

Human-Centered Artificial Intelligence in Workforce Planning for Project-Based Operations: A Conceptual Framework

Alie Fikry

School of Business & Management, Lincoln University College,
Malaysia

Abstract: Workforce demand in project-based industries fluctuates with project awards and client schedules, and organizations increasingly adopt artificial intelligence (AI) for demand forecasting, scheduling, and skills analysis. Research on human-centered and ethical scheduling remains scattered across individual methods and planning stages. The study found no framework that organizes these commitments across the strategic, tactical, and operational horizons through which project-based organizations plan their workforces. This paper addresses this gap through a structured scoping review. A Boolean search of the Semantic Scholar Academic Graph, combined with backward and forward citation tracing of ten seed studies, yielded 804 unique records, of which 53 met the inclusion criteria and venue-indexing filter. Conceptual synthesis followed three steps. First, the three planning horizons were identified from the workforce planning literature. Second, four recurring design principles were extracted from the human-centered AI (HCAI) and algorithmic management literatures. A principle was retained only when it was anchored in sociotechnical design theory and had a counterpart in an established AI-governance requirement. Third, the two dimensions were crossed into a 3×4 design space. The four principles are augmentation over automation, meaningful human control, fairness and transparency, and capability development. Each of the twelve resulting cells states one design commitment and serves as a pre-deployment checklist item. Seven propositions were then derived from the cells, with each proposition specifying a unit of analysis, a design condition, a comparison condition, an outcome variable, and a mechanism. AI-enabled workforce planning is theorized as a sociotechnical dynamic capability whose microfoundations are represented by the four principles. The framework and propositions are conceptual and remain empirically untested. Proposition 7, concerning productivity compression between novice and experienced workers, extends evidence from a customer-service setting and is advanced as an exploratory boundary condition. Validation is proposed through expert Delphi panels and multi-project field studies.

Keywords: human-centered artificial intelligence, workforce planning, project-based organizations, dynamic capabilities, sociotechnical systems, shared prosperity.

Introduction

Organizations now plan, deploy, and develop their workforce through digital systems. Artificial intelligence (AI) forecasts labor demand, optimizes schedules, and identifies skill gaps at a speed that exceeds manual planning (Verma et al., 2025; Egasmara et al., 2025). These capabilities matter most in project-based industries, where workforce demand fluctuates with project awards, sequencing decisions, and client schedules. Energy services, construction, and large-scale engineering all plan their workforces under this volatility (Biruk et al., 2022). Research on AI-enabled workforce planning has expanded, and its emphasis has moved toward human-centered and ethical concerns. A review of 50 studies found an early emphasis on efficiency outcomes such as forecasting accuracy, scheduling optimization, and cost reduction (Egasmara et al., 2025). Operations research has since reoriented toward human-centered and ethical scheduling. (Burgert et al., 2024) reviewed workforce scheduling for human-centered algorithmic management in manufacturing and proposed a conceptual optimization model. AI adoption at work also produces measurable effects on worker health and well-being (Valtonen et al., 2025). These worker-centric advances remain fragmented across separate methods and planning stages, with no integrative design framework.

Human-centered AI (HCAI) supplies the missing design vocabulary. (Shneiderman., 2020a) shows that designers can pursue high automation and high human control together. Related work specifies what this compatibility demands from designers, including reliability, transparency, and support for human oversight (Shneiderman, 2020b; Xu et al., 2022). Knowledge-management research reaches a parallel conclusion: organizations capture the most value when they treat automation and augmentation as a managed paradox rather than a choice (Raisch & Krakowski, 2021; El-Farr & Kertechian, 2024). Project management research addresses these requirements as a question of organizational readiness. That work identifies the preconditions for adopting AI (Klos et al., 2025), while the framework developed here specifies design commitments at each planning horizon. Few frameworks translate HCAI principles into the specifics of workforce planning in project-based operations (Micheli et al., 2023). Workforce planning decisions allocate people to projects, and these assignments determine incomes, skill trajectories, and career progression. When algorithms mediate such decisions without human-centered design, they intensify the contested control dynamics that algorithmic management research documents (Uddin, 2026). When organizations configure these systems around augmentation, oversight, and fairness, AI-enabled planning can

broaden access to developmental opportunities. The distinction matters for shared prosperity, because workforce planning is one of the few managerial systems that distribute work, income, and learning. This paper develops a conceptual framework that maps four HCAI design principles onto three workforce-planning horizons, from which we derive seven testable propositions. Sections 1.1 to 1.3 review the three source literatures, and Section 1.4 states the gap, the novelty claim, and the intended contribution. Section 2 details the framework development method. Section 3 presents the framework, the propositions, and their theoretical, practical, and policy implications, and Section 3.9 states the limitations. Section 4 concludes with a validation agenda.

Workforce Planning in Project-Based Operations

Workforce planning aligns labor supply with organizational demand across time horizons. Strategic human resource management theory treats this alignment as a source of competitive advantage when the resulting human capital is valuable and hard to imitate. Dynamic capability theory reframes such flexibility as an organizational ability to reconfigure resources as environments shift (Teece et al., 1997). Under this lens, workforce planning is a continuous sensing, seizing, and reconfiguring routine that operates across strategic, tactical, and operational horizons. Project-based operations strain conventional planning assumptions. Demand arrives in discrete projects with uncertain timing, and organizations assemble, deploy, and release project teams with each cycle. Researchers have not settled how organizations should embed digital technologies, and AI in particular, in that routine.

Artificial Intelligence in Workforce Planning

AI applications span the workforce planning cycle. (Egasmara et al., 2025) identified five dominant application themes across 50 reviewed studies: demand forecasting, workforce scheduling, skill gap analysis, recruitment, and retention analytics. Across these themes, AI improves predictive accuracy, resource allocation, and organizational flexibility. A multidisciplinary agenda for organizational AI reaches a broader conclusion: the central challenges are as much social and governance challenges as technical ones (Dwivedi et al., 2019). The same literature documents persistent human-centered problems. Research on the ethics of AI-enabled recruiting and selection identifies fairness and accountability as core challenges, alongside legal and reputational exposure (Hunkenschroer & Luetge, 2022). Algorithmic management extends managerial control through direction, evaluation, and discipline, and workers contest that control when the underlying criteria stay opaque (Kellogg et al., 2020; Uddin, 2026). The operations research community has begun to answer these concerns. (Burgert et al., 2024) reviewed workforce scheduling approaches and proposed a conceptual optimization model that embeds human-centered work-design criteria in the

scheduling objective. Human-in-the-loop and ethically aligned scheduling now form an active research front. Burgert et al. stop short of a design framework that organizes such commitments across strategic, tactical, and operational horizons. The unresolved problem is integration: the literature supplies human-centered methods for single planning stages but no framework that connects them into one design logic.

Human-Centered AI and Sociotechnical Systems

HCAI supplies design guidance at the general level. (Shneiderman.,2020b) rejected the assumption that more automation requires less human control and proposed a two-dimensional design space in which reliable, safe, and trustworthy systems combine high automation with high human control. Figure 1 illustrates this design space. The traditional view treats automation and control as a single trade-off along the diagonal; HCAI separates them into independent axes and makes the upper-right zone of high automation and high human control reachable (Shneiderman., 2020b).

