

ORGANISATIONAL SLACK IN WORKFORCE DESIGN AND INNOVATION OUTPUT ELASTICITY

HRM VARIABLE: DELIBERATE STAFFING SLACK ® ECONOMIC OUTCOME: ELASTICITY OF INNOVATION OUTPUT

Tests the economic case for maintaining workforce slack, comparing present conditions with lean AI-optimised operating models. Addresses a live and unresolved tension between operational efficiency and innovative capacity.

N. Subbu Krishna Sastry¹, Dr. Manjula Mallya M²

¹School of Management, CMR University, Bengaluru, Karnataka, India

²Associate Professor & Head, Dept. of Economics, Government First Grade College for Women Balmatta Mangalore, Karnataka, India. ORCID ID: 0009-0005-8812-6912,

ABSTRACT

Deliberate staffing slack – paid professional time that an employer designates as uncommitted rather than scheduled against delivery work – is one of the few human resource variables that is simultaneously a cost line and an innovation input. Artificial intelligence sharpens the conflict: it raises the marginal value of a committed hour and, for the first time, gives workforce planners the instrumentation to reclaim every uncommitted minute. This paper asks what the elasticity of innovation output with respect to deliberate staffing slack actually is, and whether the lean AI-optimised operating model now spreading across knowledge-intensive sectors is economically defensible. An analytical model is developed in which paid professional time splits between a delivery block whose productivity is AI-augmented and an innovation block in which slack must first clear a contiguity filter and then survive a governance dissipation term. The model yields a closed-form innovation peak, an elasticity function, and an elasticity rule characterising the profit-maximising slack ratio. A Monte Carlo experiment over 40,000 parameter draws calibrated to published estimates returns a median profit-maximising slack ratio of 8.9 per cent of paid professional time against an innovation-maximising ratio of 18.5 per cent, and an elasticity that peaks near 0.49 at a five per cent slack ratio before falling to 0.27 by ten per cent. A lean five per cent utilisation target under AI-augmented technology retains 99 per cent of the private value of the re-optimised policy while delivering only 87 per cent of its innovation output: the value curve is flat exactly where the innovation curve is steep, which is why the innovation loss is easy to miss and easy to authorise. In roughly one configuration in eight, no deliberate slack at all is privately optimal – slack is a lumpy commitment, not a marginal one. Once the spillover premium documented for R&D is admitted, the socially optimal slack ratio rises to between ten and fourteen per cent, so that about one per cent of firm value buys roughly eight per cent more innovation output at the margin.

KEYWORDS – Organisational Slack; Deliberate Staffing Slack; Workforce Design; Innovation Output Elasticity; Artificial Intelligence; Utilisation Targets; Knowledge Production Function; Monte Carlo Simulation; Human Resource Economics.

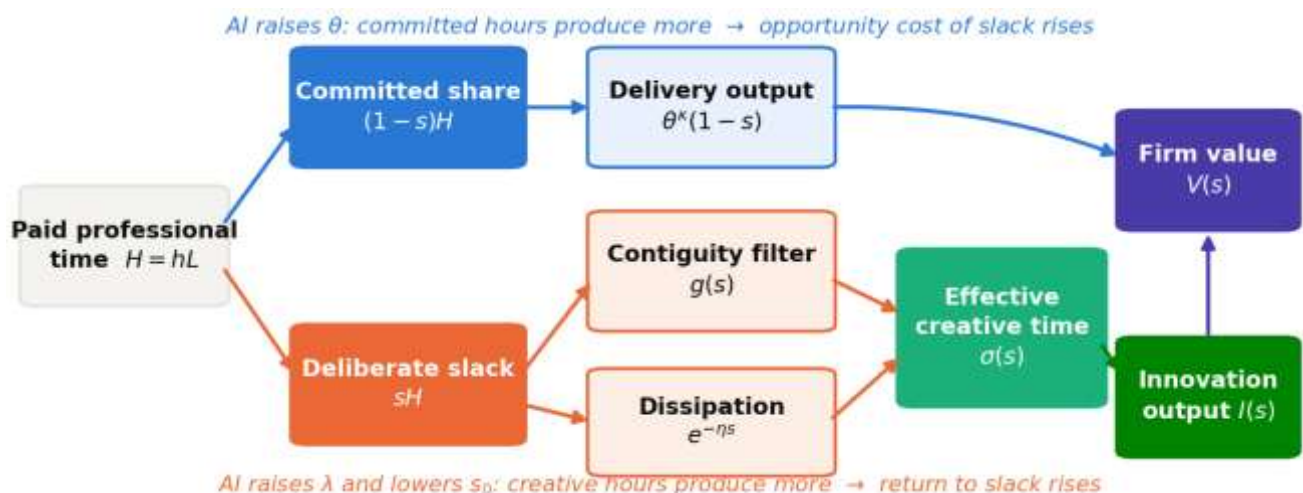


Fig. 1 Architecture of the model. Paid professional time divides into a committed block whose value AI augments and a deliberate slack block that must clear a contiguity filter and a dissipation term before it becomes effective creative time. The two annotated forces pull the optimal slack ratio in opposite directions.



I. INTRODUCTION

Every workforce plan makes an implicit statement about how much of an employee's paid time the employer expects to be able to account for. In a call centre the answer approaches all of it; in a research laboratory it does not. Between those poles sits the large and growing population of knowledge-intensive firms — software and IT services, professional services, pharmaceutical development, engineering design, higher education — where the same person is expected both to deliver scheduled output and to originate the ideas that will constitute the firm's output three years from now. For those firms the share of paid professional time deliberately left uncommitted is a genuine design variable, and this paper treats it as the human resource variable of interest under the name *deliberate staffing slack*. It is measured throughout as s , the fraction of contracted professional hours that the employer designates as unassigned, non-scheduled and not chargeable against a delivery commitment.

The question of whether such slack pays has been contested for six decades without resolution. The behavioural tradition inherited from Cyert and March [1] treats slack as the cushion that permits search, experimentation and the absorption of shocks. The economic tradition, following Leibenstein [2] and Jensen [3], treats uncommitted resources as the residue of weak monitoring and expects them to be dissipated rather than invested. Empirical work has supported both readings, and two recent syntheses of four decades of slack research conclude that the field still cannot say whether the relationship between slack and innovation is positive, negative or curvilinear, or under what conditions each obtains [4], [5].

What has changed is the urgency of the question. Artificial intelligence is being deployed precisely on the tasks that used to absorb the uncommitted margin of professional time — drafting, summarising, searching, first-pass analysis, code scaffolding — and it is being deployed alongside scheduling and workforce-optimisation systems that make previously invisible time visible and therefore reclaimable. Adoption is no longer marginal: the United States Census Bureau's Business Trends and Outlook Survey put firm-level AI use at between 17 and 20 per cent of all businesses over the six months to May 2026, rising to 37 per cent among firms with 250 or more employees and to 39.7 per cent in the information sector [6]; a Federal Reserve Board synthesis of three independent surveys places worker-level generative AI use at roughly 41 per cent [7]. At the same time the innovation aggregates are softening rather than accelerating: global research and development spending growth fell to an estimated 2.9 per cent in real terms in 2024 and a projected 2.3 per cent in 2025, the slowest rates recorded since 2010, and corporate R&D among the largest spenders grew about one per cent in real terms [8].

This juxtaposition — the most powerful productivity instrument in a generation arriving at the same moment that innovation investment decelerates — is what makes the economics of workforce slack a live rather than an academic question. If the elasticity of innovation output with respect to deliberate staffing slack is small, the lean AI-optimised operating model costs little and the efficiency case is decisive. If it is large, firms are currently authorising a quiet and cumulative reduction in their own innovative capacity, and doing so through a variable that appears on no innovation dashboard.

