

AI in 90 Minutes

*What working people
actually need to know*



Matija Vidmar

The AI Architect

AI in 90 Minutes

What working people actually need to know

© 2026 Matija Vidmar — The AI Architect

All rights reserved.

This document may be shared freely in its complete, unmodified form. It may not be sold, altered, or republished in partial form without the prior written consent of the author.

matija.vidmar@gmail.com

First edition — July 2026

Table of Contents

Why this little book exists	4
AI doesn't think (and why that's good news)	5
The golden rule: what never to paste into a chat	7
The difference is in the briefing: the four components of a request that works . .	9
Five prompts to use tomorrow morning	12
When the AI writes the prompt for you (and what comes next)	15
The ten-minute test: where your work is heading	17
This was the first step	19

What working people actually need to know

By Matija Vidmar, *The AI Architect*

Why this little book exists

For twenty years I wrote code. Ecommerce systems that cannot go down, databases that must never lose a record, servers nobody sees but that keep everything else standing. That's the perspective from which I started watching artificial intelligence, years before it became the topic of every conference. And what I see today, in 2026, is a gap that keeps widening: on one side, tools that improve every three months; on the other, most working people stuck at "I tried asking ChatGPT."

That's not a criticism. It's an observation: the foundational knowledge is missing. And the market fills that void in the worst possible way, with courses selling the twenty magic prompts for marketing and the fifteen email templates. That material works for a month. Then the model updates, the interface changes, and you're exactly where you started, holding a PDF that's already obsolete.

This little book starts from a different principle, the same one I've built all my work on: the tool depreciates, judgment doesn't. The tools you use today will be different in a year. The ability to understand what to ask, how to evaluate an answer, when to trust and when to verify: that stays, and it compounds. It's the only investment in this field that doesn't lose value with every update.

Ninety minutes won't make you an expert. But they can give you the right foundations: understanding how this technology actually works, avoiding the mistake that costs the most, learning to phrase a request that produces results instead of generic text, and walking away with five ready-to-use tools, a method for building more of your own, and an honest diagnosis of where your work is heading. That's not nothing. It's exactly the first step most people never take.

One piece of advice before you start: read with an AI tool open next to you, your phone is enough. Every time you meet an idea, test it immediately. Real understanding comes after practice, not before.

AI doesn't think (and why that's good news)

The word "intelligence" in "artificial intelligence" is the problem. It's not a technical description: it's a marketing metaphor that worked brilliantly at capturing the imagination, and just as brilliantly at generating wrong expectations in both directions. Some people expect too much and end up disappointed. Others expect too little and never start. Both make the same underlying mistake: believing "artificial intelligence" means "human intelligence replicated in a machine."

It doesn't work that way. And understanding how it actually works, even in broad strokes, radically changes how you use these tools.

Start from the simplest thing you already know: your phone's autocomplete. When you're typing a message and the phone suggests the next word, it's analyzing what you've written so far and predicting what comes next, based on millions of previously seen sequences. It doesn't understand the message. It doesn't have an opinion. It calculates a probability.

ChatGPT, Claude, and Gemini do the same thing, at a scale that makes the comparison almost comical. Their training processed an enormous share of all the text ever written: books, scientific papers, code, technical documentation, online discussions. And instead of suggesting one word at a time, they generate coherent text for thousands of words. But the mechanism stays the same: predicting the most probable sequence, fragment after fragment. Those fragments are called tokens, and they're the unit of measure for everything these systems do.

This single concept, statistical prediction, explains the three behaviors that confuse everyone who starts.

The first is hallucination. If the system generates text based on statistical plausibility, it can produce false statements that sound exactly like true ones: the confident tone, the grammatical structure, the vocabulary of an authoritative answer are all patterns it has learned and reproduces faithfully. It's not lying to you, because it has no intention. It's completing a plausible sequence, and plausibility is not the same as truth. Court rulings invented out of thin air and cited in real proceedings by real lawyers, complete with credible case references, have become their own genre of news story.

The second is inconsistency. The same question, asked twice, can produce different answers. Not because the system changed its mind, but because generation includes a controlled dose of randomness: when two words are equally plausible, it picks non-deterministically. That's a design choice; otherwise every answer would be identical and mechanical.

