

THE URBAN PROFESSIONAL IN THE AGE OF AI

*How Planning Institutions Can Use AI
Without Losing Judgement, Memory
or Public Purpose*



REMCO DEELSTRA

THE URBAN PROFESSIONAL IN THE AGE OF AI

*How Planning Institutions Can Use AI
Without Losing Judgement, Memory
or Public Purpose*

REMCO DEELSTRA

First edition, June 2026

Based on essays first published on the *Cat with Hat* Substack (catwithhat.substack.com).

Contents

Preface	4
How to Read This Collection	5
Chapter Summaries	6
Prologue: A Tuesday Morning, Five Years from Now	7
PART I. Capability and Change	9
1. The Adoption Gap	10
2. Three Paradigms, One Professional	18
PART II. The Operational Foundation	27
3. The Machine Room	28
4. The Project Memory Room	36
PART III. Options, Models and Decisions	47
5. Options and Numbers	48
6. The Digital Twin	61
PART IV. Participation and Delegation	75
7. Participation	76
8. The Burden Moves Sideways	90
PART V. Public Purpose	99
9. The model made the right choice. That was the problem.	100
Epilogue: Tuesday Morning, Revisited	109
Where to Go Next	110
Sources and Notes	113
Explanatory Glossary	122
Index	127
Editorial and AI Note	131
About the Author	132
Colophon	133

Preface

This collection began with a practical question: if AI can already perform so much of urban planning's analytical work, why has the work itself changed so little?

Many tools are already useful. They can compare policy documents, structure participation responses, test programme options, maintain project memory and make complex models accessible in ordinary language. Professional practice is more than the production of analysis. It is also the organisation of attention, the preservation of institutional memory, the negotiation of contested values and the public defence of choices that leave a physical trace.

The essays were originally published as a series. For this collection, they have been adapted in places to improve readability, strengthen the connections between chapters and remove overlap. Each chapter can still be read on its own, while the collection develops one continuous argument from capability and institutional readiness to project memory, modelling, participation, delegation and public purpose. The arguments, examples and source qualifications remain unchanged.

This collection examines the institutional consequences of AI in planning. It does not catalogue software or predict automation. The technology will date quickly. Questions of memory, judgement and accountability will endure. What should be recorded? Which assumptions may a model vary? When is participation still capable of changing an outcome? Who verifies a system that is usually right? What does a young professional still need to learn by doing?

Across the chapters, one argument returns in different forms: AI does not remove the hard work of planning. It moves it. Production becomes cheaper, while verification, judgement and accountability become more important. Information can flow faster, while the purpose of the institution has to be stated more clearly.

Urban professionals still have room to shape this transition. That requires more than adopting tools. It requires designing the conditions under which those tools enter the work, and deciding in advance which public values must remain visible when the models begin to optimise.

The central claim is simple: public planning needs a shared discipline for AI use, built around institutional memory, contestable modelling, public value protection and accountable delegation. The question is which parts of judgement, memory and democratic contestability must become stronger because AI has entered the work.

How to Read This Collection

READ THROUGH For readers who want the full argument.	USE IN PRACTICE Start with chapter openings, figures and practice blocks.	CHECK SOURCES Use sidenotes for orientation and Sources and Notes for verification.
--	---	---

The fixed right column holds brief source orientation, a single argument anchor or an occasional pull quote. It also leaves room for your own notes.

A note on evidence

This collection uses three kinds of evidence. Some sources are planning-specific. Some come from adjacent fields such as public administration, law, software engineering and knowledge work. Some are early evidence from 2025 and 2026, including preprints and foresight reports. Evidence from adjacent fields identifies mechanisms; it does not prove direct planning outcomes.

A note on the framework

The collection builds toward an operating model for AI in public planning. The model has five layers: the machine room, the project memory room, the decision package, the civic decision room and the doctrine of use. Each chapter develops one part of that model. The final chapter brings them together.

Chapter Summaries

1 The Adoption Gap

AI capability already exceeds routine use in planning. The chapter locates the gap in organisational readiness, verification capacity and the design of everyday workflows.

2 Three Paradigms, One Professional

Expert analysis, communicative legitimacy and algorithmic governance now coexist. The professional task is to combine their strengths without allowing one logic to erase the others.

3 The Machine Room

The first durable gains appear in ordinary knowledge work: notes, comparisons, briefings and handovers. Reliability depends on a clear boundary between first-pass production and final responsibility.

4 The Project Memory Room

Long projects lose reasoning faster than they lose files. A governed memory preserves decisions, assumptions, alternatives, commitments and unresolved tensions across staff and political change.

5 Options and Numbers

AI can widen the option space and expose sensitivities earlier. The decision remains political because assumptions, weights and public commitments determine what the calculations mean.

6 The Digital Twin

A useful twin works as a contestable governance workspace; an authoritative picture works against that purpose. Its value lies in making consequences, omissions and uncertainty visible before choices harden.

7 Participation

AI can widen access and accelerate synthesis, while influence still depends on institutional design. Participation matters when people can alter choices that remain genuinely open.

8 The Burden Moves Sideways

Delegation relocates work toward verification, exception handling and accountability. Each higher level therefore needs stronger rules about oversight, learning and authority.

9 The model made the right choice. That was the problem.

The closing chapter asks what the institution has instructed AI to value. A doctrine of use protects public purpose before optimisation begins and identifies who carries judgement.

Prologue

A Tuesday Morning, Five Years from Now

It is seven-thirty on a Tuesday morning. Before the first meeting of the day, a housing strategist opens a briefing that was not there when she left the office the previous evening. Overnight, her AI assistant pulled the latest transaction data from the municipal property register, cross-referenced it with the housing programme targets she uploaded three weeks ago, and flagged a signal: sales velocity in the eastern development zone has dropped by 22% over the previous six weeks, concentrated in the mid-market owner-occupied segment that carries the financial logic of the entire phase-two programme. The briefing also shows what the data changes. It runs three revised residual calculations under different assumptions about programme mix, compares them against the subsidy conditions attached to the land allocation, and drafts a one-page note for the steering group meeting at nine o'clock. She reads it on the train. By the time she arrives, she has already decided which assumption to challenge.

Across town, a planning lawyer working on the development agreement for a mixed-use waterfront project receives an alert at eight. New secondary legislation under the national spatial framework has been published, effective in thirty days, with implications for the noise contour methodology referenced in three clauses of the draft agreement. His AI assistant has already identified the affected clauses, pulled the relevant passage from the new regulation, drafted proposed revised language for each clause, and flagged one instance where the revision creates a potential conflict with a covenant agreed six months earlier with the water authority. He does not need to read the full legislative text to understand what needs to change. He needs to make one call to the water authority and sign off on the revised language before the weekend.

In the same city, a participation coordinator has 847 written responses to a neighbourhood vision consultation, submitted over the previous four weeks. She has a presentation to the district council in five days. The responses have been processed overnight: themes clustered, recurring objections separated from emotional concerns, patterns in the geographic distribution of priorities mapped against the neighbourhood boundary, and a draft summary prepared in plain language at three reading levels, one for the council report, one for the project website, one for the community newsletter. She spends her morning deciding which tensions the council needs to confront and which framing will keep the political choices visible.

None of this is science fiction. Each of these workflows is possible today, with tools that are already available, on the basis of data that most urban organisations already hold. The deliberate decision to use it is what is missing.

In early March 2026, Anthropic published a labour market study that puts a precise measure on that gap. The researchers introduce the concept of “observed exposure”: a measure that combines theoretical AI capability with actual usage data, weighting automated and work-related applications more heavily. Their conclusion is consistent across virtually all knowledge-intensive occupations: actual use remains a fraction of what is technically feasible. Legislation, software integration, verification requirements and organisational habit all slow diffusion. Capability runs ahead of adoption.¹

1. Massenkoff, 2026

Part of that lag is cognitive and epistemic as well as organisational. Planning institutions, like individuals, tend to project recent patterns forward when making decisions. When those patterns no longer hold, heuristics calibrated to a previous decade keep generating decisions for a world that has already changed. The standard professional script of the 2010s, stable career trajectories, predictable regulatory environments, incremental technology adoption, is still the implicit planning horizon for many teams and organisations. AI sits outside that script. The resulting discomfort comes from a planning apparatus facing a transition for which its existing tools were never designed.

The academic world has noticed the same pattern in urban practice specifically. A bibliometric analysis published in October 2025 found that AI applications in urban planning have passed through three distinct stages: an embryonic phase from 1984 to 2017, a development phase from 2017 to 2023, and an explosive growth phase from 2023 onwards. The number of publications in that two-year stage was 3.6 times that of the previous 39 years combined. The research is accelerating. The practice is not keeping pace.²

2. Si, 2025

The chapters that follow map that gap. They examine what AI can already do in urban practice, what the research says, where the real risks sit, and where a team should start.

There is a specific reason to read this now rather than later. The decisions being made in the next two to three years, about which data platforms to invest in, which vendors to contract with, which workflow patterns to normalise and which governance structures to establish, will be difficult to reverse. Platform choices embed themselves in procurement relationships, staff habits, data architectures and institutional dependencies. Governance frameworks, once absent, are expensive to retrofit onto systems already in use. Knowledge architectures, not set up at project initiation, cannot be reconstructed from five years of unstructured history. The organisations and professions that are building their position now, understanding the technology, shaping its governance and investing in the underlying data foundations, will operate from a significantly stronger position than those that wait for the field to settle. The transition will compound. Those who inherit a structure they did not design will find it harder to change than those who shaped it from the start.



PART I

Capability and Change

The technology is moving faster than professional practice. The first task is to understand the gap, and the longer shift in planning that gives it meaning.

ARGUMENT MAP

This part moves from capability to institutional readiness.
Read for the gap between available tools and actual use.
Key tension: speed without judgement.

CHAPTER 1

The Adoption Gap

Why the distance between what AI can do and what planning organisations actually use is now a strategic issue.

IN THIS CHAPTER

Read the adoption gap through workflow design: access, verification, ownership and the institutional capacity to maintain governed use.

USE IT FOR

Diagnosing whether an AI gap is technical, organisational or institutional.

WATCH FOR

Efficiency that outruns verification and institutional readiness.

A participation coordinator arrives at her desk on a Monday morning. She has 847 written consultation responses, a council presentation in five days, and a team of two. In most planning organisations today, she will spend the first three days coding responses manually. By the time she reaches the analysis, the presentation is already half-written in her head, shaped more by fatigue than by what the responses actually say.

That workflow is not inevitable. Every part of it that consumes her first three days can be handled overnight by tools that already exist, on data she already holds. The organisational decision to use it is what is missing.

That decision is now a strategic planning issue.

The gap in numbers

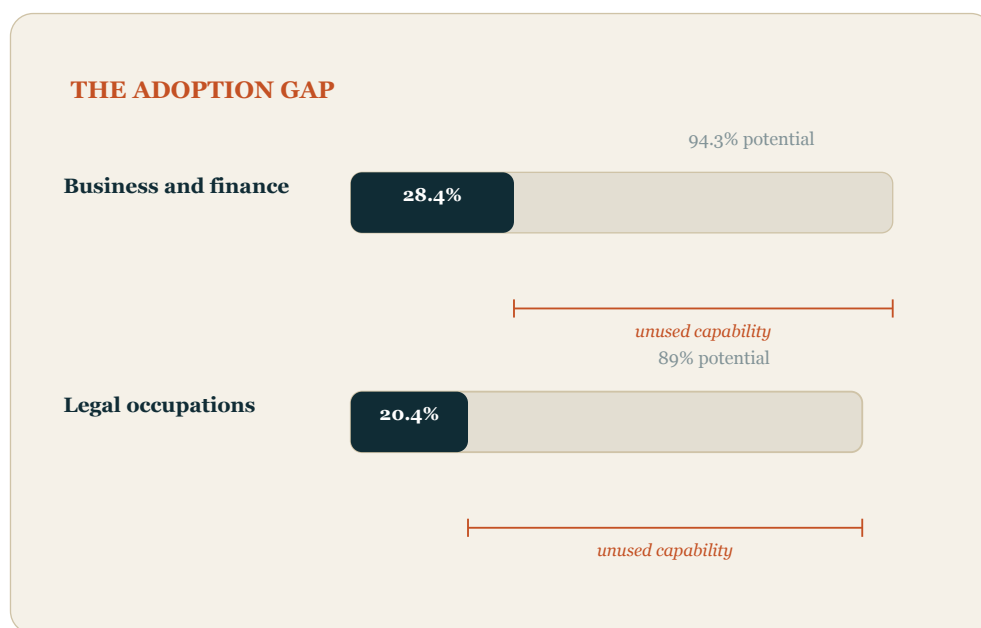


Figure 1.1. Capability and observed adoption. The bars compare theoretical task exposure with observed professional use. Source: author diagram based on Massenkoff and McCrory (2026), note 1.

In March 2026, Anthropic published a labour market study introducing a concept called “observed exposure”: a measure that quantifies which tasks AI could theoretically speed up and which are actually seeing automated usage in professional settings. For business and finance occupations, theoretical AI coverage sits at 94.3%. Observed coverage in practice is 28.4%. For legal occupations: 89% theoretically possible, 20.4% actually in use. The observed segment is much smaller than the potential segment. This capability overhang grows as capabilities advance. The unresolved question is how quickly organisations can turn potential into governed use.¹

1. Massenkoff, 2026

The adoption gap varies sharply within organisations

The gap is also unevenly distributed. OpenAI data from early 2026 shows that power users at the 95th percentile of ChatGPT usage engage the tool’s reasoning capabilities seven times more than the median paying user.² Both groups count as “AI users.” They are doing something categorically different.

2. OpenAI, 2026

The relevant adoption questions concern depth of use, task selection and the controls around each workflow.

One nuance worth naming: most of this evidence is US-based. Structural conditions may widen that lag in European public planning contexts. Procurement frameworks, data governance requirements, IT security policies and legal caution around AI-assisted decisions create real friction that the capability numbers do not capture. The gap in European municipal practice is likely wider than the aggregate figures suggest.

Research is accelerating faster than practice

In October 2025, researchers Si, Yao and Wu published a bibliometric analysis covering nearly four decades of AI applications in urban planning. The application of AI in urban planning has passed through three stages: an embryonic stage from 1984 to 2017, a development stage from 2017 to 2023, and an explosive growth stage from 2023 onwards. The number of publications in that two-year explosive growth stage was 3.6 times that of the previous 39 years combined.³

3. Si, 2025

Scientific attention to AI in planning is accelerating at a pace that has no historical precedent in the field. The practice is not keeping up. And the tools being described in that research are not abstract futures. Most of them are deployable today.

Why teams stall even when individuals are curious

Recognising the gap and closing it are different problems. Urban professionals who have tried AI and returned to their old workflows are not, in most cases, making an irrational choice. They are responding to real structural friction.

Procurement frameworks in most municipalities exclude or slow the acquisition of cloud-based AI tools. IT security policies prevent those tools from accessing the systems where relevant data lives. Budget cycles fund pilots but not the maintenance that follows. Legal departments have not yet developed clear positions on AI-assisted decisions that carry public accountability. Junior staff lack reusable prompts and verification routines. Senior staff lack the time to develop them. The data that would make AI most useful, complete, structured and queryable, is often incomplete, fragmented across legacy systems and inconsistently maintained.

Better tools leave these barriers intact. Closing the gap requires organisational decisions on procurement, access, maintenance, legal responsibility and verification. That is where the real gap lives.

The jagged frontier: where AI actually helps

The operational layer offers the clearest near-term return. The harder question is where that return stops inside a specific workflow, and why the boundary is more task-specific than most professionals expect.

Research on AI in professional settings describes what BCG and Harvard researchers call the jagged frontier. AI performs excellently on some cognitively demanding tasks and poorly on some apparently simple ones. The boundary is task-specific, not category-specific. And it shifts with every model generation.⁴

4. Described in:
Fabrizio Dell'Acqua,
2023

Return to the participation coordinator. Processing and structuring 847 texts, identifying recurring themes, mapping geographic concentrations of concern, producing summaries at multiple reading levels: AI is immediately useful in this part of the workflow. Research on large language models in participatory planning shows that NLP tools, software that reads and structures natural-language text at scale, can materially reduce the manual burden of first-pass synthesis at a scale that manual coding cannot match.⁵ What previously absorbed days of careful work can be completed overnight. That changes the sequence of the work as well as its speed.

5. Raymond et al.,
2025

But something else is also happening in those 847 responses that no AI reading them in isolation will catch. One submission reads neutrally on the surface. The coordinator who knows the project history recognises it as the voice of a community organisation that said almost exactly the same thing three years ago, before a previous round of consultation produced no visible result. That apparently neutral response carries exhaustion. The distance between those readings lies in the relationship between the text and a decade of accumulated institutional context.

This boundary separates textual analysis from institutional judgement. It also reveals one of AI's most underused capabilities. If that institutional context had been systematically logged and made queryable, a new team member arriving today could ask what was heard from that community in 2022, why it mattered, and how the project responded. AI can do that. Most urban organisations have not built the conditions for it. The gap between what AI can do for institutional memory and what teams are actually capturing is one of the most consequential vulnerabilities in area development, and it compounds every time a project manager moves on.

In this workflow, the boundary creates room for judgement, provided she actually has time for the work that is uniquely hers. The problem in most teams is that she does not, because she is still coding responses manually.

Three reasons to use AI, and why the choice matters

Before asking how to use AI, it is worth asking what for. There are three distinct answers, and they lead to fundamentally different approaches.

Doing more work: faster turnaround, higher output, more dossiers managed with the same team. This is the efficiency logic. It is the most common rationale and also the one most likely to make mediocre practice move faster.

Doing better work: stronger analysis, more thoroughly tested decisions, a wider range of perspectives brought to bear before a recommendation is made. AI can

also act as a structured challenge panel, improving the evidence and challenge applied before consequential decisions.

Doing new work: bringing disciplinary knowledge into the process that the team cannot staff directly. A municipality of 80,000 people cannot employ a behavioural economist, a climate adaptation specialist and a public health planner on every project. AI can surface relevant research from those fields, flag when a spatial proposal conflicts with established findings, and prompt better questions. This is the most underused and most interesting application.

These three goals are not equally available to every team. The structural barriers described above shape what is actually reachable. Being clear about which goal you are pursuing determines how you organise the work, what you measure and where the risks sit.

The risk that nobody explains concretely enough

01	Instruction hygiene
02	Independent AI check
03	Targeted human review

Table 1.1. Three verification layers. Each layer catches a different failure before AI-assisted work leaves the team. Source: author synthesis.

A 2026 Wharton School preprint reports controlled experiments in which participants could use AI to answer logic problems, with the AI secretly giving wrong answers half the time. Participants followed wrong AI answers 80% of the time. More significantly, their confidence went up when they had AI access, even though half the answers were deliberately wrong. High trust in AI was the strongest predictor: participants with high AI trust had 3.5 times greater odds of following a faulty answer than those with lower trust.⁶

6. Wharton, 2026, preprint

The study calls this cognitive surrender: the point at which you stop constructing the answer yourself and the AI’s output becomes your output without critical evaluation in between. Alberto Romero, writing in *The Algorithmic Bridge*, usefully frames the distinction: cognitive offloading is strategic delegation while retaining judgement. Surrender is accepting the output because it looks authoritative.⁷ What makes surrender dangerous is that it disguises itself as offloading. Most of the time AI is right, so the habit of not checking feels justified.

7. Romero, 2026

The mechanism has direct consequences in urban practice. A planning lawyer who sends out AI-drafted clause language without line-by-line verification against authoritative sources has not saved time. He has transferred risk onto a document that looks rigorous and may not be. A housing strategist who accepts an AI-generated residual land value calculation without checking the underlying assumptions has not accelerated the analysis. She has accelerated the possibility of a significant error reaching the steering group unchallenged. An urban designer who accepts an AI-generated shadow study without verifying the parameters for adjacent buildings risks an error and a legal challenge that runs for years.

Every AI-assisted workflow therefore needs a concrete verification procedure with three practical layers.

Instruction hygiene. When asking AI to analyse documents, synthesise policy or draft legal language, give it explicit constraints: *cite only from the*

documents I have provided, cite literally with page references, and flag explicitly when a question cannot be answered from the provided material.

This exposes hallucination immediately, before plausible prose can bury it.

AI-on-AI verification. A document generated in one AI session can be fed into a new session with a different prompt: check every factual claim, every citation and every source reference in this document. Do the cited sources exist? Are the quotations verbatim? Are there claims that do not follow from the provided material? The safeguard remains imperfect because models with similar training data can repeat the same errors. Independent checking still catches fabricated citations and misattributed claims that tired human review may miss.

Targeted human review. Concentrate verification time on the highest-risk categories: legal references, subsidy conditions, regulatory obligations, numerical assumptions in financial models. This is where errors are most consequential and where AI confidence combined with human fatigue creates the most exposure.

The case for waiting is not absurd

Tools are improving with every model generation, which means early adopters risk locking into platforms on 2025 assumptions. Legal liability for AI-assisted planning decisions is genuinely unresolved in most jurisdictions. The professionals who are waiting may be making a reasonable bet.

But waiting is not neutral. The decisions being made now about data platforms, workflow patterns and governance structures are difficult to reverse. Platform choices embed themselves in procurement relationships and institutional habits. Knowledge architectures not established now cannot be reconstructed from years of unstructured project history. The organisations building their foundations now will operate from a structurally stronger position than those inheriting a structure they did not design.

The question is how to move without creating dependencies you cannot audit or exit, or outrunning your own verification capacity.

Where to start: one workflow, one routine, one red flag

Pick one task you do every week that involves reading, summarising or drafting. The participation coordinator's overnight synthesis is one example. A weekly policy comparison is another. A first-draft briefing note for a steering group is a third. Do it with AI this week. Apply instruction hygiene from the start: give the tool the documents, tell it to cite literally, tell it to flag what it cannot answer. Evaluate the output against the standard required for responsible use. Run the AI-on-AI verification step on anything that will leave your desk. Adjust the prompt. Do it again next week.

The verification routine to build immediately: before sending any AI-assisted document, spend ten minutes on the three highest-risk elements. One legal reference checked against the source. One numerical assumption traced back to its origin. One claim that felt surprising verified against the primary document. Ten minutes. Every time.

The red flag that tells you cognitive surrender has begun: you stop reading AI outputs critically and start editing them. The shift from critical reader to copy editor is the point at which the AI's framing has replaced your own. When you notice it, stop. Rebuild the answer from your own reasoning first, then use AI to check and extend it. That sequence, human first, AI second, is the one that keeps your judgement intact.

Questions for Practice

- 01** Which recurring workflow offers enough value to justify a four-week trial?
- 02** Who owns the source material, the review and the final output?
- 03** Which errors would matter enough to stop or redesign the trial?
- 04** What evidence would distinguish an organisational constraint from a technical one?

CHAPTER 2

Three Paradigms, One Professional

Planning now carries three institutional logics at once: expert analysis, communicative legitimacy and algorithmic governance.

IN THIS CHAPTER

This chapter traces how the three planning logics emerged and why contemporary professionals must work across all three.

USE IT FOR

Recognising which planning logic governs a decision or workflow.

WATCH FOR

Algorithmic authority presented as a replacement for legitimacy.

A housing strategist sits across from her alderman. The programme they are reviewing has been in development for three years. The financial model is based on land values from 2022. The stakeholder analysis reflects a coalition that has partly dissolved. The scenario assumptions were built on a policy framework that was revised eight months ago. None of this is unusual. It is the normal condition of urban development work: by the time the analysis is complete, the world it was analysing has moved on.

She knows two things at once. First, that better tools now exist to reduce this lag: she has seen what AI-assisted scenario modelling can do with live transaction data. Second, that her organisation is nowhere near ready to use them in a way that she would trust in front of an alderman, a council committee, or a judicial review. The data is fragmented. The governance framework does not address AI-assisted outputs. Her team has no shared verification routine.

She is not caught between the past and the future. She is caught between three different institutional logics that her profession has accumulated over the past century, none of which has fully replaced the others, and all of which make demands on her at once.

Planning has gone through several major reorientations in what it is for, what constrains it, and what counts as professional competence. Three institutional logics now coexist inside most planning organisations, and understanding how

each emerged helps explain what the current AI moment is actually asking of professionals.

Understanding this pattern is not an academic exercise. Each major shift changed what the core constraint was, which changed what the most valuable skill was, which changed what it meant to be good at the job. Getting that analysis right is the difference between navigating the current transition deliberately and being shaped by it without realising it.

The first constraint: information lag

For most of the twentieth century, urban planning operated in what planning theorists call the rational-comprehensive mode. The logic was linear and optimistic: gather information, analyse it systematically, produce a plan that reflects the best available knowledge about how cities work and what they need. The ambition was scientific. The planner was an expert, applying rigorous methods to complex problems in the public interest.

The core constraint in this mode was information lag. By the time a plan reflected reality, reality had already moved on. Population projections went stale. Land markets shifted. Infrastructure decisions made on the basis of twenty-year forecasts collided with ten-year election cycles. The fundamental problem arose from the speed of the world relative to the planning apparatus.

The professional identity that emerged was the expert planner: trained in quantitative methods, versed in land use economics, authoritative on technical questions. The tools were projections, models, master plans. The legitimacy of the work rested on its analytical rigour.

The second constraint: coordination cost

The communicative turn of the 1980s and 1990s was a response to a double failure. First, the master plans of the post-war decades had produced environments that many communities experienced as alienating, unjust or simply wrong. Technical rationality had optimised for measurable variables and missed what actually mattered to people who lived in the places being planned. Second, the shift toward pluralistic, rights-conscious democracies made top-down expertise politically untenable in a new way.