The paper makes four contributions. First, it specifies deliberate staffing slack as an economic input rather than a residual, separating the *quantity* of designated slack from its *usability*, and modelling usability through two opposing filters: a contiguity filter, because fragmented minutes do not aggregate into creative time, and a dissipation term, because the share of designated slack that is actually converted into directed effort falls as the designated pool grows. Second, it derives in closed form the innovation-maximising slack ratio, the elasticity of innovation output with respect to slack, and an elasticity rule that characterises the profit-maximising ratio as the point at which that elasticity equals the ratio of forgone delivery value to realised innovation value. Third, it quantifies the model through a Monte Carlo experiment that compares present operating conditions with an AI-augmented technology and, within each, compares a lean utilisation target against re-optimised and innovation-maximising staffing policies. Fourth, it locates the private and social optima relative to each other and draws the resulting design and policy implications.

The remainder of the paper proceeds as follows. Section II reviews the slack literature and isolates the gap. Section III sets out the theoretical framework and the propositions the model is built to test. Section IV develops the model and its closed-form results. Section V describes the calibration and the simulation protocol. Section VI reports results. Section VII discusses their managerial and policy implications, Section VIII the limitations, and Section IX concludes.

II. LITERATURE REVIEW AND RESEARCH GAP

A. Slack in the behavioural theory of the firm

Slack entered organisation theory as an explanation of why firms do not behave as textbook optimisers. Cyert and March [1] described it as the difference between the resources a coalition commands and the payments required to hold that coalition together, and treated it as the buffer that absorbs environmental variation and finances search. Penrose [9] had already located a related idea in the managerial services that a growing firm cannot fully employ at every moment, and which therefore become the seedbed of diversification. Bourgeois [10] gave the construct its operational vocabulary, distinguishing available, recoverable and potential slack and setting the measurement conventions that still dominate empirical work.



B. The curvilinearity debate

The single most cited empirical intervention remains Nohria and Gulati [11], who tested slack against innovation inside two multinationals and reported an inverted U: too little slack prevents experimentation, too much removes the discipline that makes experimentation selective. That result has been replicated, contradicted and re-specified many times. Mellahi and Wilkinson [12] approached it from the opposite direction, showing that the level of slack reduction following downsizing is associated with subsequent innovation output — a formulation much closer to the question this paper asks, because it treats slack removal, rather than slack level, as the treatment. Two 2025–2026 syntheses now frame the state of play. Freitas and colleagues [4], mapping 340 articles across 167 journals, find the field split between a large traditional stream concerned with how slack should be allocated to raise performance and a much smaller body concerned with how firms re-orchestrate slack under environmental uncertainty. Depino-Besada and colleagues [5] are blunter: positive linear, negative and inverted-U relationships all appear in the published record; slack is assumed to be fungible when it is not; and an enlarged slack pool is as likely to finance weak projects as strong ones.

A complementary argument reaches the same conclusion from the constraint side rather than the abundance side. Sastry and Manjula Mallya [26], examining economic constraints and the status quo across Indian organisations, argue that resource tightness does not simply lower the volume of innovative activity but alters the organisation's default response to a problem, shifting it from search towards repetition of what already works. Read alongside the curvilinearity evidence, that argument implies something the inverted-U specification cannot represent: a threshold below which the marginal unit of slack buys no search at all, rather than a smooth decline towards it. The model in Section IV is built to carry exactly that feature.

C. Time as the scarce creative input

A parallel literature treats time rather than money as the binding constraint on creative work. Amabile, Hadley and Kramer [13] found that extreme time pressure depresses creative cognition and that the depression persists for days afterwards, while a sense of protected, focused time raises it. Agrawal, Catalini, Goldfarb and Luo [14] supplied the cleanest causal evidence available: using staggered college holiday breaks as exogenous variation in the availability of slack time across 5,919 United States cities and 97,423 crowdfunding projects, they report that a break raises project creation by 0.0294 projects per city-week against a baseline of 0.060 — an increase of roughly 49 per cent — with the effect concentrated in categories matching the discipline of the institution on break. Their evidence points to execution rather than ideation: the projects launched during breaks had longer development histories, implying that slack time converts existing ideas into output rather than generating the ideas themselves. That distinction matters for the present model, because it implies that the marginal product of slack time depends on a stock of unexecuted ideas and on whether the slack arrives in usable blocks.

The cost side of unprotected time is documented in the same literature. Sastry [20], in a study of leadership design against workforce burnout, locates the damage in sustained overcommitment: when every professional hour is scheduled, the discretionary effort that innovative work draws on is the first thing to be withdrawn, and it is withdrawn quietly. That observation supplies the behavioural basis for treating the conversion of designated slack into directed effort as something that can fail, rather than as an accounting identity.

D. Human resource slack specifically

Financial slack and human resource slack are not interchangeable. Rashidi and colleagues [15] argue for treating human resource slack as a construct in its own right, with its own antecedents and its own conversion mechanics. Chen, Park and Rajwani [16], studying 3,229 Chinese listed firms, find that it is specifically highly-skilled human resource slack that allows firms to convert externally acquired resources into innovation, which is consistent with the absorptive capacity argument of Cohen and Levinthal [17]: slack is only an innovation input in the hands of people who can recognise what to do with it. This is the strand the present paper extends, by asking not whether human resource slack matters but how elastic innovation output is with respect to it, and at what level of it the firm should stop.

How that capacity is released has been examined directly. Sastry and Ravish [30], [31], in two studies of delegated authority and resource provision, find that employee initiative responds to the joint presence of discretion and resource, and not to either on its own — an employee given time without authority, or authority without time, does not initiate. That joint condition is the empirical counterpart of the interaction this paper formalises between the quantity of designated slack and its usability. Sastry [19], examining departmental structure and task achievement, adds the organisational layer: the same professional hour yields different output depending on how the unit around it is arranged, which is the proposition the contiguity filter below makes precise.

E. Artificial intelligence and the lean operating model

The productivity evidence on generative AI is now reasonably well established at the task level. Brynjolfsson, Li and Raymond [18], observing 5,179 customer support agents and three million conversations, report an average 14 per cent increase in issues resolved per hour, concentrated among less experienced workers, with little measured gain for the most experienced. The significance of that number for workforce design is not the average but its arithmetic: a durable double-digit gain in the



productivity of committed hours mechanically raises the opportunity cost of any hour that is not committed. Against this, AI also augments the creative block — literature search, prototyping, synthesis — and removes some of the administrative fragmentation that prevents slack from arriving in usable blocks. Which of these effects dominates is an empirical and parametric question, and no existing study poses it in those terms.

The human resource side of that shift has been studied in this series. Sastry and Manjula Mallya [21], on analytics-based hiring, show that the composition of the professional workforce — and therefore the value of its marginal hour — is itself now an algorithmic outcome; Sastry and Manjula Mallya [23], on skill diversification under technological disruption, address how quickly that hour can be redeployed between delivery and search. Both operate on the same quantity the model treats as the opportunity cost of uncommitted time. Sastry and Manjula Mallya [29], on talent designation and social comparison, establish the point on which this paper's central construct rests: what an employer formally designates carries economic consequences beyond the resource being designated. Slack that is announced and protected is a different input from slack that merely happens to be left over, and only the first is a workforce-design variable.