The third is the most insidious: false confidence. AI uses the same assertive tone when stating something certain and when making things up. Inside the model there is no signal separating "I know this" from "this sounds right." That distinction is yours to make. And this is where the bad news becomes good news: if the added value lived inside the machine, you would be replaceable. Instead, the value lives in the judgment of whoever uses it. The machine produces; you decide what's worth keeping.

How often does it get things wrong, concretely? It depends on where you use it, a distinction almost nobody explains. On general-knowledge topics, the best models are now very reliable, with factual error rates down to a few percentage points. But in specialized domains the picture changes: in legal work, for example, even the best systems get several answers in a hundred wrong, and on precise citations of specific sources the invention rate climbs dramatically. The operating rule that follows is simple: whenever an output contains a specific piece of data you'll use in an important context, a number, a name, a date, a citation, verification is on you. It's not a defect to wait out: it's the division of labor between you and the tool.

So what good is a tool that can make things up? Flip the question: where is statistical prediction an advantage rather than a risk? The answer covers more of your daily work than you'd expect. Transforming text: a first draft to refine, a long document to summarize, a text to rewrite for a different audience. Explaining and simplifying: a contract dense with clauses, a technical study, a document in a language you barely manage. Generating options: twenty headlines to choose from, ten objections a client might raise, five different ways to open a presentation. In all these cases you're not asking the machine for truth: you're asking for material on which to exercise your judgment. And that's where it performs remarkably.

A hammer isn't defective because it doesn't work like a screwdriver. Language AI is built to generate coherent, useful text from an instruction: within that scope, it's often exceptional. Problems start when you ask it to be what it isn't, an archive of certified facts or a calculator.

One last clarification, because the market tends to muddy the waters. The most recent models "reason": before answering, they generate an internal chain of thought, explore alternative approaches, backtrack when they recognize a dead end. The right image is a mouse in a maze that, instead of running straight into every wall until it stumbles on the exit, learns to stop at the junctions and mentally explore the paths before moving. Results on complex problems, especially logic and planning, improve for real. But the underlying mechanism doesn't change: they are more capable, not more conscious. And the division of labor between you and the tool stays exactly the same.

In the full course: lesson 1.2 draws the precise map of what AI does well, what it does badly, and what it can't do at all, with hallucination-rate data by industry. It opens with the story of Sara, a lawyer, and the two nonexistent rulings she cited in a court filing. It didn't end well.

The golden rule: what never to paste into a chat

Before you even learn to write a good request, there's a rule that applies immediately, starting with the first experiment you'll run ten minutes from now. It's the rule that separates a professional from an incident waiting to happen.

Never put sensitive data into the free tiers of AI tools. Contracts with names and figures, client data, confidential financial information, employee data, documents covered by nondisclosure agreements: none of it belongs in a free chat. The reason is in the terms of service nobody reads: many free tiers reserve the right to use conversations to improve their models. Paid and business plans offer different contractual guarantees, which is exactly why they exist.

The practical rule is simple: before pasting something, ask yourself whether you'd send it to an external consultant you just met, with no signed agreement. If the answer is no, anonymize. In most cases thirty seconds are enough: "client X," "a company in industry Y," rounded figures. AI works perfectly well without knowing the real names.

A concrete example, because the rule feels abstract until you see it applied. Wrong version: "Rossi SpA owes us 47,320 euros on invoice 2026/114 and their CFO Paolo Bianchi hasn't answered in three weeks: write a payment reminder." Right version: "A long-standing client owes us a five-figure amount more than 60 days overdue and their finance contact hasn't answered in three weeks: write a firm reminder that preserves the business relationship." The quality of the email you get is identical. The exposure of your data is not. And when you need the final text with real names, you add them yourself, in your own document, outside the chat.

Note what the rule does not say. It doesn't say don't use AI at work: it says choose the right plan for the kind of data you handle. Anyone using these tools daily on company material should be working on an individual paid plan or, better, a business plan, where data handling is governed by explicit contractual terms. That's tens of euros a month against the risk of finding confidential information where it shouldn't be: the math isn't hard.