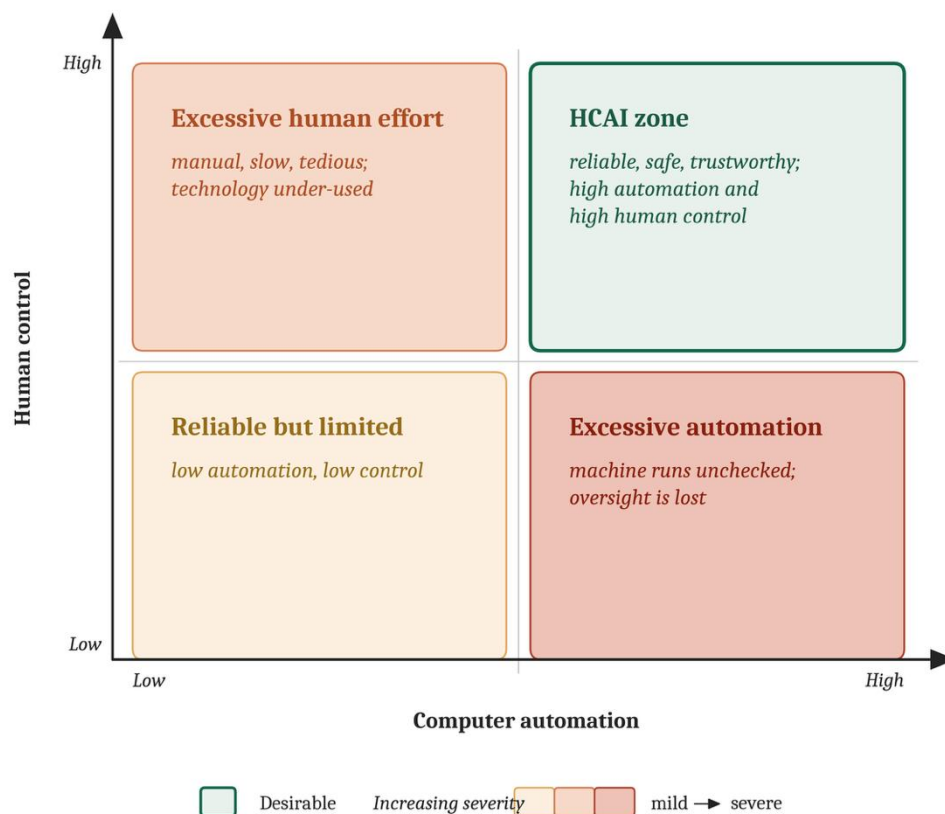


Figure 1 The two-dimensional human-centered AI (HCAI) design space. High computer automation and high human control together define the HCAI zone, where systems are reliable, safe, and trustworthy, in contrast to designs that over-rely on human effort or on unchecked automation. Adapted from (Shneiderman., 2020b).

Recent work specifies the transition challenges organizations face when they move from conventional software to AI systems (Klos et al., 2025). Knowledge-management research applies the same insight: firms that pursue automation alone trigger reinforcing problems, while firms that manage automation and augmentation together capture complementary benefits (El-Farr & Kertechian, 2024).

The sociotechnical tradition grounds these design arguments in organizational theory. (Clegg, 2000) synthesized this tradition into a set of design principles for computer-based systems, organized as meta, content, and process principles. That synthesis carries the core sociotechnical commitments: designers must optimize technical and social systems together, because optimizing the technical system alone degrades work relationships and performance; decisions should rest with human judgment wherever feasible, an idea captured in minimal critical specification; variances should be controlled at their source; and support systems should stay congruent with intended behavior. (Parker and Grote, 2022) extended the tradition to algorithmic work and showed that work design matters more, not less, as automation deepens. Recent scholarship revives this lens for AI. (Kudina and van de, 2024) analyze AI as a sociotechnical system in which technology, human behavior, and institutional rules co-produce outcomes. This perspective keeps worker development in view when algorithmic systems mediate work. Framing AI-enabled planning as a dynamic capability requires attention to its microfoundations. (Teece, 2007) argued that dynamic capabilities rest on identifiable microfoundations: the specific skills, processes, and routines through which sensing, seizing, and reconfiguring occur. This paper treats the four design principles as such microfoundations. Augmentation, human control, fairness, and capability development are the routines through which organizations enact the planning capability. The next section specifies the four microfoundations at the level where planners decide, across three horizons.

Gap, Novelty, and Contribution

Three limits recur across the reviewed work, and together they define the gap this paper addresses. First, existing guidance is stage-specific. (Burgert et al., 2024) embed human-centered work-design criteria in a scheduling objective, and (Hunkenschroer and Luetge, 2022) examine fairness and accountability in recruiting and selection. Each addresses one stage of the planning cycle. Second, existing guidance is method-specific. Optimization models, human-computer interaction design guidance, and ethics reviews each address a different class of artifact, and none states what a planning system must satisfy irrespective of the method it uses. Third, the workforce planning literature that does distinguish planning horizons carries no AI design content. (Micheli et al., 2023) set out strategic, tactical, and

operational planning for project-driven companies without addressing algorithmic mediation, whereas the HCAI literature specifies design properties without a horizon structure (Shneiderman, 2020b; Xu et al., 2022). The scoping search sharpened this gap rather than removing it. Four retrieved studies come closer than the rest. (Latos et al., 2024) derive human-centered criteria for intelligent personnel planning and treat time autonomy as a design requirement. (Gerlach et al., 2025) examine nurse perspectives on AI-based shift scheduling in terms of fairness, transparency, and work-life balance. (Gattermann-Itschert et al., 2023) embed planners' revealed preferences inside a railway crew scheduling optimizer. (Zhang et al., 2025) combine sociotechnical systems theory with strategic human resource development to bound algorithmic management. Each None of the four spans the design space. (Latos et al., 2024) cover two horizons but one design criterion. (Gerlach et al., 2025 ; Renggli et al., 2025) cover two principles at one horizon in one occupational setting. Gattermann-Itschert et al. treat human preference as an input to an optimizer rather than as a design commitment. (Zhang et al., 2025) bound algorithmic management without stating commitments by horizon. No retrieved source organizes augmentation, oversight, fairness, and development as one design logic that spans strategic, tactical, and operational planning. Table 1 records this comparison.

We do not select the four principles to cover the HCAI literature as a whole. We select them because each governs a decision that workforce planners in project-based organizations make. Augmentation over automation governs which planning tasks the system performs and which the planner performs. Meaningful human control governs who holds authority over an assignment and who answers for it. Fairness and transparency govern the criteria by which people are allocated to work. Capability development governs whether an assignment builds the skill base that it consumes. Section 2.3 states the two tests that each principle had to pass, and Section 2.4 states the boundaries between them.

The contribution differs from each source literature in a specific way. Against the HCAI literature, the framework replaces domain-general design properties with twelve horizon-specific design commitments that a planner can check before deployment. Against the workforce planning literature, it supplies the design content that horizon models omit, and it states which human-centered commitment applies at which horizon. Against the sociotechnical tradition, it treats general design principles as microfoundations of a named organizational capability and attaches a predicted, observable effect to each. The propositions carry this contribution into testable form. We claim originality for the mapping, for the derivation, and for the propositions, and not for the constituent principles or horizons, which the source literatures already establish.

Table 1 Positioning of the present framework relative to closely related studies, reviews, and conceptual models.

Study	AI focus	Workforce-planning focus	Human-centered dimension	Planning horizon addressed	Limitation for the present purpose
Egasmara et al. (2025)	Forecasting, scheduling, skill-gap analysis, recruitment, retention analytics	Full planning cycle	Identifies a human-centered shift as a research agenda	Not differentiated	Reviews applications and supplies no design framework
Burgert et al. (2024)	Optimization models for workforce scheduling	Scheduling	Work-design criteria embedded in the scheduling objective	Operational	One horizon and one method class
Micheli et al. (2023)	None	Workforce planning in project-driven companies	Not addressed	Strategic, tactical, operational	Horizon structure without AI design content
Shneiderman (2020b)	General AI system design	None	Automation and human control as independent axes	Not applicable	Domain-general; no planning specificity
Xu et al. (2022)	General AI system design	None	Opacity and loss of situational awareness	Not applicable	No workforce planning application
Hunkenschroer and Luetge (2022)	AI-enabled recruiting and selection	Recruitment	Fairness and accountability	Stage-specific (hiring)	Ethics review rather than a design framework
Kellogg et al. (2020)	Algorithmic management	Direction, evaluation, discipline	Contested control	Cross-cutting	Explains control dynamics and prescribes no design
Raisch and Krakowski (2021); El-Farr and Kertechian (2024)	Automation and augmentation as a managed paradox	None specific	Augmentation	Not applicable	Paradox-level theory without a horizon structure
Parker and Grote (2022)	Algorithmic work systems	Job and work design	Autonomy and skill use	Not applicable	Work design outcomes rather than planning-system design
Klos et al. (2025)	Organizational AI adoption	None specific	Adoption readiness	Not applicable	Preconditions for adoption rather than post-deployment design
Kudina and van de Poel (2024)	AI as a sociotechnical system	None specific	Co-production of outcomes by technology, behavior, and rules	Not applicable	General sociotechnical account
Latos et al. (2024)	Intelligent personnel	Personnel planning and shift design	Time autonomy as a human-	Tactical and operational	Criteria for one planning technology;