F. Prior work in this series and the research gap

This paper is the twelfth in a series pairing one human resource variable with one economic outcome. Earlier entries have examined departmental structure and task achievement [19], leadership design against workforce burnout [20], analytics-based hiring [21], performance appraisal and engagement [22], skill diversification under technological disruption [23], the constraints on women micro-entrepreneurs [24], and enterprise systems in higher education [25]. The common method is to take a variable that human resource management treats as an administrative parameter and ask what its economic elasticity is.

Four further studies bear on the present problem beyond those already drawn on above. Sastry [27], on fostering future leadership for innovation under global competitive pressure, and Manjula Mallya and Sastry [28], on building future-ready organisations through young-talent strategy, both locate innovative capacity in how professional attention is organised rather than in how much is spent on it, which is the premise the contiguity filter formalises. Sastry and Selvarasu [22], on performance appraisal and employee engagement, add the measurement problem that runs through this paper: what an organisation counts determines what its managers optimise, and designated slack is counted nowhere. Sastry and Manjula Mallya [32], on women employees in Bangalore's information technology start-ups and the MSME frameworks around them, and Sastry [24], on the constraints facing women micro-entrepreneurs in Karnataka, together document the operating conditions under which no discretionary allocation is affordable at all — anticipating the corner solution reported in Section VI and the cross-country pattern in Section VI-H.

Two methodological positions also carry over. Sastry and Arunachalam [33], on probabilistic safe reinforcement learning and the management of risk in dynamic environments, established the approach taken here: where a decision parameter cannot be observed directly, the honest response is to propagate its uncertainty through the model rather than fix it at a point estimate, which is why every result below is reported as a distribution rather than a single number. Vineeth Kumar and colleagues [25], on enterprise resource planning in Indian higher education, established the second: an operating technology should be evaluated not by the efficiency it delivers in isolation but by what it does to the discretion of the people working inside it — which is precisely the question this paper puts to artificial intelligence.

Set against that earlier work, the gap this paper addresses is specific and threefold. First, the slack literature almost invariably measures slack in financial ratios and innovation in patent counts, so the quantity a workforce planner actually sets — the share of paid professional time left uncommitted — is never the object of estimation. Second, the curvilinearity debate has been conducted without a formal mechanism that generates curvature, so competing empirical results cannot be reconciled. Third, no study has asked what an AI-shifted cost frontier does to the optimal slack ratio, which is the decision firms are making now. The model below is built to close all three.

III. THEORETICAL FRAMEWORK AND PROPOSITIONS

The framework rests on three claims, each of which is standard in isolation and none of which has previously been combined into a single decision problem.

The first is that designated slack and usable slack are different quantities. An employer can designate twenty per cent of a week as uncommitted and still deliver nothing usable if that time arrives as forty scattered eighteen-minute intervals. Creative work has a set-up cost, and the empirical literature on time pressure and on protected time [13], [14] is consistent with a threshold below which designated slack does not convert. The model represents this as a contiguity filter.

The second is that the conversion rate of designated slack falls as the designated pool grows. This is the economists' objection to slack, expressed as a mechanism rather than an assumption: as the uncommitted share rises, the marginal designated hour is less well directed, less visible, and more likely to fund a project that would not survive scrutiny [3], [5]. The model represents this as a dissipation term.



The third is that the firm does not maximise innovation. It maximises the sum of delivery value and the value of innovation output, and it does so on a frontier that artificial intelligence is shifting. Because AI raises the productivity of committed hours and of creative hours simultaneously, the direction of the net effect on the optimal slack ratio cannot be signed a priori; it must be computed.

Together these yield six propositions, which Sections IV and VI derive and quantify respectively.

P1. Innovation output is single-peaked in the deliberate slack ratio, with the peak generated endogenously by the interaction of the contiguity filter and the dissipation term rather than imposed by functional form.

P2. The elasticity of innovation output with respect to deliberate staffing slack is itself non-monotonic: it is small at very low slack ratios because designated time is not yet usable, reaches a maximum shortly after the contiguity threshold is cleared, and declines thereafter.

P3. The profit-maximising slack ratio is strictly below the innovation-maximising ratio, and the gap widens with the marginal delivery value of a committed hour.

P4. The private value function is flat in the neighbourhood of its optimum while the innovation function is steep, so a workforce policy that errs towards utilisation incurs a small and difficult-to-observe value penalty and a large innovation penalty.

P5. Deliberate slack is a lumpy commitment: for a non-trivial region of the parameter space the optimum is a corner at zero, and intermediate token allocations are dominated by both zero and a substantial allocation.

P6. Because innovation output carries positive spillovers, the socially optimal slack ratio strictly exceeds the privately optimal ratio, and the wedge is increasing in the spillover multiplier.

Table 1: Notation and interpretation

Symbol	Meaning	Set by
s	Deliberate staffing slack ratio: share of paid professional time designated as uncommitted	Firm (decision variable)
h, L	Contracted hours per employee; professional headcount	Workforce plan
s_0	Contiguity threshold: slack ratio at which half of designated slack arrives in usable blocks	Work design, meeting load
v	Steepness of the contiguity activation	Work design
η	Dissipation rate of designated slack	Governance, culture
$s^\dagger$	Innovation-maximising slack ratio	Implied by s_0, v, η
α	Elasticity of slack-attributable innovation with respect to effective creative input	Knowledge production function
b	Share of innovation output arising from dedicated (non-slack) channels	R&D organisation
ρ	Innovation value share: value of innovation output relative to delivery revenue at the reference slack ratio	Product strategy, market
θ	AI augmentation factor on committed delivery hours	Technology
κ	Pass-through of AI gains into marginal delivery revenue	Market structure
λ	AI augmentation factor on creative hours	Technology
μ	Spillover multiplier on the social valuation of innovation output	Policy environment
ε	Elasticity of innovation output with respect to s	Model output

Table 2: Slack constructs and their transmission to innovation output

Construct	Definition used here	Transmission channel	Sources
Financial slack	Liquid or recoverable financial resources above operating need	Funds projects; weakens capital rationing	[1], [10], [11]
Absorbed slack	Excess resources already embedded in cost structure	Slow to redeploy; can shelter inefficiency	[5], [10]
Unabsorbed slack	Uncommitted resources available for redirection	Fast redeployment to search	[4], [5]
Human resource slack	Professional capacity above the level required for delivery	Absorptive capacity; recognition of external knowledge	[15], [16], [17]
Deliberate staffing slack	Share of paid professional time designated as uncommitted	Contiguous creative time; execution of held ideas	[13], [14]; this paper

IV. THE MODEL

A. The time budget

A knowledge-intensive firm employs L professionals on h contracted hours each, giving a paid time budget H . The firm designates a fraction s of that budget as deliberate slack and the remainder as committed delivery time.

$$H = hL, \quad H_c = (1 - s)H, \quad H_s = sH \quad (1)$$

All quantities below are expressed per full-time-equivalent year and normalised so that the delivery value of one fully committed FTE-year at present technology equals unity. This removes scale and currency from the problem and makes every result a ratio.