While we're on the subject of what the tool knows about you, a quick clarification on a topic that generates confusion: memory. The model itself remembers nothing from one conversation to the next. But by 2026 all major tools have built a persistent memory layer on top of the model: ChatGPT automatically synthesizes what it learns from your conversations, Claude includes memory on every plan, Gemini has turned chat history into consultable memories. Useful, with two limits worth knowing: each tool only remembers what happened inside its own walls, and what it keeps is a selective synthesis, not an archive. If a detail matters, don't assume the tool remembers

it: repeat it.

***In the full course:** lesson 5.5 covers privacy, hallucinations, and dependence with one simple criterion, managing them without paranoia, and includes two ready-made prompts: the one-page AI policy draft for your company and the verification clause to build into sensitive processes.*

The difference is in the briefing: the four components of a request that works

In June 2026, Anthropic published a study of roughly 400,000 real AI work sessions. The question: what separates people who get excellent results from people who get mediocre ones? The answer surprised almost everyone: it's not technical skill. It's the quality of the briefing. An expert user generates, with the same tool, five times the output of a novice. And managers, not engineers, have the highest success rate of all.

Five times. Same tool, same subscription, same model. The entire difference is in how you phrase the request. And it's a skill you can learn.

The starting point is understanding why "write an email to the client" produces disappointing results. The AI doesn't know your situation: who you are, who the client is, what happened, what you want to achieve. When information is missing, it doesn't ask: it fills the gaps with statistical plausibility, which is an elegant way of saying it guesses. And it guesses in subtly wrong ways, producing the perfect answer to a generic situation that isn't yours.

The solution has a precise structure: four components that work as a checklist before you hit enter.

Context is everything the AI can't know on its own: who you are, what industry you work in, who the output is for, what circumstance led to this request. The practical test: if you sent this request to a person who has never seen your project, could they complete the job? If not, context is missing.

Role tells the AI which perspective to answer from. "You are a legal consultant specializing in commercial contracts" produces a structurally different answer than "you are a promotional copywriter." It's not make-believe: it changes how the system weighs information, the language it uses, the level of technicality. There's also a refinement almost nobody uses that pays off disproportionately: defining what the AI must not be. "You are my analyst for commercial decisions. You are not a search engine listing options without taking a position": the affirmation says where to stand, the negation says where not to go, and together they close the space that positive instructions alone leave open. One warning in the opposite direction: role works for general disciplines (lawyer, doctor, engineer, marketing manager), not for hyper-specific niches. "Software sales expert in northeastern retail" produces nothing different from "software sales expert": at that granularity, context matters, not the label.

Objective is the concrete finish line. Not "help me with this presentation," but "write the copy for five slides presenting the quarter's results to our investors." The control question: when I read the output, how will I know whether it worked?

Format is the element beginners neglect and professionals consider non-negotiable: length, structure, tone, level of formality. Without instructions, the AI chooses for you. And its default choice rarely matches what you need.

To see the progression in action, watch the same request grow one component at a time.

Minimal version: "Write a risk analysis for the project." You'll get ten pages, correct and useless.

With context: "Write a risk analysis for the implementation of a new management system in a 300-employee manufacturing company, six months from go-live, 800,000 euro budget. The main risks concern historical data migration and staff training." Already usable.

With all four components: "You are a project management consultant experienced in manufacturing system implementations. Write a risk analysis for a non-technical board of directors on the implementation described above. Format: 400 words maximum, direct tone, the three main risks with estimated probability and impact, one final recommendation." This one lands on the desk ready to use.

Each addition produces a qualitatively different output: more relevant, better calibrated to the real audience, more usable without rework. The numbers confirm what the example suggests: people who use a structure like this report improvements in perceived output quality between 40% and 60% compared to those who write on instinct.

Then there's the classic case of someone who nails the content and forgets the form. A lawyer with a small practice had started using AI for first drafts of client letters: impeccable context, frustrating results. Texts too long, structure different every time, rewriting took nearly as long as drafting by hand. The turning point was one line at the bottom of the prompt: "professional tone, 200 words maximum, structure: situation, legal position, next step." Review time was cut in half. The content had always been good: the form was creating the extra work.