Study	AI focus	Workforce-planning focus	Human-centered dimension	Planning horizon addressed	Limitation for the present purpose
	planning and scheduling		centered design criterion		no horizon-spanning design logic
Gerlach et al. (2025); Renggli et al. (2025)	AI-based shift scheduling	Nurse rostering	Fairness, transparency, preferences, work-life balance	Operational	Two principles, one horizon, one occupational setting
Gattermann-Itschert et al. (2023)	Machine learning inside a crew scheduling optimizer	Railway crew scheduling	Planners' revealed preferences as an optimization input	Operational	Human input treated as model data rather than as a design commitment
Koruca et al. (2023)	Personalized staff-scheduling algorithm	Hospital staff scheduling	Work-life balance as an explicit objective	Operational	One objective at one horizon
Hickman et al. (2025)	Machine learning in personnel assessment	Selection and assessment	Algorithmic bias mitigation	Stage-specific (selection)	Bias-mitigation methods; no planning-horizon structure
Zhang et al. (2025)	Algorithmic management	Strategic human resource development	Sociotechnical limits of algorithmic control	Cross-cutting	Bounds algorithmic management; states no design commitments by horizon
This paper	Design commitments for AI-enabled planning systems	Strategic demand planning, tactical workforce composition, operational deployment	Four HCAI principles as microfoundations	All three horizons	Conceptual; framework and propositions untested

Research Method

Literature Scoping Procedure

The synthesis rests on a structured scoping search rather than an exhaustive systematic review, because a scoping search suits a conceptual contribution. We searched the Semantic Scholar Academic Graph through its public application programming interface on 30 August 2026, and we complemented that search with backward and forward citation tracing. The query combined four concept blocks with Boolean operators: a workforce block ("workforce" OR "personnel" OR "crew" OR "staff"), a planning block ("planning" OR "scheduling" OR "allocation" OR "deployment"), an AI block ("artificial intelligence" OR "machine learning" OR "algorithmic management"), and a human-centered block ("human-centered" OR "human-centred" OR "augmentation" OR "trust" OR "fairness" OR "autonomy"). All four blocks had to match in the title or the abstract. The review period ran from 2014 to 2025. The database branch returned 607 records. Citation tracing on ten seed studies, chosen as the

retrieved records closest to the review question, returned a further 248 items, of which 202 resolved to unique records not already retrieved.

Four conditions governed inclusion. A record had to appear in a peer-reviewed journal or in conference proceedings, and to carry a publication date between 2014 and 2025. It had to treat an AI, machine learning, or optimization method applied to workforce planning, scheduling, or allocation. It also had to treat a human, ethical, or governance dimension of that method. We excluded records that reported algorithmic performance alone, records confined to consumer or patient scheduling rather than the scheduling of workers, records not available in English, records without a retrievable abstract, and records carrying a retraction notice. The search returned 32 retracted articles, all of which we excluded on that ground.

Screening proceeded in two rounds, first on title and abstract and then on the full record. One author screened every record. The single-author design makes a second independent screener unavailable, and we mitigate the resulting risk by reporting the criteria, the stage counts, and the retained set in full, so that a reader can audit any individual decision. Eligibility screening retained 79 records. The Semantic Scholar corpus indexes venues of uneven editorial quality, so we then applied a venue-indexing filter and retained a record only when its source journal appears on the CWTS Leiden Ranking core-source list as reported by OpenAlex. That filter removed 20 records. It also produced one false negative, *Oeconomia Copernicana*, which we retained after confirming Scopus and Social Sciences Citation Index coverage from the publisher. Four records sit in conference proceedings, and one sits in a journal that the open indexing sources do not resolve. No determination is available for these five, and we hold them pending confirmation against the Scopus source list. We excluded one further record at metadata verification because its issue date falls in 2026. The synthesis therefore rests on 53 sources. Table 2 reports the count at each stage, which makes the search auditable without claiming the coverage of a systematic review; Section 3.9 states the limitation that follows.

Table 2 Literature scoping: records identified, screened, excluded, and retained, by search branch. The search was executed on 30 August 2026.

Stage	Database search	Citation tracing	Total
Records identified	607	248	855
Duplicate and unresolvable records removed	5	46	51
Unique records	602	202	804
Excluded before screening (publication type, language, review period, no retrievable abstract)	79	112	191
Records screened	523	90	613
Excluded at screening, including 32 records carrying a retraction notice	395	56	451
Assessed at full-record level	128	Not applicable	128

Stage	Database search	Citation tracing	Total
Excluded at full-record level	83	Not applicable	83
Retained after eligibility screening	45	34	79
Excluded by the venue-indexing filter	19	1	20
Excluded at metadata verification	0	1	1
Held pending confirmation against the Scopus source list	4	1	5
Sources included in the synthesis	22	31	53

Two findings from this scoping motivate the framework. First, human-centered and ethical scheduling is now well established; dedicated reviews and optimization models attest to it (Burgert et al., 2024). The earlier characterization of the field as efficiency-only no longer holds. Second, these contributions remain method-specific and stage-specific. No retrieved source organized the four commitments as a single design logic spanning all three horizons. This second finding grounds the claim of underdeveloped integration, and Table 1 in Section 1.4 records the comparison study by study. Most of the retained sources address operational deployment, and rostering and shift-scheduling studies in healthcare account for a large share of them, whereas strategic demand planning and tactical workforce composition attract markedly fewer studies. That imbalance is itself evidence for the integration gap the framework addresses.

Synthesis Logic

The study developed the framework through structured conceptual synthesis. Figure 2 summarizes the three-step logic. First, we derived the planning dimension from the workforce planning literature, which distinguishes strategic, tactical, and operational horizons in project-based settings (Teece et al., 1997; Micheli et al., 2023). Strategic demand planning covers multi-quarter headcount, sourcing, and capacity decisions taken before projects are confirmed. Tactical workforce composition covers project-level staffing and competency mix, decided once a project is awarded and before mobilization. Operational deployment covers short-cycle rostering, reassignment, and shift-level allocation during execution (Micheli et al., 2023). The study uses these three labels, and no synonyms for them, throughout the paper.

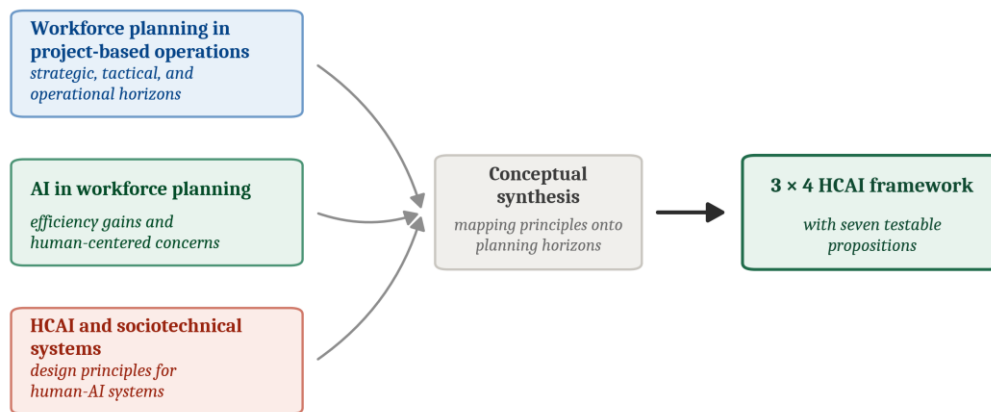


Figure 2 Framework development through conceptual synthesis of three source literatures.

Second, the study derived the design dimensions by extracting recurring prescriptions from the HCAI and algorithmic management literatures. Four principles recurred across sources: augmentation over automation (Jarrahi, 2018; El-Farr & Kertechian, 2024), meaningful human control (Santoni de Sio & van den Hoven, 2018; Shneiderman, 2020b; Xu et al., 2022), fairness and transparency (Hunkenschroer., 2022 & Luetge., 2022; Uddin, 2026), and capability development (Kudina & van de Poel, 2024; Valtonen et al., 2025). Third, we mapped each principle against each horizon and formulated propositions where the intersection generates a distinct, testable expectation. The study chose conceptual synthesis over empirical induction because the target phenomenon, HCAI-designed planning systems, remains rare in practice and cannot yet support large-sample studies. This choice limits the framework to analytical rather than statistical generalization; Section 3.9 returns to this limitation.

