propose multiple hypotheses, not just one fix

You still run the commands, confirm the environment, and decide which fix to apply.

4.4 Patterns for tests, code review, and documentation

Claude is not a substitute for your SDLC, but it can make each step more efficient.

4.4.1 Test-driven workflows with Claude

Several guides recommend pairing Claude with **TDD-style workflows**:

- Ask Claude to generate tests first, based on the spec and existing patterns.
- Review and refine the tests yourself.
- Then ask Claude to implement code that **passes those tests**.

This gives you an external oracle (the tests) to verify Claude’s implementation and reduces the chance of silently wrong behaviour.

4.4.2 AI-assisted code review

Claude can act as a second pair of eyes in code review:

- Summarising PRs in plain language (“What changed and why?”)
- Highlighting potential performance or security issues
- Suggesting additional tests or validation steps

You can either:

- Use Claude Code on the branch locally to analyse diffs, or
- Paste diffs into a chat or Project and ask for review with specific criteria (e.g., scalability, readability, security).

Claude’s comments should supplement, not replace, human review—but they often surface issues and considerations that would otherwise be missed.

4.4.3 Documentation by productising your thinking

Claude is particularly good at turning scattered notes and code into:

- Architecture overviews
- “How this module works” docs
- API reference sections based on code and comments

A simple pattern:

1. Ask Claude to read relevant code files and any existing docs.
2. Ask for a **draft doc** in your preferred template (e.g., ADR, RFC, or README section).
3. Edit for accuracy and nuance.
4. Save and, if needed, update CLAUDE.md to point to the new docs.

This is one of the highest leverage uses of Claude for long-term maintainability.

4.5 Adoption roadmap for development teams

To keep this practical, here is a simple four-week adoption roadmap tailored to engineering teams.

Week 1 – One developer, one workflow

- Choose one developer and one workflow (e.g., debugging, test generation, or small refactors).
- Set up Claude Code in their environment and create a minimal CLAUDE.md for one repo.
- Track time saved, code quality, and where human review catches issues.

Week 2 – Define standards and patterns

- Refine CLAUDE.md based on what worked and what was confusing.
- Establish **guidelines**:
 - Types of tasks allowed with Claude
 - Required review practices
 - When to avoid or limit AI (critical security modules, infra scripts, etc.)

Week 3 – Expand to a small squad

- Onboard 2–4 more developers in the same product area.
- Share a short internal guide: “Our Claude Code workflow” with examples.
- Start using Claude for more steps (planning, documentation, code review), always with human review.

Week 4 – Measure and integrate into SDLC

- Measure:
 - Lead time for changes
 - Time to understand new modules
 - PR review time
- Integrate Claude explicitly into your SDLC (definition of done, code review checklists, architecture review templates).

From here, you can extend to more repos and products, backed by the security and governance chapter later in this book.

4.6 Security & governance for developers

The developer chapter needs to reinforce that **speed never trumps safety**.

Key guardrails:

- Use Claude Code only on repos and environments that are appropriate, with **repo-scoped permissions**.
- Never paste **credentials, API keys, or secrets** into prompts or terminal sessions.
- Treat AI-generated code as you would code from a junior engineer: it must pass tests, code review, and your usual quality gates before shipping.
- For sensitive modules (security, auth, infra), consider restricting AI usage to analysis and documentation only—not direct code edits.

More detailed security, privacy, and governance practices are covered in Chapter 13, but you should already weave these principles into your engineering norms and onboarding.

4.7 Prompt patterns and examples for developers (overview)

At the end of the book, you'll find a full prompt library. For now, here are a few patterns you can introduce to your teams.

- **Codebase orientation**

"You are a senior backend engineer working on this repository. Based on the files in this directory, explain how the Checkout flow works from API request to database write. Include file paths and function names in your explanation."

- **Feature planning**

"Act as a staff engineer. Based on the attached PRD and the existing code in /services/payments, propose a step-by-step plan to add soft delete to the Payments service. Call out which files will change, what tests are needed, and any risks. Do not write code yet."

- **Implementation (phase-by-phase)**

"Implement phase 1 of the approved plan only. Show me the diffs and explain what you changed and why. Do not touch any other files or phases."

- **Debugging**

"Given this stack trace, these logs, and the code in /services/orders, list three plausible root causes. For each, show the code paths involved and propose a minimal fix plus tests to verify it."

These patterns map directly to the workflows we've discussed and keep humans firmly in control.

Chapter 5 – Claude for QA Engineers & SDETs

If developers own “what we build,” QA engineers and SDETs own “how we know it works.” In many organisations, QA is where AI has some of the fastest, most visible impact: converting requirements into tests, exploring edge cases, documenting bugs, and maintaining regression suites.

Claude—especially when used through Claude Code and Projects—can act like a **tireless QA copilot**. It will not replace human judgement about risk and coverage, but it can dramatically reduce the time spent on repetitive test design, script maintenance, and bug reporting.

5.1 How Claude fits into the QA workflow

Typical QA responsibilities include test planning, test case design, execution (manual and automated), defect reporting, regression testing, and collaboration with dev and PMs. Almost all of these activities involve language, patterns, and structured thinking—areas where Claude is strong.

Claude can help QA teams to:

- Turn requirements, user stories, and change requests into **structured test cases**
- Identify **edge cases and risk areas** that should be covered
- Generate or refactor **automation scripts** for UI, API, and integration tests
- Summarise logs and repro steps into **clear bug reports**
- Maintain regression suites and test documentation as the system evolves

The goal is not “AI that tests everything for us.” It is **AI that handles the writing, structuring, and pattern-matching**, so QA engineers can spend more time on strategy, risk assessment, and exploratory testing.

5.2 Test case generation from specs and user stories

One of the most natural places to start with Claude is **test case design**.

5.2.1 From requirements to structured test cases

Given a set of user stories, acceptance criteria, or a PRD, Claude can:

- Extract key flows and variations
- Propose positive, negative, and boundary test cases
- Organise them by priority, risk, or persona

This works even better when requirements are stored in a Project, so Claude can reference related docs (e.g., design decisions, known constraints).

For example, QA can ask:

“Based on the attached ‘Checkout v2’ PRD and user stories, generate a list of test cases for the checkout flow. Include positive, negative, and edge cases. Group them by feature (cart, payment, shipping) and mark which ones are critical for regression.”

The QA lead then reviews these cases, adjusts coverage for domain-specific risks, and adds missing scenarios.

5.2.2 Designing regression suites and smoke tests

Claude is also useful for constructing **regression suites** based on recent changes and known risk areas.

Examples:

- Given a list of recent PRs or release notes, ask Claude to identify functions and flows that are at higher risk and propose regression sets.
- Ask Claude to propose a minimal **smoke test** suite for each environment that verifies critical paths after deployment.

Again, QA remains responsible for deciding what “good enough” coverage looks like. Claude speeds up the drafting.

5.3 Using Claude Code for automation, scripts, and logs

Claude Code gives QA engineers and SDETs an extra boost when they work with automation frameworks and system logs.

5.3.1 Drafting and refactoring automation scripts

With Claude Code connected to your repos, QA can:

- Generate skeleton scripts for frameworks like Playwright, Cypress, Selenium, REST clients, or BDD tools
- Refactor flaky tests for readability and maintainability
- Add assertions, waits, or helper functions following existing patterns

Practitioners report good results when they:

- Provide **examples of existing tests** for Claude to follow
- Limit tasks to **one test or small group of tests at a time**
- Ask Claude to explain what it changed and why, then run the suite locally to validate

Claude Code is particularly useful when a QA engineer is not deeply familiar with the underlying language but understands the desired behaviour.

5.3.2 Working with logs, traces, and evidence

QA often gets buried under logs, screenshots, and traces. Claude can:

- Summarise logs around a failed scenario and highlight likely problem areas
- Extract useful error messages and categorise them (e.g., network, validation, server error, third-party)
- Turn raw notes and screenshots into structured bug reports with reproducible steps and expected vs actual behaviour

For example:

"Given these logs and my notes, draft a clear bug report. Include: title, environment, preconditions, steps to reproduce, expected result, actual result, and any suspected root causes."

This not only saves time but improves communication between QA and development.

5.4 Exploratory testing and risk-based thinking

Exploratory testing is where human curiosity and domain knowledge really shine. Claude can act as a brainstorming partner and note-taking assistant, but it cannot replace the intuition of an experienced QA engineer.

5.4.1 Generating ideas and “what-ifs”

Given a description of a feature and known constraints, Claude can:

- Propose unusual or adversarial scenarios
- Suggest persona-based tests (e.g., novice user, power user, malicious user)
- List environmental variations (browsers, devices, network conditions) worth exploring

For instance:

“Act as a senior QA for a multi-tenant SaaS app. Based on this feature description, list 20 exploratory test ideas focused on data isolation, permissions, and concurrency.”

The QA engineer can then choose the most relevant ideas and run them manually, updating notes as they go.

5.4.2 Capturing exploratory sessions

Claude can also help with **documentation** of exploratory work. QA can paste their session notes (“I tried X, saw Y, then did Z”) and ask Claude to:

- Turn them into structured session reports
- Highlight observed risks and follow-up actions
- Update test charters or checklists

This makes exploratory testing more transparent and repeatable across the team.

5.5 Building AI-assisted QA agents and pipelines

As teams mature, some move beyond interactive usage into **agentic QA workflows**.

Examples from real projects include:

- An AI QA agent built with Claude Code and Playwright that runs browsers, clicks through flows, and generates human-readable reports on every pull request.
- Pipelines that use Claude to analyse screenshots and DOMs to detect visual or layout issues beyond pixel-based diffs.
- Custom user agents that simulate novice user behaviour in web apps and log confusing areas for UX and product teams.

These setups typically require:

- Claude Code access via API or GitHub actions
- A well-designed **prompt/verification template** describing what to test on each PR
- Strong governance to ensure the agent's actions are safe and bounded (only runs in test environments, only executes approved commands)

This book does not assume every team will build such agents, but it is important to know they are possible once your basic practices are solid.

5.6 Adoption roadmap for QA and SDET teams

Like with developers, QA should start small and grow deliberately.

Week 1 – Power user pilot

- Choose one QA engineer or SDET and one product area.
- Focus on **test case generation and bug report drafting** using Claude chat or Projects.
- Measure time saved and where AI-generated test cases miss critical risks.

Week 2 – Introduce Claude Code for automation

- Connect Claude Code to your test repos.
- Use it to refactor existing tests and generate a small number of new scripts with strong human review.
- Document do's and don'ts (e.g., "Don't let AI rename core helpers without a plan.").

Week 3 – Team onboarding and shared patterns

- Onboard 2–3 more QA/SDET team members.
- Create a **QA Project** per product area with specs, past bugs, and key flows.
- Align on how to label AI-generated test cases and scripts so they can be tracked.

Week 4 – Integrate into your QA strategy

- Add clear entries in your **test strategy**:
 - Where AI is used (test design, automation, documentation)
 - What review is required for AI outputs
 - Which risk levels allow AI support vs require manual design only
- Consider piloting a simple pipeline where Claude reviews logs or drafts regression sets for each release.

The goal is to make AI a visible, governed part of your QA toolkit—not a hidden side experiment.

5.7 Security & governance notes for QA

Because QA sometimes sees **real data**, you must be careful about what is shared with Claude.

Key safeguards:

- Avoid sending **production logs or data** that contain PII, payment details, or other sensitive information. Use masked or synthetic data wherever possible.
- Use Claude Team/Enterprise with appropriate data-handling and logging settings for any work that touches internal systems.
- Ensure that AI-generated tests and scripts go through your normal review process before being added to the main suite.
- For AI QA agents, restrict them to test environments and review their prompts and permissions regularly.

Chapter 13 covers broader security and governance in detail, but QA should treat AI like any other tool that touches potentially sensitive information.