Planning theorists, drawing on Habermas and his theory of communicative action, the idea that legitimacy comes from open, undistorted deliberation rather than from expert authority, proposed a different account of what planning was actually for. Planning was recast as something other than the production of technically optimal plans: a way to facilitate legitimate collective decision-making. The planner was no longer primarily an expert who knew the right answer. The planner was a facilitator who could help diverse actors with conflicting interests navigate toward shared futures.¹

1. Healey, 1997

There is an economic logic to this shift worth making explicit. In 1937, Ronald Coase asked why firms exist at all, if markets are as efficient as economists claim. His answer: markets are expensive to use. Finding counterparties, negotiating terms, monitoring outcomes: these transaction costs are real. Organisations exist to internalise them, replacing market friction with the cheaper mechanism of direction and authority.² A similar logic applies to planning institutions. Where spatial development creates high coordination costs, institutions emerge to structure negotiation, authority and legitimacy in ways that markets cannot. The communicative paradigm did not abandon the first mode's ambitions. It recognised that the primary cost preventing good outcomes had shifted from information gathering to coordination.

2. Coase, 1937

The professional identity this produced was not soft. It centred on listening carefully, building trust, managing conflict, and translating between technical and political registers. These are skills that take years to develop. They cannot be automated. Their value becomes clearer as AI enters the work.

A candidate third logic



Figure 2.1. Three planning logics. Each responds to a different constraint and asks something different of the professional. Source: author synthesis.

A third institutional logic is taking shape, though it would be premature to declare it a settled paradigm. Algorithmic urbanism, real-time data processing, iterative simulation and AI-assisted scenario development represent a shift in the nature of the primary pressure on planning practice. Where the rational-comprehensive mode was under pressure from information lag and the communicative turn from coordination cost, the emerging pressure is data quality, data governance and the ability of planning institutions to work with

tools that reason across complexity in ways that planners cannot replicate manually.³

3. Deelstra, 2026

One caution is essential here. AI reduces some information bottlenecks. It does not remove information lag as such. In public planning, lag sits in cadastral updates, fragmented administrative systems, missing local data, legal procedure and political timing. AI cannot solve what is not digitised, interoperable or lawfully shareable. And coordination cost has not disappeared. The hard part of contested urban development is still conflict, legitimacy and political consent. AI reduces some analytical friction. It does not dissolve distributive conflict.

Two older constraints have always shaped this work: information lag and coordination cost. A third pressure now joins them, one that was not binding before because the tools that expose it did not exist. The older constraints still matter. Organisations that have not invested in clean, structured and searchable data cannot use the tools that depend on it. Governance frameworks that have not addressed AI-assisted decisions cannot legitimise the outputs of systems that produce them. Professional cultures that have not developed verification routines cannot reliably detect when those systems are wrong.

These three logics are better understood as overlapping conditions that coexist in most planning organisations today than as a clean historical sequence. Robert Goodspeed has argued that collaborative planning never achieved the uncontested paradigm status its proponents claimed, and the same will likely be true of algorithmic urbanism.⁴ Each logic adds a pressure the previous one addressed inadequately; none simply replaces its predecessor.

4. Goodspeed, 2016

More than 90 percent of state chief information officers believe generative AI can enhance the citizen experience, yet only six percent report mature, scaled implementations today.⁵ That gap is primarily a readiness problem. The tools exist. The institutional infrastructure to use them responsibly does not, in most places, yet.

5. NASCIO, 2025

Coase also observed that the right boundary between what is organised internally and what is delegated externally is not fixed. Businesspeople will be constantly experimenting, controlling more or less, and in this way equilibrium will be maintained. The same logic applies here. The boundary between internal work and AI delegation requires continuous experiment and revision. Waiting is not neutral, because others are setting that boundary in the meantime.

A third time layer

01	Infrastructure cycle: decades
02	Political cycle: years
03	AI innovation cycle: months

Table 2.1. Three planning clocks. Durable institutions now make decisions across three very different rates of change. Source: author synthesis.

Planning has always operated between two time tensions. Long infrastructure cycles, the road that takes a decade to plan and fifty years to use, sit uneasily beside short political cycles of four-year mandates and shifting priorities. Managing that tension has always been part of the work.

Current foresight work by the APA suggests how rapidly AI-related tools are entering planning practice, adding a much faster innovation cycle to the already difficult relationship between long infrastructure timelines and shorter political cycles.⁶ The tools available in 2023 were qualitatively different from those in 2021. Those available today are qualitatively different again. No infrastructure cycle and no political mandate adjusts at that pace.

6. APA, 2026

This means that the planning apparatus, designed to manage long-term commitments under conditions of medium-term political change, is now also required to make durable institutional decisions about tools that may be superseded before those decisions are fully implemented. The organisations navigating this well are not necessarily those moving fastest. They are those that understand which decisions are reversible and which are not, and that protect reversibility wherever possible.

A less discussed knowledge shift

The shift in how planning is done is visible. The shift in how planning knowledge is accessed is less discussed but worth naming.

For the first two modes, the rational-comprehensive mode and the communicative paradigm, the primary information challenge was locating relevant knowledge and evaluating its quality. The signals of quality were institutional: peer review, organisational reputation, professional endorsement. That architecture is changing. Mike Caulfield, whose SIFT framework for source evaluation has become a standard in information literacy education, identified the structural shift in an interview published in August 2025: when AI generates extensive text that no human has fact-checked, and that text is shared and acted upon, evaluating the credibility of a source no longer suffices.

The test shifts from source-level credibility to claim-level accuracy: is this specific claim correct, and how would I know? That requires claim-level checking alongside source evaluation. Tracing assertions back to original documents. AI-generated information should be treated as a provisional first pass that requires confirmation before use.⁷

7. Caulfield, 2017

This shift extends beyond the professional context. The residents, elected officials and developers that planners work with are navigating the same changing information environment. As AI-generated content floods every channel at negligible production cost, the quality signals that previously indicated reliable information lose their discriminating power. Stakeholders arrive at participation processes having formed views through information streams that are increasingly AI-curated and inconsistently verified. Planning has always required building shared understanding across different perspectives. The candidate third logic does not make that task disappear. It makes the ground under it less stable. The participation chapter returns to this question. The shift runs in both directions, changing how planners access knowledge and changing the conditions under which shared understanding must be built.

What the third logic does not make obsolete

Here is the point most commonly missed in both the enthusiast and the alarmed versions of this conversation.

The skills that the communicative turn developed were a genuine advance. The ability to listen carefully to residents who are not using planning language, to identify whose interests are structurally underrepresented in a process, to build trust with communities that have experienced planning as something done to them rather than with them: these are not skills that survive by accident. They survive by deliberate protection.

Milad Malekzadeh, writing in *Cities* in 2026, argues that planners remain central as curators, applying AI tools critically and creatively while refusing to treat them as sources of objective answers.⁸ That framing is right, but a Coasean reading sharpens what it actually means. Coase distinguished between two fundamentally different roles. Initiative means forecasting and making new institutional arrangements. Management means reacting to existing conditions and rearranging what is already there. The curator is an initiator: deciding what matters, designing how it will be assessed, determining whose voice needs amplifying, shaping the governance architecture within which AI tools operate. The operator is a manager: rearranging inputs within parameters someone else has defined. Coase located the distinction between an agent and a servant in the freedom with which the agent carries out her work. The planner who can only enter inputs that vendor software permits, whose analytical output is bounded

8. Malekzadeh, 2026

by external parameter choices, is not a curator. She is operating under direction, and the parameters determine who holds it.

There is a further tension. The depth of professional judgement that good curation requires is built through doing the work that precedes it: the first syntheses, the initial comparisons, the early legal reads, the entry-level tasks that teach you where the sharp edges are. If those tasks are delegated to AI before that grounding has formed, the curator role may persist in title while the judgement it depends on quietly erodes. Chapter 8, *The Burden Moves Sideways*, examines that delegation risk in full. This supports deliberate choices about when and how those tasks are delegated to AI.

The planner who understands what a neighbourhood organisation has been through over the previous decade, who can feel when compliance replaces genuine engagement, and who recognises that a technically optimal scenario violates something a community cares deeply about that no variable measures: that knowledge resides in professional experience. No dataset contains it. It becomes more valuable, not less, as AI handles more of the work that does not require it.

The issue is whether institutions confuse what is easier to compute with what is more important to govern. Traffic flows are easier to model than fear. Sensor-rich corridors are easier to optimise than dignity. Permitting backlogs are easier to automate than democratic conflict. AI accelerates the answer that an institution is already giving. If that answer has not been examined, the acceleration is not neutral.

Who benefits, and who pays

One dimension of the transition that deserves explicit attention is the distributional question.

The organisations best positioned to benefit are those already investing in data infrastructure, governance capacity and analytical skill. Larger municipalities with established data estates, teams already experienced in structured knowledge management, senior staff who can delegate first-pass analytical work: these are the early winners. The tools compound existing advantage.

The organisations that pay are those with fragmented legacy systems, limited IT capacity, constrained procurement frameworks and junior-heavy teams whose professional formation depends on doing the first-pass work that AI increasingly handles.

The distribution argument tends to understate a further dimension. The third pressure shifts the bottleneck from coordination cost to data quality, then from internal professional capacity to external platform dependency. As analysis comes to depend on vendor-managed models, parties outside the planning

institution increasingly define what counts as a good outcome. That brings analytical power and transfers part of the definition of optimality to a party with different interests and limited accountability. Maroš Krivý's critique of cybernetic urbanism goes further: digitalisation of planning raises a question of control as well as readiness. When optimisation displaces deliberation, the appearance of improved governance can conceal its replacement by a system serving different ends.⁹

| 9. Krivý, 2018

The transition requires explicit choices about its distributional consequences, followed by procurement that reflects them. The same planning values that the communicative turn developed, attention to whose voice is underweighted and whose interests are structurally underrepresented, apply to the profession's own internal transition. They are rarely applied there.

What this asks of professionals trained in the second logic

Return to the housing strategist who opened the chapter. She has the judgement to know when a technically clean answer is politically or socially wrong. That is the hard-won competence of the communicative turn, and it does not become less valuable in the third logic. It becomes more so, because the analytical machinery that might otherwise absorb her attention is increasingly handled elsewhere.

What she needs to add is the capacity to work in a mode that her training did not anticipate. That requires learning how data pipelines work without becoming a data engineer. It requires understanding what AI tools can and cannot do without uncritical trust in either direction. It requires developing verification routines without becoming paralysed by doubt. And it requires carrying forward the communicative turn's central insight, that legitimacy is a condition for decisions to hold, into an environment where the analysis has become much more powerful.

Urban professionals who understand all three logics, and who can work in the third without losing sight of what the second taught about inclusion and legitimacy, are the ones who will shape how this transition actually goes. Those who treat it as a purely technical upgrade will miss its institutional dimensions. Those who resist it as an existential threat will watch the transition happen around them and lose the chance to shape it.

Questions for Practice

- 01** Which planning logic currently dominates this decision?
- 02** What becomes visible under that logic, and what drops out of view?
- 03** Where does the process rely on expertise, public legitimacy or algorithmic authority?
- 04** Who has standing to challenge the balance among those three forms of authority?



PART II

The Operational Foundation

The transition begins in ordinary work: notes, comparisons, briefings and handovers. Speed matters, but only when the organisation can retain the reasoning behind its output.

ARGUMENT MAP

This part moves from daily work to long-term project memory.
Read for the hidden infrastructure of professional judgement.
Key tension: efficiency without continuity.

CHAPTER 3

The Machine Room

The fastest gains appear in the hidden operational work that structures professional attention.

IN THIS CHAPTER

The first real gains will appear in the machine room: the ordinary work of notes, comparisons, drafts and handovers.

USE IT FOR

Designing reliable first-pass and final-pass operational workflows.

WATCH FOR

Automation that removes the work through which judgement develops.

A project manager closes her laptop at 18:07. The day did not fail because the strategic questions were too hard. It failed because the operational ones consumed it entirely.

There were two meetings that needed notes and action points. A legal comment on a draft cooperation agreement that had to be checked against an older version. Three resident emails waiting for a careful response. A briefing note for tomorrow's steering group that still had to be written. A spreadsheet with comments from three departments that did not yet align. A new colleague who needed the background of a project that began before he joined.

None of this is glamorous. All of it is real work. And in most planning organisations today, it crowds out everything else.

Urban practice describes itself through its visible moments: the political choice, the design review, the participation evening, the strategic memo. But most of the profession runs on a less visible layer underneath. Notes. Summaries. First drafts. Cross-checks. Handovers. Small acts of translation between legal language, political language, technical language and plain language. This is the machine room.

And this is where AI is likely to deliver the fastest returns, at least at the level of individual tasks.¹

1. Brynjolfsson et al., 2025

The hidden economics of urban work

Urban work is expensive in a very specific way. The expense comes from the accumulation of small translation costs between one part of the process and another.

Meetings, policy documents, participation processes and legal texts all produce material that still requires conversion into a usable record, comparison, synthesis or plain-language answer. Every one of those conversions takes time. Every one interrupts something else.

The operational layer matters because it forms part of the real work. It is the layer through which the analytical, legal, political and social dimensions of urban practice are made usable by others.

The machine room is also where fragmentation becomes visible first. Badly structured data first appears in ordinary failures, such as an analyst unable to compare two policy versions quickly. Weak governance first appears when a team has no routine for checking AI-assisted output before it leaves the office. Poor institutional memory first appears when a new project lead asks why a decision was made and finds the answer only in someone's inbox.

Four ordinary cases

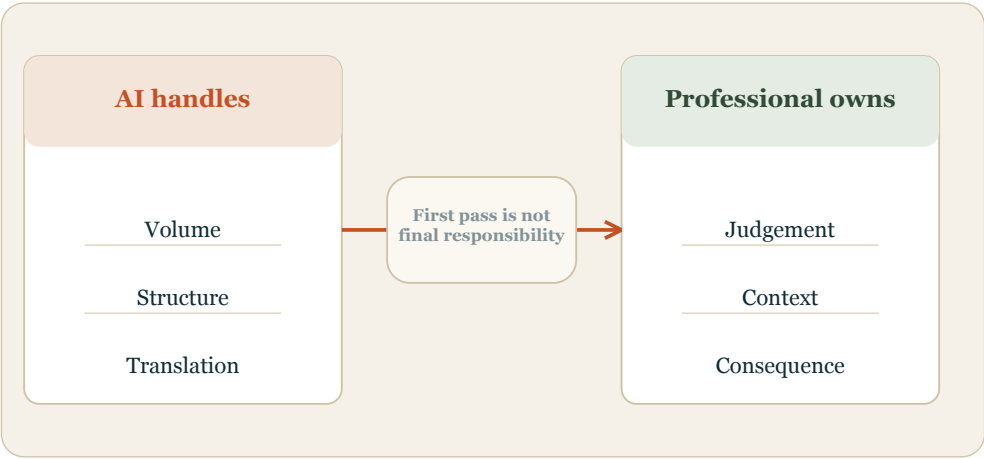


Figure 3.1. Division of labour in AI-assisted work. Professional attention remains with context, exceptions and consequences. Source: author synthesis.

These tasks vary in risk and share a family resemblance. They involve reading, structuring, comparing, extracting, drafting and translating. They are often repetitive in form, even when the consequences are not. They are exactly the kind of work where AI already performs well, provided the material is available and the verification routine is sound. The same research that documents the gains also documents what happens outside this zone: for tasks that fall beyond the frontier, AI assistance makes professionals approximately 19 percent less likely to reach the correct solution than colleagues working without it.

Professional judgement about where a task sits makes the machine room safe to operate.²

2. Dell'Acqua, 2023

Consider four ordinary cases.

A housing strategist needs a two-page note by 09:00 comparing a regional housing deal, a municipal coalition agreement and a newly revised affordability definition. The work is to identify overlap, tension, missing delivery mechanisms and political exposure across the supplied material. AI can usually produce a first structured pass in minutes, provided the documents are accessible and the task is well framed.

A planning lawyer needs to know whether a new draft clause differs materially from the agreement version discussed three weeks earlier. AI can mark the differences, group them by legal significance and flag the ones that alter obligations while separating ordinary wording changes.

A participation advisor needs to turn four hundred written responses into a usable typology of concerns, without pretending that frequency alone equals importance. AI can cluster themes, extract representative formulations and separate practical objections from trust-based ones, leaving the advisor time to judge what the themes actually mean.³

3. Raymond et al., 2025

A project team receives fifteen comments from different departments on a concept document. AI can merge, deduplicate, identify conflicts and produce a structured comment matrix instead of a mess of tracked changes.

In each case, the value lies in sequence as well as speed. Work that used to consume the first half of a process can now be moved earlier, creating more protected time for the judgement that follows. Urban professionals are rarely short of important judgement calls. They are short of the time in which to make them well.

Why efficiency is only the first gain

There is a narrow version of this argument: automate the routine, free up hours, increase output. That is true, but it misses the more important claim.

The machine room changes the quality of professional attention.

A tired professional narrows before formal compliance fails. Interpretation disappears before formal compliance does when time runs short. The second reading of a difficult response. The careful phrasing of a politically sensitive note. The effort to check whether a neutral-sounding objection is actually a signal of deeper distrust.

When operational overload dominates the day, analysis becomes thinner even if all formal tasks are completed. This is why the machine room is not a side issue. It shapes the conditions under which judgement happens.

There is a further point. The operational layer is where many organisations silently standardise their own quality. A team that produces consistent notes, clear action lists, traceable decisions and usable handovers will usually make better strategic choices than a team with the same intelligence but weak operational discipline. AI does not create that discipline by itself. It can, however, lower the cost of maintaining it.

That makes the machine room a governance issue as much as a tooling issue. The governance question is whether the organisation has decided what a meeting summary must contain: decisions taken, assumptions left open, owners, deadlines, unresolved tensions. AI can fill a template. It cannot decide what counts as a good record unless someone has done that institutional work first.

That institutional work is harder than it sounds. In most planning organisations, the operational layer has no real owner. The meeting has no assigned note-taker. The decision log has no standard. The briefing template does not exist. The absence reflects a management choice that has never been made explicitly. AI does not solve that. It inherits it.

There is also a prior condition that is easy to understate. The machine room works well when the documents it processes are accessible, structured and findable. In organisations where project files are scattered across personal drives, email threads and SharePoint folders that nobody fully navigates, The first problem is whether the material an AI tool needs exists in a usable form. Research on AI rollout readiness consistently finds a large gap between how many organisations believe they are AI-ready and how many actually are at the point of implementation.⁴ The machine room inherits whatever information hygiene the organisation already has. If the files are scattered and unstructured, the tools work on scattered, unstructured material; they organise nothing that was not organised first.

The risk of fluency without reliability

There is, however, a real risk in making this layer easier. The easier it becomes to produce acceptable-looking operational output, the easier it becomes to confuse fluency with reliability.

An AI-generated meeting summary may read more clearly than a hurried human one and still omit the single moment that matters most. A legal first draft may look rigorous and still contain a fabricated reference. The legal context makes this concrete: by 25 May 2026, Damien Charlotin's global database listed more than 1,400 court decisions addressing AI-generated

4. The readiness gap between perceived and actual AI-readiness at the point of rollout is documented across multiple enterprise AI adoption studies from 2025 to 2026, 2025

hallucinations. The tracker is a changing database with no fixed research sample, and it does not capture every filing in which fabricated material may have appeared. The pattern is consistent: the output looks authoritative, the professional does not check, and the error only surfaces when it causes damage.⁵ A resident response may be translated into calm administrative language and in the process lose the social meaning that gave it weight.

5. Charlotin, 2026

The cost of verification is real and often underestimated. Most of the measured productivity gains in the research literature are task-level gains. At organisational level, those gains are frequently offset by the time required to review, correct and re-check outputs. The net improvement is smaller and harder to see than headline figures suggest.⁶

6. The gap between task-level productivity gains and organisational-level impact is the central finding of the macro-level studies surveyed in Stanford HAI, 2026

A useful distinction helps: AI is useful for first-pass structuring; final responsibility remains human.

Let the system produce the summary, but a human checks whether the one politically decisive sentence is there. Let the system compare the documents, but a human checks the clauses that carry legal or financial consequence. Let the system cluster the consultation input, but a human checks whether the emotional centre of the material has survived the synthesis.⁷

7. Romero, 2026

AI outputs are often good enough to pass. Their fluency makes error harder to notice, and critical reading weakens when no verification routine is imposed. Cognitive surrender does not announce itself. It arrives gradually, dressed as productivity. Controlled experiments show that people follow incorrect AI answers around 80 percent of the time, and that confidence rises even when the AI is wrong.⁸ This does not prove the planning case. It identifies a mechanism planning organisations should treat as a serious risk.

8. Wharton, 2026, preprint

The apprenticeship problem

One objection deserves serious attention. Much professional judgement is built by doing exactly the first-pass work that AI now handles well. Early career staff learn by drafting notes, comparing texts, structuring comments and doing initial legal or policy reads. If those tasks disappear too quickly, what happens to judgement later?

The signals are hard to ignore. Entry-level hiring in AI-exposed professional fields dropped significantly between 2022 and 2024, as organisations found they could handle first-pass work without adding junior staff.⁹ Medical research on AI-assisted clinical work documents a parallel deskilling pattern: skills not practised atrophy, and practitioners become dependent on tools to complete tasks they previously performed independently. Neither finding proves a planning outcome. Together, they identify a professional formation mechanism that planning organisations should treat as a serious risk: you cannot learn to read a room, sense when a technically clean proposal has no political traction,

9. Stanford HAI, 2026

or recognise what a neutral-sounding objection actually signals, by reviewing AI output summaries of those situations.

There is a counterargument. Research on AI-assisted customer support found that novice workers gained the most, with 34 percent productivity improvements and suggestive evidence of faster skill transfer from senior to junior workers. In the right conditions, AI can accelerate learning; it need not obstruct it. But the conditions matter enormously. The learning benefit accrues when junior staff interact with the raw material and the AI output together, not when they only see the finished product.¹⁰

10. Raymond et al.,
2025

The real formulation is this: both dynamics are possible, and the direction depends entirely on how organisations design the work. The machine room as training ground requires deliberate construction. Junior staff should compare their own first pass against the system's, learn where the model over-compressed disagreement or missed contextual significance, and develop the habit of knowing what good looks like before they accept what fast produces.

That design choice will not happen by default. In most organisations, the path of least resistance is to give junior staff AI outputs to review and call it efficiency. That is the path toward a profession that can produce fluent documents and gradually loses the capacity to judge them.

What comes next, and why it changes the stakes

Two trajectories in model development are moving simultaneously, and they pull in different directions.

The first is improvement in grounded, document-anchored tasks. For work where the model has a source document to work from, hallucination rates on leading models have dropped sharply, from fifteen to twenty percent two years ago to below one percent on some benchmarks for summarisation and extraction tasks. This is precisely the zone where the machine room operates: processing what already exists, comparing what is in front of it, flagging differences between two known versions. That benchmark improvement is encouraging, but it does not remove the need to verify performance in the specific task and model being used. The four cases described in this chapter should become more reliable, not less, as models improve.

The second trajectory moves in the opposite direction. Models optimised for extended reasoning, the frontier systems being positioned as the next generation of professional AI, consistently hallucinate more on open-ended factual questions than simpler models do. When a reasoning model does not have a source document to anchor it, it fills the gap with statistically plausible confabulation at higher rates than its predecessors. This matters for the

machine room because the boundary between anchored tasks and open-ended tasks is not always obvious. A policy comparison is anchored. A question about whether a legal clause is consistent with emerging case law trends is not. As the machine room expands in scope and confidence, professionals may mistake the second category for the first.¹¹

11. Vectara

The shift from assistive to agentic AI compounds both trajectories. Systems that monitor documents, flag changes, draft notes and route outputs without an explicit prompt per task are already in early deployment. For the machine room, this means outputs arrive faster, at higher volume, and without the natural pause that a deliberate prompt used to create. The verification habits that are difficult to maintain today will be structurally harder to maintain when the system runs continuously. Building those habits now, while the workflow is still slow enough to make them visible, serves current practice and prepares the organisation for later autonomous operation.

For apprenticeship, the timeline is where the pressure is sharpest. The tasks disappearing from junior portfolios now are document-level: first drafts, comparisons, structured summaries. As reasoning and agentic systems mature, the next layer follows: first-pass research synthesis, initial stakeholder mapping, early scenario framing. These are precisely the tasks through which a junior planner currently learns to read a project, understand its politics, and develop the situational awareness that no amount of AI output review can substitute for. The organisations that do not deliberately redesign professional formation in the next two to three years will find the window has closed. The tools will have arrived before the pedagogy.

What a functional machine room looks like

A functional machine room has standard prompts for recurring tasks. It has templates for meeting notes, policy comparisons, first-pass legal checks, consultation summaries and handovers. It has named risk categories that always trigger human review: legal references, financial assumptions, policy compliance claims, and any statement that could materially shape a decision. It has a clear archive logic so outputs can be found, checked and reused.

Most importantly, it has a clear division of labour.

The machine handles volume, first-pass structure and repetitive translation. The professional handles judgement, exception, context and consequence.

That division sounds obvious. But most teams do not work this way. Either they keep everything manual and lose time, or they accept polished AI output too quickly because the document looks finished. The right arrangement is more disciplined than either extreme.

A planning organisation does not need a full AI strategy to begin here. It needs one workflow, one standard and one review routine. Start with something frequent and low-risk enough to learn from: a weekly policy comparison, meeting note processing, a briefing draft, a resident Q&A document. Then make one decision about quality: what must every output of this type contain? Then make one decision about review: which parts always need a human check?

That is enough to begin. But one condition has to be in place first: someone has to own the standard. Ownership must exist in day-to-day practice. The verification routine that depends on individual discipline under time pressure is the one that will not hold. The one that is assigned, scheduled and checked is the one that might.

