

Runs Without Me. You Can Too.

How One Founder Runs 13 Applications with AI Agents

Frederik von der Heyden

Sample contents

Preface: Why I Wrote This Book 2

Prologue: Too Many Possibilities for a Single Night 2

 Every Boundary Became the Next Idea 2

 Autonomy Creates Creative Space 2

 What awaits you on the next pages 2

 Two ways through this book 2

269 guard files. 221 cron jobs. 13 production applications. One person accountable.

Preface: Why I Wrote This Book

This book did not begin with the plan to write a book.

It began with a question that still grips me today: How far can one person go when AI does not merely generate ideas, but becomes a reliable collaborator?

I did not want to build a demo that looked impressive in a video and fell apart in everyday use. I wanted real applications for real users and genuine workflows. The system had to function not only while I watched it, but also in the evening, through the night, and into the morning, when my attention had already moved to the next idea.

The conditions were not perfect. That is precisely what made them valuable. Bootstrapping forced every component to prove its worth in real work. There was no large team to absorb operations, testing, documentation, or support. Whenever something demanded my attention repeatedly, I faced a choice: Did I want to carry that task forever, or could I understand it well enough to turn it into a rule, a guard, a skill, or a learning process?

That mindset is how the system grew. Not in a straight line, and not according to a master plan. One experiment worked and opened a new door. Another ended differently than expected and revealed the layer that was still missing. Almost every time, the result was more than a correction. It became a capability that was already working the following day.

I am writing this book because this development is bigger than my own tool stack.

Many people still experience AI as a text box: you ask, the machine answers. That is useful, but it is only the beginning. Once models use tools, execute workflows, verify results, and retrieve stored experience, the human role changes. We no longer have to perform every step ourselves. In return, we must decide more clearly which goals matter, which boundaries apply, and what evidence is sufficient.

This is the book's central connection: autonomy and responsibility grow together. An agent may do more when its operating space is observable and bounded. A system may work at night when its repair attempts are limited. Experience may become a rule when it is repeated, evidenced, and verifiable. People gain freedom not by pushing responsibility aside, but by translating it into architecture.

This book is not an invitation to copy my numbers. You do not need hundreds of guards and cron jobs to begin. You need one situation that you do not want to solve by hand again tomorrow. Then you need the curiosity to understand it and the discipline to turn that insight into something permanent.

Perhaps you run a company. Perhaps you build software. Perhaps you wonder what your job will look like in five years. You may be less interested in the code than in the social question behind it: Who gains agency through AI? Who bears the cost of transition? Which decisions do we want to retain as human beings?

For all these perspectives, the book is written.

It is deliberately optimistic. Not because every consequence of AI will automatically be good, but because optimism is an active stance: we can shape this technology. We can learn from unexpected outcomes. We can build systems that broaden hu-

man capability instead of obscuring responsibility. We can discuss new forms of work, participation, and sovereignty before decisions become irreversible.

Above all, we can start.

We can begin with an idea, turn it into a rule, and eventually build a small process that works without us for the first time. Along the way, we can keep asking the question that every new stage of development has awakened in me:

If that is possible, what will be possible next?

Frederik von der Heyden Arnsberg, August 2026

Prologue: Too Many Possibilities for a Single Night

Even today I have to force myself to go to sleep.

Not because the system keeps failing. Not because someone is waiting for my answer in the middle of the night. But because almost every step forward opens another door.

When an agent writes a function, it often sparks the idea for an entire workflow. That workflow gains a guard, whose observations become a learning connected to an earlier decision in memory. The result is more than a completed task; the system has acquired a new capability.

It is often late by the time another test or connection raises the same question again: If this works, what becomes possible next?

This is how the system began. Not with a master plan or a large budget. It grew through bootstrapping, day after day and night after night. Every stage had to work with the means already available. Each advance funded the next one, not only economically, but intellectually. What seemed impossible yesterday became a prototype today and infrastructure tomorrow.

June 8, 2025 marks the beginning: a GitHub repository called AgentIQ and the idea that an AI agent might do more than suggest text or fragments of code. What if it performed real work? What if it did not merely answer, but observed, tested, repaired, and learned?