5.8 Prompt patterns and examples for QA & SDETs (overview)

As with developers, a consistent set of prompts makes daily use much easier. Here are a few patterns you can teach your QA teams:

- **Test case generation**
“You are a senior QA engineer. Based on the attached user stories and acceptance criteria, generate a list of test cases for this feature. Include positive, negative, boundary, and error-handling scenarios. Group them by flow and mark which ones you consider high-risk.”
- **Regression selection**
“Given this list of recent changes and known issues, propose a regression test set for the next release. Prioritise tests by risk and impact. Explain why you selected each group.”
- **Automation support**
“Using our existing Playwright tests in /tests/ui/checkout, generate a new test that covers the ‘apply coupon’ scenario. Follow the same coding style and helper methods. Explain what you added and how to run it.”
- **Bug report drafting**
“Turn these notes and logs into a structured bug report. Include: title, environment, preconditions, exact steps to reproduce, expected vs actual results, and any hints at potential root causes.”

These patterns make Claude a predictable, reliable assistant rather than an occasional “magic button.”

Used well, Claude becomes an **extension of your QA team's intelligence**: proposing tests, clarifying behaviour, and capturing knowledge, while your engineers still decide what risk looks like and when quality is good enough. In the next chapter, we'll look at how architects and senior engineers can use Claude to review systems, document designs, and support onboarding at scale.

Chapter 6 – Claude for Architects & Senior Engineers

Architects and senior engineers sit at the intersection of **systems, teams, and time**. They care less about individual functions and more about how services interact, how data flows, how constraints shape design, and how decisions age. Claude is uniquely valuable at this level because it can read across large amounts of code and documentation and help you see patterns you might miss when you're deep in Jira and diagrams.

Used well, Claude becomes an architecture reviewer, documentation partner, and onboarding accelerator. Used poorly, it can suggest textbook patterns that don't fit your constraints or amplify design mistakes. This chapter focuses on how to keep it in the first category.

6.1 What architects really need from Claude

Most architecture work falls into a few recurring activities:

- Understanding and documenting existing systems
- Designing and reviewing new systems or major changes
- Communicating decisions to different stakeholders
- Guarding non-functional requirements (scalability, reliability, security, cost, evolution)

Claude's long context and reasoning make it especially good at:

- **Summarising complex systems** from multiple sources (code, ADRs, diagrams, tickets)
- **Challenging designs** by highlighting trade-offs and risks
- **Keeping architecture knowledge current**, instead of locked in someone's head or outdated docs

The goal is to treat Claude as a **first-pass architecture reviewer and documentation engine**, not as the final authority on what to build.

6.2 Mapping and understanding existing architecture

Before you can improve or redesign a system, you need a clear picture of what exists. This is where Claude can save days or weeks of manual digging.

6.2.1 Pulling a system narrative out of code and docs

Given a repository (or monorepo) plus a handful of key documents, Claude can:

- Identify major services, modules, and domains
- Trace end-to-end flows (e.g., “user sign-up,” “invoice generation,” “order fulfilment”)
- Highlight how data moves between components and where external dependencies sit

Architects have reported using Claude to turn multiple repos and scattered docs into a single, coherent description of a system in minutes, which would otherwise take a human hours to write.

A typical workflow:

1. Point Claude Code at the relevant repos.
2. Provide any existing architecture docs, ADRs, and sequence diagrams via a Project.
3. Ask Claude to:
 - List services and their responsibilities
 - Explain key flows in plain language
 - Flag obvious coupling or duplication

You then validate and refine this picture, but the initial mapping comes much faster.

6.2.2 Keeping architecture documentation alive

One of the hardest problems in architecture is **documentation rot**. Designs change, but docs don't. Claude can help you keep docs in sync by:

- Generating updated diagrams and descriptions after major changes
- Comparing current code to existing ADRs and calling out divergence
- Drafting new ADRs when big trade-offs are made

For example, some teams use Claude to regularly scan code and PRs to spot when implementation drifts from documented intent, prompting architects to either update the design or call for a refactor.

6.3 Using Claude as an architecture reviewer

Claude is particularly useful as a **first-pass reviewer** of proposed designs.

6.3.1 Reviewing system designs and RFCs

Given an RFC or architecture proposal, Claude can:

- Summarise the proposal and its main decisions
- Identify assumptions and missing constraints

- Evaluate against non-functional requirements like scalability, security, and maintainability
- Suggest alternative patterns (e.g., event-driven vs request-response, database per service vs shared schemas)

Real-world examples include using Claude to review design docs for data platforms, microservice decompositions, and payment systems, often highlighting issues like misuse of caches as sources of truth or inconsistent failure handling.

Claude is especially effective if you:

- Provide explicit **evaluation criteria** (e.g., “optimise for latency and operability, under budget X, for team Y’s skill level”)
- Ground it in your **current architecture docs** and constraints via a Project
- Ask for **multiple options** with pros and cons, not just a single verdict

6.3.2 Checking decisions against your principles

Many architecture teams define principles like “event-first,” “APIs before UI,” or “prefer managed services.” Claude can help enforce these by:

- Comparing a proposal against your documented principles
- Highlighting where a design deviates and whether it’s justified
- Suggesting principle-aligned alternatives

You can formalise this into a “reviewer skill” or template (some teams define a dedicated “architecture reviewer” agent configuration for Claude Code).

6.4 Supporting implementation: from design to code

Claude is not just for whiteboard work; it can help ensure the implementation stays faithful to the design.

6.4.1 From diagrams and ADRs to scaffolding

Once a design is approved, Claude Code can:

- Generate initial scaffolding for services, modules, and interfaces that match your architecture guidelines
- Insert TODOs and comments that encode the design and responsibilities
- Ensure consistent patterns across services (e.g., API structure, error handling, logging)

This is especially valuable in large teams, where multiple developers implement parts of the same design. Claude can help keep them aligned with the architecture vision.

6.4.2 Ensuring code stays aligned over time

Over time, real-world constraints and “just ship it” pressures can erode a design. Claude can act as a watchdog:

- Periodically analyse code to see whether new patterns or anti-patterns are emerging
- Compare implementation against ADRs and call out inconsistencies
- Propose refactoring plans where architecture drift has become serious

Some teams use Claude Code sub-agents configured specifically for “architecture health checks,” focusing on concerns like coupling, boundaries, and shared state.

6.5 Onboarding and knowledge transfer

Architects are often bottlenecks for knowledge. Claude can help you **scale your expertise**.

6.5.1 Creating onboarding guides

Using your docs, ADRs, and codebase, Claude can:

- Generate “System 101” guides for new engineers
- Create role-specific onboarding paths (e.g., what a new backend engineer vs SRE should learn first)
- Draft quizzes or exercises to assess understanding

These artifacts can live in Projects so new hires can ask follow-up questions and see consistent answers grounded in the same context.

6.5.2 Answering “how does this work?” at scale

With a well-maintained architecture Project and good CLAUDE.md files in key repos, Claude can answer many “how does this work?” questions that would normally land in an architect’s inbox:

- “Where is user identity validated in the checkout flow?”
- “Which services depend on the Orders DB?”
- “What happens if the payment provider times out?”

Architects still handle complex design decisions and trade-offs, but Claude takes on the repetitive explanations.

6.6 Adoption roadmap for architecture & platform teams

Just like with dev and QA, architecture adoption should be phased.

Week 1 – Single system review

- Choose one system or major feature.
- Create an **Architecture Project** with existing docs, diagrams, ADRs, and key repo links.
- Use Claude to produce:
 - A high-level system overview
 - A list of current risks and unknowns
 - A critique of an upcoming or recent design

Treat this as an experiment and note where Claude adds value vs where it misunderstands context.

Week 2 – Define architecture review patterns

- Based on week 1, design a **standard architecture review prompt** or “reviewer skill” that encodes:
 - Your principles and constraints
 - The questions Claude should always ask (scalability, security, failure modes, data ownership)
- Start using this pattern for new RFCs, while still doing human reviews as normal.

Week 3 – Embed in design and onboarding processes

- Update your **RFC template** to include a “Claude review” section, where the author or reviewer pastes the AI’s summary and critique.
- Use Claude to generate or refresh onboarding docs for at least one major system.

Week 4 – Integrate into architecture governance

- Add Claude-assisted reviews to your **architecture review board** flow, especially for medium-sized changes.
- Define clear guidelines for when Claude is advisory vs when additional human or external review is required (e.g., security-critical systems).

From here, you can scale to more systems and perhaps introduce periodic “architecture health checks” using Claude Code and Projects.

6.7 Security & governance considerations for architects

Architects often deal with **sensitive system details**, so governance matters.

Key points:

- Keep Architecture Projects **scoped**: one per system or domain, with access limited to relevant engineers and leaders.
- Avoid uploading highly sensitive config (keys, secrets, internal IPs, security playbooks); keep those in secure systems and refer to them abstractly.
- When using Claude Code, be extra cautious with infra repos and scripts; review commands and changes carefully, and prefer analysis over direct edits for critical components.
- Document where AI was used in important design decisions, so that reviewers and auditors have a trail.

Chapter 13 goes deeper into organisation-wide security and governance; as an architect, you'll likely help shape those policies.

6.8 Prompt patterns and examples for architects (overview)

Finally, here are some prompt patterns you can standardise within your architecture practice:

- **System mapping**
"Act as a principal engineer. Based on the files in this repo and the attached architecture docs, describe the current system architecture. List the main services, their responsibilities, and the key data flows between them. Call out any obvious tight coupling or single points of failure."
- **Design review**
"You are an architecture reviewer. Review this RFC against our principles (attached) and non-functional requirements. Identify risks, trade-offs, and missing considerations. Suggest at least two alternative approaches and when they might be preferable."
- **Architecture health check**
"Scan this service and its dependencies for architecture smells: shared mutable state, direct DB access from multiple services, tight coupling to external APIs, or missing failure handling. Summarise your findings and propose a phased refactor plan."

- **Onboarding pack**

“Create a ‘System 101’ onboarding guide for new backend engineers joining this team. Include overview, key services, critical paths, failure modes, and where to find up-to-date docs and runbooks.”

These patterns turn Claude into a consistent architecture assistant that complements, rather than replaces, your expertise.

Used well, Claude helps architects and senior engineers **see more, sooner**: it surfaces risks, keeps docs alive, and spreads system understanding across the team. In the next chapter, we’ll move from technical design to delivery, and explore how project managers and Scrum leads can use Claude to plan, track, and communicate work more effectively.

Chapter 7 – Claude for Project Managers & Scrum Leads

Project managers and Scrum leads live in a constant flow of inputs: stakeholder demands, changing priorities, shifting timelines, and endless meetings. A surprising amount of this work is **information choreography**—turning chaos into clear plans, decisions, and updates.

Claude is well-suited to this job. When you combine **Claude Projects** with good PM habits, the assistant can help you plan, track, and communicate more effectively—without taking responsibility away from you as the human lead.

7.1 Where Claude fits in the PM & Scrum workflow

Typical PM/Scrum responsibilities include:

- Translating strategy and stakeholder requests into backlogs and plans
- Facilitating ceremonies (stand-ups, sprint planning, reviews, retros)
- Managing risks, dependencies, and scope
- Communicating status to different audiences

Claude can support you in each of these by:

- Turning fuzzy goals into structured plans and milestones
- Summarising meetings and extracting action items
- Creating status reports tailored to executives, teams, or customers
- Helping you think through risks, trade-offs, and alternative scenarios

The idea is not to make Claude “the project manager,” but to let it handle the heavy lifting on **text, structure, and synthesis** so you can focus on relationships, decisions, and trade-offs.

7.2 Using Claude Projects as your PM workspace

For PMs and Scrum leads, **Claude Projects** should be the default way to work, not just ad-hoc chats.

7.2.1 Creating a project workspace per initiative

Best practice is to create one Claude Project for each major initiative, product, or program.

Within each Project, you can store:

- Project charter and objectives

- High-level requirements and PRDs
- Stakeholder lists and communication plans
- Backlog exports (e.g., Jira/JQL, Azure Boards, Trello)
- Risks, assumptions, issues, and decisions registers
- Templates for status reports, meeting agendas, and minutes

Guides for PMs recommend using Projects as a central AI-aware workspace where Claude always sees the same context and instructions.