And there's a side effect few people expect: Tobi Lütke, Shopify's CEO, observed that forcing himself to give AI complete context made him a better communicator across the board, sharper emails, clearer memos. Clarity of thought is the same whether the request goes to a machine or to a person. People who learn to build a good briefing for AI tend to become more precise at everything else.

You don't need all four components every time. For a simple request, context and objective are enough. The value isn't in the form-filling: it's in the mental habit of running through the four questions before writing, and deciding deliberately what to include and what to skip.

In the full course: Module 2 dedicates ten lessons to this skill, from the expert prompt (how to turn your industry knowledge into the five-times multiplier Anthropic measured) to iteration, to

meta-prompting where the AI writes the prompt for you, all the way to neutralizing AI's tendency to always agree with you.

Five prompts to use tomorrow morning

The theory of the previous chapter becomes concrete here. These five prompts cover situations almost anyone who works runs into every week. They're made to be copied, filled in where you see square brackets, and used. But their real value is something else: each one is an example of the structure you've just learned. Use them, then modify them, then write your own.

1. The email draft

```
Write an email to [name/role, e.g. "my sales manager"].  
Objective: [what should happen after they read it, e.g. "get approval for  
the trip on the 15th"].  
Tone: [formal / direct / friendly].  
Length: maximum [X] lines.  
Context: [1-2 sentences of background].
```

The field that makes the difference is the objective: not "inform," but what should happen after the reading. An email with a clear objective almost writes itself, even without AI. With AI, it arrives in thirty seconds.

2. The document query

```
I've uploaded [type of document, e.g. "a 40-page supply contract"].  
My role: [who you are in relation to this document].  
Answer these three questions, citing the exact passage each answer comes  
from:  
1. [first question]  
2. [second question]  
3. [third question]  
If an answer is not present in the document, say so explicitly instead of  
inferring it.
```

The last line is the most important: it forces the tool to distinguish between what the document says and what would merely sound plausible. It's the practical defense against the hallucination from chapter one.

3. The pre-meeting briefing

This is the documentation for the meeting on [date] about [topic]:

[paste documents, emails, previous minutes]

Prepare an operational briefing with: the current situation in 3-4 points, the open items to resolve today, the 3-5 key questions I want answered, and any tensions or risks to keep in mind.

Ten minutes before every important meeting. Whoever arrives with the right questions steers the conversation, and the right questions emerge from the documentation nobody has time to reread.

4. The reply to a difficult email

I need to reply to this email: [paste the text].

The situation: [what actually happened, 2-3 sentences, including the uncomfortable part].

What I want to achieve: [e.g. "keep the relationship but not accept the request"].

What I don't want: [e.g. "apologize more than necessary, promise dates I can't guarantee"].

Tone: professional, firm, not defensive. Maximum 10 lines.

The "what I don't want" field is the one nobody uses and it changes everything: defining boundaries in the negative spares you the AI's most typical drifts, like excessive apologies or vague promises.

5. The end-of-day review

These are my end-of-day notes: [write freely what you did, what's still open, what came up, even messy and partial]

Organize them into three blocks:

1. What I finished today
2. What's still open (with the next concrete action for each item)
3. Priorities for tomorrow morning

Short list, maximum 3-4 points per block. Operational tone, no redundancy.

Five minutes at the end of the day. The value isn't the list itself: it's that tomorrow morning you start working instead of reconstructing where you left off.

Two pointers for using these five tools well. First: don't judge a prompt by its first attempt. If the output isn't convincing, don't start over: say what's wrong ("too formal," "the renewal point is missing," "cut it in half") and let the second version build on the first. Conversation is the method, not a fallback. Second: when a modified prompt works for your case, save it somewhere, a

document, a note. After a month you'll have your own library, calibrated to your actual work. It's worth more than any collection downloaded from the internet, including this one.

***In the full course:** Appendix B contains the complete library, over forty prompts organized in nine sections, from email to data analysis to configuring a permanent assistant, with "expert" versions for lawyers, accountants, and marketing managers. And Module 4, eleven lessons, builds the method behind the prompts: email, documents, research, meetings, data analysis, all the way to using AI as a personal tutor.*

When the AI writes the prompt for you (and what comes next)

At this point you know the structure of a good request and you have five ready prompts. The next objection is predictable, because everyone raises it: "I don't have time to build a briefing like that every single time." The answer is that you don't have to: writing a prompt is itself a language task, and language tasks are exactly the territory where these tools perform best. You can delegate to the AI the job of writing the prompt for the AI.