Questions for Practice

- 01** Which first-pass task consumes attention without developing essential judgement?
- 02** What must every AI-assisted record preserve before it can leave the team?
- 03** Who owns the final pass, and what evidence must that person inspect?
- 04** Which task must remain in the junior learning path because judgement grows through doing it?

CHAPTER 4

The Project Memory Room

Long urban projects need more than files. They need a governed memory of decisions, assumptions, commitments and unresolved tensions.

IN THIS CHAPTER

Long projects fail when files survive but reasoning disappears. AI only helps if memory is governed.

USE IT FOR

Building project memory that survives staff, political and contractual change.

WATCH FOR

An archive that stores files but loses reasoning and commitments.

A new project manager joins a station area development in its ninth year.

On her first morning, she receives the usual handover package: the current programme, the latest planning schedule, the development agreement, the political decision note, the participation report and a folder of background documents that has clearly been copied forward for several years. It looks substantial.

It is not enough.

The documents tell her what the project currently says. They do not tell her why it says it. They do not tell her why the southern block lost two floors in 2021 and gained one back in 2023. They do not tell her which financial assumption the housing association challenged, or why the team decided not to change it. They do not tell her why the water authority accepted the drainage solution only on the condition that a later design phase would reopen the street profile. They do not tell her that a neighbourhood group still remembers a promise made by an alderman who left office two elections ago.

The formal record is full. The project memory is thin.

This is a normal condition of long urban development processes. Projects last longer than the teams that carry them. Councillors change. Project managers move on. Consultants rotate. Residents burn out. Developers sell their positions. Policy frameworks are revised and financial assumptions age. Participation histories compress into summaries that no longer carry the weight of what actually happened.

What remains is a strange institutional amnesia. The information usually survives somewhere. What gets lost is the meaning around it: why a choice was made, what it ruled out, and who agreed to it.

They need a project memory room.

An archive will not do it. A SharePoint folder will not do it. And a chatbot bolted onto disordered documents only makes the disorder faster to search. What the project needs is a governed, structured and queryable record: decisions, assumptions, objections, commitments, alternatives and the reasoning that connects them across time.

The point is to make deviation visible. When a project changes direction, the memory room makes explicit what is changing, which assumption no longer holds, who is affected, and which earlier commitment must now be explained. Storing everything is neither the goal nor achievable. This is a discipline of accountable adaptation, and it is the opposite of freezing the past.

Planning has a continuity problem

When planning organisations discuss AI readiness, they often describe the problem as a data problem. The data is fragmented, incomplete, outdated or locked in different systems. That is true. But in area development, the deeper problem is continuity.

The hard knowledge of a project rarely lives only in formal documents. It lives between them.¹

1. Nonaka, 1995

A council decision records that a preferred development model was approved. It may not record that two alternatives were politically impossible because they depended on land the municipality did not control. A participation report records that residents raised concerns about traffic. It may not record that traffic was the acceptable public language for a deeper distrust created by a previous process. A financial model records a residual land value. It may not record which assumptions were treated as negotiable and which were politically fixed.

This matters because area development is cumulative. Each phase rests on earlier choices. If the reasoning behind those choices disappears, later teams either treat the current position as self-evident or reopen questions without understanding why they were closed. Both outcomes are damaging. Treating inherited decisions as self-evident creates path dependency without reflection. Reopening them without understanding them creates churn, frustration and loss of trust. In both cases, the project becomes less intelligent over time, even as the file grows.

Research confirms the pattern. A 2017 comparative study of three large infrastructure projects found that knowledge in these settings tends to remain

tacit, that explicit project management knowledge is not systematically captured, and that organisational learning remains attached to individuals and absent from the institution's shared assets.² The consequence is a profession that re-learns the same lessons project after project, because no mechanism transfers them. A separate systematic review, covering 91 empirical studies on knowledge loss caused by staff turnover, found that This structural organisational problem compounds every time a senior person leaves.³

2. Aerts, 2017

3. Galan, 2023

Urban development sits at the intersection of both pressures: long-duration projects and repeated changes in the people and organisations carrying them. Knowledge loss is built into the structure.

The goal is better institutional memory, achieved through more selective and structured recording.

The difference between a file and a memory

A file answers one question: what exists?

A memory answers a different question: what mattered, why did it matter, and what follows from it?

In practice that difference decides what a successor can actually use.

A project file may contain minutes from a meeting with the housing association. A project memory should know that the meeting changed the risk allocation in the programme, that the change depended on a subsidy assumption that was never formally tested, and that the same issue will return when the next phase moves to contract.

A project file may contain a participation report. A project memory should know which concerns were structurally underrepresented, which commitments were made in response, which were later dropped, and whether the community was told why.

A project file may contain design variants. A project memory should know why a variant was rejected: too expensive, legally fragile, politically unacceptable, spatially weak, or simply premature at that moment in the process.

AI can help here, but the first design choice becomes critical. A large language model cannot reconstruct project memory from documents that never captured the relevant reasoning. It can only work with what has been recorded. That makes the founding decision more important than any technical choice about which tool to use.

Project memory has to be built from the beginning, while the reasoning is still in the room and easy to record. At project initialisation, the team should ask a simple question: what will a competent successor need to know in five years?

That question changes the structure of the work. Meeting notes should record decisions and assumptions left open. Option studies should record the selected option and the reasons for rejecting alternatives. Participation summaries should preserve frequency counts, minority positions and signals of trust or distrust.

AI can lower the cost of all of this. It can extract decisions from meetings, draft structured logs, identify recurring commitments, connect related documents and make the record queryable. But the standard has to be set by a human. The organisation has to decide what counts as memory worth keeping.

The part that gets avoided

A good project memory room preserves both the preferred account and the inconvenient record.

Old commitments that were quietly set aside. Financial assumptions that were never properly tested. Participation processes that were formally recorded but materially changed little. Conflicts managed into silence. Promises made by politicians who are no longer in office, to communities that have not forgotten them.

Forgetting is not always a failure. Sometimes it is convenient. It gives organisations room to move without explaining what shifted. A good memory room removes that room. And that is exactly why many organisations will say they want better memory while quietly preferring the flexibility of forgetting.

This is the central tension in any serious implementation.

A project memory room changes the accountability regime as well as project efficiency. When an AI assistant can surface that a commitment was made in 2022 and not followed up in 2024, the question of why becomes harder to sidestep. When a rejected alternative from 2021 becomes relevant again after a policy change in 2026, the organisation must explain why it was closed before proceeding as though it had never been considered.

Memory and forgetting are both selective. Planning institutions therefore need to replace unmanaged forgetting with governed remembering.

What a project memory room actually remembers

01	Decisions	02	Assumptions
03	Alternatives	04	Commitments
05	Unresolved tensions		

Table 4.1. Project memory in area development. These components describe area development. Other project types require different elements. Source: author synthesis.

A functional project memory room does not need to remember everything equally. It needs to remember what future judgement depends on.

Five categories matter most.

Decisions. What was decided, by whom, when, under what mandate, and with what conditions attached?

Assumptions. Which demographic, financial, legal, political or spatial assumptions supported the decision? Which were contested? Which need review when circumstances change?

Alternatives. What options were considered and rejected? Why? Are they permanently closed, or only under the conditions that applied at that moment?

Commitments. What was promised to residents, partners, landowners, housing associations, developers, water authorities or other public bodies? Was the commitment formal, political, verbal, conditional or only implied?

Unresolved tensions. Which issues were deliberately parked? Which conflicts remain unsettled beneath the management response? Which risks were accepted because no better option existed at the time?

Consider what a project director actually needs before a meeting on a revised density proposal. She wants to know which previous commitments are affected by the change, which rejected alternatives become relevant again, and which assumptions have expired since the last council decision. A well-designed memory room does not return a single answer. It returns a map: three affected commitments, two assumptions that have been overtaken by a revised subsidy framework, one unresolved issue with the water authority from 2022 that the

new design reactivates, and four source documents that require review before the meeting. This turns search results into institutional intelligence.

All of this is possible with technology that already exists. RAG-based systems, which ground AI responses in the organisation’s own documents rather than in general model memory, can already make the record queryable and source-traceable.⁴ Knowledge graphs can map connections between decisions, assumptions and documents. Decision logs and prompt logs can create the audit trail that accountability requires. The technical layer matters. But the decisive question is institutional: what is recorded, who may rely on it, and who may challenge it?

4. Klesel, 2025

How to start

01 Decision log	02 Assumption register
03 Alternative record	04 Commitment register
05 Unresolved tension log	

Table 4.2. Five project memory routines. Source: author synthesis.

Five routines are enough to begin.

A **decision log** that records decisions, the reasons behind them, the assumptions they depended on and the source documents they drew on. No material decision without a reason. No reason without a source.

A **commitment register** that tracks what was promised to residents, partners and public bodies, including the status of each commitment, who owns it and when follow-up is required. Promises made in participation are rarely legally binding, but they are politically real. Losing track of them is how trust erodes.

An **assumption register** that records the financial, demographic, legal, spatial and political assumptions on which major decisions rest, including when each assumption should be reviewed. Many project failures do not begin with a wrong decision. They begin with an old assumption that remained in use after the world around it changed. This register is also the direct foundation for serious option analysis: alternatives cannot be evaluated well until the organisation knows which assumptions are still current.

An **option history** that records rejected alternatives and the reason for rejection. Its purpose is to distinguish what was permanently ruled out from what was merely premature, without reopening every decision. That distinction matters more than most teams realise when circumstances shift later in the process.

A **handover interview protocol** for every departing project lead, policy lead or external advisor. Structured questions. Recorded answers. The tacit knowledge that would otherwise leave with the person.

AI can support all five. It can draft log entries from meeting notes, detect missing owners in a commitment register, compare current documents against earlier commitments, and prepare handover questions based on the project's recorded history. Recent research also suggests a complementary use: AI-generated output as a reflective prompt, presenting a partial or hypothetical reconstruction of a decision process to help professionals articulate reasoning they hold tacitly but have never written down.⁵ But each entry should be reviewed by a human who understands the project. The standard is simple: no AI output without a reviewer. No summary without a route back to the material it summarised.

5. Zhao, 2025

This ordinary work keeps long projects intelligent.

The governance layer

Once memory exists, who controls it matters enormously.

This distinction, between a memory room as a learning instrument and a memory room as a forensic record, is the design question that matters most. The same data can serve both purposes. A record of why a financial assumption was challenged can help a future team understand the reasoning. It can also be used, years later, to establish who knew what and when. People are aware of both possibilities simultaneously, and they will calibrate their behaviour accordingly. No governance rule resolves this tension on its own. What resolves it, partially, is organisational culture and explicit norm-setting: a shared understanding that this record exists to help future teams make better decisions, not to construct a retrospective case against past ones.

Other domains have found partial solutions worth borrowing. The FAA's Aviation Safety Action Program is a useful precedent: a safety partnership between the regulator and certificate holder that may also include an employees' labour organisation. It encourages voluntary reporting through defined enforcement-related incentives and corrective action.⁶ The result is a structured route for reporting near-misses and unsafe conditions that might otherwise remain unrecorded. Medical practice distinguishes between the clinical record and the clinician's working notes. Legal practice protects privileged communication from disclosure. These analogies map imperfectly

6. Federal Aviation Administration, 2020

onto urban development and still point toward the same principle: some knowledge can only be captured if the people who hold it trust that capturing it will not be used against them. A project memory room that wants honest input about hesitation, about rejected reasoning, about what almost went wrong, needs to build that trust into its design before it asks for that input. Planning organisations have yet to settle the practical form. The starting question remains one that most have yet to ask: what is this memory for, and what is it explicitly not for?

If the system records some voices and not others, it reproduces that imbalance. If it treats formal documents as more real than informal commitments, it favours institutional actors over residents. If it privileges frequency over significance, it mistakes organised volume for democratic weight. If it stores unresolved tensions as “closed”, it turns managerial convenience into historical fact.

The project memory room must therefore be designed with challenge built in. Users should be able to see confidence levels, inspect sources and flag errors. Public-facing summaries should be open to correction.

There is also a legal dimension. Where an application falls within the EU AI Act's high-risk regime, the regulation sets demanding requirements for logging, transparency and human oversight.⁷ Urban planning organisations should not assume that these systems will remain outside demanding governance expectations, especially when they inform decisions about housing access, permits, spatial rights or public services. Building the logging and traceability layer from project initialisation is materially easier than constructing it retroactively.

7. EU AI Act

The practical implication is concrete: the prompt, the source base, the model version, the reviewer and the resulting output may all become part of the accountability record. That discipline makes AI usable in public planning.

If an AI assistant tells a project manager that the community accepted the revised mobility plan in 2025, the project manager must be able to ask: based on which documents, which meeting notes, which participation responses, and with what level of uncertainty? If the answer cannot be traced, it is not project memory.

The result is a plausible account whose claims cannot be traced to the project record.

Bad project memory becomes more dangerous when AI makes it searchable, fluent and authoritative. Without AI, forgetting is slow and visible. With AI, it can look like institutional knowledge.

The democratic value of being remembered

Residents experience planning through time. They remember what was said and who actually listened. They notice when a concern disappears into a report and is never heard of again, and they notice when the next project manager arrives and acts as if the conversation is starting from scratch.

Institutional memory failure is often experienced as bad faith, even when no one is acting dishonestly. A promise is made in 2024. The team changes in 2026. The programme shifts in 2027. By 2028, the promise is no longer visible in the current documents. From inside the organisation, this may look like ordinary project evolution. From outside, it looks like being ignored.

That is why the project memory room should serve internal efficiency and contain a public-facing accountability layer. At minimum, residents and partner organisations should be able to see what was heard, how it was interpreted, what commitments were made, what changed later, and why. Internal negotiations require different levels of confidentiality. Any planning process that asks people for trust still owes them a traceable account of how their input moved through the project.

This matters especially when AI is used to synthesise participation. Speed can flatten meaning. A summary can be accurate in frequency and still wrong in significance. It can count the words in a hundred objections and still miss the exhaustion behind them. Cluster the themes, and the history that made those themes matter drops out of the summary. A project memory room should therefore make participation synthesis contestable. Participants should be able to ask whether this is what was said, whether this is how it was interpreted, and where their input changed the programme and where it did not.

Speed can flatten meaning.

The apprenticeship argument

A well-maintained project memory room can help junior professionals learn faster. Instead of entering a project through a chaotic folder structure and asking around until they piece together a version of the story, they can trace how a decision evolved, compare early assumptions with later outcomes, and see how legal, financial, spatial and political reasoning interacted over time.

But there is a risk in that. If AI turns project history into neat summaries, juniors may learn the conclusion without learning the process that produced it. They receive the answer and lose the chance to develop the capacity to construct one. That is a real professional formation problem, and its consequences compound over a career.

The memory room should serve as both a learning environment and a briefing tool. Junior staff should be asked to reconstruct the reasoning behind a decision from the sources, then compare their reading with the AI summary and their

supervisor's interpretation. Professional judgement develops through learning how conclusions become defensible.

The planning organisation as a remembering institution

Urban planning has always been future-oriented. It draws scenarios, tests options and imagines future places. The quality of those futures depends heavily on the capacity to remember well.

A city that cannot remember why it made a decision will struggle to change that decision well. The same failure shows up at every scale. A project that forgets what it promised loses trust. A team that forgets which assumptions were fragile mistakes habit for stability. And an organisation that forgets what it learned repeats its own mistakes, now with better software.

AI makes this problem visible precisely because it raises the standard. Once a project can be queried, the absence of memory becomes harder to excuse. The question "why did we decide this?" no longer has to disappear into inboxes, old folders and the fading recollection of people who have moved on.

But the answer will only be reliable if the memory was built deliberately. Do not start with the model. Start with the memory standard. What must be remembered? For whom? With what evidence? Under what access rules? With what right to correct?

Only then does AI become useful here.

The full project memory room described in this chapter is still unusual in urban development. Its components already exist in decision logs, knowledge systems, RAG architectures, audit trails, participation records and project learning routines. The five routines provide a practical starting point for the next project, before its first consequential decision needs to be explained years later.

Because area development is more than a design problem, a finance problem or a participation problem. Over time it becomes a continuity problem, and continuity is not something files produce on their own.

Structured memory keeps current commitments, expired assumptions and changes of direction open to explanation.

Questions for Practice

- 01** Which decisions must a successor be able to reconstruct five years from now?
- 02** Where are assumptions, rejected alternatives and unresolved tensions recorded today?
- 03** Who may correct, contest or add context to the project memory?
- 04** Which public commitments must remain visible when circumstances change?



PART III

Options, Models and Decisions

*AI widens the option space and makes assumptions visible earlier.
The question is whether those assumptions can be inspected before
preferences harden.*

ARGUMENT MAP

This part moves from option generation to spatial consequence.
Read for assumptions, sensitivities and model authority.
Key tension: better calculation without better governance.

CHAPTER 5

Options and Numbers

AI can widen the option space. The decision sits in the assumptions, sensitivities and trade-offs behind the numbers.

IN THIS CHAPTER

Build an assumption register first, then use sensitivity analysis to show which premises carry each option.

USE IT FOR

Testing options without hiding assumptions, commitments or uncertainty.

WATCH FOR

False precision and political choices hidden inside parameters.

A housing strategist receives a revised programme for a station area on Thursday afternoon.

The numbers don't work.

The change is ordinary and therefore harder to isolate: the residual land value has shifted. Construction costs have moved. The mid-market segment that carried part of the financial logic is weaker than it was eighteen months ago. The housing association can still participate, but not on the same risk allocation. The developer wants more owner-occupied housing. The alderman wants to hold the affordability promise. The neighbourhood group remembers the public courtyard commitment. The finance department wants a clearer view of exposure before the next steering group.

The team asks for options.

What follows is a familiar sequence. The financial advisor adjusts the model. The urban designer checks what extra volume might fit. The policy lead tests the affordability mix. The mobility advisor recalculates parking pressure. The project manager tries to assemble these partial views into a coherent recommendation. Two weeks later, the team has three options.

Option A is politically attractive but financially weak. Option B is financially stronger but socially harder to defend. Option C is called balanced, which usually means nobody loves it and nobody has yet found the hidden problem.

This is how option work often happens. Slowly. Sequentially. With each discipline holding part of the truth, and with the decisive trade-offs only becoming visible after preferences have already hardened.

AI changes when trade-offs become visible and how much work is required to expose them.

The option space is larger than the meeting can hold

Urban development is often discussed as if the task is to choose between a small number of alternatives. In practice, the number of possible alternatives is enormous.

Change the share of social rental. Change the phasing. Change the parking ratio. Change the block depth. Move the school. Reduce underground parking. Add a floor on the north side, remove one on the south. Change the tenure mix. Delay the public space investment. Use a different subsidy assumption. Shift risk from the public to the private side, or the other way around. Accept a lower land receipt. Change the energy concept. The list is endless.

Each adjustment affects something else. A programme mix is simultaneously a housing, financial, political, mobility and legal choice. A parking norm changes construction cost, public space, target groups, marketability and residual land value. A design variant also changes spatial conditions: it changes sunlight, noise mitigation, foundation costs, energy performance and the credibility of what was promised in participation.

This is why teams narrow the option space early. They have to. Human working memory cannot hold hundreds of combinations across ten constraints simultaneously. So teams simplify. They select a few plausible options, test those, and treat the selected set as if it represents the real decision space.

Often it doesn't.

It represents the part of the decision space the organisation had time to see.

A 2024 systematic review of computational optimisation in urban design describes precisely this gap: Urban Design Optimisation methods allow teams to explore large numbers of design solutions for a district, evaluate multiple objectives simultaneously, and select from Pareto-efficient solutions rather than from a pre-narrowed shortlist.¹

1. Tay, 2024

A Pareto-efficient solution is one where you cannot improve on one objective without making another worse. In area development terms: you cannot increase the affordable housing share without reducing the land revenue, or improve the public space quality without reducing buildable volume, unless you change the underlying conditions. There is no option that scores best on

everything. What the method produces is a frontier: a set of combinations where every step forward on one dimension requires a step back on another. The professional task is to choose the most defensible trade-off given the project's commitments, its political context, and what was promised to whom.

The tools have been available for years. The constraint has also been organisational: most planning teams have lacked the data, workflows and decision routines to use these methods in ordinary project work.

AI changes that equation by making established optimisation methods usable in the context of a real project: connected to its commitments, its political constraints, its financial assumptions, its partner agreements.

Numbers don't decide. They discipline the conversation.

There's a common mistake in discussions about AI-assisted feasibility work. The idea is simple: better models produce better decisions because they calculate more precisely.

That conclusion is incomplete.

The value of numbers in urban development is not that they decide. They discipline the conversation. They force a team to say what it's assuming. They reveal which ambition depends on which subsidy. They show when a political narrative rests on a financial condition nobody has yet tested.

A feasibility model is not reality. It's an argument in numerical form.

That argument may be useful, honest and carefully built. It is still an argument. Underneath it sit assumptions about sales values, rent levels, cost inflation, phasing, interest rates, land values, subsidy conditions and risk allocation. Some of those assumptions rest on evidence. Some were negotiated. Some were inherited from an earlier phase and never re-examined. And some are political choices wearing the clothing of technical inputs.

At this point, project memory becomes operational. A project memory room records decisions, assumptions, alternatives, commitments and unresolved tensions. It records why rejected options were rejected, and whether they were permanently closed or only impossible under the conditions of that moment. That is the direct foundation for option work.

A constraint map gives AI a basis for generating responsible variants. Without one, it can only generate variants.

An AI assistant can run twenty programme mixes in a minute. But if the affordability promise from 2023 is missing from the source base, one of those

options may look efficient precisely because it has forgotten the thing the project must still honour.

A model that forgets commitments doesn't create a better option. It creates a more elegant breach of trust.

The sensitivity analysis becomes the argument

The most useful number in a feasibility model is often not the outcome. It's the sensitivity.

Which assumption moves the result most? Which variable has become load-bearing? Which cost increase breaks the programme? Which rent assumption makes the affordable share possible? Which phasing delay creates the cashflow problem? Which public investment is treated as fixed, even though it was never formally secured?

AI makes sensitivity work faster. That is useful. But the more important shift is positional: sensitivity analysis can move from the back of a technical appendix to the centre of the decision conversation.

A steering group should see the stronger financial result for Option B and the assumptions that produce it. It should see that Option B depends on a sales value assumption based on market data from before the last interest rate shift, and that the same option reduces the public courtyard residents were told would remain. It should see that Option A has a weaker land result but is less exposed to market absorption risk and better aligned with the housing association agreement.

That choice of weighting is the decision itself, dressed up as a technical detail.

The option note of the AI era should not be organised around preferred options alone. It should be organised around load-bearing assumptions.

What must be true for this option to work? What happens if it's not? Who carries the risk if it fails? Which commitment becomes vulnerable? Which future option does this choice keep open, and which one does it close?

The last two questions are often missing. They are also the ones that matter most in long urban projects.

There is also a false precision problem. A model that reports a residual land value to the nearest euro may look more certain than the situation allows. In early option work, ranges are often more honest than point estimates. A band of uncertainty tells the steering group more than a precise number built on fragile assumptions. The reason is psychological: people trust a number with two decimals more than a range, even when the range is the more accurate

representation. AI can produce false precision faster and more convincingly than any spreadsheet. That is one more reason to keep the assumption register in the room.

What AI takes over, and what it doesn't

This is a good moment to be concrete about the division of labour. It's also where most discussions about AI in planning stay too vague.

AI can now handle some tasks in option work with limited human intervention, provided the inputs, constraints and review rules are clear. Running multiple programme variants against a fixed set of financial parameters. Generating a sensitivity table showing how the residual land value shifts under different cost or revenue assumptions. Flagging whether a proposed option may conflict with a recorded commitment in the project memory. Producing a plain-language summary of what each option costs on each criterion. These are structuring tasks. They used to consume days. They now take minutes.

Some tasks AI supports but cannot own. Translating the results of a sensitivity analysis into a recommendation that a steering group can act on. Deciding which assumptions are load-bearing and which can be loosened. Identifying which combination of conditions is politically negotiable and which is not. Judging whether a formally compliant option is actually deliverable, given the relationships, trust levels and institutional capacity in the room. These tasks require contextual knowledge that lives in the professional, not in the model.

Some decisions remain outside AI's role. Deciding what the city owes to which residents. Setting the weights: whether affordability counts more than land revenue, whether speed outweighs public space quality. Knowing when to surface a number and when to keep it off the table. Carrying the responsibility for a recommendation when the choice is genuinely hard. Those remain human. The technology is powerful enough; the problem is that authority, trust and accountability cannot be handed to a model.

The practical implication is this: a professional who used to spend two weeks producing three options can now spend two days producing a structured option field. Freed time is the point. The goal is to spend more professional attention on the options that matter, not to manufacture more of them.

But only if the organisation makes that trade consciously. If the time saved on variant generation is simply filled with more variant generation, nothing changes.

Keeping options open is a strategy, not indecision

Urban development has a cultural bias toward certainty. Decision notes prefer a clear recommendation. Political processes prefer visible progress. Project teams are rewarded for moving forward. Uncertainty is something to reduce before a decision is made.

Sometimes that's right. Often it's not.

In long projects, the best decision is not always the one that optimises the current model. It may be the one that keeps the most valuable future options open.

This idea has been formalised in climate adaptation. Haasnoot et al. (2013) developed the Dynamic Adaptive Policy Pathways approach at Deltares and TU Delft, starting from a simple recognition: many long-term decisions must be made under deep uncertainty, so planning should combine near-term action with future flexibility. The method identifies adaptation tipping points: the conditions under which a current policy will no longer meet its objectives and a new one is needed.²

2. Haasnoot, 2013

Area development needs a similar discipline.