7.2.2 Setting custom instructions for your methodology

Claude becomes much more useful when you tell it **how you manage projects**.

In each Project's instructions, you can define:

- Your methodology (Scrum, Kanban, SAFe, hybrid, or Waterfall)
- How you structure epics, stories, tasks, and bugs
- How you want status updates formatted (e.g., RAG format, bullet points vs narrative)
- The tone you need for different stakeholders (executive vs team vs client)

For example:

"You are a senior project manager working on a B2B SaaS platform. We use Scrum with 2-week sprints. Backlog items are written as 'As a [user], I want [goal] so that [value].' Status reports are in RAG format, and executives prefer concise summaries. Always ask for clarification if requirements are ambiguous."

This turns Claude into a **context-aware PM copilot** instead of a generic writing assistant.

7.3 Planning with Claude: charters, milestones, and backlogs

Claude is particularly helpful in **turning unstructured ideas into structured plans**.

7.3.1 From strategy and notes to project charters

Given raw inputs—emails, strategy decks, meeting notes—Claude can help you:

- Draft a project charter with goals, scope, success metrics, timelines, and constraints
- Identify missing information and follow-up questions
- Suggest a high-level work breakdown

For example, you can paste a CEO email plus your notes and ask:

“Based on this, draft a one-page project charter. Include: background, objectives, scope, out of scope, key stakeholders, success criteria, timeline assumptions, and high-level risks. Highlight any ambiguities that need clarification.”

You then refine this charter with stakeholders—but you no longer start from a blank page.

7.3.2 Breaking work into milestones, epics, and stories

Claude can also assist with structuring work:

- Turning objectives into milestones and deliverables
- Proposing epics and candidate stories for each milestone
- Suggesting acceptance criteria based on your templates

This works best when you:

- Provide examples of existing epics/stories in your preferred format
- Clarify constraints (team velocity, fixed dates, known dependencies)

The PM remains responsible for balancing scope, timelines, and capacity; Claude accelerates the drafting.

7.4 Running ceremonies and meetings with Claude

Many PMs find that meetings and ceremonies are where Claude saves them the most time.

7.4.1 Agendas and prep

Claude can draft and refine:

- Meeting agendas tailored to specific outcomes
- Talking points and questions for stakeholder interviews
- Refinement checklists for backlog grooming sessions

For example:

“Draft a 45-minute sprint planning agenda for a team of 8 engineers and 2 QAs. We have a 2-week sprint and a backlog attached. Include time for capacity planning, story clarification, and risk discussion. Suggest facilitation tips.”

You can iterate with Claude until the agenda fits your reality.

7.4.2 Summarising meetings into clear outcomes

If you paste meeting notes or a transcript into the relevant Project, Claude can:

- Summarise decisions and action items
- Assign owners and due dates from context

- Highlight unresolved questions and risks

This is especially useful for sprint reviews, retrospectives, and cross-team dependency meetings.

You might ask:

“Summarise this meeting in 10 bullet points: decisions, action items (with owner and due date), risks, and follow-ups. Keep it factual and neutral.”

You still sanity-check and adjust, but you don't have to build the summary from scratch.

7.5 Reporting: status, RAG, and stakeholder comms

Status reporting is a natural fit for Claude, especially when combined with Projects and connectors.

7.5.1 Generating weekly status reports

Claude can generate status reports by pulling together:

- The latest notes and meeting summaries in the Project
- Backlog exports or sprint boards (if you paste or connect them)
- Risk and issue logs

Anthropic's examples and third-party guides show how to create **custom skills** for recurring tasks like status report generation—essentially templates plus logic that you can reuse.

A typical workflow:

1. Each week, export relevant Jira or board data and paste or sync it into the Project.
2. Ask Claude to generate a status report in your template, including:
 - Overall RAG status and rationale
 - Progress against milestones
 - Scope changes
 - Updated risks and issues
 - Next steps and asks from stakeholders
3. Review and edit before sending.

This can reduce status report creation from an hour to minutes without losing nuance—especially if your inputs are consistent.

7.5.2 Tailoring communication by audience

Claude is also good at **audience-specific messaging**. From the same underlying status, it can generate:

- A concise executive summary (1–2 paragraphs, focus on outcomes and risks)

- A detailed team update (stories done, blockers, upcoming work)
- A client update (plain language, emphasising value and impact)

The key is to define your audience tone and requirements in the Project instructions, then ask Claude to produce multiple versions.

7.6 Risk management and “what-if” analysis

Good PMs are always thinking, “What can go wrong?” Claude can help structure that thinking.

7.6.1 Identifying and categorising risks

Based on your charter, backlog, and constraints, Claude can:

- Propose a risk register with likelihood/impact ratings
- Group risks by themes (scope, schedule, technical, resource, external)
- Suggest early warning indicators and mitigation actions

For example:

“From this charter and initial backlog, identify the top 15 risks. For each, give a short description, likelihood (low/medium/high), impact, early warning signs, and proposed mitigations. Follow our risk log format.”

You review, adjust ratings, and decide which risks are real vs theoretical.

7.6.2 Scenario planning

Claude can also be used for **scenario analysis**:

- “What happens if this launch date moves by four weeks?”
- “What if we lose 30% of our capacity next sprint?”
- “What if this dependency slips by one quarter?”

Claude can help you think through impacts on milestones, stakeholders, and scope, and suggest options like scope reduction, parallelisation, or descope non-critical features.

Again, you decide which options are viable—but Claude speeds up the thinking.

7.7 Adoption roadmap for PMO and Scrum teams

As with other roles, the adoption roadmap for PMs should be pragmatic.

Week 1 – One project, one PM

- Pick one active project and one PM or Scrum lead.
- Create a Claude Project with the charter, PRD, stakeholder list, and current status reports.

- Use Claude to draft one status report and one set of meeting notes that week.
- Measure time saved and quality vs your manual baseline.

Week 2 – Templates and instructions

- Refine Project instructions to reflect your PM style and templates.
- Create standard prompts for:
 - Weekly status
 - Meeting summaries
 - Risk identification

Week 3 – Extend to more ceremonies

- Start using Claude to help with sprint planning, backlog grooming, and stakeholder communication.
- Encourage the PM to capture friction points and effective patterns.

Week 4 – Share patterns and scale

- Document “How we use Claude in project management” as a short internal guide.
- Onboard 2–3 more PMs or Scrum leads with a short training and a shared set of templates.
- Start integrating Claude usage into your PMO playbooks.

7.8 Security & governance notes for PMs

PMs often see sensitive information: strategy, budgets, people issues. You need to be careful about what flows into Claude.

Key practices:

- Use Team/Enterprise deployments with **SSO and admin controls**, not personal accounts, for project work.
- Avoid uploading HR documents, performance reviews, or sensitive legal/contract details into Claude. Abstract them when possible.
- Anonymise or mask client names and confidential metrics in inputs, especially when not strictly necessary to the task.
- Keep Projects scoped: do not mix highly sensitive streams (e.g., M&A, HR) with regular delivery work in a single Project.

Chapter 13 will give you a fuller governance framework; PMs are often key partners in enforcing it.

7.9 Prompt patterns and examples for PMs & Scrum leads (overview)

Here are a few patterns you can standardise across your PM community:

- **Project charter drafting**

"You are a senior project manager. Based on these notes and this strategy email, draft a one-page project charter. Include: background, objectives, scope, out of scope, success metrics, key stakeholders, and initial risks. Highlight open questions."

- **Sprint planning support**

"Using this backlog export and team capacity (8 devs, 2 QAs, 2-week sprint), propose a draft sprint plan. Include candidate stories, estimated points, and a rationale. Flag any risks or overcommitments you see."

- **Status report generation**

"Generate a weekly status report in this template. Use the latest meeting notes, backlog export, and risk log in this Project. Include: RAG status with explanation, summary of progress, key risks and mitigations, and next week's priorities."

- **Risk register creation**

"From this charter and backlog, create a risk register with 15 items. For each, include description, category, likelihood, impact, early warning signs, and mitigation. Use our risk matrix format."

These prompts encode your methodology and expectations, turning Claude into a consistent PM partner that supports your way of working rather than imposing its own.

Used this way, Claude becomes a **force multiplier for project managers and Scrum leads**: it keeps everyone on the same page, frees you from low-value admin, and gives you more time for stakeholder alignment and decision-making.

In the next chapter, we move into more specialised territory and explore how SAP, functional consultants, and business analysts can use Claude to streamline documentation, impact analysis, and client communication.

Chapter 8 – Claude for SAP, Functional Consultants & Business Analysts

SAP consultants, functional leads, and business analysts are constantly translating between business language and system language. You interpret processes, configure systems, draft functional specs, document impacts, and explain changes to clients and stakeholders.

Claude fits naturally into this “translation” work. It can read SAP documentation, process maps, functional specs, and change requests, then help you design scenarios, impact analyses, and clear client-ready documents. Used with secure integrations—like SAP BTP, MCP connectors, or middleware—it can even pull live SAP data into that reasoning.

8.1 Where Claude fits in SAP and functional work

Typical activities for SAP/functional consultants and BAs include:

- Understanding business processes and mapping them to SAP or other systems
- Drafting functional specifications and configuration documents
- Preparing test cases and UAT scripts with business users
- Performing impact analysis for change requests and rollouts
- Writing training materials, SOPs, and client-facing documents

Practitioners already use AI (including Claude) to:

- Turn messy meeting notes or process descriptions into structured specs and user stories
- Brainstorm edge cases and variations in complex processes (e.g., order-to-cash, procure-to-pay)
- Summarise SAP configuration, tables, or Fiori apps in business language
- Generate training content and one-pagers for end users

Claude’s long-context abilities make it especially useful when you have *many* documents—blueprints, customizing guides, interface specs—that all need to be reconciled.

8.2 Drafting functional specs, config docs, and impact analyses

One of the highest-value uses of Claude for SAP/functional roles is **documentation**.

8.2.1 From process notes to functional specs

Given workshop notes, screenshots, and process maps, Claude can help you produce structured functional specs:

- Clarify scope, in and out
- Describe current vs target process flows
- List business rules and validations
- Describe integration points and data sources
- Define error handling and exception paths

You can load your notes and existing template into a Project and ask:

“Based on these workshop notes and this functional spec template, draft a functional specification for the ‘Returns & Refunds’ process in SAP S/4HANA. Use business-friendly language and clearly separate ‘as-is’ and ‘to-be’ behaviour.”

You then adjust for correctness and system details; Claude handles most of the structuring.

8.2.2 Configuration documentation and field-level descriptions

Functional consultants often need to document configuration decisions: which options are chosen, why, and how they affect the process. Claude can:

- Turn SPRO screenshots or config tables into clear descriptions
- Explain the business implications of specific config options in plain language
- Generate comparison tables for different configuration scenarios

Combined with live or exported SAP configuration data (secured via BTP, MCP, or other bridges), Claude can generate up-to-date config documentation during or after a project.

8.2.3 Impact analysis for change requests

When changes come in—new fields, new pricing rules, new approval steps—you must quickly assess impact. Claude can help by:

- Reading the change request and existing specs
- Identifying which processes, interfaces, roles, and reports may be affected
- Proposing questions to clarify the scope and risk
- Drafting an impact analysis document for review

For example:

“Based on this CR and the existing ‘Order to Cash’ functional spec, identify the impacted processes, SAP modules, Fiori apps, and integrations. Draft an impact analysis including risks, required testing, and data migration considerations.”

You validate and refine, but the initial thinking is accelerated.

8.3 Using Claude for test cases, UAT, and process scenarios

Business-facing testing—especially UAT—requires translating business flows into concrete scenarios. Claude can support this across SAP and non-SAP systems.