The technique is called meta-prompting, and it reverses the flow you've just gotten used to. Instead of showing up with the perfect request, you show up with the problem explained however it comes out, even messily, and you ask the tool to ask you the right questions and build the briefing itself. It works because the interview surfaces the context you would have forgotten to write down: the questions the AI asks you are often the ones a good consultant would ask in a first meeting.

The prompt builder

```
Help me build the best possible prompt for a task I need to complete.  
The task: [describe it however it comes out, even roughly].
```

```
Proceed like this:
```

1. Ask me the questions you need to clarify context, goal, and format (maximum five, one at a time).
2. When you have enough information, write the complete prompt, ready to copy into a new conversation.
3. Point out what I will need to add by hand (documents, data, constraints) before using it.

Try it once on a real task and you'll understand why people who work seriously with AI use this technique constantly: the prompt you get is almost always better than the one you would have written on your own, and the time invested in the interview is paid back by the first answer, which arrives usable instead of needing a redo.

But the reason this chapter exists is bigger than the technique. Meta-prompting is the first taste of a shift in posture: you stopped executing a piece of work, writing the prompt, and started having it executed, while keeping control of the result. That posture scales. The next level is permanent instructions, which you write once and which apply to every conversation. Then assistants configured for a precise role, which you don't have to reintroduce every time. Then agents: AI

systems that don't answer a question but execute an entire process, searching, comparing, filling in, producing, with you in the role of the person who assigns the work and evaluates the result. It's the passage from user to orchestrator, and it's the trajectory professional value is moving along right now: the winner isn't whoever writes the most elegant prompt, but whoever knows how to decide what to delegate, at what level of autonomy, and how to verify what comes back.

***In the full course:** lesson 2.6 develops meta-prompting into a method, 2.4 teaches iteration to refine any answer, 2.7 covers managing complex tasks in phases. And Modules 7 and 8 cover the territory of agents: the levels of autonomy, concrete cases by business function, what never to delegate, and agentic AI installed on your own computer.*

The ten-minute test: where your work is heading

So far, the tools. Now the question underneath it all, the one people ask under their breath: what about my job?

The two answers in circulation are both wrong. "AI can't do what I do, you need the human touch" is almost always hope dressed up as analysis. "In five years there will be no jobs left" is fear the data doesn't support. The measurable reality is more interesting than either.

Roles where AI amplifies human judgment without replacing it show a 22% increase in job postings. Roles where AI can execute the entire task show a 17% decline. Globally, the amplification potential is six times the potential for total substitution. But the decisive point is something else: this distinction doesn't apply to job titles. It applies to individual tasks. The same role, yours, almost certainly contains activities that are expanding and activities that are shrinking.

The real risk, then, isn't the machine replacing you. It's a colleague who uses AI well becoming more productive than you. It's not a machine against a human: it's a human with a machine against a human without one.

The useful response isn't worrying: it's measuring. This test takes ten minutes and any AI tool.

Prompt: the diagnosis of your tasks

Act as a labor market analyst.

Here is my role: [ROLE / what I do].

Here are the 5 tasks I spend the most time on: [TASK LIST].

For each task, tell me whether it requires non-routine judgment (AI amplifies me)
or is codifiable (AI can replace it), and why. Then tell me: overall, is my mix
of tasks expanding toward more judgment or shrinking toward execution?
Close with 3 concrete signals I should monitor over the next 6 months.

Fill in the two fields honestly, including the less glamorous parts of the job that take up more time than you'd admit. The picture you get isn't a verdict: it's a map. It tells you where to invest (the judgment activities, to amplify with AI) and where it's time to delegate (the codifiable activities, which someone will delegate anyway, with or without you).