B. The contiguity filter and the dissipation term

Designated slack becomes usable creative time only to the extent that it arrives in blocks long enough to sustain non-routine work. This is modelled as a smooth activation function with half-activation at the contiguity threshold s_0 and steepness v .

$$g(s) = \frac{s^v}{s^v + s_0^v} \quad (2)$$

Against this, the share of designated slack actually converted into directed creative effort declines as the designated pool grows, at rate η . Effective creative intensity per unit of paid professional time is therefore the product of the designated quantity, the activation and the dissipation.

$$\sigma(s) = s g(s) e^{-\eta s} \quad (3)$$

Equation (3) is the analytical core of the paper. The first two factors rise in s and the third falls, so σ is single-peaked without any peak being assumed. Because $\log \sigma$ has derivative $(v+1) - v g(s) - \eta s$ with respect to $\log s$, and g is strictly increasing, that derivative is strictly decreasing: the peak is unique. Setting it to zero at the desired peak location $s^\dagger$ gives the dissipation rate implied by that peak, which is how the model is calibrated.

$$\eta = \frac{(v+1) - v g(s^\dagger)}{s^\dagger} \quad (4)$$

C. Innovation output

Slack-attributable innovation output is a constant-elasticity function of effective creative input, augmented by the AI creative factor λ and normalised to unity at a reference slack ratio s_r taken as ten per cent of paid professional time.

$$I_s(s) = \left[\frac{\lambda \sigma(s)}{\sigma(s_r)} \right]^\alpha \quad (5)$$

Not all innovation comes from slack. A share b of innovation output at the reference point originates in dedicated research organisations, formal project pipelines and external acquisition, and is unaffected by the general workforce's slack ratio. Total

innovation output is therefore a convex combination, which also prevents the model from making the untenable claim that a firm with no designated slack innovates not at all.

$$I(s) = b + (1 - b)I_s(s) \quad (6)$$

D. Delivery value and the firm's problem

Committed hours produce delivery value at a rate augmented by AI. The augmentation reaches the firm's marginal revenue only to the extent that competition does not compete it away, which the pass-through exponent κ captures: $\kappa = 1$ corresponds to elastic demand in which every productivity gain is sellable, κ near zero to a market in which gains are dissipated in price.

$$D(s) = \theta^\kappa (1 - s) \quad (7)$$

Writing the innovation value share at the reference point as ρ , the firm's objective per FTE-year is the sum of the two blocks.

$$V(s) = \theta^\kappa (1 - s) + \rho (1 - s_r) I(s) \quad (8)$$

E. The elasticity of innovation output

Differentiating (6) in logarithms gives the quantity named in the title of the paper: the elasticity of total innovation output with respect to deliberate staffing slack.

$$\epsilon_{I,s}(s) = \frac{(1 - b)I_s(s)}{I(s)} \alpha [(\nu + 1) - \nu g(s) - \eta s] \quad (9)$$

Equation (9) has an important structural property. The bracketed term depends only on the work-design and governance parameters ν , s_0 and η and on the knowledge-production elasticity α ; it does not contain θ or κ . Artificial intelligence therefore does not shift the elasticity curve through the delivery channel at all. It shifts the curve only if it changes how slack converts — by lowering the contiguity threshold, by raising λ , or by tightening governance — and it shifts the firm's chosen point on that curve through the opportunity cost of a committed hour. Separating these two kinds of effect is what allows Section VI to attribute the change in the optimal slack ratio to specific mechanisms.

F. The elasticity rule

The first-order condition of (8) can be written entirely in elasticity terms. At an interior optimum the elasticity of innovation output with respect to slack equals the ratio of forgone delivery value to realised innovation value, scaled by the ratio of slack to committed time.

$$\epsilon_{I,s}(s^*) = \frac{D(s^*)}{\rho(1 - s_r)I(s^*)} \cdot \frac{s^*}{1 - s^*} \quad (10)$$

Equation (10) is the paper's decision rule and it is directly usable. It says that a firm need not know the whole innovation production function to test its own staffing policy; it needs only an estimate of the local elasticity and of the ratio between the value of its innovation output and its delivery revenue. Because the right-hand side is increasing in s and the left-hand side is decreasing beyond the contiguity threshold, the solution is unique wherever an interior solution exists.

Two boundary results follow. First, since the marginal value of innovation output is zero at the innovation peak $s^\dagger$ while a committed hour still has positive value there, the profit-maximising ratio is strictly below $s^\dagger$ for every finite parameter set — Proposition P3. Second, inverting (10) at any candidate lean target s_ℓ gives the break-even innovation value share below which that lean target is genuinely optimal.

$$\rho^c = \frac{\theta^\kappa}{(1 - s_r)I'(s_\ell)} \quad (11)$$

G. Spillovers and the social optimum

Innovation output generates returns that the innovating firm does not capture. Bloom, Schankerman and Van Reenen [34] find that the gross social return to R&D is at least double the private return; Jones and Summers [35] compute an average social benefit-cost ratio of about 13.3 for innovation investment, with a conservative floor near 4.9 once long payoff lags are imposed. Introducing a spillover multiplier μ on the social valuation of innovation output gives the social objective, from which the socially optimal slack ratio follows in the same way.

$$V_{\text{soc}}(s) = \theta^\kappa (1 - s) + \mu \rho (1 - s_r) I(s) \quad (12)$$

V. CALIBRATION AND SIMULATION DESIGN

A. Parameter anchors

The model is not estimated on firm data; it is calibrated to published parameter ranges and then solved numerically across the resulting uncertainty. This is the appropriate design here because the decision variable — the share of paid professional time an employer designates as uncommitted — is not systematically reported in any firm-level dataset, so an estimation exercise would have to proxy it and would inherit the proxy's error. A calibrated model states its assumptions explicitly and lets the reader vary them.

Every parameter is drawn from a triangular distribution over a plausible range, except the pass-through exponent, which is drawn uniformly because no central value is defensible. Table 3 sets out the anchors. The knowledge-production elasticity α is centred at 0.55, within the band that firm-level patent-to-research-input elasticities have occupied since Griliches [36] and consistent with the diminishing research productivity documented by Bloom, Jones, Van Reenen and Webb [37]. The innovation peak $s^\dagger$ is centred at 18 per cent, matching the discretionary-innovation-time allowances long reported at research-intensive manufacturers and technology firms, which cluster between fifteen and twenty per cent of paid time. The AI delivery augmentation θ is centred at 1.14, taken directly from the 14 per cent average productivity gain measured by Brynjolfsson, Li and Raymond [18].

Table 3: Calibration of the parameter space

Parameter	Distribution	Central	Anchor
α creative-input elasticity	Triangular(0.35, 0.55, 0.80)	0.55	Patent–research-input elasticities [36], [37]
s_0 contiguity threshold	Triangular(0.03, 0.06, 0.10)	0.06	Protected-time and time-pressure evidence [13], [14]
v activation steepness	Triangular(2.0, 3.0, 5.0)	3.0	No direct anchor; carried in sensitivity
$s^\dagger$ innovation peak	Triangular(0.13, 0.18, 0.25)	0.18	Reported discretionary-time allowances, 15–20%
b non-slack share of innovation	Triangular(0.30, 0.50, 0.70)	0.50	Dedicated R&D vs distributed innovation [16], [17]
ρ innovation value share	Triangular(0.10, 0.30, 0.60)	0.30	New-product vitality in R&D-intensive firms [8]
θ AI delivery augmentation	Triangular(1.05, 1.14, 1.35)	1.14	14% measured productivity gain [18]
κ pass-through	Uniform(0.30, 1.00)	0.65	Market structure; no central estimate
λ AI creative augmentation	Triangular(1.00, 1.12, 1.30)	1.12	Task-level generative AI evidence [18]
Δs_0 AI fragmentation relief	Triangular(0.00, 0.30, 0.50)	0.30	Automation of administrative interruption
ζ governance tightening	Triangular(0.00, 0.15, 0.35)	0.15	Algorithmic scheduling and utilisation tracking

B. Simulation protocol

Forty thousand parameter vectors are drawn. For each vector the model is solved twice: once under present technology, in which $\theta = \lambda = 1$ and the contiguity threshold takes its drawn value, and once under an AI-augmented technology, in which θ , κ and λ take their drawn values and the contiguity threshold is reduced by the drawn fragmentation relief. Within each technology the objective is evaluated on a grid of 1,751 slack ratios between zero and 35 per cent, and the global maximum is taken, so that corner solutions at $s = 0$ are found rather than assumed away. Three staffing policies are then compared within each technology: a lean utilisation target fixed at five per cent designated slack, the profit-maximising ratio, and the innovation-maximising ratio. A variant of the lean policy adds the governance tightening term ζ , representing the algorithmic scheduling that typically accompanies a lean AI operating model.