8.3.1 Scenario and test case generation

From requirements, process maps, and system descriptions, Claude can:

- Generate test scenarios and UAT scripts in business language
- Suggest positive, negative, and edge cases
- Map steps to specific SAP transactions, Fiori apps, or fields (when given that context)

For example:

“Using this functional spec and process map for ‘Procure to Pay’, generate UAT scenarios for business users. For each scenario, include preconditions, SAP roles, transactions or Fiori apps, step-by-step actions, expected results, and data variations.”

This is especially helpful when you need to prepare materials for non-technical business testers.

8.3.2 End-to-end process coverage

Claude can also help ensure you don’t miss **cross-module** scenarios:

- Order-to-cash with credit limits, discounts, and returns
- Procure-to-pay with multiple approvals and GR/IR variances
- Hire-to-retire with edge cases in HR or payroll

By feeding in process documentation and related specs, you can ask Claude to list complete end-to-end flows and highlight where additional tests are needed.

8.4 Integrating Claude with SAP and enterprise data (safely)

Beyond static docs, many SAP teams are starting to connect Claude to live or near-real-time SAP data through secure bridges.

8.4.1 SAP BTP and AI Core

SAP Business Technology Platform (BTP) and SAP AI Core allow enterprises to access models like Claude from within the SAP ecosystem.

Typical patterns include:

- Building custom Fiori apps that call Claude to summarise purchase orders, sales orders, or invoices for users.
- Using Claude to analyse SAP data in Datasphere or other warehousing layers via secure connectors.
- Generating draft narratives for management reports directly from SAP data.

Because this happens inside SAP's managed environment, you benefit from existing controls, logging, and compliance frameworks.

8.4.2 Model Context Protocol (MCP) and desktop workflows

Anthropic's **Model Context Protocol (MCP)** and third-party MCP servers for SAP allow Claude (especially the desktop app) to query SAP systems from your laptop in a governed way.

Examples:

- Asking Claude: "Show me the last 5 sales orders for Customer X and summarise any credit or delivery issues," with data pulled live from S/4HANA.
- Exploring SAP Datasphere tables and views in natural language, then having Claude draft transformation logic or SQL for downstream analysis.
- Building AI agents that read SAP maintenance data and suggest work orders or optimisations.

These patterns require careful setup, but they allow business analysts and consultants to interact with SAP data conversationally while still respecting governance.

8.4.3 Middleware and low-code connectors

For teams without deep SAP AI expertise, middleware platforms provide connectors for both Claude and SAP.

Common use cases:

- When a new invoice or order is created in SAP, send it to Claude for validation or summarisation, then post a summary in Slack or Teams.
- Automatically generate exception reports or escalation emails based on SAP events.

These workflows are typically orchestrated in tools like Make, Workato, or similar, with Claude providing the language and reasoning layer.

8.5 Business analysis beyond SAP: insights, scenarios, and communication

Not all functional work happens inside SAP. Business analysts often work with spreadsheets, BI tools, and planning systems.

Claude can help BAs to:

- Summarise large datasets or dashboards into narratives and talking points
- Identify trends and anomalies described in plain language (after you provide relevant data summaries or exports)
- Map business processes and create visual workflows to explain them
- Prepare stakeholder-specific decks, one-pagers, and FAQs

Courses and guides on “Generative AI for business analysts” emphasise exactly these skills: using AI to automate repetitive analysis and documentation so analysts can focus on **interpretation and decision support**.

8.6 Adoption roadmap for SAP/functional and BA teams

Because these roles are close to both business users and core systems, adoption should be phased and governed.

Week 1 – Document-only pilots

- Start with **static documents**: existing specs, process descriptions, training materials.
- Use Claude Projects to draft or update functional specs, UAT scenarios, and training docs.
- Avoid any live SAP or sensitive data in this phase.

Week 2 – Structured analysis and test design

- Introduce prompt templates for:
 - Functional spec drafting
 - UAT script generation
 - Impact analysis for change requests
- Get 1–2 consultants/BAs comfortable with these patterns and collect feedback.

Week 3 – Exploring secure integrations (with IT/CoE)

- With your AI CoE and SAP/BTP teams, evaluate **BTP AI Core, MCP, or middleware** options for secure integration.
- Run a small POC: e.g., ask Claude to summarise a limited set of SAP orders or tickets via a secure connector.

Week 4 – Team training and governance

- Document “How we use Claude in SAP/functional/BA work,” including allowed and prohibited data types and use cases.
- Train a small group of consultants and BAs on both the opportunities and the guardrails.

From here, you can expand to more use cases as governance matures.

8.7 Security & governance notes for SAP and functional work

Because SAP systems often hold **highly sensitive business and personal data**, governance is critical.

Key points:

- Use enterprise integrations (SAP BTP, secure MCP servers, or managed connectors) rather than ad-hoc exports with raw production data.
- Avoid putting PII, financial details, or confidential contracts directly into chats; when needed, anonymise or mask data.
- Make sure your use of Claude aligns with both your organisation’s and SAP’s data privacy and AI policies.
- Log and monitor AI interactions that touch live SAP data, so that you can audit and improve over time.

Chapter 13 extends these ideas into a full governance framework; SAP and BA teams will be key stakeholders.

8.8 Prompt patterns and examples for SAP consultants & BAs (overview)

Finally, here are prompt patterns tuned to SAP/functional and BA work:

- **Functional spec drafting**

"You are a senior SAP SD functional consultant. Based on these workshop notes and the attached template, draft a functional specification for the new discount approval workflow. Clearly describe as-is vs to-be, affected transactions/Fiori apps, business rules, and error handling."

- **Impact analysis**

"Act as a lead SAP FI/CO consultant. Given this change request and the existing 'Order to Cash' spec, identify impacted processes, modules, interfaces, and reports. Draft an impact analysis including risks, required configuration changes, testing scope, and data migration considerations."

- **UAT scenarios**

"You are a business analyst preparing UAT for a new procurement approval process in S/4HANA. Using this process description and config overview, generate UAT scripts in business language. Include roles, steps in SAP, expected results, and variants (e.g., rejection, escalation, exception cases)."

- **SAP data Q&A via secure integration**

"Using the connected SAP Datasphere data, analyse the last 90 days of late deliveries. Summarise top 5 root causes, affected customers, and suggest follow-up analysis for the supply chain team."

These patterns turn Claude into a practical copilot that respects SAP realities and business constraints, while dramatically reducing the time you spend on heavy documentation and repetitive analysis.

Used thoughtfully, Claude helps SAP and functional consultants and business analysts **bridge business and systems faster**—without sacrificing accuracy, governance, or client trust.

Chapter 9 – Claude for Sales & Business Development

Modern B2B sales teams juggle research, outreach, discovery, proposals, and follow-up across dozens of accounts. Much of this work is **information heavy**: understanding companies, tailoring messaging, and keeping playbooks consistent.

Claude can act as an **AI sales analyst and copy partner**—researching accounts, drafting personalised outreach, preparing call briefs, and helping you keep your sales process structured. You still own targeting, judgment, and relationships; Claude takes over the heavy reading and writing.

9.1 Where Claude fits in the sales workflow

A typical B2B sales cycle involves:

- Account and contact research
- Qualification (BANT, MEDDIC, or similar frameworks)
- Outreach sequences across email, LinkedIn, and calls
- Discovery and demo preparation
- Proposals, follow-ups, and renewals

Sales teams are already using Claude and similar tools to:

- Turn raw research into concise account briefs
- Qualify leads against frameworks like BANT or MEDDIC
- Generate personalised outreach sequences at scale
- Prepare discovery call questions and objection-handling scripts
- Capture and summarise notes from calls and meetings

Claude's long-context and Projects features make it particularly useful when you want your entire sales motion—ICP, messaging, playbooks—to be **baked into** your AI assistant.

9.2 Building a sales “brain” with Claude Projects

For sales and BD teams, the most powerful pattern is to create a dedicated **Sales Project** (or multiple Projects by segment/region) that becomes your shared sales “brain.”

9.2.1 What to put into your Sales Project

Sales leaders and operators who've done this successfully typically load their Project with:

- ICP and segment definitions

- Product one-pagers and feature sheets
- Battlecards and competitor comparisons
- Past successful emails and sequences
- Call scripts, qualification frameworks (BANT, MEDDIC, SPICED, etc.)
- Case studies, testimonials, and pricing guidelines (in high-level form)
- Brand and tone of voice guidelines

You then set **Project-level instructions** so Claude knows it is “part of” your sales team: how you talk, who you sell to, what problems you solve, and what you never promise.

This means that when reps ask Claude for an email, objection response, or call plan, it responds in your voice, informed by your playbook—not generic sales advice.

9.2.2 Encoding your process and frameworks

You can also encode your sales methodology:

- How you define a qualified opportunity
- Your definitions of stages and exit criteria
- Questions you want to cover for BANT or MEDDIC
- How you handle pricing and discount conversations

Some teams even create **Claude Skills** (reusable prompt files) like “Sales Playbook Skill” or “Account Research Skill” that automate repeatable steps such as creating one-page account playbooks and status snapshots.

9.3 Account and contact research

Claude is not a replacement for data providers (LinkedIn, Clearbit, Databar, etc.), but it can **make sense** of research data and turn it into actionable insights.

9.3.1 Summarising accounts and finding angles

Given raw inputs—website copy, blog posts, job listings, funding news—Claude can:

- Summarise what the company does, who its customers are, and how it makes money
- Infer likely tech stack, GTM motion, or pain points from public signals
- Identify “hooks” and trigger events (new funding, hiring, product launches, leadership changes)

There are Claude Code skills and open-source tools that orchestrate this: they crawl sites, pull job postings, and produce a structured company profile that sellers can quickly scan.

A typical workflow:

1. Use your research tool (or a Claude Code skill) to gather raw company data into a file or CRM.

2. Ask Claude in your Sales Project to produce a concise account brief with:
 - Company overview and ICP fit
 - Key initiatives and likely KPIs
 - Hypotheses about pains you can solve
 - Recommended messaging angles

This reduces hours of scattered research into one structured brief you can trust and refine.

9.3.2 Qualification and opportunity notes

Claude can help you structure your thinking against frameworks:

- Turn rough notes into BANT or MEDDIC-formatted fields
- Highlight missing information you still need to gather
- Suggest follow-up questions to ask in the next meeting

For example:

“Using these call notes and our MEDDIC template in this Project, fill out a MEDDIC summary: Metrics, Economic buyer, Decision criteria, Decision process, Identified pain, Champion. Highlight any gaps and suggest questions for the next call.”

You remain responsible for correctness and judgment; Claude keeps your documentation tight and consistent.

9.4 Personalised outreach at scale

One of the most visible uses of Claude in sales is **cold outreach and follow-up sequences**. Done badly, this becomes spam. Done well, it becomes **highly personalised at scale**.

9.4.1 Building your own template engine

Rather than feeding Claude a single prompt and hoping for the best, teams that see strong results build a **template engine** with Claude Code and Projects:

- Create a small set of base templates (problem-opening, trigger event, competitor displacement, referral follow-up, webinar follow-up, etc.).
- Define your voice and tone guidelines (words to use/avoid, formality, length).
- Give Claude structured context: company data, persona, trigger event, product value proposition.

Then instruct Claude to:

- Fill templates with context-specific details
- Enforce quality checks (no clichés, no false claims, no over-promising)
- Output CSVs or structured files that plug into your sequencing tool (Apollo, Instantly, HubSpot, etc.).

This is the approach described in step-by-step guides where Claude Code builds hundreds of customized emails from a contact list, while still keeping quality review in the loop.

9.4.2 Multi-channel sequences and follow-ups

Claude can generate not just email copy, but full **multichannel sequences**:

- First touch: email + LinkedIn connection message
- Second touch: follow-up referencing a recent event or content
- Third touch: short value nugget or case study
- Final touch: polite breakup email

Guides for Claude-based outreach automation emphasise:

- Using **real signals** (funding, hiring, tech stack, content engagement) to drive personalisation
- Keeping messages **short, clear, and not overly “AI”** in tone
- Avoiding sending anything without at least spot-checking and A/B testing performance

Claude can also create variations for A/B tests and help you interpret results.