2.3 Justification for the Four Design Principles

A conceptual framework must justify why it selects the dimensions it does. Two tests governed selection here. The study retained a principle only when it had an anchor in classical sociotechnical design theory and a counterpart in an established AI-governance requirement. The sociotechnical anchor guards against ad hoc invention. (Clegg., 2000) supplies it: minimal critical specification implies leaving decisions to human judgment where possible, which motivates augmentation over automation; control of variance at its source motivates meaningful human control; the human-values principle motivates fairness and transparency; and the treatment of design as continuous and incomplete motivates capability development.

The external counterpart guards against parochialism. (The European Commission High-Level Expert Group on Artificial Intelligence., 2019) defined seven requirements for trustworthy AI. The four principles map onto the subset of those requirements that a planning

system can design for, while the remaining requirements sit at the organizational or infrastructural level. Two principles, augmentation and human control, operationalize different facets of one requirement, human agency and oversight; the external test therefore confirms relevance to trustworthy-AI goals rather than establishing four independent requirements. Table 3 records both mappings. We kept the set to four for parsimony, because a design checklist with too many dimensions is not usable at the point of decision. Requirements such as technical robustness and privacy are essential but belong to systems engineering and data governance, not to the planning-design choices modeled here.

Table 3 Derivation of the four design principles from sociotechnical design theory (Clegg, 2000) and the EU trustworthy AI requirements (European Commission High-Level Expert Group, 2019).

Design principle	Sociotechnical anchor (Clegg, 2000)	EU trustworthy AI requirement (2019)
Augmentation over automation	Minimal critical specification	Human agency and oversight
Meaningful human control	Control of variance at source	Human agency and oversight
Fairness and transparency	Human values; congruence of support systems	Transparency; diversity, non-discrimination and fairness
Capability development	Continuous, incomplete design	Societal and environmental well-being

Conceptual Boundaries Among the Four Design Principles

Two pairs of principles sit close enough to require an explicit boundary. Augmentation over automation and meaningful human control both concern human agency, and Section 2.3 records that both map onto one EU requirement. The two principles govern different objects. Augmentation governs the division of cognitive labor: it fixes which tasks the system performs and which the planner performs, and designers set it before deployment. Meaningful human control governs authority over the resulting decision: it fixes who can interrogate an output, who can override it, on what grounds a worker can contest it, and to whom the outcome traces (Santoni de Sio & van den Hoven, 2018). A system can therefore augment and still escape control, when planners hold the decision but cannot inspect the reasoning behind a recommendation. A system can also stay under control and still automate, when planners may override outputs but the system now performs the judgment tasks that planners once performed. Fairness and transparency form the second pair, and this paper treats them as one principle for a stated reason. Fairness is a property of the allocation rule, and it concerns the distribution of assignments, hours, and development exposure that the criteria produce. Transparency is a property of the criteria as knowledge, and it concerns whether affected

workers can see, understand, and contest them. In workforce planning the two fail together. A fairness rule that workers cannot inspect does not reduce contestation, because workers judge the outcomes they observe rather than the rule they cannot see (Kellogg et al., 2020; Uddin, 2026). Visible criteria that encode an unfair rule make the unfairness legible without correcting it, and they expose the organization to the legal and reputational consequences that (Hunkenschroer and Luetge ., 2022) identify. We therefore retain the pair as a single principle with two design objects. Table 4 records the boundaries across the full set.

Table 4 Conceptual boundaries among the four design principles.

Design principle	Governing question	Design object	Failure mode when the principle is absent	Boundary with the adjacent principle
Augmentation over automation	Which planning tasks does the system perform, and which does the planner perform?	The division of cognitive labor, fixed at design time	Skill erosion among the planners who must handle exceptions (El-Farr & Kertechian, 2024)	Fixes the allocation of tasks; it does not fix who holds authority over the resulting decision
Meaningful human control	Who decides, who can override, on what grounds can a worker contest, and to whom does the outcome trace?	Interrogation, override, contestation, and audit mechanisms attached to each decision	Loss of situational awareness and decisions that no person can be held to (Xu et al., 2022)	Fixes authority over the decision; it does not fix which party performed the underlying task
Fairness and transparency	On what criteria are people allocated to work, and can affected workers see those criteria?	The allocation rule, and the visibility of that rule to affected workers	Contested control, grievances, and legal exposure (Hunkenschroer & Luetge, 2022; Uddin, 2026)	Fixes the criteria and their visibility; it does not fix who applies them or how tasks are divided
Capability development	Does the assignment build the skill base that it consumes?	Development exposure as an explicit objective in the planning model	Depletion of the skill portfolio on which later planning cycles depend	Fixes the effect of an assignment on future capability; the other three fix the design of the present decision

Derivation of the Propositions

The twelve cells state design commitments. They do not by themselves make predictions, and the step from a cell to a proposition requires a stated rule. We formulated a proposition only where a cell, or a row of cells sharing one principle, satisfied three conditions. First, the cell contrasts two design options that a planner can choose between, so that the proposition carries a comparison condition rather than an ideal. Second, the predicted outcome can be measured in a project-based organization, either from records the organization already holds, such as revision logs, override records, grievance records, and mobility records, or from a short survey of planners and affected workers. Third, a mechanism named in the reviewed literature connects the design condition to the outcome. The seven propositions therefore do not stand in one-to-one correspondence with the twelve cells. Cells that met the first condition but not the third yielded design guidance without a proposition.

Five propositions sit at the principle level, where the prediction holds across horizons. Proposition 4 is the exception within that group, since it contrasts two horizons directly. Proposition 6 sits at the framework level, because its condition is the sustained use of the whole framework rather than any single principle. Proposition 7 sits at the workforce level, because its outcome is distributional across worker cohorts rather than a property of the planning system. Section 3.5 closes with a summary table recording the unit of analysis, the design condition, the comparison condition, the outcome variable, the mechanism, and the evidential status of each proposition. We state each proposition as an expectation to be tested and not as a finding, and that table marks which propositions rest on mechanisms documented in prior studies and which rest on conceptual conjecture.

Results and Discussion

Framework Structure

The framework crosses four HCAI design principles with three planning horizons, producing a 3×4 design space. Figure 3 presents the full space, and each cell states one design commitment for that principle at that horizon. Table 5 records how each cell was derived, naming the design commitment, the theoretical rationale, and the source for each. We state the propositions that follow at the principle level, where they hold across horizons. Proposition 4 is the exception, since it contrasts the operational and strategic horizons. The twelve cells function as a pre-deployment design checklist. Horizon-differentiated prediction remains part of the validation agenda, and Proposition 4 is the only proposition that addresses it.

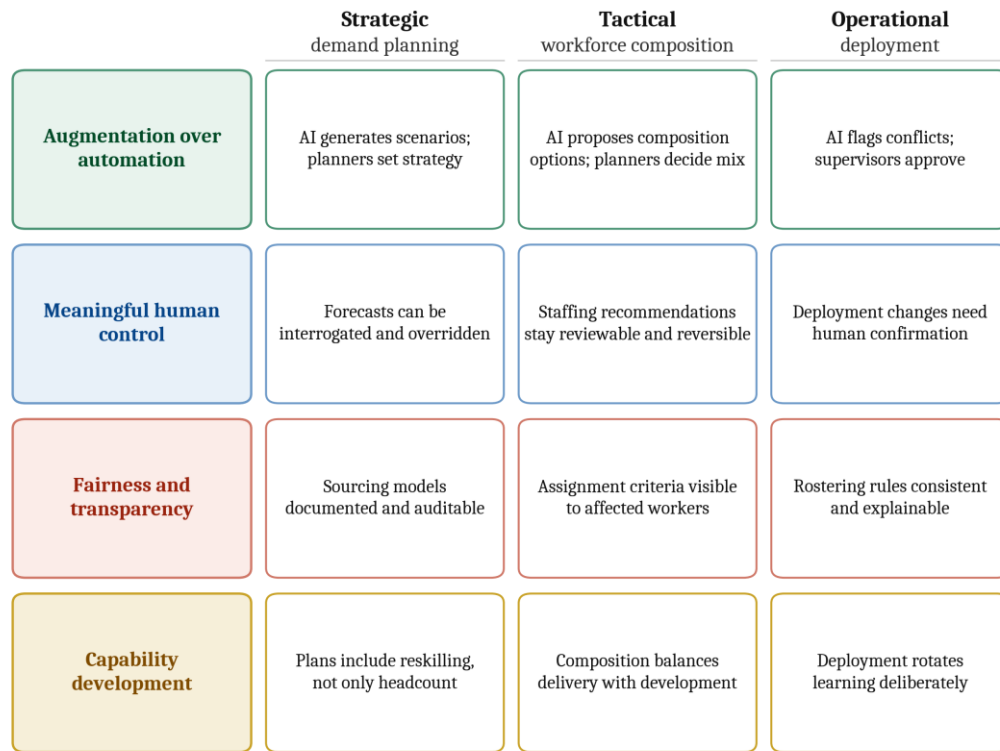


Figure 3 The HCAI workforce planning framework: four design principles across three planning horizons.