All innovation and value quantities are indexed so that the profit-maximising policy under present technology equals 100. Sensitivity is assessed by Spearman rank correlation between each input parameter and each output of interest, which is robust to the non-linearity of the mapping. The simulation is implemented in Python with a fixed random seed and is fully reproducible.

VI. RESULTS

A. The two filters and the shape they generate

Figure 2 shows the mechanism at the central calibration. The contiguity filter rises steeply through the threshold region and is effectively complete by twelve per cent, while dissipation falls monotonically from the outset. Their product, effective creative intensity, rises to a peak at $s^\dagger$ and declines thereafter. Nothing about the peak is assumed: it is the point at which the gain from having more designated time is exactly offset by the loss from directing it less well.

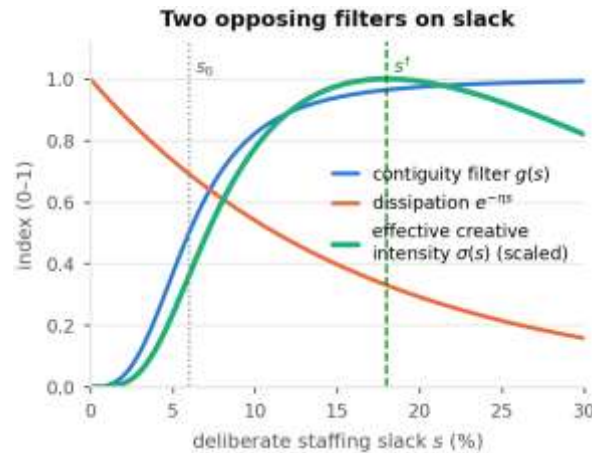


Fig. 2 The contiguity filter, the dissipation term and the effective creative intensity they generate. The peak of $\sigma(s)$ is endogenous, not imposed.

Figure 3 places innovation output and firm value on the same index. The visual result is the paper's central empirical claim and is stated formally as Proposition P4: through the region between five and twelve per cent designated slack the value curve moves by half a per cent while the innovation curve rises by 37 per cent. A firm optimising on observable near-term value faces a nearly flat objective in precisely the range where its innovative capacity is most sensitive. The local minimum visible near two per cent is the lumpiness result of Proposition P5: a token slack allocation is worse than either no allocation or a serious one.

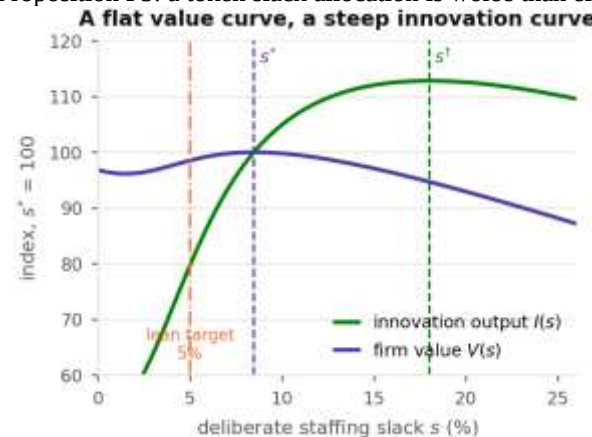


Fig. 3 Innovation output and firm value against the deliberate slack ratio at the central calibration, both indexed to the profit-maximising point. The value curve is flat where the innovation curve is steep.

B. The elasticity profile

Figure 4 and Table 4 report the elasticity of innovation output with respect to deliberate staffing slack across the simulated parameter space. The profile confirms Proposition P2. At a two per cent slack ratio the median elasticity is only 0.16, because designated time has not yet cleared the contiguity threshold and is being spent on fragments. It rises to a median of 0.49 at five per cent, falls to 0.38 at eight per cent and 0.27 at ten per cent, and is statistically indistinguishable from zero by fifteen per cent. At the profit-maximising ratio the median elasticity is 0.29.

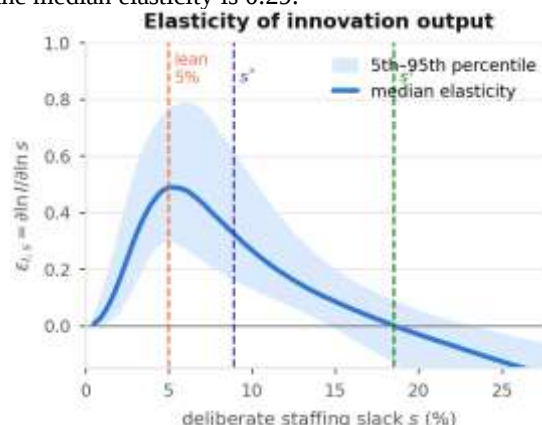


Fig. 4 Elasticity of innovation output with respect to deliberate staffing slack, median and 5th–95th percentile band across 40,000 parameter draws.

Table 4: Elasticity of innovation output with respect to deliberate staffing slack

Slack ratio	5th pct	25th pct	Median	75th pct	95th pct
2%	0.05	0.11	0.16	0.22	0.33
5% (lean target)	0.30	0.40	0.49	0.59	0.76
8%	0.20	0.29	0.38	0.49	0.70
10% (reference)	0.13	0.20	0.27	0.35	0.52
15%	-0.01	0.05	0.08	0.12	0.20
At profit-maximising s^*	0.00	0.24	0.29	0.34	0.44

Two features of Table 4 deserve emphasis. First, an elasticity between 0.3 and 0.5 is economically large for a variable that most firms treat as an administrative residual: it implies that halving the designated slack ratio from ten to five per cent reduces innovation output by roughly a quarter. Second, the elasticity is not a constant. Any empirical study that estimates a single linear coefficient on slack is estimating a local average over an unknown mixture of positions on this curve, which is a plausible explanation for the contradictory findings that the two recent syntheses [4], [5] document.

C. Where the private optimum lands

Across the simulated space the profit-maximising slack ratio under present technology has a median of 8.9 per cent of paid professional time, with an interquartile range of 7.3 to 10.3 per cent. The innovation-maximising ratio has a median of 18.5 per cent. The wedge between the two — a little under ten percentage points of paid time — is the quantitative content of Proposition P3: a profit-maximising firm should be expected to run at roughly half the slack ratio that would maximise its own innovation output, and this is a rational choice, not a failure of management.

Under the AI-augmented technology the median falls to 7.4 per cent, with the interquartile range moving down to 6.0 to 8.7 per cent. In 13.0 per cent of present-technology draws and 9.9 per cent of AI-augmented draws the optimum is a corner at zero designated slack, which is the spike at the left of Figure 6 and the formal content of Proposition P5.