9.5 Call preparation, discovery, and follow-up

Claude is also strong at supporting **live conversations**—before and after calls.

9.5.1 Call prep and discovery frameworks

Given an account brief and opportunity context, Claude can:

- Draft call agendas with clear objectives
- Propose discovery questions tailored to the prospect’s role and industry
- Suggest how to position your product against their likely pains

You might ask:

“For this upcoming discovery call with a VP of Sales at a Series B SaaS company (see account brief and product info in this Project), propose a 30-minute agenda, 15 discovery questions, and 3 ways to frame our value proposition based on their hiring and tech stack signals.”

You refine this with your own knowledge, but you don’t start from scratch.

9.5.2 Summarising calls and defining next steps

After the call, Claude can help you quickly turn transcripts or notes into:

- CRM-ready summaries
- Key points by MEDDIC/BANT fields
- Action items, owners, and timelines
- Follow-up email drafts recapping what you heard and next steps

This reduces time spent on admin and ensures better alignment across sales, CS, and product teams.

9.6 Customer success, renewals, and expansions

Sales does not stop at the closed-won stage. Claude can help with **ongoing account management**.

9.6.1 QBR and renewal preparation

Given usage reports, support tickets, and prior notes, Claude can:

- Summarise value delivered
- Identify expansion opportunities based on usage patterns and feedback
- Draft QBR decks or talking points
- Suggest scripts for addressing known risks before renewal

This mirrors how business analysts use Claude; here, the focus is on **value storytelling and strategy**.

9.6.2 Risk detection and playbooks

With structured data and notes, Claude can help you:

- Flag accounts at risk based on usage, sentiment, or support patterns (when integrated via secure connectors and tools).
- Suggest playbook steps based on your retention strategies (e.g., stakeholder mapping, ROI review, executive alignment).

Again, Claude is not deciding which customers are at risk; it's **surfacing patterns** and suggesting actions that you validate.

9.7 Adoption roadmap for sales and GTM teams

To avoid “AI spray and pray,” roll Claude into sales in phases.

Week 1 – Sales leader + one segment

- Sales leader and one AE/BDR pick a single segment (e.g., mid-market SaaS).
- Create a Sales Project and load ICP, messaging, and example emails.
- Use Claude to generate account briefs and draft outreach for a **small batch** of accounts.
- Manually review every output and send only what you’re confident in.

Week 2 – Template engine and repeatable workflows

- Build base templates and a voice guide (or adapt existing ones).
- Define standard prompts/skills for:
 - Account research
 - Discovery prep
 - Sequence generation
- Integrate Claude output with your sequencing tool via CSV or simple scripts.

Week 3 – Enable more reps

- Train a small group (3–5 reps) on the Sales Project and workflows.
- Establish **QA rules**: what must be reviewed manually, what can be lightly spot-checked only.
- Start capturing performance metrics (opens, replies, meetings booked) per template.

Week 4 – Scale and refine

- Iterate on templates and prompts based on performance and rep feedback.
- Consider more automation (Claude Skills, Claude Code agents, CRM integrations) for steps that are working well.
- Document your “AI in Sales” guidelines so new reps know how to use Claude responsibly.

9.8 Security, compliance, and ethics in sales

Sales uses a lot of **personal and company data**, so you must use Claude responsibly.

Key considerations:

- Avoid putting sensitive personal information (non-public details, HR information) into prompts. Work with publicly available data and CRM data your company has a right to use.
- Do not fabricate or exaggerate case studies, results, or customer names; instruct Claude explicitly not to invent proof points or references.
- Make sure your use of AI in outreach respects regulations (GDPR, CAN-SPAM, local privacy laws) and your company's legal guidance.
- Keep your Sales Projects scoped by team/segment and control access via your enterprise settings, not personal accounts.

Ethical AI in sales is partly about compliance—but also about **maintaining trust** with prospects and customers.

9.9 Prompt patterns and examples for sales & BD (overview)

Here are some prompt patterns you can standardise:

- **Account brief**
“You are an SDR for a B2B SaaS company. Based on this website, LinkedIn page, and recent news (attached), create an account brief: who they are, ICP fit, likely pains we solve, relevant trigger events, and 3 messaging angles. Keep it to one page.”
- **Outreach sequence generation**
“Using our ICP, voice guide, and templates in this Project, write a 3-email outbound sequence for this contact. Each email must reference a real signal from the account data, stay under 100 words, and avoid clichés and false claims.”
- **Discovery prep**
“For this upcoming discovery call (see account brief and product overview), propose a 30-minute call plan: agenda, discovery questions tailored to a VP of Operations, and 3 ways to frame our value proposition based on their recent hiring and tech stack.”
- **Post-call follow-up**
“Turn these call notes into: (1) a concise CRM summary, (2) an internal recap for

the team, and (3) a follow-up email to the prospect recapping their pains, what we proposed, and next steps.”

These patterns give reps a consistent, powerful assistant embedded in your sales motion rather than a one-off copy generator.

Used this way, Claude helps sales and business development teams **do more of the right work**—better research, sharper messaging, cleaner follow-ups—without sacrificing authenticity or control.

Chapter 10 – Claude for Your Career Growth

Claude is not only a tool for your company; it can also be a career coach, tutor, and practice partner. Used intentionally, it helps you plan your path, learn faster, build a stronger portfolio, and prepare for interviews—without losing your own judgment or voice.

In this chapter, we'll treat Claude as a personal teammate: a system you can configure once and then reuse as you grow, rather than a one-off "homework helper."

10.1 Setting up Claude as your personal career coach

The first step is to give Claude enough context to help you **as an individual**, not just as a generic professional.

10.1.1 Create a personal Career Project

Anthropic's own guidance recommends creating a dedicated Project for your career planning. In that Project, upload:

- Your current resume or CV
- Links or exports from LinkedIn and any portfolio sites
- A few job descriptions that represent your target roles
- A list of your skills, certifications, and experience (even if it's rough)

Within the Project instructions, tell Claude who you are, what you've done, and what you're aiming for—role, industry, seniority, and time horizon.

This means every conversation in that Project automatically has your background and goals in mind; you don't have to re-explain yourself each time.

10.1.2 Define how you want Claude to behave

Claude becomes more helpful when you define a **style** for its responses. Anthropic suggests you can configure it to behave like a career counselor, coach, or mentor.

For example:

"In this Project, act as my career coach. Ask clarifying questions before giving advice. When I propose an idea, help me think through trade-offs. When I feel stuck, suggest 2–3 specific experiments I can run over the next week."

You can adjust this over time (more direct, more challenging, more encouraging) based on what actually helps you move forward.

10.2 Planning your career path with Claude

With your context loaded, Claude can help you **analyse where you are and where you want to go**, then turn that into a concrete plan.

10.2.1 Analysing your current profile against target roles

Claude can compare your resume and skills to target job descriptions and produce:

- A gap analysis: skills you have vs skills you need
- A prioritised list of capabilities to build in the next 6–12 months
- Suggestions for how to demonstrate those skills (projects, contributions, certifications)

Anthropic's "Plan your career path" workflow does exactly this: you share your resume and target jobs, and Claude creates both an action-oriented roadmap and a skills tracking system.

You might ask:

"Given my resume and these 3 target job postings, identify the top 10 skills and experiences I need to strengthen. Group them into: (1) technical skills, (2) domain knowledge, (3) soft skills. Then propose a 6-month learning and project plan."

10.2.2 Creating a living action tracker

Claude can then help you design a **skills log and action tracker**:

- A simple table of skills, current level, target level, and evidence
- Specific actions (courses, books, projects, contributions, networking) with due dates
- A way to check in weekly or monthly and update progress

You can keep these as spreadsheets or documents in the same Career Project, so Claude can refer to them when giving advice later.

10.3 Using Claude to learn faster and build portfolio projects

Claude is not a replacement for doing the work, but it is a very effective **tutor and project co-designer**.

10.3.1 Learning new topics with AI support

Anthropic's own teams report using Claude heavily for learning codebases and fixing errors; individuals do the same for general learning.

You can use Claude to:

- Explain concepts at the level you need ("explain this as if I'm a junior developer / new BA / PM")
- Recommend learning paths and resources tailored to your goals
- Quiz you on key ideas and give feedback on your explanations

A useful pattern is to:

1. Read or watch a resource.
2. Summarise what you understood to Claude.
3. Ask it to challenge your understanding, correct mistakes, and give next-step questions.

This makes learning more active and less passive.

10.3.2 Designing and executing portfolio projects

Career guides emphasise "AI-enhanced portfolios": real projects that show how you use tools like Claude in practice.

Claude can help you:

- Brainstorm project ideas that align with your target roles
- Scope the project into phases and milestones
- Draft documentation, readme files, or case studies as you go

For example, if you're a developer or data analyst, Claude Code can help you build a small app or analysis project faster, while you still own architecture and decisions. If you're a PM, BA, or marketer, Claude can help you structure case-study documents or mock project artifacts.

You can then ask Claude to help you turn these into **portfolio narratives**: why you did the project, what you learned, what the outcomes were.

10.4 Building your resume, LinkedIn, and personal brand

Claude is very strong at **structured writing and rewriting**, which is exactly what you need for resumes and online profiles.

10.4.1 Resume and LinkedIn optimisation

With your target roles in mind, Claude can:

- Rewrite your resume bullets to be more impact-focused (metrics, outcomes)
- Tailor your resume to specific job postings while preserving truth
- Suggest LinkedIn headline and About section variations aligned to your positioning

Anthropic's career templates and community "Career Helper" skills show how Claude can run through this systematically for job seekers.

The safe pattern is:

- You write the initial content.
- Claude refines phrasing, structure, and alignment to roles.
- You verify accuracy and adjust to keep your own voice.

10.4.2 Portfolio and personal website content

Creators are already using Claude to help build portfolio websites and write case-study-style project descriptions.

Claude can:

- Propose a site structure (sections, pages) for your portfolio
- Draft project pages given your notes and code or artifacts
- Help you keep tone consistent across your site and profiles

You still select which projects to highlight and which stories to tell; Claude handles the heavy writing and organisation.

10.5 Preparing for interviews with Claude

Interview preparation is one of the clearest personal use cases for Claude.

10.5.1 Analysing job descriptions and company context

Before an interview, Claude can:

- Analyse the job description and highlight what the company really cares about
- Suggest how your experience maps to those needs
- Help you research the company, product, and industry trends

Anthropic's career path guide explicitly recommends using Claude's web search and research capabilities to ground your preparation in up-to-date company and industry information.

You can ask:

"Given my resume and this job description, identify the top 5 themes they care about. For each, suggest one story from my experience I should be ready to tell."

10.5.2 Mock interviews and practice

People are already using Claude for **mock interviews**, including technical and behavioural sessions.

Effective patterns include:

- Sharing the job description, your resume, and any recruiter notes.
- Asking Claude to behave like an interviewer and ask one question at a time.
- Answering verbally (using voice input) to make practice more realistic.
- Asking Claude to critique your answers: clarity, structure (e.g., STAR), relevance, and missing details.

For technical roles, you can also use Claude Code as a practice partner for coding problems: it can generate questions, review your solutions, and suggest improvements.

10.6 A 30–60–90 day personal adoption plan

To make Claude a sustainable part of your career growth, treat it as a **long-term habit**, not a one-off gadget.

First 30 days – Setup and experiments

- Create your Career Project and upload your core materials.
- Do a first **career gap analysis** and action plan with Claude.
- Use Claude for at least one learning session per week and one document (resume, LinkedIn, or portfolio) update.

Days 31–60 – Projects and interview prep

- Start or continue one **portfolio-grade project** where you consciously use Claude as a helper.
- Use Claude for deeper learning (explanations, quizzing, reviewing your notes).
- If you're job-searching, do 2–3 mock interviews and refine your stories.

Days 61–90 – Integrate into your routine

- Set a weekly check-in with Claude in your Career Project: review progress, update your skills log, adjust your plan.