The reason it transfers is structural: both housing programmes and flood risk strategies involve long-lived investments, uncertain futures and decisions that lock in before the uncertainty is resolved.

A project may not know whether the mid-market rental segment will recover in three years. It may not know whether a mobility hub will reduce car ownership enough to justify a lower parking ratio. It may not know whether a subsidy scheme will survive the next cabinet. It may not know whether a later climate standard will make today's public space design insufficient.

The answer is not to wait for certainty. Certainty arrives too late.

The answer is to make decisions that are explicit about thresholds. If sales velocity falls below this level, revisit the tenure mix. If parking pressure exceeds this level after phase one, activate the mobility reserve. If construction costs rise above this band, reopen the phasing sequence. If the subsidy is not secured by this date, shift to the fallback programme.

AI can help maintain this adaptive logic. It can monitor indicators, update scenarios, flag expired assumptions and show which pathway the project is drifting toward. But the adaptive strategy itself remains a governance choice. The organisation must decide which uncertainties matter, which signals count, and what it's willing to change when conditions shift.

That adaptive response depends on institutional discipline.

The metrics trap

There's a danger here, and it's serious.

Once options can be scored quickly, the score starts to acquire authority.

This is partly structural and partly psychological. What is measured is easier to defend. What is harder to measure becomes easier to treat as secondary. A number produced by a model often gains extra authority over a colleague's estimate. A colleague's estimate can be questioned directly. A model's output first requires understanding what the model assumed, which most people in the room do not have time for. So the number stands. Often it stays in place for one reason only: nobody knows how to push back without looking like they are rejecting evidence.

AI intensifies this tendency because it can produce more scores, more quickly, with more apparent precision.

A 2026 study in *npj Urban Sustainability* examines this directly. Based on 28 AI implementations selected from 157 deployments across six urban domains, Moghaddam and Cao document cases in which technical accuracy did not translate into the intended policy result. One documented example: an acoustic detection system reached 97 percent accuracy on its own metric while yielding only 9.1 percent effectiveness on the actual policy objective it was deployed to address.³

3. Moghaddam & Cao

The example is not imported here because gunshot detection resembles area development. It does not. It matters because the failure mode is transferable: a system can perform well against its own technical metric and still fail against the public purpose for which it was deployed.

For option work in area development, the pattern is close at hand.

The objective determines what disappears from view. Housing totals say little about household formation, land revenue can displace long-term public value, green-space area does not show whether teenage girls, older residents or children without private gardens can use it, and average accessibility conceals who remains disconnected.

Biased data is one cause. A clean dataset can still produce a bad decision when the model works exactly as designed against the wrong objective.

This is the part of options and numbers that people would rather not name. Better calculation doesn't remove politics. It moves politics upstream into the choice of variables, weights and thresholds.

Who decides whether affordability counts more than land revenue? Who decides whether a public space commitment is a hard constraint or a soft

preference? Who decides when a model is allowed to say that a politically desired option is not financially feasible?

These are governance questions with numbers attached, even when they arrive dressed as technical ones.

This is why weighting should not be hidden inside the model. If affordability counts twice as much as land revenue, say so. If delivery speed is allowed to outweigh public space quality, say so. If political feasibility is included as a criterion, name who assessed it and on what basis. A weighting framework is a political document written in numerical form.

Goodhart's Law describes a recurrent governance problem: once an indicator becomes a target, behaviour adapts to the indicator and weakens its value as a measure of the underlying objective. In option work, this means that the moment a score becomes the goal, the underlying purpose it was meant to serve starts to recede. In area development, this is why a housing number can start to crowd out housing need, why a green space norm can crowd out actual use, and why a financial score can crowd out public value. AI accelerates an existing risk.

What the decision package should contain

The most valuable output from AI-assisted option work is a better decision package, not a beautiful image or a ranked list. A strong package holds at least six things.

01 Constraint map	02 Assumption register
03 Commitment check	04 Sensitivity summary
05 Option history	06 Recommendation Separate calculation from judgement.

Table 5.1. Decision package. Source: author synthesis.

A **constraint map**: what is legally fixed, financially constrained, contractually committed, spatially bounded, politically preferred or self-imposed? Many planning processes collapse this distinction, treating soft preferences as hard constraints and ignoring actual hard constraints until they cause a problem.

An **assumption register**: which numbers drive the result, where do they come from, who owns them, and when do they expire? No number without an assumption. No assumption without an owner.

A **commitment check**: which promises, participation outcomes or partner agreements are touched by each option? The project memory room and the option field connect at this point. An option that performs well against financial criteria but contradicts a formal commitment to a neighbourhood group is not a free choice. It has a cost.

A **sensitivity summary**, written in plain language and placed in the main decision package. The purpose is to show the steering group where the project is fragile, and why.

An **option history**: what was tested, what was rejected, and under what conditions could it return?

A **recommendation that separates calculation from judgement**. The model may show that an option performs better on selected criteria. The professional still has to explain whether that option is defensible.

AI becomes genuinely useful in a planning organisation when it makes the choice harder to obscure and leaves the choice with the organisation.

The new professional skill: specifying the question

Less time on	More time on
producing options	understanding them
building models	knowing which assumptions to challenge
writing the trade-off summary	explaining it to people who will resist parts of it

Table 5.2. The shift in option work. Less production time creates more room for interpretation and explanation. Source: author synthesis.

The work becomes more interpretive and more political, not less.

The model changes the urban professional's work; the role remains.

The tasks that used to define option work move partly to AI: running the model, generating variants, building the comparison, writing the first draft of the trade-off summary. What remains is harder to delegate and harder to learn from a textbook.

The most important skill becomes the ability to specify the question precisely enough that the output is actually useful.

A weak prompt asks: make three options for this site.

A stronger prompt asks: *generate programme variants for 450 to 650 dwellings, with at least 35 percent affordable housing, no loss of the public courtyard commitment, no net increase in public parking pressure, and a maximum public investment exposure of X. Use current cost assumptions but test sensitivity for construction cost increases of 5, 10 and 15 percent. Flag options that depend on assumptions not yet formally approved. Identify which earlier commitments each option affects. Do not rank the options until the constraints and assumptions have been shown.*

That is better planning, not just better prompting.

The professional who can write that specification understands the project. She knows which constraints are real, which assumptions are fragile, which promises matter and which numbers are politically loaded. AI makes her more effective because she already knows what the work requires.

The opposite is equally true. A professional who can't specify the question will receive outputs that look useful and may be fundamentally misdirected.

This is also where the character of professional development shifts. Junior professionals used to learn option work by doing option work: building models, running variants, seeing what changed when you moved a parameter. That learning path is partly disrupted when AI handles the mechanical steps. The risk is that something important is lost before something better replaces it. The skill of reading a model critically, of knowing which number to distrust, of recognising when a result is too clean, develops through practice with the model itself. If that practice disappears, the professional may be able to specify questions but unable to evaluate answers.

This supports deliberate choices about what junior professionals still do by hand.

Where conflict must remain visible

There's a temptation to make option work smoother. AI can help produce cleaner dashboards, clearer scorecards and more confident recommendations. That's useful. But some discomfort must remain.

The purpose of option analysis is not to remove discomfort. It's to locate it accurately.

If all options look good, the model is probably hiding the conflict. If one option clearly wins, the evaluation framework may be too narrow. If the numbers settle the matter, the wrong questions may have been excluded.

Good option work should reveal where the project is genuinely constrained. It should show the alderman that keeping the affordability promise requires either more subsidy, lower land revenue, higher density, reduced parking cost or a different phasing strategy. It should show the developer that viability depends on public choices that cannot be treated as technical inputs. It should show residents when a promise is being changed, and why. It should show the project team which decision can be made now and which one should remain conditional.

And it should show where the institutional capacity to make those decisions is thin.

Which option is best is only half the question. The other half is whether the organisation has the governance capacity, the political will and the partner trust to carry it out. An option that depends on a decision the organisation cannot make will not survive contact with reality.

AI can make the conversation more precise. It cannot take the pain out of it.

How to start

A planning organisation doesn't need a full digital twin to begin with options and numbers.

It needs one live feasibility model, one structured assumption register and one repeatable option template.

Start with a project where a programme decision is coming. Take the existing model. Identify the ten assumptions that matter most. Record their source, owner, date and review condition. Then ask AI to help generate a structured option field around those assumptions: a map of feasible combinations and exposed trade-offs. The final recommendation remains a separate act of judgement.

Use the first round internally. Don't present it as truth. Present it as a challenge to the team's existing understanding.

Which option did the team expect to work, but doesn't? Which rejected option becomes viable under a changed assumption? Which preferred option depends on an assumption nobody owns? Which political promise has a financial consequence that hasn't yet been named?

That's enough to begin.

The next step is to connect the option field to the project memory room. Every option study should leave a trace: what was tested, what was rejected, why it was rejected, which assumptions were used, which commitments were affected, and what should be reviewed if conditions change.

That creates cumulative intelligence. The next time the market shifts, the team doesn't start again. It asks which assumptions have expired and which alternatives should return.

What this changes

The machine room chapter argued that AI first changes the operational conditions of urban work. The project memory chapter showed why long-running projects need a governed memory: decisions without memory become brittle, and organisations that cannot remember cannot adapt responsibly.

Once the machine room works and the project memory exists, option generation becomes different. The team is no longer limited to the three variants it had time to prepare. It can explore a wider field, test assumptions earlier, expose trade-offs more honestly and keep future pathways visible.

The hardest constraints remain: land ownership, finance gaps, political courage, legal procedure, public trust, implementation capacity. AI doesn't dissolve any of them.

Back to the housing strategist on Thursday afternoon. The numbers don't work. That hasn't changed. But what changes is what she can do with that problem before the next steering group. She can run the option field against the actual constraints. She can test which assumptions are load-bearing and which can be moved. She can check which commitments are affected by each variant. She can show the team the contours of the decision space, explain why three options are too narrow, and state the conditions under which each path holds.

She still has to decide which path to recommend.

AI leaves judgement in the work and changes the material it operates on.

The promise of options and numbers was never certainty. What it offers is a wider decision space, a clearer view of what each path costs, and a better record of why one path won out over another.

Questions for Practice

- 01** Which assumptions carry the preferred result?
- 02** What changes when each influential assumption is varied within a credible range?
- 03** Which public commitments constrain the option space before ranking begins?
- 04** Does the decision package separate calculation, interpretation and recommendation?

CHAPTER 6

The Digital Twin

A digital twin is a governance workspace where assumptions become visible and contestable.

IN THIS CHAPTER

A digital twin is useful when it makes consequences visible before decisions harden.

USE IT FOR

Using digital twins as contestable decision spaces, with safeguards against authoritative pictures.

WATCH FOR

Visual certainty that exceeds the model's knowledge.

A planning team opens the digital twin of a station area on the morning of a steering group meeting.

The model is impressive.

The station is there. The tracks are there. The planned housing blocks are there, coloured by phase. The public space layer shows the new square, the cycling routes and the line of trees along the main street. The mobility layer shows parking demand. The climate layer shows heat stress. The water layer shows where heavy rain collects. The energy layer shows the grid constraint that everyone knows about but no one wants to make central yet.

The team loads three programme options.

Option A adds density close to the station. The model shows a better housing yield and stronger land revenue, but also more shadow on the schoolyard in winter and a sharper peak in electricity demand during phase two. Option B protects more public space and reduces construction risk, but weakens the affordability mix. Option C keeps the political promise on affordable housing, but only if the parking ratio is lowered and the mobility hub is delivered before the first residents arrive.

For the first time, the team can see the problem together.

That is progress.

It is also the risk.

A digital twin makes a contested future look present. It turns assumptions into images. It gives uncertain relationships the appearance of a system. It lets councillors, residents, developers and public agencies point at the same screen and believe they are looking at the same problem.

The twin has not decided.

A digital twin is not a better picture

01	Model: represents selected relationships
02	Map: locates information
03	Dashboard: monitors indicators
04	Digital twin: connects changing data, models and decisions to a real system

Table 6.1. Model, map, dashboard and digital twin. The categories can overlap, but only the twin is designed to remain dynamically connected to the represented system. Source: author working definition based on Batty (2024), note 1.

The first mistake is to treat the digital twin as a visualisation tool.

That is understandable. The visual layer is what people notice first. Buildings can be rotated. Flood depths can be coloured. Shadow can move across a public square. Traffic flows can be animated. A future neighbourhood can be made visible before it exists.

But if that is all the twin does, it remains a more expensive version of the presentation.

A serious urban digital twin links place, data, assumptions and decisions. Visibility is a byproduct. Consequence is the point.

Michael Batty describes digital twins in city planning as covering a wide range of models, from aggregate economic and behavioural processes to local simulations and detailed spatial models. The central issue is how scientists, policymakers and planners interact with real cities through these models: to understand, predict and design.¹

1. Batty, 2024

A useful minimum definition follows from that. An urban digital twin is a digital representation of a physical urban system, made up of entities and relationships over time, used to simulate possible changes and inform decisions. A representation that lacks relationships between layers, a time dimension and a defined decision purpose remains a model, a map or a dashboard. Useful in its own right; not the same thing.

An urban digital twin is a digital representation of a physical urban system, made up of entities and relationships over time, used to simulate possible changes and inform decisions.

Figure 6.1. Working definition of an urban digital twin. Source: author working definition based on Batty (2024), note 1.

The distinction matters in practice. A digital model is updated manually. A digital shadow receives data from the physical world automatically, but mainly in one direction. A digital twin goes further: it connects data, models, assumptions and decisions in ways that allow change in one layer to affect another. In urban planning, that bidirectional flow rarely means a machine acting on the city. More often it means a decision being made to confront its own consequences.

If we add 1,200 homes to this station area, what happens to school capacity, mobility demand, shadow, heat, water storage, energy load, construction phasing and public space quality? If we lower the parking ratio, what has to be true for the mobility system to hold? If the housing association can no longer carry the same share of risk, which programme commitments become vulnerable? If the water authority updates its climate scenario, which earlier design choices must return to the table?

The digital twin becomes useful when it can hold those questions in relation to each other.

The twin creates a room in which the answer becomes harder to hide.

The twin is where layers meet

01	Technical twin: asset and system performance
02	Project twin: options, constraints and delivery
03	Civic twin: consequences, participation and public choice

Table 6.2. Three versions of a digital twin. The versions differ in purpose, users and governance; technical sophistication alone does not distinguish them. Source: author synthesis.

Urban development is full of layers that are usually handled separately.

The financial model sits with the planeconomist. The mobility assumptions sit with the traffic advisor. The phasing logic sits with the project manager. The legal commitments sit in agreements and council decisions. The public space ambition sits in design documents. The participation history sits in summaries. The energy constraint sits partly with the grid operator, partly with the municipal energy team, partly in everyone’s anxiety.

Each layer is rational on its own terms. The problem is that the decision is not made in one layer.

A higher density option changes the financial result, the construction sequence, the energy demand, the sunlight impact, the school pressure and the political story. It is also a programme and governance move. A lower parking ratio changes target groups, marketability, public space, residual land value and the credibility of the commitment to prevent parking spillover into adjacent streets. It is also a financial, spatial and distributional choice. A climate adaptation measure changes maintenance costs, public space design, heat comfort, biodiversity, land allocation and sometimes the number of homes that can be built. It looks like a water choice while spanning many parts of the project.

The promise of the digital twin is earlier visibility of these relationships, while accepting that its knowledge remains partial.

A massing change can trigger a first daylight check. A phasing change can trigger a first infrastructure dependency check. A programme change can trigger a first school capacity and energy load check. A public space reduction can trigger a commitment check against the participation record.

The early result is a warning system, not a final truth.

AI matters because it makes the twin’s structured data and models queryable. It lets a project director ask in plain language: show me which assumptions change if we move 300 homes from phase three to phase two. It lets a councillor ask: which option best protects the affordable housing commitment, and what

does that cost elsewhere? It lets a resident group ask: how does this revised mobility plan affect walking routes to the primary school?

A twin operated only by specialists keeps access tied to technical skill. AI shifts the question from who knows how to run the model to who is allowed to question it.

From specialist instrument to governance workspace

The GIS team builds them. The data team maintains them. The consultants demonstrate them. The project team uses screenshots in presentations. Councillors see the output, not the working environment. Residents see a polished view, not the assumptions behind it.

That model has limits.

It keeps the twin technically impressive and institutionally weak. The people who most need to understand the trade-offs remain outside the actual model. They receive a translation of the twin and remain outside the decision space itself.

The smart city tradition often prizes speed: real-time sensing, instant response, automated adjustment. Urban planning needs a different ideal. Many planning choices are slow for a reason. They involve rights, losses, future obligations and democratic accountability. The useful twin is not always the fastest twin. Sometimes its value is that it slows the room down enough to see the consequences before a decision hardens.

There are three versions of this. A technical twin is used by specialists. A project twin is used by teams and partners to test choices. A civic twin goes further: assumptions can be challenged, scenarios contested and decisions traced by actors with real standing. Governance conditions mean that most twins will remain below that level, appropriately so. The direction still matters.

Open in this context does not mean everyone can change everything. That would be irresponsible. It means different actors can ask questions, see the assumptions relevant to their role, understand the evidence base and challenge what the model is doing.

A water authority can add a revised climate assumption and show its effect on the design. A housing association can test whether a programme option still meets its financing and management constraints. A mobility department can show the dependency between parking reduction, public transport frequency and the timing of a mobility hub. A neighbourhood group can see whether an earlier commitment on public space has been preserved, changed or quietly

traded away. A councillor can see that the option labelled “balanced” is only balanced because three risks have been pushed into later phases.

The twin does not remove disagreement. It gives disagreement a more precise object.

The European Commission’s Local Digital Twin Toolbox is moving in this direction by presenting local digital twins as modular, standards-based tools that help cities simulate, analyse and plan urban environments, with an emphasis on interoperability, openness and scalability.² Interoperability determines whether the twin can function as a public decision room across systems. A closed visual platform may look like a digital twin. Institutionally, it may function like a vendor-controlled presentation layer.

2. European
Commission, 2025

A political distinction, not a technical one

Who can ask questions of the model? Who can see the assumptions? Who can challenge the data? Who decides which indicators matter? Who benefits when the model becomes the accepted version of the project?

The digital twin becomes a decision room only when those questions are answered explicitly.

Precision can be unwelcome. A digital twin that makes trade-offs visible also removes the room that ambiguity provides. A commitment that was never quite defined cannot be broken. A risk that was never quite located cannot be owned. Some actors benefit from that fog because ambiguity may be the condition holding a coalition together, even without bad faith.

The harder critique is this.

Some argue that visible trade-offs can depoliticise planning instead of democratising it. When a councillor can see that option C only works if the mobility hub is delivered early, the political decision has already been displaced into the question “what does the model say?” Complex choices become technical parameters. The room changes, but the power over what goes into the model, whose data counts and which scenarios are run does not. On this reading, the civic digital commons is a more sophisticated form of managed consent, not a step toward shared governance.

Otherwise it becomes a very beautiful black box.

The authoritative picture problem

Intentional omission bounded, documented and open to challenge	Accidental omission missing, stale or invisible to decision makers
--	--

Table 6.3. Two kinds of omission. A partial model can be useful when the boundary is explicit. Hidden absence produces false authority. Source: author synthesis.

There is a particular danger in urban digital twins: they look more certain than they are.

A spreadsheet can be challenged because it looks like a spreadsheet. A policy note can be challenged because everyone understands that language frames reality. A map can be challenged because maps have visible omissions.

A high-resolution digital twin is harder to resist. It feels complete. It feels neutral. It feels like the place itself has spoken.

It has not.

Every twin is selective. It includes some relationships and excludes others. It has a boundary, a time step, assumptions about behaviour, climate, mobility, costs, demographics, public space use and institutional capacity. It has data that is current and data that is stale. Trust, fear, informal use, social safety, attachment, loss, memory, political exhaustion: these are not easily captured by sensor data or 3D geometry.

The persuasive surface of a twin creates a burden of proof. Whoever challenges the model has to prove that the model is wrong. Whoever built the model does not have to prove that it is right. That asymmetry is institutional, not mathematical. It shapes who speaks first in the room, whose objections are taken seriously, and which concerns get translated into a parameter and which get written off as sentiment.

Stefano Moroni’s 2025 work on the limitations of city-scale digital twins makes this point from a planning theory perspective. Urban knowledge is dispersed, partly tacit and partly local. It cannot be fully captured in a technical representation, however advanced the model becomes.³

3. Moroni, 2025

That is exactly why the digital twin must be connected to the project memory room.

A twin without memory can show the spatial effect of a new proposal. It may not show that the same corner of the site carries a broken promise from 2022. A twin without participation history can show that a public square still meets a

technical standard. It may not show that residents were told it would be the social heart of the neighbourhood and now experience the revision as a loss. A twin without institutional context can show that a phasing change is efficient. It may not show that it depends on an agreement with a partner whose trust has already been weakened.

The model can be accurate and still be incomplete.

But incompleteness takes two forms, and the difference matters. Some omissions are intentional: the model abstracts from reality because a useful twin is not a full copy of the world. Other omissions are accidental: missing data, measurement errors, outdated inputs, human mistakes, relationships that were never modelled because no one thought to include them. Governance starts with knowing which is which. What has been deliberately excluded because it is outside scope? And what has been excluded because the organisation does not know, measure or maintain it? A useful abstraction is different from a blind spot. That distinction must be visible.

That is the authoritative picture problem: the more persuasive the twin becomes, the more discipline it takes to remember that it is still an argument about the city.

The participation promise is real, but not automatic

01	Not possible: a legal or physical constraint
02	Not chosen: an available option was rejected
03	Not funded: the option remains possible but lacks an approved budget

Table 6.4. Three statuses that must remain distinct. Clear status labels prevent preference and budget pressure from being presented as necessity. Source: author synthesis.

Digital twins are often presented as tools for participation.

There is good reason for that. A twin can make complex spatial choices easier to understand. It can allow people to explore scenarios before they respond to finished plans. It can show the consequences of different choices in ways that written reports rarely achieve.

A 2025 systematic review on digital twins and citizen participation in land use agenda-setting found that citizen-centred twins can support engagement by helping citizens frame problems, convince others and propose alternatives.⁴

4. Adade & de Vries, 2025

That is promising. But access to the model is not the same as influence over the decision.

Public input that cannot change the model, the decision or the project memory turns the twin into participation theatre with better graphics. People are invited to interact after the option space has closed. The model displays the proposed public space while concealing the financial trade-off that reduced it. That is managed attention, not participation.⁵

5. Cardullo, 2019

A digital twin can also create a new participation divide: between people who can work through the model and people whose knowledge remains outside it. If access and routes to challenge are left undesigned, the twin raises barriers as easily as it lowers them.

For a digital twin to support participation, it needs access and contestability. One of the most useful things a twin can do is separate three statements that planning processes often blur: not possible, not chosen and not funded. Many conflicts become toxic precisely because those three are collapsed into one administrative answer. A twin that keeps them distinct gives residents, partners and councillors better questions to ask. Better questions are more valuable here than a better visualisation.

What AI adds to the twin

A digital twin represents the urban system. AI helps people question it.

Many twins still depend on expert mediation. Someone has to open the right layer, change the right parameter, run the right simulation and explain the output. AI can lower that access cost. A project lead can ask which option creates the greatest exposure to delayed grid capacity. A councillor can ask which assumptions changed since the last decision. A resident can ask how a revised mobility plan affects the route to the primary school.

The system should answer while showing its sources, assumptions and uncertainty.

But AI cannot own final responsibility. It may explain the model, generate scenarios and surface conflicts. Decisions about legitimate trade-offs, politically honest scenarios and institutional conflict remain human.

That boundary is what keeps the tool in its place.

The twin needs rules before authority

01	Interoperability and semantics
02	Infrastructure
03	Data acquisition and actuation
04	Data quality and harmonisation
05	Modelling, simulation and decision support
06	Visualisation and display
07	Human and capital resources
08	Governance, organisational and social issues

Table 6.5. Eight digital twin bottlenecks. The bottlenecks show why a digital twin is an institutional programme, not a software purchase. Source: author table based on Weil et al. (2023), note 6.

Planning organisations spend considerable time on the technical architecture of a digital twin. The governance architecture gets far less attention. That imbalance is the wrong way round, because a twin that becomes influential before its rules are settled will be used in ways no one formally decided. At minimum, those rules have to answer these questions:

1. What decision is the twin meant to support?
2. What geography, time horizon and policy questions are inside the model, and which are deliberately outside?
3. Where does each dataset come from, when was it updated and who owns it?
4. Which assumptions drive the results, and which of those are evidence, which are expert estimates and which are political choices?
5. Who may view, query, edit, export or challenge the model?
6. Which scenarios were tested, by whom, and how did they influence a decision?
7. How can a resident, partner, councillor or professional challenge the data, interpretation or use of the model?

These controls form the safety system for the twin.

The list (see Table 6.5) should warn every planning organisation against treating the digital twin as a software purchase.⁶

6. Weil, 2023

The hard problem is institutional fit.

A city can buy a platform. It cannot buy trust, data discipline, shared standards, interdepartmental coordination or political clarity. Those have to be built.

It also cannot buy the capacity to keep the data current.

A digital twin is only as reliable as its least-recently updated layer. An energy constraint from 2022 is not a minor detail. A mobility assumption that predates a new development decision is not a footnote. A climate scenario that the water authority has since revised is not background noise. Each of these turns the twin from a decision support tool into a source of false confidence. An outdated layer turns neglect into a political decision.

This is the data maintenance problem that most twin initiatives underestimate. The recurring cost lies in keeping the twin current and reliable after the initial build. The investment in building the model is visible. The investment in keeping it honest is structural, recurring and unglamorous. Planning organisations that treat the twin as a one-off project will discover within two years that it must operate as a continuing service. The wider digital twin field repeatedly runs into what practitioners call “death by pilot scheme”: successful demonstrations that fail to become operational services because governance, data ownership, financing and maintenance were never settled before the demonstration ended.