Figure 3 HCAI workforce planning framework across three planning horizons and four design principles. Each combination of principle and horizon translates into a distinct planning cell, which further elaborates the framework. These twelve cells are summarized in Table 5 and the design commitment, theoretical rationale, and source behind each are displayed.

Table 5 Derivation of the twelve framework cells, with the design commitment, the theoretical rationale, and the source for each.

Cell (principle × horizon)	Design commitment	Theoretical rationale	Source
Augmentation × strategic demand planning	AI generates demand scenarios; planners set the strategy	Analytical breadth suits AI, whereas reading a client's real schedule intent is equivocal and suits human judgment	Jarrahi (2018); Raisch and Krakowski (2021)
Augmentation × tactical workforce composition	AI proposes composition options; planners decide the mix	Minimal critical specification leaves the composition decision to human judgment	Clegg (2000); Jarrahi (2018)
Augmentation × operational deployment	AI flags scheduling conflicts; supervisors approve the resolution	Exception handling preserves the judgment skills that full automation erodes	El-Farr and Kertechian (2024)
Meaningful human control × strategic demand planning	Forecasts can be interrogated and overridden	The tracking condition requires the system to respond to the reasons of those who deploy it	Santoni de Sio and van den Hoven (2018)

Cell (principle × horizon)	Design commitment	Theoretical rationale	Source
Meaningful human control × tactical workforce composition	Staffing recommendations stay reviewable and reversible	Reversibility keeps the composition decision attributable to a named human	Santoni de Sio and van den Hoven (2018); Shneiderman (2020b)
Meaningful human control × operational deployment	Deployment changes require human confirmation	Control of variance at source, where deployment errors cascade into project delay	Clegg (2000); Dwivedi et al. (2019)
Fairness and transparency × strategic demand planning	Sourcing models are documented and auditable	Accountability for allocation criteria set before individual workers are identifiable	Hunkenschroer and Luetge (2022)
Fairness and transparency × tactical workforce composition	Assignment criteria are visible to affected workers	Contested control declines when evaluative criteria become legible	Kellogg et al. (2020); Uddin (2026)
Fairness and transparency × operational deployment	Rostering rules are consistent and explainable	Rostering determines income for day-rate and rotational workers, which raises the fairness stakes	Uddin (2026); Parker and Grote (2022)
Capability development × strategic demand planning	Plans include reskilling and not only headcount	Design is continuous and incomplete, so the plan must renew the skill base it draws on	Clegg (2000); Teece (2007)
Capability development × tactical workforce composition	Composition balances delivery with development	Firms resolve the automation-augmentation paradox by developing worker knowledge alongside the technology	El-Farr and Kertechian (2024)
Capability development × operational deployment	Deployment rotates learning deliberately	Assignment is the mechanism through which learning opportunities are distributed across a workforce	Kudina and van de Poel (2024); Valtonen et al. (2025)

Table 5 also fixes the status of the cells. Each commitment states what one principle requires at one horizon, and each cites a source that supports the underlying mechanism. None of the twelve has been tested as a construct. We present them as design commitments and not as findings, and Section 3.9 returns to this limitation.

Augmentation Over Automation

The first principle assigns AI and humans complementary roles at every horizon. (Jarrahi, 2018) argued that AI outperforms humans on analytical breadth while humans retain advantages in intuition and equivocal judgment. Workforce planning contains both task types. Demand forecasting across many projects is analytically broad, whereas judging a client's real schedule intentions is equivocal. Organizations that pursue automation alone generate reinforcing problems, including skill erosion among the humans who must handle exceptions (El-Farr & Kertechian, 2024).

Proposition 1. In project-based planning units, planning systems designed for augmentation, in which AI generates ranked options and a named planner records the decision, will show

fewer within-cycle planning reversals and fewer post-hoc exception overrides than systems designed for end-to-end automation with exception-only human involvement. Planning reversals denote assignment decisions rescinded within the same planning cycle. The mechanism is complementarity: the design assigns analytically broad tasks to AI and equivocal judgment tasks to planners (Jarrahi., 2018), so equivocal cases are resolved before commitment rather than after it.

Forecast utilization, the share of AI-generated forecasts that planners adopt without material revision, presupposes that planners review forecasts. The measure is therefore undefined under end-to-end automation. We retain it as a within-condition diagnostic for augmentation designs and not as a comparative outcome across the two conditions.

Meaningful Human Control

The second principle operationalizes (Shneiderman's., 2020b) two-dimensional design space. The term originates in the ethics of autonomous systems, where control requires that a system respond to the moral reasons of the humans who deploy it and that its outcomes trace back to an identifiable human (Santoni de Sio & van den Hoven, 2018). High automation in forecasting and scheduling is desirable, but only when planners can interrogate, override, and audit system outputs. (Xu et al., 2022) identified opacity and loss of situational awareness as central risks when humans transition to AI-mediated work. In workforce planning, this loss is costly because deployment errors cascade into project delays and worker hardship (Dwivedi et al., 2019).

Proposition 2. In project-based planning units, planning systems that provide interrogation, override, and audit trails at all three horizons will sustain higher planner trust and shorter time to recovery after a planning error than systems that issue outputs without recourse. The mechanism is the pair of conditions that meaningful human control requires. The system responds to the reasons of the planners who deploy it, and each outcome traces to an identifiable human (Santoni de Sio & van den Hoven., 2018). Both conditions preserve the situational awareness that AI-mediated work otherwise erodes (Xu et al., 2022).

Fairness and Transparency

The third principle addresses the contested control dynamics that algorithmic management research documents. Workers resist algorithmic direction and evaluation when the underlying criteria stay opaque (Uddin., 2026). Fairness failures in HR algorithms also carry legal and reputational consequences beyond worker resistance (Hunkenschroer & Luetge., 2022). In project-based operations, assignment decisions determine income for day-rate and rotational workers, which raises the fairness stakes further. Fairness in scheduling carries a cost: an

allocation that satisfies an explicit fairness rule trades some efficiency for equity. Planners can design for that trade-off explicitly, and they can retain human oversight at the assignment decisions that most affect workers (Shneiderman., 2020b).

Proposition 3. In project-based planning units, making assignment criteria visible and auditable to affected workers, and setting an explicit fairness rule, will reduce grievance frequency and perceived assignment injustice relative to units that apply opaque, efficiency-only criteria. The mechanism is legibility: workers contest algorithmic direction when the criteria stay opaque, and contestation falls once the criteria become inspectable (Kellogg et al., 2020; Uddin, 2026).

Proposition 4. For the same planning unit and the same fairness rule, the association between transparency and trust will be stronger at the operational horizon than at the strategic horizon. The mechanism is proximity: operational decisions assign identifiable individuals to identifiable work within the current pay cycle, whereas strategic decisions set aggregate capacity before individual workers are identifiable.

Capability Development

The fourth principle connects planning design to shared prosperity. AI-enabled planning can build worker capability alongside short-term output. Skill gap analytics can identify development needs, and deployment algorithms can treat competency building as an assignment criterion in its own right. Knowledge-management research shows that firms resolve the automation-augmentation paradox by developing worker knowledge alongside the technology (El-Farr & Kertechian, 2024). Treated this way, capability development becomes a microfoundation of the planning capability (Teece, 2007).