- Incrementally improve your public presence (portfolio, LinkedIn, GitHub, case studies) with Claude's help.
- Start thinking about how to showcase your **AI-augmented workflows** in interviews and performance reviews.

By the end of 90 days, Claude should feel like a **reliable, configured career tool** you regularly return to, not just something you tried once.

Used carefully, Claude can significantly accelerate your career development—by making planning, learning, portfolio building, and interview prep more structured and less overwhelming, while you stay in control of decisions and direction

PART III

Scaling Safely: Adoption, Governance, and Checklists



Using Claude alone is easy; using it across a company is hard. Part III gives leaders, platform teams, and the AI Centre of Excellence a concrete roadmap for rolling Claude out across the organisation—safely, measurably, and without chaos.

You'll define adoption waves, write clear AI usage policies, and put security, privacy, and governance controls in place. This is where Claude stops being a side project and becomes part of “how we work here,” with checklists you can drop straight into your internal wiki.

Chapters in this Part

- Chapter 12 – Building a Team Adoption Roadmap
- Chapter 13 – Security, Data Privacy, and Governance for Claude
- Chapter 14 – Implementation Checklists and Templates

Chapter 11 – Comparing AI Workflows: Standard Prompting vs Projects vs Claude Code

By this point, you have seen Claude used in many ways: quick chats, long-running Projects, and Claude Code in real repos. Under the hood, these are different workflows, each with strengths and limitations.

If you treat everything as “just chat,” you either underuse Claude or overload it with the wrong jobs. This chapter helps you choose **the right mode for the right task**, and shows how to combine them into end-to-end flows.

11.1 Three modes, three mental models

You can think of the main Claude modes like this:

- **Standard prompting (chat)** – “Ask a smart colleague a one-off question.”
- **Projects** – “Work in a shared, long-running workspace with all the relevant docs.”
- **Claude Code** – “Pair with a coding coworker inside your filesystem and tools.”

Each is built on the same underlying models, but they are optimised for different types of work.

11.1.1 Standard prompting (chat)

Standard prompting is what most people start with:

- A free-form chat interface
- You paste text or files and ask questions
- Claude responds based on that context and your instructions

It is ideal for:

- Brainstorming, idea generation, and quick questions
- Drafting or editing small pieces of content
- Simple code snippets or explanations

Limitations:

- Context is relatively short-lived and fragmented
- You repeat instructions and re-upload files
- Harder to build repeatable, team-wide processes

Standard prompting is the “**scratchpad**” of AI work: flexible, fast, but not where you run your core workflows.

11.1.2 Projects: persistent context and shared knowledge

Projects are persistent workspaces that store **files, instructions, and chats** around a specific initiative.

They excel when:

- You have a body of work (product, client, codebase, program) that you keep returning to
- Multiple people need consistent answers from the same context
- You want Claude to “remember” domain knowledge and instructions between sessions

Think of a Project as:

- A **knowledge base** for a product, team, or use case
- A place to encode background (docs, style guides, specs) and **how you want work done** (instructions, templates)

Compared to chat, Projects:

- Reduce hallucinations by grounding Claude in your data
- Make outputs more consistent across time and people
- Support more complex, long-running work (e.g., an entire product launch or architecture phase)

11.1.3 Claude Code: working where the code lives

Claude Code brings Claude into your editor/terminal and lets it see your **files, repo, and commands**.

It is designed for:

- Writing and refactoring code
- Running tests and commands
- Navigating and understanding large codebases
- Building automations and agents that interact with tools and APIs

Claude Code is where the AI stops being just a writer and starts being a **doer** in your development environment. It complements Projects by grounding in *runtime context* (files, outputs) rather than just static docs.

11.2 When standard prompting is enough

There are many tasks where simple chat is still the best option.

11.2.1 Low-stakes, one-off problems

Use standard prompting when:

- The task is small and self-contained (e.g., “rewrite this email,” “explain this error message,” “summarise this article”).
- You don’t expect to revisit the context often.
- The risk of a mistake is small and easy to catch by eyeballing the output.

For example:

- A PM asking for alternative phrasings of a slide title
- A BA asking for a quick explanation of a statistical concept
- A developer asking why a regex behaves a certain way in a small snippet

11.2.2 Fast exploration and “rough drafts”

Chat is also perfect for **exploratory thinking**:

- Brainstorm risks or ideas
- Explore pros and cons of options before deeper work
- Draft a rough outline before moving into a Project for refinement

As soon as something becomes **reusable, long-running, or multi-stakeholder**, it belongs in a Project or Claude Code flow instead.

11.3 When you should move into Projects

Projects are the backbone of serious, repeatable Claude usage.

11.3.1 Signals you need a Project

You probably need a Project when:

- You keep uploading the same documents.
- Multiple people ask Claude similar questions about the same topic.
- You have a long-running initiative where context keeps expanding (product, client, program).
- You care about consistency of voice, structure, and policy more than one-off creativity.

Common examples:

- “Product X – Roadmap & Docs” (PM + Eng + QA)
- “Client Y – SAP Rollout” (SAP + BA + PM + QA)

- “Sales – ICP + Messaging + Sequences” (Sales + Marketing)

11.3.2 What Projects give you that chat cannot

Compared to one-off prompts, Projects give you:

- **Persistent context:** you can load many files (docs, slides, specs, code snippets) and keep them around.
- **Shared configuration:** instructions, templates, and glossary terms apply to all chats in that Project.
- **Reduced friction:** you no longer waste time “reminding” Claude who you are and what you’re doing.

This transforms Claude from “a clever scratchpad” into a **team workspace**.

11.3.3 Projects vs Skills (process vs knowledge)

Anthropic and community guides emphasise that Projects and Skills do different jobs:

- **Projects** store *knowledge* for a specific area of work (docs, code, client info).
- **Skills** store *processes*—reusable instruction sets that tell Claude how to do something the same way every time, across any chat or Project.

Example:

- A Project for “Q2 Product Launch” with all briefs and plans.
- A “Status Report Skill” that knows exactly how to turn Project context into your standard report format.

We focus on Projects in this book, but Skills are the next layer when you want **full reuse** of your best prompts and workflows anywhere in Claude.

11.4 When Claude Code is the right tool

Claude Code is overkill for simple prompts—but irreplaceable when the work is code-centric.

11.4.1 Characteristics of Claude-Code-worthy tasks

Claude Code shines when:

- The work happens in a codebase or data folder, not just prose.
- You need to read or edit multiple files or run commands.
- The task can be described as “spec a behaviour, get an implementation,” especially with clear tests.
- You want repeatable workflows (routines, skills) that operate on your project structure.

Examples from practitioners include:

- Generating and refactoring tests in your project's framework
- Large-scale refactors across services
- Building internal tools and automations
- Setting up and maintaining AI agents that run on repos and external tools

11.4.2 When chat or Projects are not enough

You know you're misusing chat/Projects instead of Claude Code when:

- You keep copy-pasting large code snippets instead of letting Claude see the repo.
- You struggle to keep context in sync with the actual files.
- You want to run commands (tests, builds, linters) based on Claude's suggestions.

In those cases, moving to Claude Code gives you:

- Direct access to your files and directory structure
- The ability to apply diffs and run commands
- Better support for complex, multi-step coding tasks

11.5 Choosing the right workflow by role and task

Rather than thinking "which Claude feature do I like?", start from **role + task** and choose the minimal tool that works.

11.5.1 Developers & QA

- Quick code questions: **chat**
- Understanding a module/service, drafting docs: **Project + Claude Code**
- Implementing features, refactors, tests: **Claude Code**, optionally backed by a Project with PRDs and ADRs
- Debugging complex issues: **Claude Code** (with logs and notes as Project docs)

11.5.2 Architects

- Reviewing RFC drafts: **Project** (architecture docs + proposal)
- Checking code vs design: **Claude Code**, guided by instructions and ADRs referenced via Project
- Onboarding docs and system 101s: **Project**, with Claude pulling from repos and docs

11.5.3 PMs & Scrum leads

- Quick email rewrites: **chat**
- Full initiative workspace: **Project** (charter, backlog exports, notes)
- Status reports and risk logs: **Project + Skills** (reusable templates)

11.5.4 SAP, functional, BA

- One-off “explain this” Q&A: **chat**
- Client or process workspaces: **Projects** with specs, process docs, and test materials
- App- or data-level integration: **Claude Code + MCP/BTP** in advanced setups

11.5.5 Sales & GTM

- Rough copy brainstorming: **chat**
- ICP, messaging, and playbook brain: **Projects**
- Automated outreach and research pipelines: **Claude Code + Projects + Skills**

11.6 Example “day in the life” workflows

To make this concrete, here are two illustrative end-to-end flows.

11.6.1 Developer “day in the life”

1. **Morning** – Reads a new ticket in Jira. Uses chat to clarify language and draft questions for the PM.
2. **Design** – In the product’s Project, asks Claude to summarise the relevant PRD and ADRs and outline an implementation plan.
3. **Coding** – Opens Claude Code in the repo, loads CLAUDE.md, and asks it to implement phase 1 of the plan with tests.
4. **Review** – Uses Claude Code to explain the diff and highlight potential risks, then submits a PR for human review.
5. **Docs** – Asks Claude in the Project to update a small design note or README based on the PR.

Same model, three modes, each doing what it does best.

11.6.2 PM “day in the life”

1. **Morning** – In the initiative’s Project, asks Claude to summarise yesterday’s stand-up notes and Slack thread.
2. **Planning** – Drafts a sprint planning agenda and checklist based on the backlog export.

3. Reporting – Uses a Status Report Skill in the Project to create a weekly update, then edits and sends.

4. Risk review – Asks Claude to refresh the risk log given latest changes.

They might occasionally use chat separately (e.g., “rewrite this message to be more concise”), but most of the important work happens in the Project.

The rest of this book assumes you are comfortable picking the right workflow: chat for small, one-off tasks; Projects for persistent context; and Claude Code for work that lives in your repos and tools.

Chapter 12 – Building a Team Adoption Roadmap

You now have role-by-role patterns for using Claude. The remaining challenge is **how to roll this out across your organisation** without chaos, hype cycles, or hidden risk.

Successful enterprises treat AI adoption as a phased transformation: start with focused pilots, prove value, build governance, then scale in deliberate waves. This chapter gives you a practical roadmap to do exactly that.

12.1 Principles for AI adoption that actually sticks

Before diving into phases, it helps to anchor on a few principles that recur across AI adoption playbooks and maturity models:

- **Business first, tools second** – Every initiative must tie to a clear business outcome (revenue, cost, risk, speed, quality).
- **Start narrow, then scale** – Begin with a small number of high-impact, feasible use cases; expand only when they prove value.
- **Central guardrails, local innovation** – An AI Centre of Excellence (CoE) provides platforms, governance, and patterns; teams experiment within those boundaries.
- **Measure and iterate** – Each phase has explicit metrics and feedback loops, not just anecdotal “it feels faster.”

We’ll structure the roadmap around four phases you can loop through for each wave of adoption.

12.2 Phase 1 – Identify high-impact use cases

Most failed AI programs start with “Let’s use AI everywhere.” Most successful ones start with **three to five very specific workflows**.

12.2.1 Criteria for picking use cases

AI adoption guides and maturity models recommend focusing on use cases that:

- Have **clear business value** (time saved, errors reduced, revenue influence)
- Are **feasible** with current data, tools, and skills
- Are **repeatable** (daily or weekly tasks, not edge-case fire drills)
- Are **safe enough** to experiment on (non-production or low-risk to start)

Within this book’s scope, good first Claude use cases include:

- Developers: test generation, small refactors, documentation
- QA: test case design, bug report drafting
- PMs: status reports, meeting summaries
- SAP/BA: functional spec drafting, UAT scripts
- Sales: account briefs, first-pass outreach drafts

12.2.2 Involving stakeholders early

Phase 1 should include voices from:

- Business owners (who own the KPI)
- Team leads (who know day-to-day work)
- Security/compliance (to flag constraints)
- AI/IT platform owners (to assess tooling feasibility)

The output of Phase 1 is a **shortlist** of use cases, each with:

- Business outcome (what to improve, and by how much)
- Current baseline (time, cost, error rate, etc.)
- Owner and initial team
- Constraints and risks

12.3 Phase 2 – Design and run pilots

With your use cases chosen, you move into **pilot mode**: small, time-boxed experiments designed to learn, not just “try things.”