UN-Habitat’s people-centred smart cities work makes a similar governance point from the rights and inclusion side.⁷ The questions are concrete. Can older residents use the interface? Can people without technical skills understand what the model is doing? Can communities see how their data is used? Can the city prevent vendor lock-in? Can the organisation explain why one scenario was selected over another? Can the council still make a political decision without pretending the model has decided for it?

7. UN-Habitat,
2022

A democratic twin gives councillors better questions, not better images. Which assumptions changed? Which options were not tested? Which public commitment was traded away? Which layer is outdated?

These are not side questions. They define whether the digital twin strengthens public planning or weakens it.

The real test of a digital twin is whether the organisation is willing to own what its layers reveal. A model that exposes an unfunded mobility promise, an outdated energy assumption or a broken participation commitment does not merely improve analysis. It creates responsibility. The hardest question concerns accountability: who is prepared to own what the twin makes visible?

The twin may meet resistance even when it is right. A correct model can expose that a coalition was held together by incompatible assumptions. Resistance can be evidence that the twin has found the real conflict.

The same question returns after the decision. Many planning twins are built for the moment before approval: to compare options and support a choice. But urban places live much longer than the decision that created them. If the twin dies after the zoning decision, it was only a tool. If it continues into construction, maintenance, energy operation, mobility management and public space stewardship, it becomes part of the operating memory of the place. That asks for different rules. Who may use the model before the decision? Who maintains it after? Who updates it when reality changes? Who pays for that work over time?

Where to start

A planning organisation does not need a complete urban digital twin to begin.

A planning organisation needs one live decision environment for one consequential project. The logic is not to build the whole city first and then use it. Start with a decision that is complex enough to expose interdependencies, and build the minimum twin that makes those interdependencies visible.

A station area is often the right starting point. The layers already collide there in ways that are politically real: housing, mobility, climate, energy, finance and public commitments. The conflicts are visible. The stakeholders are identifiable. The decisions have consequences that stretch across disciplines and across time.

The Dutch experience with local digital twins suggests a similar logic. The municipality of Amersfoort, together with RIVM and Province Utrecht, has been integrating health data into spatial planning through a dedicated twin, making the health consequences of design choices visible before decisions are locked.⁸ The work began with one set of questions in one area, together with the partners who carry the relevant data. That is the right scale to begin.

8. Future City
Foundation, 2021

The first question is not “what can we build?” It is “which decision is actually in front of us, and which layers does it touch?” That question defines the scope of the twin. It also defines the assumption register that needs to exist before any scenario is run. Every number that matters should have a source, an owner, a date and a trigger for review. Without that register, the twin is only as reliable as its least-examined input.

From there, the twin can be opened to partners with controlled access. The water authority, the housing association, the mobility team and the energy team each carry assumptions that affect the others. Letting each test the assumptions

that fall within their responsibility is governance and makes the model more trustworthy. It is how the model becomes trustworthy.

Public use comes last and should begin with a clear explanation before any polished presentation of what the model can show, what it cannot show, which choices are still open and how input can alter the decision.

The first goal is a better decision room.

What this changes

A digital twin is a partial and governed representation of the city. It supports political judgement and participation while substituting for neither.

It is a shared environment in which options can be tested against spatial, physical, financial, social and institutional consequences.

That is powerful. It is also dangerous if the model becomes more trusted than the process around it.

The hardest constraints remain: land ownership, finance gaps, energy capacity, legal procedure, political courage, public trust, implementation capacity. A digital twin does not dissolve them. It makes their interaction harder to ignore.

Consider what happens when the twin reveals that two political promises cannot both be kept. The public space commitment requires one footprint. The housing target requires another. Both were made. Both are in the project memory. The twin exposes the conflict without resolving it. Something must give.

That is the moment the twin earns its place.

The twin earns its place by removing the final technical argument for postponing the political choice.

Back to the planning team on the morning of the steering group.

The model is impressive. But that is no longer the point.

The point is that the team can now show why option C only works if the mobility hub is delivered early, why option A creates an energy exposure that cannot be solved inside the project alone, and why option B weakens a public space commitment that residents will recognise immediately.

The twin has not decided.

It has changed the room in which the decision is made.

It has changed the room in which the decision is made.

Questions for Practice

- 01** Which concrete decision should the twin help people examine?
- 02** Which consequences, uncertainties and omissions must remain visible?
- 03** Who can inspect, challenge and update its assumptions?
- 04** How will the team distinguish what is impossible, unchosen and unfunded?



PART IV

Participation and Delegation

More access, faster synthesis and higher delegation do not automatically produce legitimacy. They move the burden of judgement, verification and influence.

ARGUMENT MAP

This part moves from access to influence, and from delegation to accountability.
Read for participation, verification and burden shifting.
Key tension: more interaction without more power.

CHAPTER 7

Participation

AI can make participation broader and more legible. Organisations still decide whether it becomes influential.

IN THIS CHAPTER

Follow one contribution from its first expression to the point where the project records what changed, what stayed fixed and why.

USE IT FOR

Designing participation that can still alter programme, phasing or design.

WATCH FOR

More interaction without more public influence.

A participation lead opens the digital twin of a station area in a community centre on a wet Tuesday evening.

Forty residents sit in the room. Another thirty have joined online. The project team has brought three development options, all revised after the last steering group. The model shows the station square, the housing blocks, the schoolyard, the mobility hub, the public courtyard that has become contentious, and the phasing dependency that makes everyone in the project team slightly uneasy.

A resident asks what happens if the parking ratio is lowered further.

The mobility layer updates. The financial layer shifts. The model shows less underground parking, more room for trees, a small improvement in affordability, and a new dependency on the mobility hub opening earlier than currently funded.

Another resident asks whether the courtyard could be enlarged without losing affordable homes.

The system runs through the available options. It shows that this is possible in one variant, but only by moving volume toward the schoolyard and accepting more winter shadow. The AI assistant turns the output into plain language, with the uncertainty clearly marked. A planner adds that one of the assumptions is contested and will need separate verification.

For a moment, the meeting feels different.

Agreement remains absent and the model leaves the trade-off unresolved. What has changed is the timing: the questions now reach the decision space while it remains open. Residents now test the structure of the choice instead of merely reacting to a finished proposal.

Then someone at the back raises a hand.

“Which of these things can still actually change?”

The room goes quiet.

This scene describes what AI enabled participation can look like when the process is well designed, the project team is honest, and the option space is still genuinely open. Most processes are not like that. The tools are real; the conditions are not guaranteed. The decisive issue is whether the organisation is willing to let participation alter the work.

*The room goes
quiet.*

AI cannot create that willingness. But where the willingness exists, it can change the old economics of participation: more explanation, more follow-up, more iterative reasoning, at a depth that was previously too labour-intensive to offer at scale.

AI can make participation broader, faster and more legible.

It can also make participation more convincing while leaving power exactly where it was.

Participation is where the whole argument is tested

Participation is meaningful when public input can change the decision before the decision is made, and when that change is traceable afterward.

Figure 7.1. Working definition of meaningful participation. Source: author working definition based on Arnstein (1969) and later participation scholarship, notes 1 and 2.

Participation data became processable in the machine room, while project memory made commitments traceable. Options and numbers widened the decision space, and the digital twin made that space shared and spatial. Influence still depends on a political design choice.

If AI improves internal project work but leaves participation as a late stage reaction exercise, the operating system becomes more capable without becoming more legitimate. The organisation will see more, calculate more, remember more and simulate more, while citizens still enter after the meaningful choices have narrowed.

Leaving participation at that late stage would reproduce an old planning failure with better tools.

Sherry Arnstein's critique from 1969 still matters because it names the issue directly. Power defines participation: who can shape the decision, who can only respond, and who is present mainly to validate a process already under way. Who can shape the decision, who can only respond to it, and who is present mainly to validate a process that has already moved on.¹ A working definition follows from that, though it is a normative one rather than a logical necessity: participation is meaningful when public input can change the decision before the decision is made, and when that change is traceable afterward. Arnstein's framework addresses redistribution; recent scholarship extends it by adding recognition and representation as equally necessary conditions: whose knowledge is treated as valid, and who bears the consequences of choices made.²

1. Arnstein, 1969

2. Blue et al., 2019

AI does not make that question obsolete.

It makes it harder to avoid.

AI cannot create willingness

Most participation does not fail because the tools are bad. It fails because serious influence is organisationally inconvenient and politically risky. An open option space slows projects down. Real influence can threaten financial feasibility. A developer wants to reduce risk, not distribute it. An elected official does not want to revisit a decision that was politically difficult to reach. An administrative organisation does not want a participation memory that later reveals which warnings were ignored.

Many organisations say they want better participation. What they often want is better acceptance. AI will not close that gap. It may make it harder to hide.

The question is whether organisations choose to go further.

The old participation problem included more than exclusion

Urban professionals often describe the weakness of participation as a problem of reach.

The same people come to evening meetings. Written consultations attract the highly motivated. Digital platforms bring in more responses but not necessarily a more representative public. People with less time, less confidence, lower literacy, language barriers or distrust of institutions remain underrepresented. Those are real problems.

But there is a second problem, and it receives less attention.

Planning organisations often do not know what to do with participation once they have it.

Hundreds of written responses arrive. Meeting notes pile up. Workshop post-its are photographed and stored. Online comments are exported into spreadsheets. The formal report appears weeks later, often organised into themes that are accurate enough and politically unusable. “Traffic”, “green space”, “safety”, “housing mix”, “identity of the area”. Everything has been heard. Little has been made sharper.

The deeper tension disappears in the act of summary.

Was the traffic concern about congestion, child safety, construction nuisance or distrust that mobility promises will be dropped later? Was the concern about green space mainly ecological, recreational, aesthetic, or a proxy for fear of losing a familiar place? Was the objection to density an objection to height, to social change, to pressure on services, or to the sense that the conclusion had been reached before the meeting began?

Participation data is not self-interpreting. It carries conflict, history, tone and strategic behaviour. It also carries weak signals that matter precisely because they are not the dominant theme.

AI can help immediately with this processing burden. It can process large amounts of input, cluster recurring issues, compare consultation rounds, surface contradictions, map concerns geographically, produce first summaries in plain language and support translation across languages. OECD's 2025 work on AI in civic participation identifies precisely these functions: processing citizen input, reducing implementation costs, improving access to information and lowering language barriers. A 2025 Delphi study among Finnish planners found especially strong agreement around AI's usefulness for analysing participatory data.^{3,4}

3. OECD, 2025

The saved time has a professional use. A participation advisor who spends less time coding responses manually has more time to interpret what those responses mean, check whether minority positions are being flattened and ask what the project should do with the signal.

4. Raymond et al., 2025

The first gain is more professional attention for the part of participation that cannot be automated.

AI can widen the entrance, and lower the cost of going deeper

Serious participation is expensive. The cost sits in everything the listening commits you to: explanation, follow-up and interpretation all take time. A participation advisor can facilitate one good conversation, but cannot have five hundred of them. Informational evenings are time-bounded. Project websites are formally public and practically inaccessible. The result is that most participation formats are designed around what is affordable, not around what actually produces the best public knowledge.

A conversational AI layer changes what a resident can do with a plan before a meeting even starts.

It can offer layered explanations: a short version for orientation, a deeper version for those who want detail, a source-linked version for those who want to verify. It can turn technical plan language into questions residents can actually judge. It can help people understand the consequences of their own preferences before they commit to them.

At its best, this does more than make government documents readable. It gives residents a patient first interlocutor. One that can explain the plan, answer follow-up questions, ask what matters underneath a first objection, and help turn an inchoate concern into a contribution that the planning process can actually work with. That kind of exchange was previously only possible at scale

if you had a very large facilitation team. AI changes that constraint without replacing what the facilitator does in the room.

Consider a resident whose evening shifts keep her from the Tuesday meeting. At 23:40, on her phone, she can ask what the plan means for the playground behind her block, why the parking street is being moved, and whether the green route shown in the visualisation is funded or merely indicative. She can be asked what matters most in her objection, and receive a response that helps her turn a concern into a usable contribution. That is a different threshold of access than a PDF and an evening in a community centre.

None of that replaces human facilitation. It extends the reach of good facilitation to more people, more questions, more moments in the project timeline.

A 2025 systematic review of citizen-centred digital twins in land use agenda setting found that such tools can help residents frame problems, convince others and propose alternatives, provided the systems are understandable, interactive, inclusive and privacy oriented.⁵

5. Adade & de Vries,
2025

Rotating a model on a screen is not the point. The value is that residents can get inside the logic of the decision early enough to contest it. That is a different kind of access than a PDF and a Tuesday evening meeting, and it forces the next question: was the option space still open when they arrived?

The option space must be open at the moment of participation

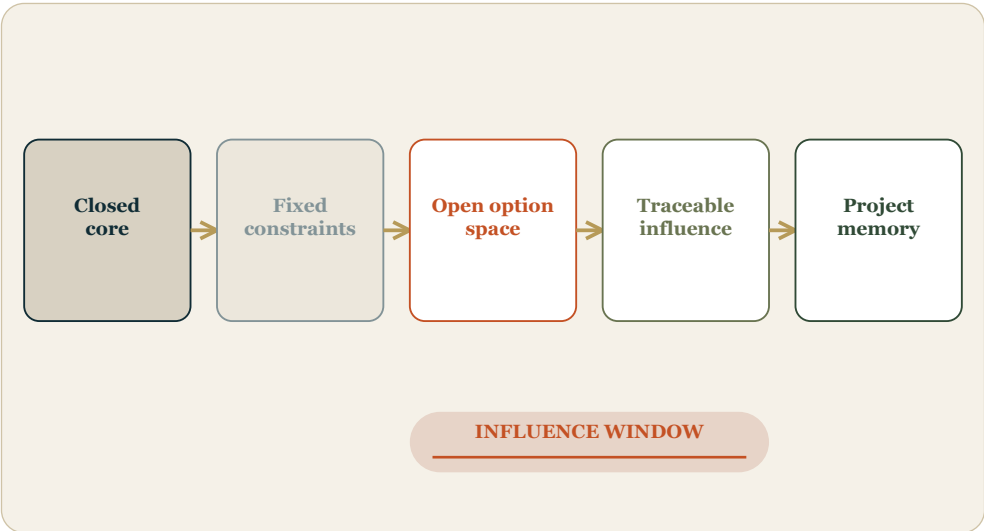


Figure 7.2. The influence window. Participation can influence only choices that remain genuinely open. Source: author synthesis.

The polished version of AI enabled participation looks compelling.

Residents test variants on a model. The process feels active. The municipality can document meaningful public access. Everyone leaves with the impression that planning has become more responsive.

But there is one prior question.

Was the option space still open?

If the financially and politically viable range had already been fixed before the session, the model may allow interaction without allowing influence. People can move trees, widen a pavement and discuss roof gardens while the housing mix, phasing logic and public space footprint are no longer genuinely in play. They are participating in surface variation around a closed core.

That is a process design choice, not a technology failure.

The distinction between not possible, not chosen and not funded matters here. Participation becomes poisonous when those three are blurred. Residents ask why something cannot happen. The organisation answers with the authority of necessity, even where the real answer is that it was not chosen, or not chosen under current budget assumptions, or not chosen because another public goal was prioritised.

AI can help separate those categories. It can show which constraints are legal, which are physical, which are financial, and which are project-specific assumptions that could be revisited. Doing so requires care. AI can identify distinctions in text. Institutional knowledge must decide whether a constraint is political or genuinely impossible. The organisation must do that analytical work first. AI can then help make it visible.

But only if the organisation wants that visibility.

A digital twin without an open option space is a theatre set. AI makes the lighting better.

Synthesis is where power hides

AI is often presented as a neutral summariser of public input.

Every synthesis contains analytical choices.

Any summary makes choices. It decides what counts as a theme, what gets grouped together, what is treated as repetition, what is labelled as an outlier and what disappears because it does not fit the dominant categories. Human analysts do this too. The difference is that AI can do it at a scale and speed that makes the politics of compression less visible.

A 2025 study on AI in participatory urban planning warns that analytical systems can make parts of the population invisible and that generative systems

raise deeper questions about which public is being represented in a process.⁴ That concern should be read as a design requirement, not as a reason to avoid the tools altogether.

4. Raymond et al.,
2025

Participation synthesis in the AI era should therefore be inspectable. AI does not provide that inspectability automatically. Most large language models are less transparent about their clustering logic than a careful human analyst with a visible coding scheme. Inspectability is a design requirement: the raw input must be stored and accessible; the theme labels must be documented with enough reasoning to be challenged; minority positions must be surfaced through an explicit protocol before the synthesis is accepted as the official account; and a human must sign off on the result before it leaves the organisation. None of that is technically difficult. All of it requires deliberate process design.

This is also where quantitative instincts can mislead. Frequency matters. It does not settle significance. A concern raised by three people may be marginal. It may also identify a future failure that the majority cannot yet see. A single disability access concern can reveal a design blind spot. A small group's distrust can indicate a long project history that newer staff do not know. A recurring irritation expressed in restrained language can matter more than a louder objection that has little bearing on the final decision.

Good participation analysis does not confuse prevalence with importance.

AI can help reveal both. It should not be allowed to collapse one into the other.

From consultation to iterative reasoning

The deepest change AI makes possible in participation is that it can make it iterative.

Conventional participation is mostly one-directional. A resident states a preference. The project team notes it. The formal report places it in a category. The process moves on. The resident rarely discovers what happened to what they said, or what would change if they had said something different.

AI, linked to a shared scenario environment, can work differently. A resident who says “the courtyard should be larger” can be asked: why does that matter to you? What are you worried would happen without it? If making it larger means accepting a different building height elsewhere, is that a trade off you could accept? What if the enlarged courtyard comes with a reduced green buffer on the other side?

Those questions do not manipulate. They deepen. They help people move from a position to a reasoning, from a preference to an interest, from a reaction to a considered judgement. It draws on a Delphi-like logic without claiming to be a formal Delphi method: iterative clarification, structured reflection, and the

possibility of revising a judgement in light of consequences the participant had not yet seen.

This is not a description of AI as autonomous facilitator. The section on the AI penalty documents a real finding: people are less willing to join deliberative processes when they know AI is running them. The iterative model only works if AI operates as an instrument under human direction, not as the face of the process. A human facilitator leads the session. The AI deepens the background work: it runs the follow-up, tracks what has been said, shows the trade off consequences, feeds the summary back for human review. The resident experiences a richer conversation. They do not necessarily experience a machine.

But the shadow of this possibility must be named immediately. More depth in the conversation is not automatically more influence in the decision. An AI that asks good follow-up questions, listens carefully and produces a sophisticated synthesis can still be embedded in a process where none of that input reaches the decision. The iterative layer can become the most sophisticated form of participation theatre yet invented: residents feel genuinely heard, reflect carefully, receive nuanced responses, and change nothing.

The test is whether the decision changed.

That check must be external to the AI system. It must be institutional.

Participation must enter the memory

A participation process usually leaves behind two records.

The first is the formal report. It says how many people attended, how many comments were received and which themes emerged.

The second is the actual memory of the process. That record is messier and more valuable. It includes which residents felt that a promise was made. Which subject kept returning despite never ranking highly by volume. Which objection was technically answered but relationally unresolved. Which group had a different reading of the place than the official project team. Which answer created relief and which answer deepened distrust.

That does not mean storing every raw emotional response forever. It means preserving the parts of participation that remain relevant for future decisions: commitments, unresolved tensions, minority positions, contested assumptions, accepted trade offs and rejected alternatives. If the project later changes direction, the organisation should be able to see what participation had already made visible and whether that knowledge was honoured, revised or quietly dropped.

AI can help build that record. It can compare participation rounds, trace when a concern first appeared, detect whether a promise made in one phase has disappeared from later documents, and draft a transparent change log: this concern was raised, this assumption was tested, this element changed, this element did not change, and this is why.

Many organisations use that phrase, or some version of it, with good intentions. But urban development is rarely so direct. A resident says one thing, several interests interact, a decision shifts partly, a later constraint reopens the question. A better format would be:

The familiar formula "you said, we did" is too simple for urban development. A stronger record would show what was heard, what changed, what did not change, why, and what remains open.

That last item matters most. A process that has no unresolved questions has not been honest. It has been managed.

Visualisation helps people speak. It can also over-persuade.

Generative images are beginning to appear in participatory planning. Their appeal is obvious. Many residents struggle to respond to abstract plan language, but they can react immediately to an image of a greener street, a widened pavement or a redeveloped square. Images can unlock conversations that technical drawings do not.

A 2025 study in *Landscape and Urban Planning* explored generative AI images in community participation and proposed a useful middle ground: “controlled imperfect” images. The useful range lies between polished pseudo-final renderings that imply false certainty and rough outputs too vague to support discussion. Images should be concrete enough to provoke response, but open enough to invite correction.⁶

| 6. Awashra, 2025

Participation images should not make a future look decided before it is. A visually compelling AI-generated streetscape can quietly privilege one option, make trade offs disappear and pull discussion toward aesthetic approval, pushing substantive choice aside. The more beautiful the image, the more disciplined the process around it needs to be.

What assumption does this image contain? What is missing from it? What would change if the budget, density or management model changed? Is this showing an option, or selling one?

Visual participation becomes valuable when images are treated as thinking instruments.

It becomes dangerous when they become soft persuasion.

The AI penalty is real

People may not want AI in the room.

A 2025 preregistered survey experiment with a representative German sample found that people were less willing to join deliberative processes when they were told AI would facilitate them, and expected those processes to be of lower quality than human-led formats. The authors call this an “AI penalty” and warn that it could create a new deliberative divide, based on attitudes toward AI rather than education or income.⁷

7. Jungherr, 2025

The penalty compounds when institutional trust is already low. Research on trust in AI-assisted government consistently finds that confidence in the technology is not independent of confidence in the institution using it: trust in high-risk AI systems depends on the technology's features and on trust in the institutions under which they operate.⁸ In neighbourhoods where the municipality has a history of broken commitments or top-down decisions, AI in the participation room may read as another layer of managed distance. Residents may experience it as a tool for managing them more efficiently.

8. Pande et al., 2025

That finding should not be overgeneralised. A virtual focus group study published the same year found that disclosed AI summaries were positively received by participants, especially when transparency was clear. Participants also expressed concern that AI might miss emotion and nuance in sensitive discussions.⁹

9. Wang, 2025

AI may be widely accepted as a support tool behind the process: translation, first summaries, source retrieval, question handling, comparison of consultation rounds. It may be more contested when it acts as moderator, facilitator or simulated conversational partner. A technical housing workshop is not the same as a painful redevelopment process in a neighbourhood with a history of distrust.

Participation depends on legitimacy as much as functionality.

Simulated publics are useful only when they remain simulated

Before a meeting, a project team could ask a model to test how different residents, entrepreneurs, housing advocates and accessibility groups might respond to a proposed scenario. It could surface weak arguments, missing questions and possible conflict points. A 2026 article in the *Journal of Deliberative Democracy* argues that generative AI simulations may complement human deliberation by helping train facilitators and providing rapid exploratory consultation under time pressure.¹⁰

10. Rountree, 2026

Used that way, simulation can sharpen preparation.

But it cannot substitute for public participation.

A simulated resident can predict concerns and help rehearse political questions, but it has no standing, lived stake or exposure to the consequences of the decision.

Simulated participation can help prepare for real participation. Where it replaces real participation, the organisation has not saved time. It has produced a decision without a public, and called it consulted.

What meaningful AI supported participation would look like

01	Participation should begin before the core option space has closed.	02	The process should distinguish fixed constraints from open choices.
03	AI assisted synthesis should be reviewable.	04	Participation should feed the project memory.
05	Visual and conversational AI should be treated as instruments for exploration, not tools of persuasion.	06	The role of AI should be disclosed in language people understand.

Table 7.1. Six disciplines for participation. Source: author synthesis.

By that point, six disciplines define the standard. Technology remains secondary.

First, participation should begin before the core option space has closed. If a project team can only explain, not change, that should be said plainly.

Second, the process should distinguish fixed constraints from open choices. Legal impossibility, financial pressure, political preference and project habit are not the same category.

Third, AI assisted synthesis should be reviewable. Raw input, coding choices, summary logic and minority positions should not vanish behind a clean report.

Fourth, participation should feed the project memory. Commitments, contested assumptions, unresolved tensions and reasons for rejecting proposals should remain queryable later.

Fifth, visual and conversational AI should be treated as instruments for exploration, not tools of persuasion. Their outputs should show uncertainty and invite correction.

Sixth, the role of AI should be disclosed in language people understand. Disclosure should give a clear account of what the system does, what remains outside its role and where human judgement carries responsibility.

Where to start

Most planning organisations should not begin with an ambitious civic digital twin and an AI facilitator in every workshop. That is too large a leap and, in many cases, a solution in search of a process.

Start smaller.

Choose one live project with a real but manageable participation stream. Use AI to structure written responses, but keep the analysis reviewable: store the raw input, document the clustering choices, flag the minority positions explicitly before finalising the summary. That last step takes ten minutes. It is the difference between a synthesis and an official account. Produce a participation log that separates themes, minority positions, commitments, open questions and decisions. Compare the result with the traditional summary and ask which one better supports an honest project conversation.

| *Start smaller.*

In a second step, connect that participation log to the project memory. When a new option is developed, make visible which earlier concerns or commitments it touches. When an option is rejected, record why.

Only after that does it make sense to bring participation into a shared scenario environment. By then the organisation has learned the harder discipline of making public input alter the work in a traceable way.

What this changes

The resident at the back has asked the essential question: “Which of these things can still actually change?”