To see how to read the result, take a typical case: an administrative manager who lists payment reconciliation, monthly reporting, supplier replies, handling contract anomalies, and training the new colleague. The test will tell her the first three are largely codifiable, and the last two are her real profession: judgment on ambiguous cases and transferring experience. The operational conclusion isn't "you're at risk." It's: use AI to compress the first three and free up hours for the two no machine will take from her. Almost everyone who runs this exercise discovers that the irreplaceable part of their job exists, but occupies a minority share of the week. Reversing that proportion is the professional project of the coming years, and it pays to start before someone else decides for you.

And there's one final data point that works as a compass. A second Anthropic study, 9,700 real users cross-referenced with their actual usage data, measured something counterintuitive: people who delegate more to AI are systematically more optimistic about their professional future, across every dimension measured, from pay to job security. Not out of naivety: because they've already seen the results firsthand, and AI has stopped being an abstract threat and become a tool they've built something with. Fear is highest where experience is lowest. The fastest way to reduce the first is to increase the second.

In the full course: lessons 1.4 and 1.5 develop this data in full, from the measured decline in junior roles to the domain-expertise multiplier, with the honest picture neither the doomsayers nor the minimizers will give you.

This was the first step

If you've made it here and run even one of the experiments along the way, you're already ahead of most people in your industry: the ones who stopped at the half-watched webinar or the newsletter skimmed on the train. I'm not saying that to flatter you. I'm saying it because it's useful to know where you stand.

Before we part, a taste of the format you'll find in every lesson of the full course. It's called a Reality Check, and it exists because the AI market systematically exaggerates in both directions.

REALITY CHECK

The most effective way to use AI today is this: use it to draft what you will then revise, use it to process text that would otherwise cost you hours, use it to generate options you will choose from. Don't use it as an oracle of facts, don't use it for critical calculations without verification. Keep this division of responsibility in mind and you'll rarely be disappointed.

This little book was the first step. The full journey is called "From User to Orchestrator," and it's built on the trajectory the data in these pages already pointed to: value is shifting from those who execute to those who know how to direct execution, by people and by AI systems, with the right level of control.

The numbers first, then the map: over sixty lessons in eight modules plus an orientation module, five operational appendices, and a library of over forty ready prompts with versions by profession. Every lesson closes with a Reality Check like the one above and an exercise to do right away. Here's what it contains, module by module:

- **Module 1 — Foundations.** The complete data behind the numbers in this little book: what AI does well, what it does poorly, who's actually winning with it. For those who want solid foundations instead of opinions.
- **Module 2 — Talking to AI.** Ten lessons that turn the briefing chapter into a complete method: the expert prompt, iteration, meta-prompting, and how to neutralize the AI that always agrees with you.
- **Module 3 — The tools.** ChatGPT, Claude, Gemini, and the rest: how to choose with criteria that outlive the rankings of the month, and when the free plan is truly enough.
- **Module 4 — Everyday productivity.** The practical core: eleven lessons on email, documents, research, meetings, data analysis, voice and images, all the way to AI as a personal tutor.

- **Module 5 — The language of business.** Value calculation, use cases that work and overrated ones, risks, copyright, and the diagnostic to locate where your organization stands.
- **Module 6 — AI in the team.** Delegating with the three supervision models, introducing AI to a group that resists, redesigning processes instead of bolting it onto the old ones.
- **Module 7 — Agents.** The territory almost nobody covers for non-technical readers: AI that doesn't answer but acts, the levels of autonomy, orchestration, what never to delegate.
- **Module 8 — AI on your own computer.** Agentic AI that works on your files and in your programs, with a guided path through Claude Cowork, one of the most widely used agentic suites of the moment: how to set it up without writing a line of code.

And the five appendices aren't footnotes but tools: a glossary of 35 terms explained in plain language, the complete prompt library, the resources to stay current, five rules for reading AI numbers without being fooled, and the sixteen-question self-assessment to diagnose your organization.

If you run a company, modules 1, 5, and 6 pay for themselves. If you want immediate results in daily work, modules 2 and 4 are the core. If you coordinate a team, start with 5, 6, and 7.

The course comes in two complete, identical versions, Italian and English, and it's updated regularly: as you read these lines it's already in its second edition, because this field doesn't stand still, and an AI course frozen in 2025 is already a historical document.

You'll find it here: ai.evolbot.com

The tool depreciates. Judgment doesn't. Start building yours today.

Matija Vidmar, The AI Architect