12.3.1 Pilot design: scope, metrics, and duration

AI adoption playbooks suggest pilots should:

- Be short (typically 4–8 weeks)
- Have a **narrow scope** (one or two teams, one or two key workflows)
- Be tied to **measurable targets** (e.g., “reduce time to produce test cases by 40%,” “cut weekly status report prep from 2 hours to 30 minutes”)

For each pilot, define:

- **Who** – which roles and specific people will participate
- **What** – the exact workflows they will use Claude for
- **How** – which Claude interfaces (chat, Projects, Claude Code) and any integrations
- **Measures** – baseline and target metrics for business impact and user satisfaction

12.3.2 Guardrails and governance from day one

Even in pilots, you need **minimum viable governance**:

- Access via Claude Team/Enterprise with SSO and logging, not unmanaged accounts
- Clear data classification rules: what can and cannot be shared with Claude
- Simple usage policies for participants (e.g., “no production credentials,” “no client PII,” “human review required for external outputs”)

Security guides for Claude emphasise treating security as a **stack of decisions**—identity, data handling, sandboxing, cost controls—rather than a single switch.

12.3.3 Training and change management in pilots

Pilots are not just about tools; they are **change experiments**. Effective programs include:

- Short role-based training sessions (“Claude for Developers,” “Claude for PMs”)
- Example prompts and patterns tailored to your use cases
- A clear way for participants to share feedback, blockers, and wins

There are even Claude Skills dedicated to **change management** that walk leaders through assessing impact, planning communications, and monitoring adoption.

12.4 Phase 3 – Evaluate pilots and decide what to scale

Not every pilot deserves to scale. Before you expand, you must **pressure-test** whether the experiment delivered real value.

12.4.1 Measuring impact

Enterprise workbooks and roadmaps recommend looking at both **technical** and **business** metrics:

- Technical: model performance, error rates, stability, latency
- Business: time savings, cost reduction, quality improvements, throughput, satisfaction

For Claude pilots, you might measure:

- Developer: time to implement features or write tests, PR review quality, defect rates
- QA: number and quality of test cases per week, bug report clarity, regression failures caught
- PM: time spent on reporting, meeting prep, risk visibility
- SAP/BA: spec turnaround time, UAT readiness
- Sales: meetings booked, reply rates, time spent on research and writing

You don’t need perfect measurement, but you do need **directionally clear evidence** that the new workflow is better than the old one.

12.4.2 Learning from failures and partial successes

The most mature organisations treat pilots as **learning labs**, not pass/fail tests.

Ask:

- Where did Claude clearly help?
- Where did it cause confusion or extra work?
- Were issues due to tooling, training, governance, or use case selection?

Use these insights to refine both your **workflows** and your **selection criteria** for future pilots.

12.4.3 Scale, adjust, or stop

For each pilot, you decide:

- **Scale** – results are strong, risks manageable, users positive → move to Phase 4.
- **Adjust** – results mixed, but promising → refine scope, prompts, or training and rerun.
- **Stop** – results poor or risk too high → stop and explicitly share why with stakeholders.

Stopping a pilot is not failure; it prevents wasted investment and builds credibility.

12.5 Phase 4 – Scale in waves, not everywhere at once

When a pilot works, you might feel tempted to “roll it out to everyone.” The better pattern is to **scale in concentric rings**.

12.5.1 Selecting expansion targets

Use pilot data to select where to expand next:

- Teams with similar workflows and readiness
- Use cases with clear, repeatable templates and guardrails
- Areas where champions are eager to mentor others

Each expansion wave should have its own **90-day loop**: Assess, Prioritise, Execute, Measure, Adapt.

12.5.2 Role-based adoption roadmaps

You can stitch the role-specific guidance in this book into an overall expansion plan:

- **Wave 1** – Developers and QA in 1–2 product areas (Claude Code + Projects)
- **Wave 2** – PMs, architects, and SAP/BA in those same areas (Projects + Skills)
- **Wave 3** – Sales and customer success for selected segments (Sales Projects + Skills + integrations)
- **Wave 4** – Cross-cutting functions like Finance, Operations, HR where appropriate

AI adoption roadmaps for 2026 emphasise exactly this pattern: start with clear value functions and expand into adjacent functions once patterns are proven.

12.5.3 Building champions and communities of practice

Scaling is not just copying playbooks; it is building **people networks**:

- Identify early champions and give them explicit time and recognition to mentor others.
- Set up regular forums (guilds, brown bags) where teams share prompts, pitfalls, and wins.
- Codify what works into Skills, templates, and internal docs as living assets.

This keeps your adoption from fragmenting into disconnected experiments.

12.6 The role of the AI CoE and platform teams

Across all phases, an AI Centre of Excellence (or equivalent) plays a central role.

12.6.1 Responsibilities of the AI CoE

Maturity models suggest CoEs should:

- Define enterprise AI strategy and align it to business objectives
- Select and manage core platforms (e.g., Claude Team/Enterprise, Claude Code, gateways)
- Establish security, data, and usage policies in partnership with InfoSec and Legal
- Provide enablement: training, templates, Skills, and support for pilots and scale-ups
- Monitor adoption, performance, and risk across the organisation

In this book's context, the CoE is also the natural home for **Claude Projects and Skills standards**.

12.6.2 Platform and governance stack for Claude

Security and platform guides for Claude recommend a layered approach:

- **Identity & access** – SSO, role-based access, workspace separation
- **API and credential hierarchy** – central provider keys, scoped project keys, budget and rate limits
- **Data controls** – DLP, PII filtering, retention policies, workspace scoping
- **Observability** – logging, tagging requests, cost reporting, model performance monitoring
- **Guardrails** – input/output validation, content safety, prompt injection protections

The CoE and platform teams implement this stack once, so individual teams can focus on workflows, not plumbing.

12.7 Putting it together: a 12-month Claude adoption arc

Finally, here is a simple 12-month arc you can adapt.

Months 1–3 – Foundation and first pilots

- Establish your AI CoE mandate and platform decisions.
- Run 3–5 pilots across Dev/QA, PM, and one business function (SAP/BA or Sales).
- Put basic governance and cost controls in place.

Months 4–6 – Evaluate and first scale wave

- Evaluate pilots; choose 1–2 to scale.
- Document patterns as Projects + Skills, with clear guides.
- Expand to adjacent teams in the same domains (e.g., more product squads).

Months 7–9 – Second scale wave and deeper integration

- Introduce more automation and integrations (e.g., MCP, data connectors, CI/CD hooks) where pilots have proven value.
- Scale to new functions that can reuse patterns (e.g., from one SAP rollout to multiple clients).

Months 10–12 – Optimisation and portfolio planning

- Review adoption metrics, costs, and risk incidents; adjust governance.
- Identify the next wave of higher-ambition use cases (agentic workflows, cross-system automations) based on maturity.
- Integrate AI metrics into regular business reviews and budgeting.

At the end of this arc, Claude is no longer an experiment. It is part of your **operating model**—governed, measured, and continuously improved.

Chapter 13 – Security, Data Privacy, and Governance for Claude in the Enterprise

Claude comes with strong, enterprise-grade security on Anthropic's side. But your biggest risks don't live inside the model; they live in **how your organisation uses it**—what data you send, who has access, what gets logged, and how outputs are reviewed.

This chapter lays out a practical governance stack for Claude: security capabilities you can rely on from Anthropic, and the policies, controls, and checklists you must implement yourself to safely support developers, QA, PMs, SAP teams, and sales.

13.1 Claude Enterprise security and compliance: what you get out of the box

Anthropic provides a strong baseline for Claude Enterprise, Claude for Work, and Claude Code deployments.

Key capabilities include:

- **Infrastructure security** – Encryption in transit (TLS 1.2+), encryption at rest (AES-256), hardened cloud environments.
- **Compliance** – SOC 2 Type II, ISO 27001, and related certifications, with details available in the Anthropic Trust Center.
- **Enterprise features** – SAML/SSO, private networking options, Zero Data Retention (ZDR) for regulated workloads, and audit logging across products.
- **New security integrations** – As of 2026, Claude integrates with dozens of security and compliance platforms to improve visibility, data loss prevention, and policy enforcement.

These controls mean Claude can meet enterprise-grade expectations—but they **do not replace** your own governance. Anthropic's security covers their infrastructure; you remain responsible for how your people access and use Claude and what data they send to it.

13.2 Your side of the shared responsibility model

Security and AI governance experts emphasise Claude's **shared responsibility model**:

- Anthropic secures the platform, models, and core services.
- Your organisation must:
 - Decide who can use Claude and how
 - Control what data is permitted in prompts
 - Govern API keys and integrations
 - Monitor usage and handle incidents

Think of your responsibilities in five layers:

1. **Policy** – What is allowed and why
2. **Identity & access** – Who can use Claude, where, and for what
3. **Data** – What can reach the model and how it is logged or retained
4. **Operations** – How you monitor, respond, and improve
5. **Culture** – How teams think about risk, review, and accountability

The rest of this chapter walks those layers from top to bottom.

13.3 Defining clear AI usage policies

A good Claude governance program starts with written policies that are **specific, simple, and enforceable**.

13.3.1 Data classification and allowed use

AI governance and privacy roadmaps for 2026 strongly recommend aligning AI use with your existing data classification schemes.

You should explicitly define:

- **Allowed data** – Public information, synthetic data, non-sensitive internal docs, non-prod configs.
- **Restricted data** – Internal confidential docs, non-prod logs with some identifiers, anonymised metrics (allowed only in specific, controlled contexts).
- **Prohibited data** – Secrets and credentials, payment details, regulated personal data (e.g., unmasked PII subject to GDPR or HIPAA), highly sensitive legal or HR content.

Valence Security's guidance for Claude stresses that these rules should apply **consistently across all AI tools**, not just Claude.

Policies should be written in plain language and heavily used in training so every developer, QA, PM, SAP consultant, and salesperson knows **what they can and cannot paste into Claude**.

13.3.2 Usage policies by role and risk

Beyond data types, you should define **role-based AI usage policies**:

- Developers: allowed to use Claude Code for non-prod code; production-critical modules require senior review and stricter guardrails.
- QA: allowed for test generation and bug reports, but only on masked or synthetic data.
- PMs and architects: allowed for docs and planning; prohibited from exposing confidential HR or M&A details.
- SAP/BA: must use enterprise integrations (BTP, MCP, secure connectors) and avoid raw SAP exports with live PII.
- Sales: allowed for public and CRM data; forbidden to expose confidential contracts or non-public personal data outside approved channels.

These policies should tie back to your data classification and regulatory obligations.

13.4 Identity, access, and environment controls

Security guidance for Claude emphasises **access control and workspace scoping** as first-line defences.

13.4.1 SSO, role-based access, and workspace separation

Best practices include:

- Enabling **SAML-based SSO** for Claude Enterprise and Claude for Work so access follows your identity provider's policies.
- Limiting Claude access to users and teams that have a documented business need; avoid "enable it for everyone and see what happens."
- Separating **workspaces or organisations** by sensitivity or region if needed (e.g., EU vs non-EU, regulated vs non-regulated units).

Workspace scoping should mirror how data is isolated in your existing SaaS estate.

13.4.2 Governing API keys, service accounts, and MCP

Enterprise security guides repeatedly warn that API keys and service accounts are high-value targets.

Controls should include:

- A **credential hierarchy** – central platform keys, project-scoped keys, no raw provider keys in developer laptops or CI logs.
- Key management – central tracking, rotation, and revocation of Claude API keys and MCP servers, with least privilege access.