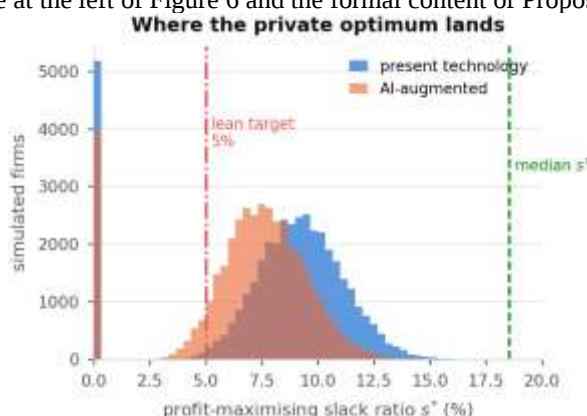


Fig. 6 Distribution of the profit-maximising slack ratio under present and AI-augmented technology across 40,000 draws. The mass at zero is the corner solution.

D. The lean AI-optimised operating model

Table 5 and Figure 5 compare the three staffing policies within each technology. The comparison is the paper's answer to the question posed in the title.

Table 5: Innovation output and value created by technology and staffing policy (medians; present-technology profit-maximising policy = 100)

Technology	Staffing policy	Slack	Innov.	Value
Present	Lean target	5.0%	78.6	97.9
Present	Profit-maximising	8.9%	100.0	100.0
Present	Innovation-maximising	18.5%	111.4	94.8
AI-augmented	Lean target	5.0%	87.2	107.6
AI-augmented	Lean + tight scheduling	5.0%	88.1	107.8
AI-augmented	Profit-maximising	7.4%	100.0	109.0
AI-augmented	Innovation-maximising	18.5%	113.7	102.2

Three readings follow. The first is that the lean five per cent target is not, on this model, a trade-off at all: it is dominated. Moving from the lean target to the re-optimised ratio raises both innovation output and firm value, and does so in 99.9 per cent of draws. The lean policy is not buying efficiency at the price of innovation; it is giving up innovation for almost nothing.

The second is the magnitude of the asymmetry. In median terms the re-optimised policy is worth 1.1 index points of value and 12.9 index points of innovation output relative to the lean target: the innovation gain is roughly twelve times the value gain. Expressed as a ratio, the lean policy retains 99 per cent of the private value available under AI-augmented technology while delivering 87 per cent of the innovation output. This is Proposition P4 in numbers, and it explains how a rational, well-run firm arrives at a policy that is wrong: the loss it can see is a rounding error and the loss it cannot see is not.

The third is that AI-augmented technology raises firm value under every policy, by between 7 and 10 index points, and that the lean policy captures most of that gain. The efficiency case for AI adoption is therefore correct on its own terms. What is not correct is the inference that the freed time should be reclaimed rather than redesignated.

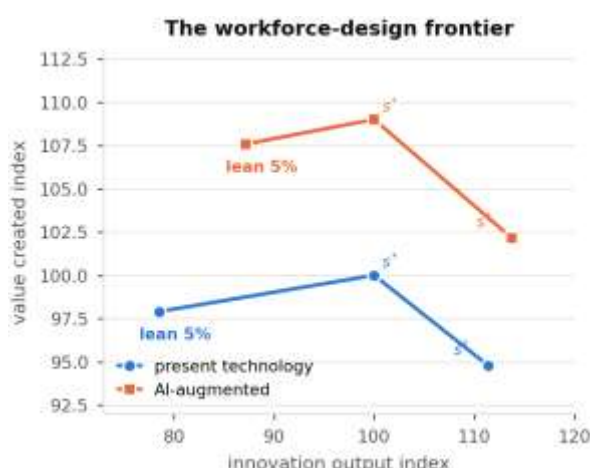


Fig. 5 The workforce-design frontier. Each line traces one technology through the three staffing policies; the lean target lies inside the frontier under both technologies.

E. Which AI channel moves the optimum

Table 6 decomposes the movement in the optimal slack ratio at the central calibration by switching on one AI channel at a time. The result is not what the public debate assumes.

Table 6: Channel decomposition of the AI effect at the central calibration

Step	s*	Innov.	Value	ϵ at s*
A. Present technology	8.48%	0.953	1.173	0.330
B. + delivery gain $\theta = 1.14$, full pass-through	7.96%	0.933	1.301	0.360
C. + creative gain $\lambda = 1.12$	8.20%	0.971	1.309	0.357
D. + fragmentation relief, so falls 30%	6.78%	0.940	1.317	0.305
E. + algorithmic scheduling, η rises 15%	6.66%	0.942	1.318	0.299

The delivery-productivity channel — the one that dominates discussion of AI and staffing — moves the optimal slack ratio by only about half a percentage point, from 8.48 to 7.96 per cent, because it raises the opportunity cost of an uncommitted hour but leaves the innovation block untouched. The creative-augmentation channel pushes back in the opposite direction. The largest single mover is the fragmentation-relief channel, which cuts the optimum by 1.4 percentage points — and it does so for a reason that is easy to misread. When AI lowers the contiguity threshold, less designated time is needed to produce the same usable creative blocks, so the firm rationally designates less. That is a genuine efficiency gain in the innovation block, not a reduction in innovative capacity: innovation output at step D is barely below step A despite a 20 per cent smaller slack allocation. Algorithmic scheduling adds only a further 0.12 percentage points.

The implication for practice is direct. A firm that reduces designated slack because AI has genuinely removed fragmentation is behaving optimally. A firm that reduces designated slack because AI has raised billable throughput is trading roughly twelve units of innovation for one unit of value. Distinguishing the two requires knowing which channel is actually operating, and utilisation dashboards cannot make that distinction.

F. The private-social wedge

Because innovation output spills over, the socially optimal slack ratio exceeds the private one. Table 7 and Figure 8 quantify the wedge at the central AI-augmented calibration and across the simulated space. At the conservative multiplier of two implied by Bloom, Schankerman and Van Reenen [34], the socially optimal ratio rises from 6.7 to 9.0 per cent; at a multiplier of three

it reaches 10.4 per cent; at five, 12.0 per cent. In the full Monte Carlo the corresponding medians are 10.2, 11.8 and 13.6 per cent.

Table 7: Private and social optima under alternative spillover multipliers (central AI-augmented calibration)

Spillover multiplier μ	Optimal slack	Innovation	Private value
1.0 (private optimum)	6.66%	0.942	1.3184
2.0 (Bloom et al. floor)	9.04%	1.015	1.3109
3.0	10.40%	1.038	1.3017
5.0 (Jones–Summers floor)	12.00%	1.056	1.2883
Innovation-maximising	18.00%	1.067	1.2230

The economically important number in Table 7 is not the level but the exchange rate. Moving from the private optimum to the multiplier-three social optimum costs 1.27 per cent of firm value and buys 10.2 per cent more innovation output — an exchange rate of about eight to one. Moving further, all the way to the innovation-maximising ratio, costs a further 6.0 per cent of firm value for only 2.8 per cent more innovation, an exchange rate of less than one to two. There is therefore a well-defined intermediate region, roughly nine to twelve per cent of paid professional time, in which additional slack is cheap in value terms and productive in innovation terms, and it lies entirely above where firms are currently being pushed.

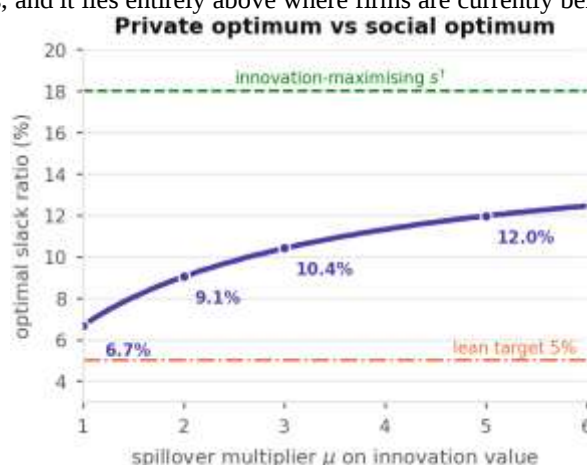


Fig. 8 The socially optimal slack ratio as a function of the spillover multiplier, against the private optimum, the lean target and the innovation-maximising ratio.