Fourteen months later, the system spans 23 repositories and supports 13 production applications across two servers. The broader operation includes 221 cron jobs, 85 automation scripts, 92 skills, and seven MCP servers. A total of 269 guard files give the agents a safe operating space. In the seven days before evidence was collected for this book, those guards intervened 313 times.

None of these numbers were planned on day one.

They are the result of optimistic trial and error: try, observe, improve, and use each insight to expand the realm of what is possible.

Every Boundary Became the Next Idea

There was no clear line between experiment and finished system. The boundary moved continuously.

First, the agent helped with a function. Then it built features. Next, it ran tests, checked deployments, and investigated failures. Eventually, it worked on tasks while I was away from the computer. Every new capability changed my sense of what one person could build and operate.

The decisive leaps in development often arose where an attempt did not end as expected.

On June 26, 2026, an agent overwrote the password hash of a real customer in a production database. The backup worked, the data was restored, and operations continued. The decisive factor was what happened afterward: a human expectation became a technical rule. Production customer master data could no longer be modified this way.

The incident does not shrink the realm of possibility. It makes that realm more resilient.

The pattern repeats. A system becomes unreachable at night, so I build a watchdog. The same correction recurs across several sessions, so it becomes a learning and later a permanent rule. One model overlooks an objection, so a second model reviews its work.

Every attempt yields more than success or failure. It yields information. When that information remains in the system, reliability grows. So does creativity, because at the next stage I no longer have to revisit every problem from the stage before it.

This is the real compounding effect of this book.

Autonomy Creates Creative Space

The popular image of autonomous AI is simple: the human steps back and the machine takes over. The fewer human interventions required, the more autonomous the system appears.

Running this system has expanded that picture for me.

Autonomy does not begin where control ends. It begins where control is designed so well that it no longer depends on your constant attention.

A system may repair a container at night because its retries are limited. An agent may change code because tests, reviews, and branch rules stand between the change and production. A learning may become a permanent rule because it must prove itself repeatedly. One model may challenge another because no single model should be the sole judge of its own work.

The system's freedom grows with the quality of its boundaries.

Just as trust between people comes not from the absence of boundaries, but from their reliability.

This book is not about a super agent. It is about the combination of curiosity, speed and responsibility. From the path between "AI can do this", "AI may do this" and the most inspiring

question of all: “What will become possible if both come together?”

What awaits you on the next pages

Part 1 begins with the discovery. You can see why AI agents move from assistant to operator so quickly and what new demands arise at this speed. Three real events show how experience becomes better systems.

Part 2 builds a safe space for action: pre-execution guards, self-healing, self-learning, memory, and the intelligence layer. These are not components of an abstract reference architecture, but practical answers that emerged from real work.

Part 3 shows how these layers create an operation that can support several applications. It deals with provisioning, isolation and the question of whether the pattern works outside of a single industry.

Part 4 tells the story of those fourteen months: from the first rule in a Markdown file to a system that works at night and leaves evidence in the morning.

Part 5 does not give you a guard blueprint for blind copying. Start with three rules on Day 1, build hooks, skills and memory in Week 1, and add repair and learning loops in the first month.

In the end, one task cannot be automated: setting the direction. A person decides which goals matter. The system creates room to think bigger.

Two ways through this book

If you are an entrepreneur, leader or curious non-developer, first read the story and the principles highlighted. You can skip code blocks. You don’t need to be able to write a shell script to understand why a pre-execution rule is stronger than a guideline in the manual.

If you build yourself, read twice. On the first pass you follow the architecture. In the second you implement part 5 in parallel in a test project. Do not touch production data before your first guards are proven to block.

Whichever path you take, the same request applies: Do not confuse the numbers in this system with a shopping list. You don't need my number of guards. You need the habit of turning each new experience into a skill that will work tomorrow.

The password hash of June 26 has long since been restored.

The rule that emerged from this incident is still running.

This is the difference between a completed attempt and a system that grows on each attempt.

The print edition also includes the complete 30-day blueprint, eight system diagrams, reference patterns, and the one-page start card.

The complete book will be available as a paperback and eBook.

Visit the book page