Proposition 5. In project-based planning units, planning systems that include development exposure as an explicit optimization objective will improve internal mobility and twelve-month retention relative to systems that optimize utilization alone, while holding short-term utilization within its pre-deployment range. The mechanism is knowledge renewal: development exposure builds the worker knowledge that the technology alone cannot supply, which is how firms resolve the automation-augmentation paradox (El-Farr & Kertechian., 2024).

Proposition 6. Organizations that treat HCAI-designed planning as a dynamic capability, refined continuously across all three horizons, will outperform organizations that deploy equivalent AI tools as one-off automation projects on three dimensions: forecast-to-actual staffing accuracy, elapsed time to re-plan after a change in project schedule, and planner

retention. The mechanism is accumulation: repeated sensing, seizing, and reconfiguring cycles build routines that a single deployment cannot (Teece., 2007).

Capability development carries distributional consequences that the framework must account for. Evidence on generative AI at work shows that productivity gains concentrate among less experienced and lower-skilled workers, while the most experienced gain little (Brynjolfsson et al., 2025). That evidence comes from a customer-service setting. Its transfer to judgment-intensive project work is a conjecture we advance, not an established result. In project teams, this productivity compression would narrow the performance gap between novice and experienced workers. The effect can raise aggregate output and help newer workers. Yet it can also erode the relative value, status, and bargaining position of highly experienced workers. The system now partly encodes and redistributes their tacit expertise. Left unmanaged, this dynamic threatens the trust that the human-control and fairness principles protect. Designers must therefore build the capability-development principle to preserve the distinct value of experienced team members. For instance, deployment algorithms can route experienced workers toward judgment-intensive and mentoring roles that AI does not compress.

Proposition 7 (exploratory). Where AI compresses productivity differences between novice and experienced workers, planning systems that redeploy experienced workers into judgment-intensive and mentoring roles will sustain the trust and retention of those workers better than systems that treat all workers as interchangeable beneficiaries of AI assistance.

We label Proposition 7 exploratory for the reason stated above: its antecedent has not been observed in project-based work. Testing the proposition therefore requires establishing that antecedent first. A study would need to measure the productivity gain from AI assistance separately for novice and experienced planners or engineers in project settings, and to demonstrate compression, before the design prediction becomes meaningful. Should compression fail to occur in judgment-intensive project work, the proposition does not apply, and that result would identify the boundary of the distributional argument advanced here. Table 6 summarizes the structure and the evidential status of all seven propositions.

Table 6 Structure of the seven propositions, with the unit of analysis, design condition, comparison condition, outcome variable, mechanism, and evidential status of each.

Proposition	Unit of analysis	Design condition	Comparison condition	Outcome variable	Mechanism and evidential status
P1	Planning unit	Augmentation: AI generates ranked options, a named planner decides	End-to-end automation with exception-only human involvement	Within-cycle planning reversals; post-hoc exception overrides	Complementarity (Jarrahi, 2018; El-Farr & Kertechian, 2024). Mechanism documented; prediction untested

Proposition	Unit of analysis	Design condition	Comparison condition	Outcome variable	Mechanism and evidential status
P2	Planning unit	Interrogation, override, and audit trails at all three horizons	Outputs issued without recourse	Planner trust; time to recovery after a planning error	Tracking and tracing (Santoni de Sio & van den Hoven, 2018; Xu et al., 2022). Mechanism documented; prediction untested
P3	Planning unit	Assignment criteria visible and auditable; explicit fairness rule	Opaque, efficiency-only criteria	Grievance frequency; perceived assignment injustice	Legibility of evaluative criteria (Kellogg et al., 2020; Uddin, 2026). Mechanism documented; prediction untested
P4	Planning horizon within a unit	Transparency at the operational horizon	Transparency at the strategic horizon	Strength of the transparency-trust association	Proximity of the decision to the individual worker. Conceptual conjecture of the authors
P5	Planning unit	Development exposure as an explicit optimization objective	Utilization optimized alone	Internal mobility; twelve-month retention; short-term utilization held in range	Knowledge renewal (El-Farr & Kertechian, 2024). Mechanism documented; prediction untested
P6	Organization	HCAI-designed planning refined continuously across all three horizons	Equivalent AI tools deployed as one-off automation projects	Forecast-to-actual staffing accuracy; time to re-plan; planner retention	Accumulation of sensing, seizing, and reconfiguring routines (Teece, 2007). Mechanism documented; framework-level prediction untested
P7 (exploratory)	Worker cohorts within a workforce	Redeployment of experienced workers into judgment-intensive and mentoring roles	All workers treated as interchangeable beneficiaries of AI assistance	Trust and retention of experienced workers	Productivity compression, evidenced in a customer-service setting (Brynjolfsson et al., 2025). Antecedent unobserved in project work; exploratory conjecture

Figure 4 maps Propositions 1 to 5 onto their design principles and predicted planning effects. The map shows principle-level relationships only, and Propositions 6 and 7 sit outside it because they operate at different levels of analysis. Proposition 6 is framework-level: its condition is the sustained use of all four principles across all three horizons, so no single principle can anchor it. Proposition 7 is workforce-level: its outcome is the distribution of standing and retention across worker cohorts rather than an effect on the planning system. Placing either proposition inside the principle-level map would misstate what it predicts.

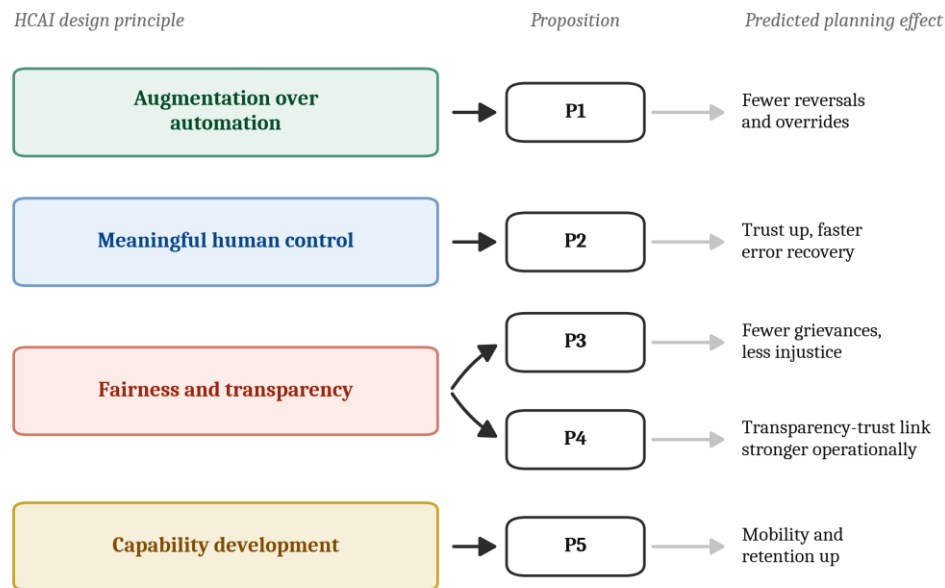


Figure 4 Proposition map linking each HCAI design principle to its predicted planning effect.

Figure 4 summarizes the proposed associations between the HCAI design principles and their expected impact on workforce planning outcomes. These propositions form the basis for a theoretical interpretation of the framework and for connecting the design principles to well-known ideas in the literature.

Theoretical Implications

The framework repositions AI-enabled workforce planning within dynamic capability theory at the microfoundational level. Prior work treats dynamic capabilities as the ability to reconfigure resources (Teece et al., 1997), and Teece (2007) locates that ability in specific microfoundations. The framework names four such microfoundations for the planning capability and specifies the human-AI division of labor through which each operates. It also provides a planning-specific expression of sociotechnical systems theory (Clegg, 2000), and it converts human-centered concerns into design variables with predicted effects, which moves the conversation from documenting problems to testing solutions.

The framework also occupies a distinct position relative to research on organizational readiness for AI. That research diagnoses the preconditions for adoption in projects (Klos et al., 2025), whereas the present framework begins where readiness ends. It assumes deployment proceeds and asks how to design the planning system, cell by cell, to keep automation and human control jointly high. Readiness diagnostics and design frameworks therefore work in sequence: an organization first resolves its readiness tensions, then applies the horizon-specific design commitments specified here. The framework is a design instrument, testable at the level of individual planning decisions.