A weak process gives a reassuring answer. “We are here to listen.” That may be true, but it avoids the point.

A stronger process answers concretely.

The housing target is fixed by a council decision. The exact block distribution is not. The affordability commitment remains in place, but the risk allocation is under review. The mobility hub remains unfunded, so any option that depends on it carries an unresolved condition. Enlarging the courtyard creates

consequences elsewhere. The public space quality standard is open to refinement. The phasing sequence is still under discussion. These three issues raised tonight will be tested against the model and added to the project memory. The team will show what changed, what did not and why.

That answer does not manufacture agreement.

It creates a more honest conflict.

The promise of participation in the AI era is a more consequential and traceable conflict. Harmony or instant consensus would be a false measure.

A better decision room is one in which public knowledge enters earlier, is developed more carefully, stays visible longer and changes the options before the options become fate.

Questions for Practice

- 01** Which choices remain genuinely open when participation begins?
- 02** Can a resident trace what happened to a contribution after the meeting ends?
- 03** How are minority positions preserved when AI groups or summarises responses?
- 04** Who records whether public input changed an option, a condition or the final decision?

CHAPTER 8

The Burden Moves Sideways

Five levels of AI delegation in planning practice.

Higher delegation relocates verification, judgement and accountability.

IN THIS CHAPTER

This chapter maps five delegation modes and identifies the new burden created at each level.

USE IT FOR

Choosing delegation levels and locating the burden they create.

WATCH FOR

Delegation presented as progress while burden relocation goes unseen.

She almost misses it.

The revised programme matrix for the station area looks harmless on her screen: density targets met, financial model balanced, project risks marked amber rather than red. One relationship sits slightly out of place. The subterranean parking ratio has moved; the land value assumption has not. To protect the developer's residual margin, the underlying optimisation model has compressed the long-term infrastructure buffer and reduced the spatial footprint of the public courtyard promised to the neighbourhood group in 2024.

She didn't catch this through manual spreadsheet tracking. Her AI assistant flagged the deviation by cross-referencing the draft agreement against the project's recorded commitment register: the project memory layer where council conditions, partner agreements, and participation promises were kept.

She knows two things at once. Large parts of the bureaucratic pipeline can now be accelerated. Making that pipeline fast is not the same as making it governable. The tool has exposed a compromise hidden inside an automated calculation. It hasn't solved the conflict; it's changed the room in which the tension must be faced.

The illusion of the ladder

Planning organisations use different delegation modes across fragmented tasks; no single ladder describes the movement.

In public planning, the levels are delegation modes inside specific, fragmented tasks. A planning department is not a single organism climbing a ladder. A single municipality may run advanced agentic workflows for participation synthesis while handling land acquisition, site negotiations, and field observations almost entirely by hand.

The levels matter for a particular reason. Each one moves the burden somewhere else. The levels do not mark progress up a ladder. Understanding that sideways movement is the difference between governing a transition and losing contact with the material of the craft.

Put differently: the levels don't tell you how capable a planning organisation is. They tell you where attention has been relocated, and where it may disappear if the organisation does not design for it.

Adapting an early automation taxonomy used by the National Highway Traffic Safety Administration (NHTSA) in 2013 for autonomous vehicles gives a useful scaffold for thinking about delegation.¹ But we should resist treating it as an institutional roadmap.

1. NHTSA

Where the burden moves

Level	Role of AI	Role of professional	Core risk	Where the burden moves
Lo Manual	None	Produce and document	Knowledge loss	Production
L1 Assistive	Isolated tasks	Review output	Fluent but unreliable output	Output checking
L2 Co-creative	Iterative co-creation	Choose and calibrate	Prompt dependence	Choice and calibration
L3 Oversight	Agents and monitoring	Verify and remain accountable	Cognitive surrender	Review and accountability
L4 Governed optimisation	Governed scenarios	Set institutional rules	Hidden political choices	Institutional design and democratic accountability
Higher delegation shifts attention toward verification, institutional design and democratic accountability.				

Figure 8.1. Five levels of AI delegation and where the burden moves. Source: author adaptation of an early NHTSA automation taxonomy (2013), note 1.

Higher delegation changes the work without guaranteeing professional advancement; lower delegation can also waste scarce professional attention. A municipality that keeps all participation synthesis manual while its neighbours use well-governed agentic workflows may not be cautious. It may simply be burning coordinator time on first-pass clustering and comparison, leaving less capacity for the judgement that actually matters.

Urban planning is structurally different from legal discovery or software engineering in one specific way: it's bound to democratic procedures and political cycles in which the ground rules themselves are contested. Planning goals are frequently ambiguous. Public values shift across time, place, and political mandate. A planning decision also leaves a physical and legal trace: land is allocated, rights are fixed, public space is shaped, and losses are distributed unevenly. Because of that, maintaining a lower level of delegation is

often a deliberate protective choice: keeping the professional in direct contact with local friction that no parameter can fully encode.

The Coasean shift: moving the friction

Building on the Coasean reading introduced in Chapter 2, the transition between these modes can be understood through the lens of transaction costs. Ronald Coase observed in 1937 that organisations arrange their internal structures to manage the costs of negotiation, direction, and information processing.² Recent economic analysis has extended that insight directly to AI agents, arguing that as AI reduces the cost of coordination inside firms, it shifts which activities get organised internally and which don't.³ AI relocates these costs instead of abolishing them.

2. Coase, 1937

3. Coase, 1937

For planning organisations, this points to three relocations.

Search costs collapse. At Levels 1 and 2, the assistive and co-creative modes in Figure 8.1, locating policy precedents, comparing municipal zoning updates, or scanning regional housing deals moves from a sequential paper chase to near-instant retrieval.

Coordination costs drop partially. Disciplines that historically operated in silos, urban design, traffic engineering, land economics, can increasingly be brought into the same modelling or scenario workspace. Running iterative variants costs a fraction of what it once did.

Verification and accountability costs rise sharply. This is the hidden price of efficiency. When production costs fall, more outputs get generated, and more choices need to be audited, tracked, and defended before a council committee. The professional's core task shifts from finding and producing information to verifying, legitimising, and governing it.

From flow to fatigue

Level 2, the co-creative phase in Figure 8.1, is the seductive one. An urban designer adjusting parameters in a parametric model, a policy lead co-writing a brief with an LLM: the human remains the director of the loop. It feels highly productive, and it often is.

At Level 3, the oversight phase in Figure 8.1, systems run continuously behind the scenes to process inputs or update models, and the professional enters a fundamentally different role: the reviewer.

Evidence from general knowledge work identifies a mechanism that planning organisations should take seriously, while leaving the planning-specific outcome unproven. A 2026 BCG study published in Harvard Business Review surveyed 1,488 full-time US workers and found that those overseeing multiple

AI agents, rather than using AI to replace discrete tasks, were most at risk for acute cognitive fatigue.⁴ The researchers call this “AI brain fry”: mental exhaustion from holding multiple streams of half-trusted output in sustained attention. Workers described mental fog, slower decision-making, and difficulty concentrating. Brain fry was specifically linked to oversight and multitasking across AI tools, not to AI use as such.

4. BCG and HBR, 2026

That 14% figure deserves attention. The study covered general knowledge workers. It identifies a plausible mechanism for planning without proving a planning-specific outcome. Oversight becomes mentally taxing because it demands continuous vigilance against a system that is fluent, confident, and usually correct. When an automated analysis looks rigorous most of the time, the cognitive cost of line-by-line verification starts to exceed the time the organisation has scheduled for it.

Whether this fatigue is transitional or structural remains genuinely open. Professionals who grow up with AI as a standard tool may calibrate differently. What is less likely to disappear is the underlying vigilance asymmetry: a system that is correct most of the time is harder to check than one that fails frequently. Even as AI improves, the oversight burden persists and concentrates around quiet failures.

Separate experimental research from Wharton found that participants often deferred to an AI assistant even when it was deliberately wrong. The researchers call this cognitive surrender, the pattern described in Chapter 1, where people accept AI output without engaging deliberative judgement.⁵ This is a preprint, not yet peer-reviewed, and the experimental conditions involved structured reasoning tasks rather than complex planning decisions. The mechanism is still worth noting. In public administration, cognitive surrender transfers liability while creating the appearance of efficiency. An error in an environmental contour or a land allocation covenant doesn’t disappear when it passes unchecked through an automated summary. The accountability stays with the professional and the institution they represent.

5. Wharton, 2026, preprint

The limits of the specification

Level 4, governed optimisation in Figure 8.1, is often framed as the professional destination: the planner as specification writer who defines the rules, weights, and constraints, then lets the machine generate compliant pathways. That model works well when constraints are explicit and static. Checking a site against formal floor area ratios or building-height limits is genuinely a Level 4 task. It should be automated.

The problem starts when that same logic extends to choices that aren’t formally bounded.

Consider a standard Level 4 constraint: “Prioritise long-term infrastructure debt minimisation at 60% and delivery speed at 40%.” Looks precise. But it conceals a series of questions no technical system can answer. Who legitimised that specific weighting? The municipal council? The housing association? The existing neighbourhood residents? An alderman will routinely choose a development variant that is politically explainable and locally tolerable over one that a model declares mathematically optimal.

A parameter compresses a political settlement into a technical instruction. It shapes which option is preferred and which options become visible.

The professional at Level 4 can’t just be a technical writer of prompts. They need to be a designer of the rules that govern what counts as a valid constraint: which weights can be freely varied, which public values must remain explicit in the register, and which choices must stay visible to the council regardless of what optimisation the model prefers. Call it institutional design: deciding what the machine is and isn’t allowed to optimise away.

This is where specification becomes politics. Setting the weight between affordability and land revenue is not a modelling choice; it is a redistributive decision. Choosing whether a scenario can trade public space against density is not a parameter; it is a spatial value judgement that belongs in a council chamber, not a prompt. At Level 4, the professional who writes the constraint set is exercising a form of power that is largely invisible precisely because it is technical in appearance. The open parameter register exists to make that power visible.

But the register only works if planners understand what they are actually doing when they write the constraint: encoding politics through software configuration.

The loss of guided friction

Higher delegation has a slow cost that the efficiency numbers don’t show: it disrupts how junior professionals learn to plan.

Labour market data from Massenkoff and McCrory (2026) shows no detectable unemployment increase among highly AI-exposed workers since late 2022.⁶ Broad displacement hasn’t happened. The absence of broad displacement does not mean the training structure is intact. A different question is opening up, one the aggregate data isn’t designed to catch: what happens when the routine tasks through which judgement is formed are automated before institutions have built a replacement training path?

6. Massenkoff, 2026

Historically, professional judgement in urban planning is formed through tedious, repetitive work: manually coding hundreds of participation responses, drafting cross-checks of policy variants, building spreadsheets cell by cell.

Machine room names this operational layer. Project memory shows how the same work generates the institutional memory a project carries across years and personnel changes. This operational work forms part of the training ground in which professional judgement develops. It's where guided friction occurs. It's how a junior planner learns pattern recognition, builds familiarity with the legal edges of agreements, and develops an intuitive sense for what data actually represents on the ground. Many of the tasks now being automated were not incidental; they formed part of the substrate on which judgement was built. Evidence from wider labour markets does not prove this planning outcome. It identifies a formation mechanism planning organisations should address before automating the training ground away.

The harder institutional question follows immediately: who has an incentive to rebuild that training path? A planning bureau that invests in junior development while its competitors don't is absorbing a cost the market doesn't reward. A municipality under budget pressure has no short-term reason to maintain the pipeline either. This is a collective action problem, not a training design problem. No individual organisation can solve it fully on its own.

A useful comparison: GIS, CAD, and spreadsheets all automated tasks that were previously done by hand. Planners who grew up with those tools still learned to read the terrain behind the interface. What's different now is the scope and the speed. Large parts of the operational layer are becoming automatable at the same time, faster than many institutions can build new training paths to replace them.

Microsoft's 2025 Work Trend Index points in the same direction: as organisations move toward agent-based work, professionals need domain expertise and the ability to work with and manage AI-enabled workflows.⁷ That distinction matters. Experienced professionals can calibrate AI outputs against prior knowledge of how the work is normally produced. A junior planner reviewing outputs they've never learned to produce doesn't have that reference point. They lose contact with the material of the craft.

7. Microsoft, 2025

Rules for an accountable workspace

To scale delegation without succumbing to system blindness, planning organisations need explicit institutional design principles. Three can anchor the practice.

The pedagogical mirror. Organisations should reject passive review workflows for junior staff. When evaluating complex project materials, a junior professional should conduct a targeted manual analysis of one component before accessing the AI-generated summary. Judgement is formed in the conscious comparison between human interpretation and machine output, not in approving the output alone.

The democratic boundary. Some choices cannot be converted into optimisation tasks. Direct democratic conflict, the reconciliation of minority positions in participation, the judgement of trade-offs between public values: these require human mediation because legitimacy is not a system output. The system can log the arguments. It can't grant legitimacy to the settlement.

The open parameter register. All weights, thresholds, and objective functions used to drive scenario variations must be pulled out of the software's prompt history and recorded in an explicit, accessible register. If a model assumes a specific inflation rate, ranks delivery speed over public space quality, or treats a promised courtyard as adjustable project space, that choice must be auditable by council members, partners, and residents. A register entry needs more than "public space quality: 0.35." It should record who approved the weight, what promise or policy it reflects, and when it must return to council if changed. The risk here is compliance theatre: a register that exists but is never read, never contested, and never revised. Documentation is not the same as accountability. The register only works if there is a political moment in which its contents are actually put to question. The public authority that must defend the choices should own the register, with the vendor and project analyst serving defined supporting roles.

What this changes

The profession isn't climbing toward a hands-free horizon; it's moving sideways into a more contested ecology of human judgement, machine-generated options, institutional incentives, and legal exposure.

The risk isn't that professionals fail to climb. It's that organisations mistake specification for judgement.

Back to the housing strategist on Thursday afternoon. The programme matrix remains on her screen.

The anomaly flagged by her assistant is real: the model has optimised for density and land receipts by treating the courtyard promise as adjustable project space and stripping it of the status of public commitment. The calculation is technically accurate on its own metrics. It is blind to the loss of trust because that loss was never encoded as a constraint the model was allowed to respect.

She doesn't accept the variation. She doesn't rewrite the document by hand either. She rejects the automated optimisation, overrides the parameter constraint to protect the public promise, and forces the system to run the financial numbers again.

She can do this because she knows the history. She was in the room when the promise was made. But the harder question is what happens when no one in the

chain was in that room: when the parameter was inherited from a previous project, the model configured by a vendor, and the output approved in a meeting that ran over. The accountability still stays with the institution. Nobody inside it can reconstruct where the choice was made, by whom, or with what democratic mandate. That absence of reconstructability leads directly to the final question: what must the institution be able to explain in public?

AI hasn't made her job simpler. It's stripped away some administrative noise, leaving the political and normative choice where it belongs: visible, contested, and owned by the institution that must defend it.

*The
accountability
still stays with
the institution.*

Questions for Practice

- 01** At which delegation level does each recurring workflow currently operate?
- 02** Who specifies, reviews, overrides and signs the output at that level?
- 03** Where has saved production time created new verification or exception work?
- 04** Which judgement still needs to be learned through first-hand practice?



PART V

Public Purpose

The final question concerns what the institution has told AI to value, protect and ignore.

ARGUMENT MAP

This part moves from tools to public purpose.

Read for what institutions instruct AI to optimise, protect or ignore.

Key tension: operational clarity without democratic reduction.

CHAPTER 9

The model made the right choice. That was the problem.

The decisive question concerns what planning institutions instruct AI to optimise, protect or trade away.

IN THIS CHAPTER

This chapter turns the operating model into a doctrine for deciding which values AI may optimise and which commitments remain binding.

USE IT FOR

Writing a doctrine of use around public purpose and accountability.

WATCH FOR

Operational clarity that reduces democratic choices to optimisation.

The station area case from the previous chapter returns here, but the question is no longer how the model reduced the courtyard. It is what the institution owed before the model ever ran.

The calculation was coherent on its own terms. The courtyard promise remained adjustable because the institution had recorded it without protecting it in the optimisation logic. No rule was broken and no technical error occurred. The failure came earlier, when the organisation left a public commitment outside the system's hard constraints.

The project director therefore faces an institutional question: what had the organisation told AI to value?

AI capability makes institutional instruction decisive. Planning institutions must name what the system may not optimise away before the model runs. A wrong choice can be challenged. A hidden institutional choice, made easier to defend by the model, is harder to see.

The argument has moved from capacity to consequence. AI can do more of the analytical work of planning than most teams are currently using it for. But each gain moves the decisive question one step earlier: what has the institution instructed the system to value? Across the preceding chapters, operational

friction gave way to questions of memory, weights, model boundaries, civic influence and delegated burdens.

What is the city for, and who gets to decide?

AI accelerates the answer already present in the institution.

AI does not arrive in a neutral institution. It enters organisations with existing incentives, blind spots, political pressures and habits of avoidance.¹ Where the institutional logic is sound, it compounds good judgement. Where it is weak, it compounds avoidance. A municipality that treats housing delivery as the only real target will use AI to accelerate delivery. A development process that treats participation as risk management will use AI to manage that risk more efficiently.

1. UN-Habitat,
2022

A model may improve housing totals, participation volume and delivery speed while overlooking housing need, actual influence and the long-term cost of distrust.

Better calculation does not settle the question of what should count. That question has to be asked before the tool enters the room.

The decisive prompt is never in the window

Prompting skill is useful and insufficient on its own.

The more important prompt is the one the institution gives, usually without writing it down.

Maximise delivery. Minimise risk. Avoid political exposure. Protect land value. Keep the process moving. Make the proposal defensible. Reduce conflict before the council meeting.

None of these instructions needs to appear in a prompt window. They live in templates, performance indicators, budget cycles, legal advice and meeting habits. They shape what the model is asked to do and what its output will be used to justify.

The prior question is whether the institution's instructions are legitimate, alongside whether the system follows them.

A planning model can execute an illegitimate goal with great consistency, optimising land revenue at the expense of affordability, reducing complaint volume by narrowing participation, or producing an answer that is legally defensible yet publicly unrecognisable.

Cugurullo and Xu (2025) call this the oracle problem: AI systems in urban governance that generate authoritative-looking outputs without exposing the value assumptions built into them.² The oracle does not need to lie. It selects. And the selection happens before the steering group meets.

2. Cugurullo & Xu, 2025

The APA Foresight Trend Report 2026 points to a governance reframe already underway: the question has shifted from “can this work technically?” to “under what rules, with what oversight, for what public purpose, and with what protection of rights?”³ That shift places the governance problem inside the institution operating the algorithm.

3. APA, 2026

A public planning institution cannot treat public purpose as an internal preference. It works with democratic mandates, legal rights, distributional consequences and places people cannot simply leave.

The city is a shared condition governed through public mandates; a client brief cannot carry that role.

That is why alignment in planning is political before it is technical. There is no model that can rank competing public values once and for all. Housing delivery matters, and so does who can actually afford to live in what gets built. Climate adaptation matters, and so does whether children can still move freely through the neighbourhood. Public trust matters, and so does the right of existing residents not to be treated as an obstacle to a future drawn without them.

That is why the professional role does not disappear. It changes. The planner becomes less a producer of all analysis and more the person who ensures the questions, assumptions and value choices behind the analysis can be inspected, contested and defended.

The counterargument deserves a harder hearing

The capacity version is familiar. Many planning organisations are overloaded. Housing targets are urgent. Staff are tired. Residents often receive poor explanations because teams lack time, even where exclusion is absent from the intention. AI can protect planning values when it reduces administrative overload and gives smaller municipalities access to expertise they cannot hire.

AI does not make these problems worse by default. It makes them faster, harder to see and more difficult to contest. Slow, manual planning at least created friction. Sometimes that friction was the only moment the wrong answer could be questioned.

Capacity without direction produces faster drift and no public value. The question is what the institution spends that capacity on, and who inside the institution has the standing to ask that question before the model runs.

The city as optimisation problem

A city can be described as an optimisation problem. Land is scarce, public money is scarce, political attention is scarce. Every plan allocates scarcity, and every allocation creates winners, losers and future constraints. Optimisation has a place here. A planning organisation that refuses it because it sounds technocratic will make worse decisions.

The city exceeds an optimisation problem.

It is also a place where people form attachments, remember promises, avoid certain streets, claim dignity, grow old and disagree about what matters. Some of these things can be measured, some can only be named, and some only appear when the process has made room for them.

Optimisation has already entered planning.⁴ The danger begins when it becomes the grammar through which all other values must speak.

4. Sanchez et al.,
2025

A tree matters because it reduces heat. A courtyard matters because it raises social interaction metrics. A quiet corner matters because the public health literature supports it. Those arguments may be useful. But something is lost when every urban value needs an instrumental justification before it can enter the model. Some values need to be named as values.

Public value must be made operational, but not reduced

A value that is never operationalised will lose every time it meets a cost estimate or a legal deadline. Affordability becomes a rounding error. Trust disappears from the model entirely. But the moment values become operational, they risk being reduced: affordability to a percentage, dignity to an accessibility score, spatial quality to a checklist.

Each translation is useful. Each is incomplete. The work of governance is to keep both truths visible at the same time, which means every public value that enters an AI-assisted workflow needs two records: the measurable proxy and the reason the proxy is not enough. Someone must own that difference, or the proxy will quietly become the policy.

Take affordable housing. A model can track tenure mix, rent levels and eligibility categories. That is necessary. But the public question is wider: who can actually live here, under what security, with what access to schools and care? A technically compliant affordability score may still fail the urban question.

The professional question changes

The old professional question was: what is the right plan? The communicative question added: who must be heard before a plan can claim legitimacy? The algorithmic question now adds: what has the system been instructed to see, optimise, ignore and remember?

Urban professionals need enough technical literacy to recognise when a model exceeds its proper model boundary, where the source base is thin and which choices must remain under human authority.

The profession must develop the competence and secure the institutional standing to use it. Without that standing, professional judgement becomes a review task inside someone else's system.

The dominant promise of AI is speed. In many parts of planning, speed is welcome. But democratic planning also needs deliberate friction. Time to notice that a clean option depends on a broken promise. Time to let conflict surface before it becomes litigation, protest or quiet withdrawal.

A profession that cannot slow the tool down will eventually be governed by its pace.

The operating model

The five layers share one task: each keeps a different part of public judgement visible as AI moves through the workflow. Their force lies in sequence and connection, not in separate compliance checks.

The machine room reduces operational friction, but only when final responsibility remains human.

The project memory room preserves decisions, assumptions, commitments and unresolved tensions across time.

The decision package turns model output into a governable choice by separating constraints, assumptions, commitments, sensitivities, option history and judgement.

The civic decision room makes options, consequences and participation traceable before choices harden.

The doctrine of use names what AI may not trade away, who carries judgement and where freed capacity must go.

Together, these layers form a public planning discipline for AI use, grounded in institutional memory, democratic contestability and professional judgement.

The institution needs a doctrine of use

01 Before the model runs

Name what it may not trade away.

02 Before output reaches a decision maker

Name who carries judgement.

03 Before capacity is freed

Name where the attention goes.

Table 9.1. Doctrine of use. Source: author synthesis.

A doctrine can be stated in three rules. Before the model runs, name what it is not allowed to trade away. Before the output reaches a decision maker, name who carries the judgement. Before the capacity is freed, name where the attention goes.

On the first rule: every AI-assisted workflow requires a prior public value statement that names which commitments are hard constraints, not variables. In the station area case, that means naming the courtyard promise before the optimisation runs, not discovering its status afterwards. The EU AI Act establishes traceability and human oversight as legal baselines for high-risk applications;⁵ a doctrine of use makes those baselines operational in the specific context of urban planning.

5. EU AI Act, 2024

On the second: AI may support analysis, generate options and surface patterns. It may not decide which public promise is negotiable. In practice, this means a named project leader, not “the AI workflow”, signs for each consequential judgement. Goldsmith and Yang (2025) argue that AI in urban governance redistributes discretion without always redistributing accountability.⁶ The

6. Goldsmith & Yang, 2025, preprint

doctrine closes that gap by making accountability visible before a decision, not auditable after.

On the third: capacity freed by AI should be reinvested in judgement. When a participation synthesis falls from three weeks to three days, the recovered time belongs with the minority positions that the synthesis flagged without resolving. The APA Foresight Trend Report 2026 points to a risk already visible across municipalities: AI-governance quality is uneven, which means citizen protection has started to depend on which city someone lives in.³ Reinvesting freed attention is how a public institution refuses to let that divergence determine what counts as legitimate planning.

3. APA, 2026

The strongest critique accepts the tools

The strongest critique is that the tools may become useful too quickly for the institutions around them to adapt. Weak processes start to look competent. Thin participation looks responsive. Contested priorities look technically balanced. Political choices become harder to see because they are embedded earlier, in cleaner systems, behind better interfaces.

There is a harder possibility. The same instruments proposed here could make planning less contestable if they are used only to document decisions more elegantly. A project memory room can become a defensive archive. A parameter register can become compliance theatre. A digital twin can make a narrow option space look open. A doctrine of use can become another checklist that nobody is empowered to invoke.

The test is not whether the institution has records, models or rules. The test is whether someone outside the immediate project logic can challenge them before they harden into a decision.

That challenge right must be institutional, not personal. It cannot depend on a brave project manager, a technically literate councillor or an unusually well-organised residents' group. The system only becomes publicly governable when challenge is designed into the workflow: who can pause an optimisation, who can demand that a commitment is treated as hard, who can contest a weight, who can require an alternative scenario, and where that challenge is recorded.