- MCP and integration governance – only approved MCP servers (for SAP, databases, storage, etc.) deployed; each with constrained permissions and logging.

TrueFoundry's governance guide for Claude and related assessments emphasise failing securely: if an integration misbehaves or a key leaks, the blast radius should be small and detectable.

13.5 Governing the data layer: keeping sensitive data out (or under control)

Security practitioners are clear: **the biggest risks aren't in the model; they're in the data environment around it.**

13.5.1 Restricting what reaches the model

Key recommendations from security and trust-and-safety guidance:

- Configure **Zero Data Retention (ZDR)** where possible for regulated workloads.
- Use **data loss prevention (DLP)** and SaaS security tools integrated with Claude to scan and block sensitive content from leaving your environment.
- Automatically detect and redact secrets, tokens, and credentials before they are sent—through pre-prompt hooks, proxy services, or DLP integrations.
- Treat prompts that include external content (like web pages or emails) as potential vectors for prompt injection; sanitise or constrain what gets forwarded.

These measures are in line with 2026 data governance roadmaps that emphasise **semantic layers, access controls, and auditability** as foundations for trustworthy AI.

13.5.2 Logging, retention, and audit trails

For privacy and audit readiness, you must control and document data flows:

- Decide which logs are kept where: Claude workspace logs vs your own observability stack.
- Define retention periods for AI-related logs, aligned to your data privacy and security policies.
- Ensure you have **audit trails** of who accessed what, which Projects and prompts might have touched sensitive data, and how outputs were used.

TrustArc and similar privacy frameworks highlight that regulators expect not just secure design, but **documented controls** and the ability to reconstruct what happened when something goes wrong.

13.6 Managing model risks: hallucinations, misuse, and safety

Beyond data, you must govern **how** Claude is used and how its limitations are handled.

13.6.1 Hallucinations as a governance issue

As discussed in Chapter 3, hallucinations—confident but wrong answers—are inherent to LLMs. Governance frameworks like NIST’s AI Risk Management Framework and emerging AI regulations treat reliability and accuracy as key risk dimensions.

Organisational controls include:

- Defining tasks where **human review is mandatory** (production code, legal text, customer-facing comms, design decisions in regulated contexts).
- Requiring grounding in internal data (Projects, RAG, MCP) for high-stakes answers instead of open-ended questions.
- Mandating “show your work” and citation prompts in Projects where decisions depend on Claude’s outputs.

These are not just UX choices; they are part of your **risk controls**.

13.6.2 Misuse, abuse, and content safety

Anthropic implements its own safety guardrails, but organisations should still:

- Configure **content filters and allow lists** where available (e.g., blocking certain prompt types in shared environments).
- Monitor for suspicious or non-compliant usage—prompts that suggest scraping internal systems, exfiltrating data, or generating disallowed content.
- Include AI usage in your **acceptable use** and disciplinary policies so misuse is handled consistently.

Security and governance guidance stresses that Claude should be treated like any other powerful SaaS: high trust, but **not blind trust**.

13.7 Putting it together: an AI governance checklist for Claude

Drawing from enterprise security guides, AI adoption workbooks, and privacy roadmaps, you can turn this chapter into a practical checklist.

Policy and scope

- Document an AI usage policy that covers Claude across roles and use cases.
- Map AI usage to your data classification scheme: allowed, restricted, prohibited.
- Identify regulatory obligations (GDPR, HIPAA, sectoral rules, EU AI Act readiness) that apply to Claude usage.

Identity & access

- Enforce SSO for Claude Enterprise / Workspaces; disable unmanaged accounts for work purposes.
- Limit user access by role and business need; avoid blanket enabling for the whole company.
- Separate workspaces or organisations where necessary (e.g., geography, sensitivity).

API, keys, and integrations

- Use a central credential management pattern—no raw provider keys on developer machines.
- Track, rotate, and revoke Claude API keys and MCP servers; implement least privilege.
- Approve and document all integrations (MCP, SAP, data warehouses, ticketing tools).

Data protection and privacy

- Enable ZDR and private networking where appropriate.
- Integrate DLP and SaaS security tools that cover Claude usage; block secrets and sensitive PII from prompts.
- Define logging and retention strategy for prompts, responses, and integration activity.

Operational governance

- Set up monitoring dashboards for usage: costs, volumes, top Projects, and high-risk patterns.
- Include Claude-related incidents in your security incident response plan.
- Run periodic reviews of Projects and Skills to prune stale data and update guardrails.

Culture and training

- Train each role (Dev, QA, PM, SAP, Sales) on both **how to use Claude** and **how to use it safely**.
- Communicate that Claude is a **tool, not a decision-maker**; humans own outcomes.
- Encourage reporting of near-misses or concerns to continuously improve governance.

This checklist aligns with broader AI governance roadmaps for 2026 that call out data governance, access controls, transparency, and risk management as the foundations of trustworthy AI.

Used this way, Claude becomes not just powerful but **trustworthy**: embedded in your engineering, product, SAP, and sales workflows under clear guardrails, with privacy and security built in by design—not bolted on later.

Chapter 14 – Implementation

Checklists and Templates

This chapter turns the ideas in the book into **practical checklists and templates** you can copy into your internal wiki, policy docs, and onboarding material. Use them as starting points and adapt to your organisation's specifics, regulations, and risk appetite.

14.1 Claude security & data privacy policy checklist

Use this to draft or review your Claude usage policy.

1. Scope and purpose

- Define why your organisation is using Claude (e.g., productivity, quality, faster delivery) and where it is in scope (Dev, QA, PM, SAP, Sales).
- Map Claude to your existing security, data governance, and acceptable use policies.

2. Data classification and handling

- Align Claude usage with your data classification scheme (public, internal, confidential, restricted).
- For each class, decide: allowed, allowed with controls, or prohibited for prompts and Projects.
- Explicitly forbid secrets, credentials, and regulated PII from being pasted into Claude, unless via approved, governed integrations.

3. Roles and responsibilities

- Assign clear owners: AI CoE, security, legal/privacy, HR, IT, and business unit leaders.
- Define who approves new Claude Projects, Skills, and integrations.
- Document who handles incidents related to AI misuse or data exposure.

4. Access and environment

- Require SSO for Claude Enterprise / Workspaces; block unmanaged work use of consumer accounts.
- Limit access to Claude Code and Cowork to approved roles and pilot groups; expand via defined rollout phases.

- Separate workspaces by department, sensitivity, or region when needed.

5. Logging, monitoring, and retention

- Decide what is logged where (Claude logs vs your SIEM/observability).
- Define log retention and access controls for AI-related logs.
- Integrate Claude with your security stack (DLP, CASB/SSPM, SIEM) to monitor and alert on policy violations.

6. Governance lifecycle

- Establish a review cadence for your AI policy (e.g., quarterly) aligned with regulatory changes and Anthropic updates.
- Require periodic access reviews and Project/Skill audits to remove stale or risky configurations.
- Plan for continuous training updates as features and policies evolve.

14.2 Role-based AI usage policy templates

Below are concise templates you can adapt for each role.

14.2.1 Developers

Allowed

- Code understanding, refactoring proposals, test generation, documentation drafts using Claude Code and Projects.
- Non-prod logs and configs that do not contain secrets or regulated PII.

Required controls

- All AI-generated code must pass tests and human code review before merging.
- Use repo-scoped credentials; never paste secrets or production configs into prompts.

Prohibited

- Using Claude to modify production environments directly.
- Sharing customer or employee PII outside approved integrations.

14.2.2 QA & SDETs

Allowed

- Test case design, regression planning, bug report drafting, and log analysis on masked/synthetic data.

Required controls

- Clearly label AI-generated tests; require review before adding to regression suites.

Prohibited

- Raw production logs containing PII or secrets in prompts.

14.2.3 PMs & Scrum leads

Allowed

- Using Projects for charters, roadmaps, status reports, and meeting notes.

Required controls

- Review and approve all external-facing content derived from Claude.

Prohibited

- HR performance reviews, salary discussions, legal/M&A docs in Claude.

14.2.4 SAP, functional, BA

Allowed

- Spec drafting, process documentation, UAT scripts using static docs and governed SAP connectors (BTP, MCP).

Required controls

- Use approved integration paths for live SAP data; avoid ad-hoc exports with sensitive data.

Prohibited

- Pasting raw SAP screens or dumps with unmasked PII into generic chats.

14.2.5 Sales & GTM

Allowed

- Account briefs, outreach drafts, call prep using ICP docs, public data, and CRM info.

Required controls

- Do not fabricate case studies or client names; clearly mark AI-drafted content for review.

Prohibited

- Non-public personal data or confidential contract details in prompts.

14.3 Risk assessment worksheet for AI use cases

Use this worksheet for each proposed Claude use case, as suggested in AI readiness and risk-mitigation checklists.

1. Basic information

- Use case name:
- Owner:
- Teams involved:
- Claude modes: chat / Project / Claude Code / Skills / MCP / other

2. Business value

- Problem we are solving:
- Expected benefits (time saved, cost reduced, quality improved, risk reduced):
- KPI(s) to measure success:

3. Data profile

- Data types used (public, internal, confidential, PII, regulated):
- Systems integrated (e.g., GitHub, Jira, SAP, CRM, Data warehouse):
- Is sensitive or regulated data required? If yes, how will it be masked or protected?

4. Risk classification

Rate each 1–5 (low to high):

- Impact of incorrect AI output:
- Likelihood of exposure of sensitive data:
- Operational impact if the workflow fails:

5. Controls and mitigations

- Required human review steps:
- Data minimisation and masking strategy:
- Access control and logging requirements:
- Fallback if Claude is unavailable or fails:

6. Approval

- Security review complete (Y/N):
- Legal/privacy review complete (Y/N):
- Approved for: pilot only / limited production / broad rollout

This format aligns with modern AI governance and enterprise risk practices.

14.4 Data classification template for AI tools

Adapt this simple table to your environment (drawn from AI readiness and governance resources).

Data class	Examples	Claude usage policy
Public	Marketing site, public docs, blog posts	Allowed in all Claude modes
Internal	Non-sensitive internal docs, non-prod logs	Allowed in Projects and Claude Code with SSO
Confidential	Internal roadmaps, non-public metrics, client IDs	Allowed only in scoped Projects with DLP & logs
Restricted	PII, payment data, health data, sensitive legal/HR	Prohibited except via approved, governed flows

Make this visible wherever people work with Claude so decisions are quick and consistent.

14.5 Sample 90-day Claude adoption plan

Finally, here is a concise 90-day plan you can embed in your internal docs, aligned with broader AI adoption and readiness checklists.

Days 1–30 – Foundation and pilots

- Stand up Claude Enterprise with SSO, basic policies, and monitoring.
- Choose 3–5 pilot workflows across one or two domains (e.g., Dev/QA + PM).
- Train participants on role-specific usage and safety.
- Start measuring baseline and early impact.

Days 31–60 – Refine and expand pilots

- Adjust prompts, Projects, and Skills based on feedback.
- Introduce simple integrations (e.g., Git/CI, Jira exports, SAP test data) where needed.
- Document successful patterns in an internal “Claude playbook.”

Days 61–90 – Decide and prepare to scale

- Evaluate pilots against KPIs; decide what to scale vs pause.
- Harden governance (DLP integrations, key management, periodic audits).
- Plan the next wave (more teams, more use cases) and update training accordingly.

At the end of these 90 days, you should know **where Claude drives real value, how to keep it safe, and how to scale it with confidence**

Continue Your AI Journey With Paramount



This guide is just the beginning. As AI continues to transform how organizations build, collaborate, innovate, and grow, success will depend on adopting these technologies with the right balance of strategy, governance, and practical execution.

Explore more technology insights, workforce solutions, digital transformation initiatives, and emerging technology expertise at [Paramount Software Solutions](#) and discover how your organization can unlock greater productivity, efficiency, and business value through the strategic use of technology and AI.

Paramount

Technology. Talent. Transformation.

Follow Us



© 2026. Paramount Software Solutions