The Planning Capability and Its Microfoundations

Naming the four principles as microfoundations requires three specifications: the capability they underpin, the resources that capability reconfigures, and the routine through which each principle operates. We define the capability as AI-enabled workforce planning capability. It is the organization's ability to sense shifts in project-driven labor demand, to commit people to projects on that reading, and to renew the skill base that later commitments draw on. The resource base it reconfigures has three elements: the deployable skill portfolio, the assignment routines that allocate that portfolio to projects, and the planning information system through which planners manage both. Each principle operates on a different phase of the capability. Augmentation over automation operates on sensing. The division of cognitive labor determines which demand signals the organization notices, because AI scans breadth while planners interpret the equivocal signals that clients send about schedule intent (Jarrahi, 2018). Meaningful human control operates on seizing. Committing people to a project requires an accountable decision-maker who can override the recommendation and to whom the outcome traces (Santoni de Sio & van den Hoven, 2018). Fairness and transparency also operate on seizing, through legitimacy, because a reconfiguration that workers contest does not execute as designed, and contestation rises when the criteria stay opaque (Kellogg et al., 2020; Uddin, 2026). Capability development operates on reconfiguring. It changes the resource base itself, because assignments distribute the learning that determines what the organization can sense and seize in the following cycle (Teece, 2007). These microfoundations differ in kind from general AI design principles and from AI governance requirements. A governance requirement states a property that a system must hold, and an assessor evaluates the system against that property at a point in time (European Commission High-Level Expert Group, 2019). A microfoundation is a routine that identified people perform repeatedly, and its test is whether the resource base changes across successive planning cycles (Teece, 2007). Table 3 maps each principle to the governance requirement it addresses, and the distinction drawn here explains why that mapping does not exhaust the principle. Augmentation over automation satisfies part of the human agency and oversight requirement, and it also constitutes the sensing routine through which the planning capability operates. The framework claims the second role, which the governance literature does not.

Shared Prosperity as a Downstream Outcome

Shared prosperity follows from the framework as a downstream outcome rather than as a separate argument. Two of the four principles produce distributional effects directly. Fairness and transparency determine which workers receive which assignments, and assignments determine income for day-rate and rotational workers. Capability development determines which workers receive developmental exposure, and that exposure determines future earning

capacity. The framework therefore predicts a distributional consequence, and not only an efficiency consequence, from the same design choices. Propositions 3, 5, and 7 carry that prediction: Proposition 3 concerns perceived justice in allocation, Proposition 5 concerns internal mobility, and Proposition 7 concerns the relative standing of experienced workers. Section states the policy implications that follow. The study present shared prosperity as a theorized outcome of the four principles and not as an established result.

Practical Implications

Workforce planners in project-based industries can use the framework as a design checklist. Each of the twelve cells in Figure 3 poses one question: does the current or planned system satisfy this principle at this horizon? Organizations adopting AI forecasting or scheduling tools can apply the checklist before deployment, which is less costly than remediating trust failures afterward. The propositions also give managers measurable expectations. Planning reversals, override frequency, grievance rates, and internal mobility are all observable in routine workforce data, so organizations can evaluate their own designs without research infrastructure. The twelve cells convert into procurement and configuration questions, and Table 5 supplies the wording. At the tactical horizon under fairness and transparency, for example, the commitment that assignment criteria stay visible to affected workers becomes two configurable requirements. The system exposes the criteria that generated a given assignment, and it records those criteria in a form that a worker can request. An organization unable to answer one of the twelve questions has identified a design gap before deployment rather than after a trust failure.

Policy Implications and Shared Prosperity

Workforce planning is a concrete transmission channel between human-centered AI and shared prosperity. Planning systems distribute work, income, and learning across a workforce. Evidence that AI adoption affects worker stress and well-being (Valtonen et al., 2025), and that algorithmic control provokes contestation (Uddin, 2026), implies that organizations cannot leave distribution design to default settings. The framework suggests two policy-relevant levers. Policymakers can write fairness and transparency requirements into procurement standards for workforce management systems. Organizations can write capability development objectives into workforce planning governance, so that productivity gains from AI fund reskilling rather than displacement.

Shared prosperity meets friction in distribution. The same productivity compression that benefits novices can devalue experienced workers (Brynjolfsson et al., 2025), so prosperity that is shared in aggregate may still be contested in distribution. An experienced team leader

whose hard-won expertise now sits in a recommendation engine may experience the gain as a loss of standing. Policy that pursues shared prosperity through AI must therefore address relative position as well as average productivity. This means pairing reskilling for lower-skill workers with recognition, role redesign, and reward structures that keep experienced workers invested. Without this balance, a fair distribution of gains can coexist with real erosion of trust among the workers a project most depends on.

Limitations

Four limitations bound the framework. First, it is conceptual. The seven propositions remain untested, the twelve cells are analytically derived design commitments rather than validated constructs, and we claim analytical generalization and not statistical generalization. Second, we assembled the evidence base through a structured scoping search rather than a systematic review. The search drew on the Semantic Scholar Academic Graph rather than on Scopus or Web of Science, and that corpus indexes venues of uneven editorial quality: it returned 32 retracted articles, and a venue-indexing filter was needed to remove a further 20 records published in journals outside recognized indexing. A search executed in Scopus would return a different record set, and Table 2 reports the counts so that the difference can be assessed. One author screened every record, and five records remain held pending confirmation of their indexing status. The search may also have missed relevant studies in practitioner outlets and in languages other than English. Third, the synthesis privileged project-based operations, so the horizon structure may transfer poorly to continuous-process or retail workforces, where demand does not arrive in discrete projects. Fourth, we extracted the design principles from a literature that is itself young. Recent work documents a shift toward human-centered issues in AI-enabled workforce planning (Egasmara et al., 2025), so the set of four principles may require revision as evidence accumulates. Proposition 7 carries the additional limitation stated in Section 3.5, since its antecedent has not been observed in project-based work.

Conclusions

This paper asked how organizations can embed human-centered AI principles in workforce planning for project-based operations. The answer developed here is a 3×4 framework that maps augmentation over automation, meaningful human control, fairness and transparency, and capability development onto strategic demand planning, tactical workforce composition, and operational deployment. The framework makes three contributions. It converts domain-general HCAI design properties into twelve horizon-specific design commitments that planners can check before deployment. It theorizes the four principles as microfoundations of an AI-enabled workforce planning capability and specifies the phase of that capability on

which each principle operates. It states seven propositions, each carrying a unit of analysis, a comparison condition, an observable outcome, and a named mechanism.

The study distinguish these contributions from empirical findings. The framework rests on a synthesis of prior studies, and the mechanisms it invokes are documented in that literature. The twelve cells and the seven propositions are the authors' conceptual claims, and no part of the framework has been tested. Proposition 7 rests on evidence drawn from a different work setting and remains exploratory. For practice, the twelve cells give workforce planners a pre-deployment checklist, and Table 5 supplies the wording for each check. The propositions give managers measurable expectations, because planning reversals, override frequency, grievance rates, and internal mobility all appear in routine workforce data. Validation should proceed in two stages. A Delphi study with workforce planning experts in project-based industries can assess the completeness of the twelve cells, refine their wording, and rank the propositions by expected effect size. Multi-project field studies can then test the propositions using the indicators identified in Section 3.7, comparing planning units that differ in design within the same organization. Establishing whether AI compresses productivity differences in judgment-intensive project work is the prior question for Proposition 7. The aim is to move the field from cataloging human-centered concerns toward planning systems whose design commitments can be specified, checked, and tested.