The better response is institutional seriousness. Use the tools, but with a theory of public value. Use them with records that can be inspected and professional formation designed into the workflow. Use them with enough scepticism to know when the output is fluent and wrong. A sociotechnical audit of six commercial LLMs in urban infrastructure decisions found that 51.3% of cited regulations, codes and technical sources were fabricated or unverifiable, and that model confidence correlated negatively with reasoning quality: the worst-performing models answered with the greatest certainty.⁷ The audit

7. Poudel et al., 2026

shows why fluent output cannot serve as evidence of competence. Use them with enough ambition to know when manual work is no longer defensible.

Urban professionals protect their future role by shaping the systems entering planning practice.

What this means on Monday morning

The lesson from Option D is that revising a public promise requires authority, explanation and memory. Conditions change, and promises can be revised through a legitimate process. The model may identify the trade-off. It cannot decide whether the institution is allowed to make it.

Before using AI to draft a note, ask what the note must not flatten. Before summarising participation, ask whose position is most likely to disappear. Before generating options, ask which commitments are hard constraints and which are merely preferences. Before ranking scenarios, ask who approved the weights.

If the time saved is spent producing more output, the profession becomes faster. If it is spent checking assumptions, widening options, improving participation and confronting political choices earlier, the profession becomes better.

The institution around the model determines which of those outcomes follows.

The urban operating system is institutional

Every planning organisation already runs on an operating system, even if no one calls it that. It consists of laws and habits, templates and political routines, professional norms and the informal knowledge that never gets written down. AI does not replace that operating system. It runs on top of it. And when the operating system is weak, AI does not fix it. It exposes the weakness, scales it, or makes it harder to notice.

What does the organisation remember, and what does it forget? Who can challenge the model? Who owns the judgement? Who benefits from speed, and who pays for error? What kind of city is being made easier to build?

Back to the station area

The project director rejects Option D despite its technical strength.

She rejects it because the system has treated a public promise as a negotiable variable without asking whether the institution had the authority to do that.

The team reruns the model.

The new options are less elegant. More expensive. Harder to explain. One requires a political choice about the capital plan. One requires renegotiating the phasing with the developer. One requires the council to decide whether the housing target or the public space promise carries more weight under the changed financial conditions.

The conflict is difficult and unavoidable.

It is also the work.

For a moment, the AI system succeeds by forcing the institution to confront a choice it was about to hide from itself.

This model makes the political shortage of housing, municipal fiscal limits and conflicts between public values harder to hide behind speed, fluency or technical authority. That narrower contribution still matters.

The question behind the tools was never whether AI can help urban professionals. It can.

The question is whether urban professionals, and the institutions they work in, can become clear enough about public purpose to govern what AI makes possible.

Public values, evasions and compromises will shape the AI-mediated city through whatever the tools are allowed to accelerate.

Questions for Practice

- 01** Which public commitments operate as hard constraints before optimisation begins?
- 02** Who may change an objective, weight or constraint, and under what mandate?
- 03** What must the institution explain when the model's preferred option conflicts with an earlier promise?
- 04** Which decisions remain outside the system because they require democratic authority?

Epilogue

Tuesday Morning, Revisited

The station area has been under development for eleven years. The current project manager arrived eight months ago. She did not start from zero.

On her first week, she asked the project's AI assistant to brief her on the history of the southern residential block, which is currently stalled. Within an hour she had a structured account of the seven iterations the programme had gone through since 2021, the reasons each version was revised, the financial model assumptions that had been contested and why, and the three outstanding commitments made to the community association in 2023 that the previous team had not yet translated into binding contract language. She had the names of the relevant people on every side, the key dates in the permit track, and a note flagging that the water authority's consent conditions from 2022 contained a clause that had never been formally agreed by the developer.

That afternoon she met the community association. They asked her whether the green roof commitment from 2023 was still in the programme. She knew the answer before they finished the question. More importantly, she could show them the record of where that commitment came from, what it was tied to, and why it had survived three programme revisions. The meeting was an hour shorter than it might have been. It ended in trust; suspicion had receded.

The following week, new housing market data showed that demand for the mid-market rental segment had softened by 14% across the region over the previous two quarters. Her AI assistant ran the revised residual calculations overnight, flagged three programme scenarios that remained financially viable under the new assumptions, and drafted a note for the steering group identifying the decision that needed to be made and the conditions under which each option was defensible. She spent her morning writing the recommendation. The model did not need recalculation.

The project is not finished. It will not be finished for another six years. There will be more team changes, more market shifts, more legislative updates, more community objections that need to be heard and addressed. The knowledge built over eleven years remains queryable, auditable and accessible to everyone who needs it, under terms agreed at the start and maintained throughout.

The technology became dependable only after the institution defined how it could be used, checked and challenged.

The first briefing came from the machine room, while project memory recovered the old commitment. The decision package made the renewed trade-off defensible, and the civic decision room kept the discussion traceable. The doctrine of use set the boundary: changed conditions gave the system no licence to erase that commitment.

Where to Go Next

Figure 9.1 brings the argument of this collection into one connected operating model. The machine room, project memory room, decision package and civic decision room describe different places where professional work now changes. AI makes some acts of production cheaper and faster. The work does not disappear. It moves toward specification, interpretation, verification, institutional memory and public explanation. The practical question is where an organisation wants that burden to land.

The four rooms depend on one another. A fast machine room without project memory produces fluent work that cannot explain its own history. A project memory room without a decision package preserves material without turning it into inspectable choices. A decision package without a civic decision room can expose trade-offs while leaving those affected without standing to challenge them. And participation without a reliable operational foundation asks people to enter a process that may be unable to remember what they said. These are institutional capacities, not a sequence of software purchases.

A bounded workflow is still the most useful place to begin. It makes ownership, sources, review points and consequential errors concrete. Yet the boundary of the trial should never become the boundary of the thinking. From the first use, the team needs to decide what evidence remains recoverable, which assumptions must be visible and who carries the final judgement. A small workflow can teach an organisation a great deal, provided the learning is recorded and shared beyond the small group that ran the experiment.

Memory changes what later decisions can become. Urban projects outlast appointments, political terms, procurement cycles and software contracts. Their intelligence depends on whether a successor can reconstruct why a choice was made, which alternative was rejected and what was promised along the way. AI can make that record easier to query, but queryability is only useful when the underlying record has been governed. The institution must still decide who may add context, correct an error, contest a summary or explain a deviation from an earlier commitment.

The same discipline belongs in models and digital twins. A useful model helps a decision confront its consequences. It reveals the assumptions carrying a result, preserves credible alternatives and distinguishes what is impossible from what was unchosen or unfunded. Figure 9.1 places the decision package between memory and civic influence for a reason. Calculation becomes public value only

when the route from evidence to recommendation can be inspected and when participation arrives while the option space is still open.

Professional development has to be designed into this operating model. Junior staff learn judgement by making first passes, noticing exceptions, comparing imperfect sources and receiving correction from people who understand the context. AI can produce that first pass faster, but speed is a poor substitute for formation. Some tasks should remain human long enough for people to learn what a plausible answer conceals. Others can be delegated when the organisation has created a stronger route for learning from review, challenge and consequence.

The doctrine of use governs all four rooms. It names the public commitments that a model may not trade away, the people authorised to change objectives or weights, the evidence required before an output travels and the explanation owed when a recommendation conflicts with an earlier promise. Such a doctrine cannot be written once and filed. Tools change, laws develop, responsibilities move and experience reveals new failure modes. The doctrine has to remain a living public commitment, revised openly and tested against real decisions.

The next move depends less on technological ambition than on institutional character. Planning organisations will acquire more capable tools. The decisive issue is what they build around those tools: memory strong enough to survive turnover, judgement confident enough to resist fluent error, participation early enough to alter choices and accountability clear enough to remain visible when the model is right on its own terms. That responsibility belongs to the institution before it belongs to the system.

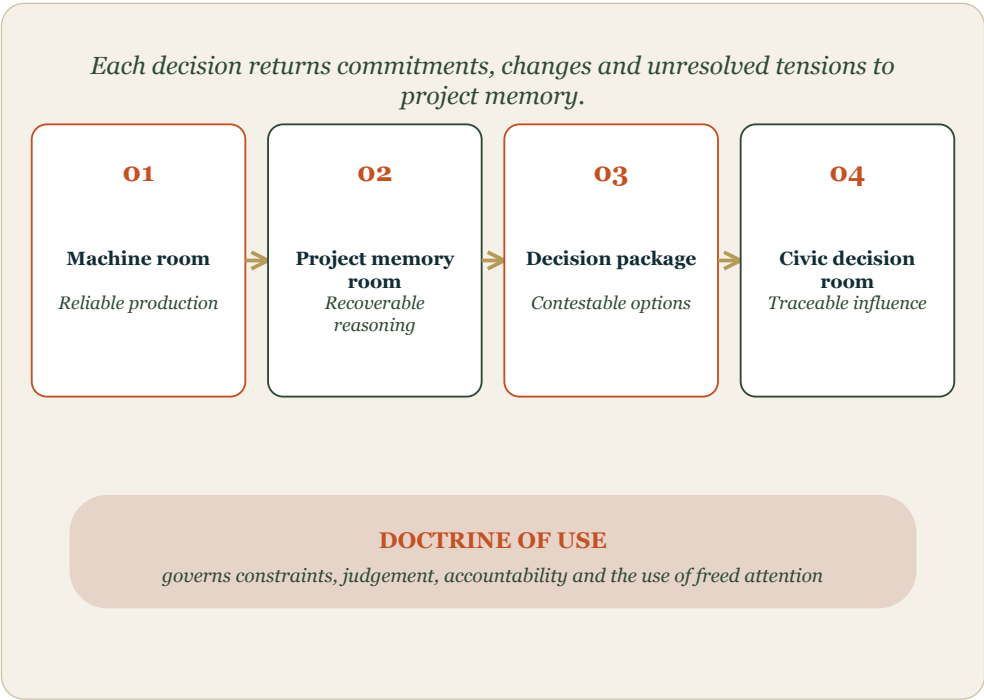


Figure 9.1. The connected operating model. Four workspaces carry planning from daily production to public influence, governed throughout by the doctrine of use. Source: author synthesis.

Sources and Notes

Short source labels appear on the page where a claim is cited. Full references, links and source qualifications are collected here.

Prologue

A Tuesday Morning, Five Years from Now

- 1 Massenkoff, M. & McCrory, P. (2026). *Labor market impacts of AI: A new measure and early evidence*. Anthropic Economic Index.
<https://www.anthropic.com/research/labor-market-impacts>. Used for the observed-exposure measure and the gap between technical exposure and observed professional use.
- 2 Si, S., Yao, Y. & Wu, J. (2025). *Artificial Intelligence in Urban Planning: A Bibliometric Analysis and Hotspot Prediction*. Land, 14(11). <https://www.mdpi.com/2073-445X/14/11/2100>. Used for the three-stage trajectory and the 3.6-times publication-volume finding.

Chapter 1

The Adoption Gap

- 1 Massenkoff, M. & McCrory, P. (2026). *Labor market impacts of AI: A new measure and early evidence*. Anthropic Economic Index.
<https://www.anthropic.com/research/labor-market-impacts>
- 2 OpenAI (2026). *Ending the Capability Overhang*.
<https://cdn.openai.com/pdf/openai-ending-the-capability-overhang.pdf>
- 3 Si, S., Yao, Y. & Wu, J. (2025). *Artificial Intelligence in Urban Planning: A Bibliometric Analysis and Hotspot Prediction*. Land, MDPI, 14(11). <https://www.mdpi.com/2073-445X/14/11/2100>
- 4 Described in: Fabrizio Dell'Acqua et al. (2023). *Navigating the Jagged Technological Frontier*. Harvard Business School Working Paper.
Limitation: Preprint; used cautiously, not planning-specific.
- 5 Raymond, C.M. et al. (2025). *Uses, opportunities and risks of artificial intelligence in participatory urban planning*. Discover Cities, Springer. <https://doi.org/10.1007/s44327-025-00137-4>.
Othengrafen, F., Sievers, L. & Reinecke, E. (2025). *From Vision to Reality: The Use of AI in Different Urban Planning Phases*. Urban Planning, 10.
<https://www.cogitatiopress.com/urbanplanning/article/view/8576>
- 6 Wharton School, University of Pennsylvania (2026). *Thinking, Fast, Slow, and Artificial: How AI is Reshaping Human Reasoning and the Rise of Cognitive Surrender*.
https://papers.ssrn.com/sol3/papers.cfm?abstract_id=6097646
Limitation: Preprint; used cautiously, not planning-specific.
- 7 Romero, A. (2026). *A New Wharton Study on AI Warns of a Growing Problem: Cognitive Surrender*. The Algorithmic Bridge. <https://www.thealgorithmicbridge.com/>
Limitation: Commentary; used as conceptual framing, not primary evidence.

Chapter 2

Three Paradigms, One Professional

- 1 The communicative turn in planning theory is developed most fully in: Healey, P. (1997). *Collaborative Planning: Shaping Places in Fragmented Societies*. Macmillan. The original theoretical grounding draws on Habermas, J. (1984). *The Theory of Communicative Action*. Beacon Press.
- 2 Coase, R. (1937). The nature of the firm. *Economica*, 4(16), 386-405. The transaction cost argument, the initiative/management distinction, and the direction/freedom definition of the employment relationship are all drawn from this source. The application to planning institutions is an analytical analogy, not a claim Coase himself makes.
- 3 The three-mode framework and the shift toward algorithmic urbanism are developed in: Deelstra, R. (2026). *The Urban Operating System*. <https://remcodeelstra.substack.com/p/the-urban-operating-system>. Earlier author essay used as background for the three-mode framework; not used as independent external evidence.
- 4 Goodspeed, R. (2016). The death and life of collaborative planning theory. *Planning Theory & Practice*, 17(2), 275-281.
- 5 NASCIO & Accenture (2025). *Harnessing GenAI to Elevate the Citizen Experience*. September 2025. <https://www.nascio.org/press-releases/new-report-reveals-how-genai-can-elevate-citizen-experience-in-state-government/>. The survey covers US state chief information officers, not planning organisations specifically. It is used here as an indicator of institutional readiness patterns in the public sector more broadly.
- 6 APA (2026). *2026 APA Foresight Trend Report for Planners*. American Planning Association. The “third time layer” framing is the author’s synthesis of current AI tool cycles and planning practice timelines, informed by the report’s treatment of AI as a fast-moving planning-relevant trend.
- 7 Caulfield, M. (2017). *Web Literacy for Student Fact-Checkers*. The SIFT framework is the foundational work. His adaptation for AI-generated content is discussed in: Panke, S. (2025). Interview with Mike Caulfield about Deep Background, AI literacy and future skills. *AACE Review*, August 2025. <https://aace.org/review/caulfield-ai>. The “epistemological shift” framing in this section is the author’s interpretation of Caulfield’s argument applied to planning practice; it should not be read as Caulfield’s own claim about planning specifically.
- 8 Malekzadeh, M. (2026). Urban planners should not be afraid of AI: the planner as curator. *Cities*, Elsevier, 168, article 106497. <https://doi.org/10.1016/j.cities.2025.106497>
- 9 Krivý, M. (2018). Towards a critique of cybernetic urbanism: the smart city and the society of control. *Planning Theory*, 17(1), 8-30.

Chapter 3

The Machine Room

- 1 Brynjolfsson, E., Li, D. & Raymond, L. (2025). *Generative AI at Work*. Quarterly Journal of Economics, 140(2), 889-942. NBER Working Paper 31161. <https://www.nber.org/papers/w31161>. Average 14% productivity gain across 5,179 customer support agents; 34% for novice workers; near-zero for the most experienced. Aldasoro, I. et al. (2026). *AI Adoption and Firm Productivity*. CEPR VoxEU / BIS. Cited in Stanford HAI, *AI Index Report 2026*, Ch. 4, Fig. 4.4.28. Study of 12,000 European firms: +4% labour productivity from AI adoption, amplified by training. Yotzov, I. et al. (2026). *Firm Data on AI*. NBER Working Paper 34836. <https://www.nber.org/papers/w34836>. Nine in ten executives across 6,000 firms report no measurable productivity impact despite 69% active AI use.

Limitation: Preprint; used cautiously, not planning-specific.

- 2 Dell’Acqua, F. et al. (2023). *Navigating the Jagged Technological Frontier*. Harvard Business School Working Paper 24-013. hbs.edu/faculty. Within the frontier: +12.2% task completion, +25.1% speed, +40% quality. Outside the frontier, a separate study found the inverse: Becker, M. et al. (2025). *Measuring the Impact of AI on Software Engineers’ Productivity*. METR. Cited in Stanford HAI, *AI Index Report 2026*, Ch. 4, Fig. 4.4.27. Experienced developers became 19% slower with AI assistance, with a gap between perceived helpfulness and actual performance.

Limitation: Preprint; used as mechanism evidence, not planning-specific.

- 3 Raymond, C.M. et al. (2025). *Uses, opportunities and risks of artificial intelligence in participatory urban planning*. Discover Cities, Springer. <https://doi.org/10.1007/s44327-025-00137-4>. Othengrafen, F., Sievers, L. & Reinecke, E. (2025). *From Vision to Reality: The Use of AI in Different Urban Planning Phases*. Urban Planning, 10. <https://www.cogitatiopress.com/urbanplanning/article/view/8576>. Both studies are the empirical basis for the participation advisor case. Raymond et al. find that planners see AI as most useful for summarising, reporting and analysing feedback.

- 4 The readiness gap between perceived and actual AI-readiness at the point of rollout is documented across multiple enterprise AI adoption studies from 2025 to 2026, including findings cited in the AI Index Report 2026 (Stanford HAI) and enterprise data governance research. The consistent finding: organisations overestimate their data and infrastructure readiness before implementation.

- 5 Charlotin, D. (2026). *AI Hallucination Cases Database*. HEC Paris / Sciences Po. <https://www.damiencharlotin.com/hallucinations/>. The live database listed more than 1,400 court decisions on 25 May 2026 and explicitly notes that it tracks judicial decisions discussing alleged or established AI hallucinations, not every problematic filing. Illinois Courts guidance confirms that lawyers retain professional responsibility for verifying cited cases: <https://www.illinoiscourts.gov/News/1665/Paste-in-Haste-The-Fallout-of-AI-Hallucinations-in-Court-Filings-and-the-New-ARDCs-Guide-to-Implementing-AI/news-detail/>

- 6 The gap between task-level productivity gains and organisational-level impact is the central finding of the macro-level studies surveyed in Stanford HAI, *AI Index Report 2026*, Ch. 4, Fig. 4.4.28. Yotzov et al. (2026) found that nine in ten executives reported no measurable productivity impact despite active AI use, consistent with Brynjolfsson’s “J-curve” hypothesis that organisations absorb adoption costs before gains materialise. See note 1 for full citations.

- 7 Romero, A. (2026). *A New Wharton Study on AI Warns of a Growing Problem: Cognitive Surrender*. The Algorithmic Bridge. <https://www.thealgorithmicbridge.com/>. Romero’s distinction between cognitive offloading (strategic delegation while retaining judgement) and cognitive surrender (accepting output because it looks authoritative) is the conceptual basis for the first-pass / final-pass distinction in this chapter.

Limitation: Commentary; used as conceptual framing, not primary evidence.

- 8 Wharton School, University of Pennsylvania (2026). *Thinking, Fast, Slow, and Artificial: How AI is Reshaping Human Reasoning and the Rise of Cognitive Surrender*. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=6097646. 1,372 participants; incorrect AI advice followed in 79.8% of cases; confidence increased by 11.7 percentage points even when the AI was wrong; high AI trust associated with a 3.5-fold increase in the likelihood of following a faulty answer.

Limitation: Preprint; used cautiously, not planning-specific.

- 9 On entry-level employment: Brynjolfsson, E. et al. (2025). Cited in Stanford HAI, *AI Index Report 2026*, Ch. 4, Figs. 4.4.29-30. Employment for workers aged 22-25 in AI-exposed occupations fell approximately 16-20% from 2022 peaks, while headcount for older groups continued to grow. Massenkoff, M. & McCrory, P. (2026). *Labor market impacts of AI*. Anthropic Economic Index. <https://www.anthropic.com/research/labor-market-impacts>. Complementary signal: -14% job-finding rate for the same age group. On learning penalties: Shen, C. & Tamkin, A. (2025). Cited in Stanford HAI, *AI Index Report 2026*, Ch. 4, Fig. 4.4.27. Software engineers who relied heavily on AI for learning showed no measurable speed improvement and faced measurable learning penalties.

Limitation: *Used as mechanism evidence, not planning-specific.*

- 10 Brynjolfsson, E., Li, D. & Raymond, L. (2025). *Generative AI at Work*. Quarterly Journal of Economics, 140(2), 889-942. NBER Working Paper 31161. <https://www.nber.org/papers/w31161>. The 34% productivity gain for novice and low-skilled workers is the empirical basis for the counterargument. The study also provides suggestive evidence that AI disseminates the tacit knowledge of more experienced workers to less experienced colleagues, but only when the less experienced worker is actively engaged with the task, not when they are passively receiving completed outputs.

Limitation: *Preprint; used cautiously, not planning-specific.*

- 11 Vectara Hallucination Leaderboard (November 2025 update). On grounded summarisation tasks, leading models now achieve below 1% hallucination rates. On open-ended factual tasks, reasoning models perform substantially worse: OpenAI o3 hallucinated at 33-51% on SimpleQA/PersonQA, compared to 16% for its predecessor o1; DeepSeek-R1 hallucinated at 14.3% versus 3.9% for its base model. The divergence between anchored and open-ended performance is documented in: Graffius, S.M. (2026). *Are AI Hallucinations Getting Better or Worse?* <https://www.scottgraffius.com/blog/files/ai-hallucinations-2026.html>.

Chapter 4

The Project Memory Room

- 1 Nonaka, I. & Takeuchi, H. (1995). *The Knowledge-Creating Company: How Japanese companies create the dynamics of innovation*. Oxford University Press. The SECI model describes four modes of knowledge conversion: socialisation (tacit to tacit, through shared experience and proximity), externalisation (tacit to explicit, through articulation and reflection), combination (explicit to explicit) and internalisation (explicit to tacit, through learning by doing). The central observation relevant here: the most strategically valuable organisational knowledge resides in the socialisation mode, shared in corridor meetings, informal practice and observation. It is structurally difficult to capture through documentation alone. Nonaka and Takeuchi note that Western organisations tend to over-rely on explicit-to-explicit conversion, systematically underinvesting in the conditions under which tacit knowledge is shared.
- 2 Aerts, G., Dooms, M. & Haezendonck, E. (2017). *Knowledge transfers and project-based learning in large scale infrastructure development projects: an exploratory and comparative ex-post analysis*. International Journal of Project Management, 35(3), 224-240. <https://doi.org/10.1016/j.ijproman.2016.10.010>. Comparative case study of three projects: Gaaspedammer Tunnel (Netherlands), Hong Kong-Zhuhai-Macau Bridge (China) and Crossrail (UK). Core finding: knowledge in large infrastructure settings tends to remain tacit, explicit project management knowledge is not systematically captured, and learning stays attached to individuals and does not become an institutional asset.

- 3 Galan, N. (2023). *Knowledge loss induced by organizational member turnover: a review of empirical literature, synthesis and future research directions*. Parts I and II. *The Learning Organization*, 30(2), 117–136 and 137–161. <https://doi.org/10.1108/TLO-09-2022-0107> and <https://doi.org/10.1108/TLO-09-2022-0108>. Systematic review based on 91 empirical studies. Central finding: knowledge loss through staff turnover is a structural organisational problem with documented performance effects at unit and organisational level, not primarily an HR issue.
- 4 Klesel, M. & Wittmann, H.F. (2025). *Retrieval-Augmented Generation (RAG)*. *Business & Information Systems Engineering*, 67(4), 551–561. Frankfurt University of Applied Sciences. <https://doi.org/10.1007/s12599-025-00945-3>. RAG architectures ground LLM responses in an organisation's own documents, reducing hallucination risk and enabling source-traceable answers. Key limitations for public-sector use: data quality requirements are high; the “blinkered chunk effect” limits comprehensive understanding in standard implementations; skewed source bases introduce new bias effects.
- 5 Zhao, A. et al. (2025). *Identifying, Capturing, and Reusing Tacit Knowledge in Creative Domains with Generative AI*. Adjunct Proceedings of the 38th Annual ACM Symposium on User Interface Software and Technology (UIST). <https://dl.acm.org/doi/10.1145/3746058.3758467>. Research on using generative AI to surface tacit knowledge in design workflows. Core mechanism: AI-generated output presented as a partial or hypothetical reconstruction of a decision process prompts professionals to reflect on and articulate the implicit reasoning behind their choices. The research domain is graphic and UI/UX design. The mechanism is transferable to professional domains where expertise is tacit and handover is structurally incomplete.
- 6 Federal Aviation Administration. (2020). *Aviation Safety Action Program*, Advisory Circular AC 120-66C. <https://www.faa.gov/regulationspolicies/advisorycirculars/index.cfm/go/document.information/documentID/1037363>. ASAP encourages voluntary safety reporting through a partnership between the FAA and certificate holder that may include an employees' labour organisation. An FAA programme overview lists 262 operators and 767 programmes: <https://www.faa.gov/newsroom/aviation-voluntary-reporting-programs-1>. The analogy supports protected learning and reporting, while its application to planning remains an author interpretation.
- 7 EU AI Act, Articles 12 and 13. Article 12 requires high-risk AI systems to maintain logs enabling post-hoc risk identification and conformity evaluation; Article 13 requires transparency enabling deploying organisations to interpret and act on outputs responsibly. AI Act Service Desk, European Commission. <https://ai-act-service-desk.ec.europa.eu/en/ai-act/article-12>. Urban planning organisations should not assume that AI systems informing decisions about housing access, spatial rights, permits or public services will remain outside demanding governance expectations as the regulatory framework develops.