G. What actually drives the optimum

Figure 7 and Table 8 report Spearman rank correlations between each parameter and the profit-maximising slack ratio under AI-augmented technology. The ordering is instructive and partly counter-intuitive.

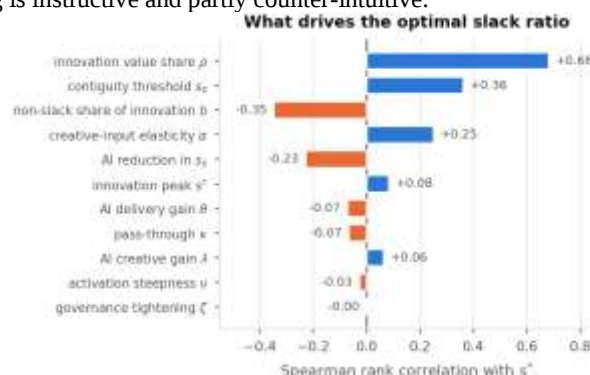


Fig. 7 Rank-correlation sensitivity of the profit-maximising slack ratio to each model parameter, AI-augmented technology.

Table 8: Sensitivity of the profit-maximising slack ratio (Spearman rank correlation)

Parameter	Correlation	Reading
ρ innovation value share	+0.68	Dominant driver
s_0 contiguity threshold	+0.36	Work design matters as much as strategy
b non-slack share of innovation	-0.35	Strong central R&D substitutes for slack
α creative-input elasticity	+0.25	Knowledge production function
Δs_0 AI fragmentation relief	-0.23	Efficiency in the innovation block
$s^\dagger$ innovation peak	+0.08	Weak
θ AI delivery gain	-0.07	Weak
κ pass-through	-0.07	Weak
λ AI creative gain	+0.06	Weak
v activation steepness	-0.03	Negligible
ζ governance tightening	-0.00	Negligible

The size of the AI productivity shock is almost irrelevant to how much slack a firm should designate. What determines the answer is how much the firm's innovation output is worth relative to its delivery revenue, how fragmented its professionals' time already is, and how much of its innovation comes from a dedicated function rather than from the general workforce. This is a useful corrective: the workforce-design decision that AI appears to force is in fact governed by variables that predate AI entirely, and a firm that reasons from the size of its AI productivity gain to its staffing slack policy is reasoning from the wrong parameter.

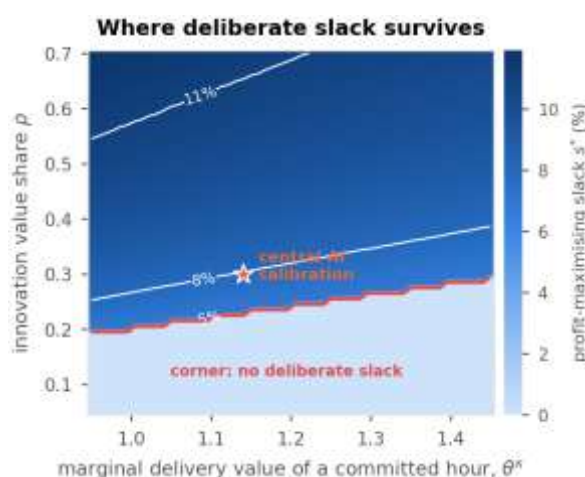


Fig. 9 Profit-maximising slack ratio over the innovation value share and the marginal delivery value of a committed hour. Below the red boundary the optimum is a corner at zero.

Figure 9 makes the corner result visible. Below an innovation value share of roughly 0.20 the optimum collapses discontinuously to zero: there is no gradual decline through one and two per cent, because a token allocation cannot clear the contiguity threshold and therefore returns nothing while still costing delivery time. The boundary shifts upward as the marginal value of a committed hour rises, which is exactly what AI does. Firms sitting slightly above the boundary today can be pushed across it by a productivity improvement that in every other respect is good for them.

H. A cross-country illustration

The innovation value share ρ is the dominant parameter, and it varies systematically across economies with the intensity of research activity. Figure 10 and Table 9 map observed 2024 research inputs into the model by scaling ρ with each economy's R&D expenditure as a share of GDP, holding all other parameters at their central values. This is an illustration of the model's logic, not an estimate of any country's actual slack ratio, and it should be read as such.

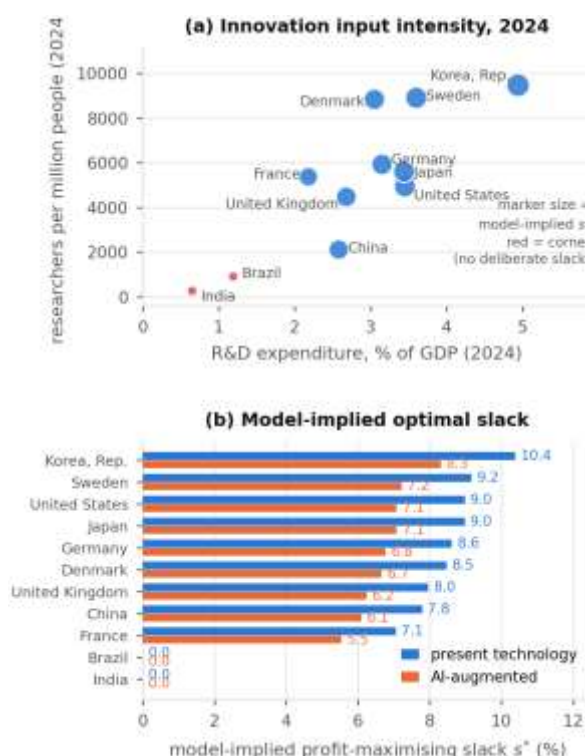


Fig. 10 (a) Research inputs across eleven economies, 2024. (b) Model-implied profit-maximising slack ratio under present and AI-augmented technology. Illustrative calibration; ρ scaled by R&D intensity. Data: UNESCO Institute for Statistics via Our World in Data [38].

Table 9: Cross-country illustrative calibration (2024 research inputs; ρ scaled by R&D intensity)

Economy	R&D/GD P	Researchers per mn	ρ	s^* now	s^* AI
Korea, Rep.	4.94	9,472	0.49	10.4%	8.3%
Sweden	3.60	8,911	0.35	9.2%	7.2%
United States	3.45	4,937	0.34	9.0%	7.1%
Japan	3.44	5,609	0.34	9.0%	7.1%
Germany	3.15	5,926	0.31	8.6%	6.8%
Denmark	3.05	8,838	0.30	8.5%	6.7%
United Kingdom	2.68	4,473	0.26	8.0%	6.2%
China	2.58	2,107	0.25	7.8%	6.1%
France	2.18	5,369	0.21	7.1%	5.5%
Brazil	1.19	903	0.12	0.0%	0.0%
India	0.65	259	0.06	0.0%	0.0%

The pattern is a slack trap. In economies where research intensity is high, the private return to designated slack is high enough that firms will fund it themselves, and the model puts the optimum between eight and ten and a half per cent of paid professional time. In economies where research intensity is low — the operating condition documented for resource-constrained technology and MSME firms [32] — the private return falls below the threshold at which any allocation clears the contiguity filter, and the privately rational allocation is zero. The mechanism is circular: low aggregate research intensity depresses the value of firm-level innovation output, which removes the private case for the slack that would raise it. AI widens the gap, because it lowers the implied optimum everywhere and therefore pushes marginal economies further into the corner. This is precisely the configuration in which the private–social wedge of Section VI-F becomes a policy argument rather than an observation, and it is worth noting that the fastest-growing patent jurisdictions are not the ones with the highest research intensity: India recorded 19.1 per cent growth in patent applications by applicant residence in 2024, the strongest among the top twenty origins, from the lowest research-input base in Table 9 [39].