References

- Biruk, S., Jaśkowski, P., & Maciaszczyk, M. (2022). Conceptual framework of a simulation-based manpower planning method for construction enterprises. *Sustainability*, 14(9), 5341. <https://doi.org/10.3390/su14095341>
- Brynjolfsson, E., Li, D., & Raymond, L. (2025). Generative AI at work. *The Quarterly Journal of Economics*, 140(2), 889–942. <https://doi.org/10.1093/qje/qjae044>
- Burgert, F. L., Windhausen, M., Kehder, M., Steireif, N., Mütze-Niewöhner, S., & Nitsch, V. (2024). Workforce scheduling approaches for supporting human-centered algorithmic management in manufacturing: A systematic literature review and a conceptual optimization model. *Procedia Computer Science*, 232, 1573–1583. <https://doi.org/10.1016/j.procs.2024.01.155>
- Clegg, C. W. (2000). Sociotechnical principles for system design. *Applied Ergonomics*, 31(5), 463–477. [https://doi.org/10.1016/S0003-6870\(00\)00009-0](https://doi.org/10.1016/S0003-6870(00)00009-0)
- Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., Duan, Y., Dwivedi, R., Edwards, J. S., Eirug, A., Galanos, V., Ilavarasan, P. V., Janssen, M., Jones, P., Kar, A. K., Kizgin, H., Kronemann, B., Lal, B., Lucini, B., ... Williams, M. D. (2019). Artificial intelligence (AI): Multidisciplinary perspectives on emerging challenges,

- opportunities, and agenda for research, practice and policy. *International Journal of Information Management*, 57, 101994.
<https://doi.org/10.1016/j.ijinfomgt.2019.08.002>
- Egasmar, F., Rahayu, A., Wibowo, L. A., Rofaida, R., Sofia, A., & Fauziyah, A. (2025). Workforce planning optimization utilising AI to improve firm performance: A systematic literature review using VOSviewer. *Journal of Work-Applied Management*.
<https://doi.org/10.1108/JWAM-06-2025-0102>
- El-Farr, H., & Kertechian, K. S. (2024). Knowledge management and knowledge leadership in the fourth industrial revolution: Resolving the automation-augmentation paradox. *IntechOpen*. <https://doi.org/10.5772/intechopen.1005236>
- European Commission High-Level Expert Group on Artificial Intelligence. (2019). Ethics guidelines for trustworthy AI. European Commission.
- Gattermann-Itschert, T., Poreschack, L. M., & Thonemann, U. W. (2023). Using machine learning to include planners' preferences in railway crew scheduling optimization. *Transportation Science*, 57(3), 796–812. <https://doi.org/10.1287/trsc.2022.1196>
- Gerlach, M., Renggli, F. J., Bieri, J. S., Sariyar, M., & Golz, C. (2025). Exploring nurse perspectives on AI-based shift scheduling for fairness, transparency, and work-life balance. *BMC Nursing*, 24(1), Article 1161. <https://doi.org/10.1186/s12912-025-03808-0>
- Hickman, L., Huynh, C., Gass, J., Booth, B. M., Kuruzovich, J., & Tay, L. (2025). Whither bias goes, I will go: An integrative, systematic review of algorithmic bias mitigation. *Journal of Applied Psychology*, 110(7), 979–1000. <https://doi.org/10.1037/apl0001255>
- Hunkenschroer, A. L., & Luetge, C. (2022). Ethics of AI-enabled recruiting and selection: A review and research agenda. *Journal of Business Ethics*, 178(4), 977–1007. <https://doi.org/10.1007/s10551-022-05049-6>
- Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making. *Business Horizons*, 61(4), 577–586. <https://doi.org/10.1016/j.bushor.2018.03.007>
- Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Academy of Management Annals*, 14(1), 366–410. <https://doi.org/10.5465/annals.2018.0174>
- Klos, T. B., Zarccone, A., Rentrop, C., Riedinger, C., Prinz, N., & Huber, M. (2025). Entangled by design: A structured overview of management challenges concerning AI adoption in organizations. *Annals of Computer Science and Information Systems*, 43, 705–713. <https://doi.org/10.15439/2025F1129>

- Korucu, H. İ., Emek, M. S., & Gülmez, E. (2023). Development of a new personalized staff-scheduling method with a work-life balance perspective: Case of a hospital. *Annals of Operations Research*, 328(1), 793–820.
<https://doi.org/10.1007/s10479-023-05244-2>
- Kudina, O., & van de Poel, I. (2024). A sociotechnical system perspective on AI. *Minds and Machines*, 34(3), Article 21. <https://doi.org/10.1007/s11023-024-09680-2>
- Latos, B. A., Buckhorst, A. F., Kalantar, P., Bentler, D., Gabriel, S., Dumitrescu, R., Minge, M., Steinmann, B., & Guhr, N. (2024). Time autonomy in personnel planning: Requirements and solution approaches in the context of intelligent scheduling from a holistic organizational perspective. *Zeitschrift für Arbeitswissenschaft*, 78(3), 277–298. <https://doi.org/10.1007/s41449-024-00432-7>
- Micheli, G. J. L., Martino, A., Porta, F., Cravello, A., Panaro, M. A., & Calabrese, A. (2023). Workforce planning in project-driven companies: A high-level guideline. *Frontiers in Industrial Engineering*, 1, 1267244. <https://doi.org/10.3389/fieng.2023.1267244>
- Parker, S. K., & Grote, G. (2022). Automation, algorithms, and beyond: Why work design matters more than ever in a digital world. *Applied Psychology*, 71(4), 1171–1204. <https://doi.org/10.1111/apps.12241>
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation–augmentation paradox. *Academy of Management Review*, 46(1), 192–210. <https://doi.org/10.5465/amr.2018.0072>
- Renggli, F. J., Gerlach, M., Bieri, J. S., Golz, C., & Sariyar, M. (2025). Integrating nurse preferences into AI-based scheduling systems: Qualitative study. *JMIR Formative Research*, 9, Article e67747. <https://doi.org/10.2196/67747>
- Santoni de Sio, F., & van den Hoven, J. (2018). Meaningful human control over autonomous systems: A philosophical account. *Frontiers in Robotics and AI*, 5, Article 15. <https://doi.org/10.3389/frobt.2018.00015>
- Shneiderman, B. (2020a). Human-centered artificial intelligence: Three fresh ideas. *AIS Transactions on Human-Computer Interaction*, 12(3), 109–124. <https://doi.org/10.17705/1thci.00131>
- Shneiderman, B. (2020b). Human-centered artificial intelligence: Reliable, safe & trustworthy. *International Journal of Human-Computer Interaction*, 36(6), 495–504. <https://doi.org/10.1080/10447318.2020.1741118>
- Teece, D. J. (2007). Explicating dynamic capabilities: The nature and microfoundations of (sustainable) enterprise performance. *Strategic Management Journal*, 28(13), 1319–1350. <https://doi.org/10.1002/smj.640>

- Teece, D. J., Pisano, G., & Shuen, A. (1997). Dynamic capabilities and strategic management. *Strategic Management Journal*, 18(7), 509–533. [https://doi.org/10.1002/\(SICI\)1097-0266\(199708\)18:7<509::AID-SMJ882>3.0.CO;2-Z](https://doi.org/10.1002/(SICI)1097-0266(199708)18:7<509::AID-SMJ882>3.0.CO;2-Z)
- Uddin, A. F. M. A. (2026). Trust under the algorithm: Employee perceptions of control, fairness, and autonomy in algorithmic management. *Saudi Journal of Business and Management Studies*, 11(2), 40–52. <https://doi.org/10.36348/sjbms.2026.v11i02.003>
- Valtonen, A., Saunila, M., Ukko, J., Treves, L., & Ritala, P. (2025). AI and employee wellbeing in the workplace: An empirical study. *Journal of Business Research*, 199, 115584. <https://doi.org/10.1016/j.jbusres.2025.115584>
- Verma, D., Kumar, K. S., & Bhanot, S. (2025). Workforce management in the era of generative AI: Insights and research agendas. *European Economic Letters*, 15(1). <https://doi.org/10.52783/eel.v15i1.2359>
- Xu, W., Dainoff, M. J., Ge, L., & Gao, Z. (2022). Transitioning to human interaction with AI systems: New challenges and opportunities for HCI professionals to enable human-centered AI. *International Journal of Human-Computer Interaction*, 39(3), 494–518. <https://doi.org/10.1080/10447318.2022.2041900>
- Zhang, P., Dillard, N., & Cavallo, T. (2025). Navigating the limitations of algorithmic management: An integrative framework of sociotechnical systems theory (STS) and strategic HRD. *Human Resource Development Review*, 24(3), 279–303. <https://doi.org/10.1177/15344843251320252>