Chapter 5

Options and Numbers

- 1 Tay, J. Z., Ortner, F. P., Wortmann, T. & Aydin, E. E. (2024). *Computational Optimisation of Urban Design Models: A Systematic Literature Review*. *Urban Science*, 8(3), 93. The review covers research from 2012 to 2022 and defines Urban Design Optimisation as a way to explore many design solutions for a district, evaluate multiple objectives simultaneously, and select from Pareto-efficient options. Pareto-efficient here means: combinations where improving one objective requires accepting a cost on another. <https://www.mdpi.com/2413-8851/8/3/93>

- 2 Haasnoot, M., Kwakkel, J. H., Walker, W. E. & ter Maat, J. (2013). *Dynamic adaptive policy pathways: A method for crafting robust decisions for a deeply uncertain world*. Global Environmental Change, 23(2), 485–498. The DAPP method, developed at Deltares and TU Delft, introduces the concept of adaptation tipping points: conditions under which a current policy no longer meets its objectives and new action is required. The Deltares website provides a current overview of the method and its applications. <https://www.deltares.nl/en/expertise/areas-of-expertise/sea-level-rise/dynamic-adaptive-policy-pathways>
- 3 Mashhadi Moghaddam, S. N. & Cao, H. (2026). *The metrics trap: how technical sophistication masks social harm in urban AI systems*. npj Urban Sustainability. The authors analyse 28 AI implementations selected from 157 deployments across six domains and show how technical accuracy can mask policy failure. <https://www.nature.com/articles/s42949-026-00394-1>

Chapter 6

The Digital Twin

- 1 Batty, M. (2024). *Digital twins in city planning*. Nature Computational Science, 4, 192–199. Batty frames urban digital twins as a broad family of models, from aggregate economic and behavioural processes to local simulations, and focuses on how scientists, policymakers and planners interact with real cities through these models. <https://www.nature.com/articles/s43588-024-00606-7>
- 2 European Commission. (2025). *EU Local Digital Twin Toolbox*. The toolbox is described as a modular, standards-based suite of tools for European cities and communities, built around interoperability, openness and scalability, to support simulation, analysis, policy testing and sustainable development strategies. https://regions-and-cities.ec.europa.eu/cities-portal/eu-initiatives-cities/eu-local-digital-twin-ldt-toolbox_en
- 3 Moroni, S. (2025). *Insurmountable limitations of city-scale digital twins? On urban knowledge and planning*. Computational Urban Science, 5, article 17. Moroni's argument is stronger than this chapter's use of it suggests. He contends that city-scale digital twins face not merely practical limitations but in-principle constraints: urban systems are complex; calling them merely complicated understates the constraint, and the dispersed, tacit, dynamic nature of urban knowledge cannot be resolved by better data or stronger models. This chapter draws on Moroni's diagnosis while taking a less conclusive position on the remedy: the twin remains useful as a partial and contestable representation, not as a complete one. <https://link.springer.com/article/10.1007/s43762-025-00174-0>
- 4 Adade, D. & de Vries, W. T. (2025). *A systematic review of digital twins' potential for citizen participation and influence in land use agenda-setting*. Discover Sustainability, 6, article 354. The review identifies 34 studies and argues that citizen-centred digital twins can support agenda setting by helping citizens frame problems, convince others and propose alternatives, while also warning about access, privacy and institutional barriers. <https://link.springer.com/article/10.1007/s43621-025-01204-x>
- 5 Cardullo, P. & Kitchin, R. (2019). Being a 'citizen' in the smart city: up and down the scaffold of smart citizen participation. *Urban Geography*, 40(1), 103–126. Cardullo and Kitchin show that participation in smart city environments is often structured from above: citizens are invited to engage within parameters set by technology and institutions, without real influence over the choices that matter. Their concept of "citizenship from above" is the empirical counterpart to the managed attention pattern described here. <https://doi.org/10.1080/02723638.2018.1473436>
- 6 Weil, C., Bibri, S. E., Longchamp, R., Golay, F. & Alahi, A. (2023). *Urban digital twin challenges: A systematic review and perspectives for sustainable smart cities*. Sustainable Cities and Society, 99, 104862. The review identifies eight categories of bottlenecks across technical, organisational and governance dimensions. <https://www.sciencedirect.com/science/article/pii/S2210670723004730>

- 7 UN-Habitat. (2021). *Centering People in Smart Cities*. The playbook argues for digital governance centred on people, human rights, digital equity, responsible data governance, security and public capacity, including open standards, shared data ownership, interoperability and protection against vendor lock-in. <https://unhabitat.org/centering-people-in-smart-cities>
- 8 Future City Foundation / Data- en Kennishub Gezond Stedelijk Leven. (2021-2025). *Innovatielijn Digital Twin: koppeling gezondheidsdata aan 3D-planomgevingen*. A consortium of RIVM, municipality of Amersfoort, municipality of Zwolle, Province Utrecht and Future City Foundation coupled health calculation tools (including RIVM's Groene Batenplanner) to 3D planning environments such as Tygron. The methodology makes the health effects of spatial choices, including green scenarios, air quality and heat stress, visible before decisions are locked. Practical tests were conducted in Amersfoort (Schuilenburg, Soesterkwartier and Langs Eem en Spoor). The GGO Digital Twin was subsequently developed by Province Utrecht in collaboration with municipality of Amersfoort, Urban Sync and Tygron as a replicable planning support instrument. <https://future-city.nl/groene-baten-planner-gekoppeld-aan-digital-twin> and <https://gezondstedelijklevenhub.nl/tools/handboek-ggo-digital-twin/>

Chapter 7

Participation

- 1 Arnstein, S. R. (1969). *A Ladder of Citizen Participation*. Journal of the American Institute of Planners, 35(4), 216-224. <https://doi.org/10.1080/01944366908977225>. Arnstein's core claim is that participation is meaningful to the extent that it redistributes influence over decisions.
- 2 Blue, G., Rosol, M. & Fast, V. (2019). *Justice as Parity of Participation: Enhancing Arnstein's Ladder Through Fraser's Justice Framework*. Journal of the American Planning Association, 85(3), 363-376. <https://doi.org/10.1080/01944363.2019.1619476>. The article extends Arnstein's redistribution focus through Fraser's linked dimensions of recognition and representation.
- 3 OECD. (2025). *AI in civic participation and open government*, in *Governing with Artificial Intelligence: The State of Play and Way Forward in Core Government Functions*. OECD Publishing. <https://doi.org/10.1787/795de142-en>. The chapter identifies practical uses of AI for participation while stressing transparency, accountability, testing and evaluation.
- 4 Raymond, C. M. et al. (2025). *Uses, opportunities and risks of artificial intelligence in participatory urban planning*. Discover Cities. <https://doi.org/10.1007/s44327-025-00137-4>. A Delphi study with Finnish planners maps opportunities and risks across analytical, generative and humanised AI.
- 5 Adade, D. & de Vries, W. T. (2025). *A systematic review of digital twins' potential for citizen participation and influence in land use agenda-setting*. Discover Sustainability, 6, article 354. <https://doi.org/10.1007/s43621-025-01204-x>. The review finds that citizen-centred digital twins can strengthen agenda setting and proposal formation under accessibility, interactivity, privacy and institutional conditions.
- 6 Awashra, I. et al. (2025). *Generative AI text-to-image for community participation in landscape planning*. Landscape and Urban Planning, 264, 105464. <https://doi.org/10.1016/j.landurbplan.2025.105464>. The authors use the term "controlled imperfect" for images intended to support discussion without implying a settled future.
- 7 Jungherr, A. & Rauchfleisch, A. (2025). *Artificial Intelligence in Deliberation: The AI Penalty and the Emergence of a New Deliberative Divide*. Government Information Quarterly, 42(4), 102079. <https://doi.org/10.1016/j.giq.2025.102079>. A preregistered survey experiment with a representative German sample found lower willingness to participate and lower expected process quality in AI-facilitated deliberation.

- 8 Pande, D., Khanal, S. & Taeihagh, A. (2025). *Conceptualising Public Trust in High-Risk AI Policymaking: A Comparative Analysis of Three Asian Cities*. *Journal of Comparative Policy Analysis: Research and Practice*, 27(5-6), 578-602. <https://doi.org/10.1080/13876988.2025.2543546>. The study uses 3,600 survey responses from Seoul, Singapore and Tokyo and finds that institutional trust shapes trust in high-risk AI.
- 9 Wang, Y., Chen, H., Wei, X., Chang, C. & Zuo, X. (2025). *Trusting the machine: Exploring participant perceptions of AI-driven summaries in virtual focus groups with and without human oversight*. *Computers in Human Behavior: Artificial Humans*, 6, 100198. <https://doi.org/10.1016/j.chbah.2025.100198>. Disclosed AI summarisation was accepted by participants, alongside concerns about accuracy, authenticity and emotional nuance.
- 10 Rountree, J. & Gastil, J. (2026). *The Case for Using Generative AI to Run Deliberation Simulations*. *Journal of Deliberative Democracy*, 21(1), 1-14. <https://doi.org/10.16997/jdd.1625>. The authors present simulations as preparation and research tools that should complement rather than replace human deliberation.

Chapter 8

The Burden Moves Sideways

- 1 National Highway Traffic Safety Administration (NHTSA). (2013). *Preliminary Statement of Policy Concerning Automated Vehicles*. An early public automation classification, used here as an analytical framework for task delegation.
- 2 Coase, R. (1937). The nature of the firm. *Economica*, 4(16), 386-405.
- 3 Shahidi, P., Rusak, G., Manning, B. S., Fradkin, A. & Horton, J. J. (2025). The Coasean Singularity? Demand, supply, and market design with AI agents. NBER Working Paper No. 34468. Used here for the argument that AI agents may reduce transaction costs and reshape what is organised inside firms and markets.

Limitation: Preprint; used cautiously, not planning-specific.
- 4 Bedard, J., Kropp, M., Hsu, M., Karaman, O.T., Hawes, J. & Rosen Kellerman, G. (2026). When using AI leads to “brain fry.” *Harvard Business Review*, March 2026. BCG / University of California, Riverside. Survey of 1,488 full-time US workers. Brain fry defined as acute cognitive fatigue from excessive use or oversight of AI tools beyond one’s cognitive capacity; highest risk found among those overseeing multiple AI agents simultaneously. Used here as context for oversight-specific cognitive fatigue; adapted for the planning domain.

Limitation: Commentary; used as conceptual framing, not primary evidence.
- 5 Shaw, S.D. & Nave, G. (2026). Thinking, fast, slow, and artificial: How AI is reshaping human reasoning and the rise of cognitive surrender. Wharton School Research Paper (preprint). Three pre-registered experiments, N=1,372, 9,593 individual trials. Note: at time of writing, this paper has not yet completed peer review; the experimental conditions involved structured reasoning tasks and did not reproduce complex professional planning decisions.

Limitation: Preprint; used cautiously, not planning-specific.
- 6 Massenkoff, M. & McCrory, P. (2026). Labor market impacts of AI: A new measure and early evidence. *Anthropic Economic Index*, March 2026. The study found no detectable unemployment increase for highly AI-exposed workers through late 2025. Used here as broader labour-market context, not as direct evidence for the apprenticeship claim.

- 7 Microsoft. (2025). 2025 Work Trend Index: Microsoft and LinkedIn. Used here as broader context on changing AI skills, the rise of agent-based work, and the growing distinction between using AI, working with AI, and managing AI-enabled workflows.

Limitation: *Commentary; used as conceptual framing, not primary evidence.*

Chapter 9

The model made the right choice. That was the problem.

- 1 UN-Habitat. 2022. Artificial Intelligence and Cities: Risks, Applications and Governance. <https://unhabitat.org/ai-cities-risks-applications-and-governance>. Used for the foundational observation that AI embeds value choices unconsciously unless institutions deliberately govern toward the public interest, and that optimisation can sideline justice, equity and lived experience.
- 2 Cugurullo, F. and Xu, Y. 2025. When AIs become oracles: generative artificial intelligence, anticipatory urban governance, and the future of cities. *Policy and Society*, 44(1), 98–115. <https://doi.org/10.1093/polsoc/puae025>. Used for the oracle problem: AI systems in urban governance that generate authoritative-looking outputs without exposing their underlying value assumptions.
- 3 American Planning Association. 2026. *Foresight Trend Report for Planners*. <https://www.planning.org/publications/document/9283891/>. Used for: the governance reframe from technical capability to rules, oversight and public purpose; and the observation that local AI-governance quality is already uneven across municipalities.
- 4 Sanchez, T.W., Brenman, M. and Ye, X. 2025. The Ethical Concerns of Artificial Intelligence in Urban Planning. *Journal of the American Planning Association*, 91(2), 294–307. <https://doi.org/10.1080/01944363.2024.2355305>. Used for the argument that bias, opacity and accountability gaps in AI planning tools are not hypothetical but already present in deployed systems.
- 5 European Union. Regulation EU 2024/1689, the Artificial Intelligence Act. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689>. Used for the legal baseline on traceability, human oversight and risk management in high-risk AI applications, as context for the doctrine of use argument.
- 6 Goldsmith, S. and Yang, J. 2025. *AI and the Transformation of Accountability and Discretion in Urban Governance*. Preprint, arXiv:2502.13101. <https://arxiv.org/abs/2502.13101>. Used for the argument that AI can redistribute discretion and therefore requires corresponding accountability arrangements.
Limitation: *Preprint; used cautiously, not planning-specific.*
- 7 Poudel, A., Barrios, C., De La Torre, P., Ton, H., Surface, T., Mehta, V. and Silwal, S. 2026. *Governance risks of AI reasoning in urban infrastructure through Delphi audit of human and large language model judgment*. Discover Cities, 3, article 81. <https://doi.org/10.1007/s44327-026-00268-2>. A Delphi-derived audit of six commercial LLMs found that 51.3% of model-cited sources were unverifiable or fabricated and that self-reported confidence correlated negatively with reasoning quality ($r = -0.23$).

Explanatory Glossary

Terms are defined as they are used in this collection. The definitions describe a working vocabulary for practice. They do not claim to form a universal taxonomy.

**adaptation
tipping point**

The condition or moment at which a current policy no longer meets its objectives and a different action or pathway becomes necessary.

adoption gap

The distance between what available AI can do and what organisations use in governed professional workflows.

AI penalty

A decline in performance that occurs when people use AI outside the model's reliable capability frontier or fail to apply the judgement needed to detect a poor answer.

**algorithmic
governance**

The rules, authority and review arrangements through which institutions direct and control algorithmic systems used in public decisions.

**algorithmic
urbanism**

An emerging planning logic built around real-time data, iterative simulation and AI-assisted analysis. It changes how options are produced without settling who should decide among them.

**assumption
register**

A maintained record of the assumptions behind an option or calculation, including their source, owner, confidence and effect on the result.

**authoritative
picture**

A representation that is treated as the correct view of reality because it looks complete and persuasive, even though it rests on hidden assumptions, omissions and value choices.

brain fry

Practitioners' informal term for the fatigue produced by sustained AI-assisted work, especially repeated prompting, checking and correction.

**capability
overhang**

The gap between capabilities already available in AI systems and the narrower set that organisations use in practice.

**citizenship from
above**

A form of participation designed and bounded by institutions, where authorities determine who may participate, through which categories and on what terms.

**civic decision
room**

The setting in which evidence, options, public input, authority and commitments meet while a decision can still change.

civic twin	A twin in which actors with real standing can challenge assumptions, contest scenarios and trace how evidence and participation affect decisions.
Coasean shift	A relocation of production, coordination and verification costs after AI makes some acts of production cheaper.
cognitive offloading	The transfer of part of a mental task to an external aid, such as notes, software or AI, while the person retains enough understanding to inspect and use the result.
cognitive surrender	Acceptance of AI output without constructing or testing the reasoning needed to judge it.
communicative action (Habermas)	Action oriented toward mutual understanding through reason-giving among participants. In planning, it supports legitimacy through deliberation, with technical authority alone insufficient.
constraint map	A structured view of legal, physical, financial, political and public-value conditions that shape which options remain available.
contestability	The practical ability to inspect, question, pause and change a model, assumption, weight or decision before it becomes authoritative.
controlled imperfect principle	A design principle for AI in participation in which the system supports human deliberation while keeping its limits and failure modes visible.
DAPP (dynamic adaptive policy pathways)	An approach that keeps several future pathways open and links later actions to monitored conditions, preserving the ability to revise a forecast.
data governance	The allocation of responsibility for data quality, access, provenance, security, maintenance and lawful use.
decision package	A set of materials that exposes the option space, assumptions, constraints, sensitivities, commitments and recommendation so decision-makers can inspect how the conclusion was reached.
digital model	A digital representation that describes or simulates a system without an automatic flow of observed data back into it. It may support analysis while remaining a static or manually updated artefact.
digital shadow	A digital representation updated by data flowing from the physical system into the model. Information moves from reality to the representation, without the model automatically acting back on reality.

digital twin	A digital representation of a physical urban system, made up of entities and relationships over time, used to simulate change and inform a defined decision. Its value depends on visible assumptions, omissions and routes for contesting the result.
doctrine of use	A public institution's explicit rules for what AI may optimise, which commitments constrain it, who may change objectives and how outputs must be reviewed and explained.
EU AI Act	Article 12 addresses record-keeping for high-risk AI systems. Article 13 addresses transparency and the information deployers need to interpret and use those systems appropriately.
false precision	The appearance of certainty created by exact numbers that exceed the quality of the underlying evidence or assumptions.
five levels of delegation (Lo to L4)	Five modes ranging from manual production to governed optimisation. Each level changes the role of AI and relocates work toward review, calibration, oversight or institutional design.
Goodhart's Law	When a measure becomes a target, behaviour adapts to the measure and can weaken its value as an indicator of the underlying public objective.
initiative versus management (Coasean)	Coase's distinction between acting under one's own initiative and acting under another person's direction within an organisation. AI can alter where that boundary is practical without removing questions of authority.
institutional memory	Knowledge that remains available to an organisation across staff changes, project phases and political terms.
jagged frontier	The uneven boundary of AI capability, where a model may perform strongly on one demanding task and fail on a nearby task that appears simpler.
knowledge graph	A structured network of entities and relationships that makes connections among decisions, actors, assumptions, places and documents queryable.
ladder of participation (Arnstein)	A framework that distinguishes degrees of public power, from manipulation and tokenism to partnership, delegated power and citizen control.
LLM	A large language model trained on large text collections to generate and transform language. Fluency does not establish that an answer is accurate, complete or grounded in the relevant sources.

machine room	The operational layer of planning where notes, comparisons, drafts, translations and handovers are produced.
managed attention	An institutional approach that directs scarce professional attention toward interpretation, exceptions and consequences after routine production becomes cheaper.
meaningful participation	Participation that begins while consequential choices remain open and gives people a traceable route to influence the decision.
metrics trap	The pattern in which quantitative scores crowd out the values they were meant to represent, especially when models optimise toward measurable proxies and lose sight of public purpose.
model boundary	The explicit limit of what a model represents, knows or may decide. Work outside that boundary remains subject to other evidence and human judgement.
NLP	Natural language processing methods classify, compare, extract or summarise patterns in written or spoken language.
observed exposure	The share of work in which people are actually observed using AI. The larger technical exposure measure covers tasks a model could theoretically perform.
open parameter register	An accessible record of the weights, thresholds and objective functions used in a model, including who approved them and when they may change.
oracle problem	The use of authoritative-looking AI output without visibility of the assumptions and value choices that shaped it.
Pareto-efficient	A state in which one objective or party cannot be improved without making another worse off. It identifies a trade-off boundary and does not decide which point on that boundary is publicly desirable.
project memory room	A governed, queryable record of decisions, assumptions, alternatives, commitments and unresolved tensions across the life of a project.
project twin	A twin used by a project team and its partners to compare scenarios, expose dependencies and support a shared decision process.
public purpose	The democratic mandate and public values that direct what a planning institution should protect, pursue and explain.

public value	The benefit a public institution is expected to create and protect through legitimate action, including outcomes that cannot be reduced to a single metric.
RAG	Retrieval-augmented generation retrieves material from a selected document collection before a language model produces an answer. It can improve traceability when the sources, retrieval rules and citations remain inspectable.
recognition and representation	Recognition concerns whose identity, experience and claims are treated as legitimate. Representation concerns who is present, who speaks and whose interests carry standing in the process.
residual land value	The amount left for land after expected revenues are reduced by construction costs, finance, risk, fees and other development costs. Its apparent precision rests on assumptions that can change the result materially.
SECI model	Nonaka and Takeuchi's model of knowledge creation through socialisation, externalisation, combination and internalisation. It describes movement between tacit and explicit knowledge; simple document storage captures only part of that process.
sensitivity analysis	Testing how an outcome changes when influential assumptions or input values change within credible ranges.
simulated publics	AI-generated proxies for resident views, useful for rehearsing arguments and surfacing concerns, but without democratic standing or lived stake.
tacit versus explicit knowledge	Tacit knowledge is embodied in experience, judgement and relationships; explicit knowledge has been articulated in words, diagrams, data or procedures. Projects need both, and lose capacity when they confuse documentation with complete understanding.
technical twin	A digital twin used mainly by technical specialists to inspect system behaviour, data or design performance.
traceability	The ability to reconstruct which evidence, assumptions, changes and approvals produced an output or decision.
transaction cost	The effort required to find information, coordinate actors, negotiate terms, verify performance and manage an exchange or decision process.
verification routine	A repeatable procedure for checking sources, calculations, assumptions and consequential claims before AI-assisted work is used.

Index

Page references lead to the first substantial discussion of each term.

A

Adade	81
adaptation tipping point	53
adoption gap	10
AI	10
AI penalty	29
AI-assisted participation	80
Aldasoro	32
algorithmic governance	20
algorithmic urbanism	20
APA	102
Arnstein	78
assumption register	56
authoritative picture	67
Awashra	85

B

Batty	62
Becker	29
Bedard	93
Blue	78
brain fry	93
Brynjolfsson	33

C

capability overhang	11
Cardullo	69
Caulfield	22
citizenship from above	69
civic decision room	104
civic twin	65
Coase	20
Coasean shift	93
Coasean Singularity	93

cognitive offloading	15
cognitive surrender	94
communicative action (Habermas)	19
constraint map	55
contestability	43
controlled imperfect principle	85
Cugurullo	102
D	
DAPP (dynamic adaptive policy pathways)	53
data governance	11
de Vries	81
decision package	55
Dell'Acqua	12
digital model	63
digital shadow	63
digital twin	63
doctrine of use	105
E	
EU AI Act	105
F	
false precision	51
five levels of delegation (Lo to L4)	91
G	
Gastil	86
Goldsmith	105
Goodhart's Law	55
H	
Habermas	19
Healey	19
I	
initiative versus management (Coasean)	23
institutional memory	38
J	
jagged frontier	12
Jungherr	84

K

Kitchin	69
Klesel	41
knowledge graph	41

L

ladder of participation (Arnstein)	78
LLM	13

M

machine room	28
Malekzadeh	23
managed attention	69
Mashhadi Moghaddam	54
Massenkoff	11
McCrory	11
meaningful participation	78
model boundary	64
Moroni	67

N

NLP	13
Nonaka	37

O

observed exposure	11
OECD	80
open parameter register	95
option space	49
oracle problem	102

P

Pareto-efficient	49
Poudel	106
project memory room	37
project twin	65
public purpose	102
public value	54

R

RAG	41
---------------	----

Rauchfleisch	84
Raymond	82
recognition and representation	78
residual land value	48
Rountree	86
S	
Sanchez	103
SECI model	37
sensitivity analysis	51
Shahidi	93
Si	12
simulated publics	87
Stanford HAI	32
T	
tacit versus explicit knowledge	42
Taeihagh	86
technical twin	65
traceability	105
transaction cost	20
U	
UN-Habitat	71
V	
verification routine	17
W	
Wang	86
Weil	70
Wittmann	41
X	
Xu	102
Y	
Yotzov	32

Editorial and AI Note

This collection was written, edited and designed as a human-led argument supported by AI-assisted workflows.

The essays in this collection were originally published as a series and have been edited into one continuous argument. Repeated series navigation, subscription messages, article-level disclosures and duplicated introductory material have been removed. Short transitions and the preface were written for this edition. The original arguments, examples and source qualifications have been retained where relevant.

AI supported research, drafting, critical review, editorial consolidation, layout production and the creation of the original illustrations. The ideas, choices and conclusions remain the author's. The illustrations were generated as a coherent editorial series using an image supplied by the author as a visual reference for palette, texture and atmosphere, not as a composition to reproduce.

Remco Deelstra

June 2026

About the Author

Remco Deelstra is a strategist in housing and urban development at the City of Leeuwarden, the Netherlands. His work sits at the intersection of long-term urban change, public institutions and the practical choices through which policy becomes place.

This collection grew from an essay series written for urban planning professionals in a Dutch and European context. The series is grounded in a founding essay, *The Urban Operating System* (March 2026), which set out the framework that the nine chapters develop in turn. The essays began as an attempt to understand a practical gap: tools were becoming capable faster than planning organisations were becoming ready to use them. The collection follows that gap from everyday knowledge work to project memory, modelling, participation and public purpose.

He publishes on the Cat with Hat Substack (<https://catwithhat.substack.com>) and on LinkedIn.

Colophon

First edition, June 2026

Text, editing and art direction: Remco Deelstra

Original editorial illustrations developed with AI-assisted image generation

Copyright © 2026 R. Deelstra. Licensed under Creative Commons Attribution 4.0 International (CC BY 4.0), unless otherwise noted.

You may share and adapt this publication with attribution. Third-party source material, quoted excerpts and externally referenced works remain subject to their own rights and licences.

Suggested attribution: Remco Deelstra, *The Urban Professional in the Age of AI*, 2026.