VII. DISCUSSION

A. Why the innovation loss is invisible

The dominant finding of this paper is not that slack is valuable. It is that the private objective function is flat in the region where the innovation function is steep. A chief operating officer who reduces designated slack from the profit-maximising ratio to a lean five per cent target sees a value change of about one per cent — well inside the noise of any quarterly measurement — and

an innovation change of about thirteen per cent that will not be measurable for several years and will never be attributable to the staffing decision when it is. The asymmetry is not a failure of rationality; it is a failure of instrumentation. Utilisation is measured weekly with three significant figures. The elasticity in Table 4 is measured nowhere.

This reframes the practical problem. The remedy is not to argue that slack is good — the literature has done that for sixty years without changing behaviour — but to make the innovation consequence of a staffing decision visible at the moment the decision is taken. Equation (10) is designed for exactly that purpose: it requires only a local elasticity and the ratio of innovation value to delivery revenue, both of which a firm can estimate from its own new-product revenue history.

B. Slack is lumpy, not marginal

The corner result deserves separate emphasis because it contradicts how slack policies are usually introduced. Organisations that want to encourage innovation frequently begin with a small allowance — half a day a month, five per cent of time — on the reasonable assumption that a small commitment is a safe first step. The model says the opposite. A token allocation sits below the contiguity threshold, so it produces fragmented time that yields nothing while still costing delivery hours; it is dominated by both zero and a serious allocation. The finding aligns with the evidence that initiative responds only to discretion and resource together [31], and with the observation that constraint changes an organisation's default setting rather than merely its output level [26]. It is also consistent with the widely observed pattern in which discretionary-time programmes are introduced, produce nothing, and are withdrawn as evidence that slack does not work. The model's reading is that they were never above the threshold.

It follows that the contiguity threshold s_0 is a first-order design variable, and Table 8 confirms it: its rank correlation with the optimum is 0.36, five times that of the AI productivity gain. A firm that reduces meeting load, consolidates interruption, and protects whole days rather than scattered hours lowers s_0 and can then achieve the same innovation output with less designated slack. Work design and slack policy are substitutes, and the cheaper lever is usually work design.

C. Design rules for slack that survives an AI-instrumented workforce plan

Table 10 translates the model's parameters into levers. The logic throughout is that slack policy should be set on the parameters that actually govern the optimum, not on the utilisation target that is easiest to observe.

Table 10: From model parameters to managerial and policy levers

Parameter	Managerial lever	Policy lever
ρ innovation value share	Measure and publish revenue from products introduced in the last three years; make it a board metric	Innovation-output disclosure standards
s_0 contiguity threshold	Consolidate meetings; protect whole days; route interruption through AI triage rather than people	Working-time regulation on fragmentation
η dissipation rate	Light-touch accountability for slack output — a visible artefact, not a timesheet	—
b non-slack innovation share	Decide explicitly whether innovation is a function or a distributed capability; do not fund both weakly	Public research funding mix
s slack ratio	Set a floor, not a target; never introduce slack below the contiguity threshold	Innovation-time tax credit on designated hours
μ spillover multiplier	—	R&D tax credit extended from spend to designated innovation time

Three rules follow directly from the results. First, set slack as a floor rather than a target: the cost of exceeding the optimum by two percentage points is a fraction of the cost of falling two points short, because the value curve is flat above the optimum while the innovation curve is steep below it. Second, never introduce a slack allocation below the contiguity threshold; if the



organisation cannot commit to a serious allocation, the model's advice is to commit to none and concentrate innovation in a dedicated function instead. Third, when AI frees professional time, make the redesignation decision explicitly rather than by default, because the default in an instrumented workforce plan is reclamation.

D. Reconciling the curvilinearity debate

The model offers a reconciliation of the contradictory empirical record that both recent syntheses report [4], [5]. If the true relationship is the one in Figure 3, then a study sampling firms clustered below the contiguity threshold will find slack insignificant or negative; a study sampling firms in the eight-to-twelve per cent region will find it strongly positive; a study sampling firms above the innovation peak will find it negative; and a study pooling all three will find an inverted U whose turning point is an artefact of its sample composition. The disagreement in the literature may therefore be less a disagreement about the world than an artefact of estimating a single coefficient on a variable whose elasticity ranges from 0.16 to 0.49 to below zero across the relevant domain.

This also suggests the empirical design that would test the model most sharply: not a cross-section of slack levels against innovation counts, but a panel that observes changes in designated time policy — the kind of exogenous variation that Agrawal and colleagues [14] exploited through academic calendars and that Mellahi and Wilkinson [12] exploited through downsizing events — and estimates the elasticity locally rather than globally.

VIII. LIMITATIONS AND FUTURE RESEARCH

Four limitations bound the claims made here. The first is that the model is calibrated rather than estimated. Its parameter ranges are anchored in published work, and the Monte Carlo design propagates the resulting uncertainty into every reported quantity, but no result in this paper is an estimate of any actual firm's elasticity. The appropriate reading of the numbers is as the implication of a stated mechanism under stated ranges, not as measurement.

The second is that the innovation block is deliberately reduced. Effective creative time enters through a single index; the model does not distinguish exploration from exploitation in March's sense [40], does not model the stock of unexecuted ideas that Agrawal and colleagues [14] argue slack time draws down, and treats the knowledge-production elasticity as constant. Extending the model to a two-stage structure in which slack first accumulates and then executes an idea stock would allow the timing of innovation output to be modelled and would probably raise the optimal ratio, because option value would be added to expected output.

The third is that the AI channels are represented parametrically. The delivery augmentation θ is anchored on a single well-identified study in a customer-support setting [18]; the creative augmentation λ and the fragmentation relief have no comparably clean anchor. As firm-level evidence accumulates — and the microdata now being released from the Business Trends and Outlook Survey [6] is the most likely source — these should be replaced by measured values.

The fourth is that the model treats the firm as choosing a single slack ratio for a homogeneous professional workforce. In practice slack is distributed unevenly, and Chen, Park and Rajwani's finding [16] that it is specifically highly-skilled slack that converts implies that the allocation of slack across skill groups is itself a design variable. A heterogeneous-agent version of this model would be a natural next step, and would speak to whether concentrated slack among senior professionals dominates uniform slack across the workforce.

Beyond these, three empirical projects follow directly. First, the elasticity in Table 4 can be estimated on firms that have changed a formal discretionary-time policy, using the policy change as the treatment. Second, the contiguity threshold s_0 is measurable from calendar and workflow telemetry, which many organisations now hold, and its estimation would be a contribution independent of the rest of the model. Third, the cross-country mapping in Table 9 should be replaced by direct measurement of designated non-chargeable professional time, which some professional-services and IT-services firms already report internally.

IX. CONCLUSION

This paper set out to test the economic case for maintaining deliberate staffing slack against the lean, AI-optimised operating model now spreading through knowledge-intensive sectors. The answer the model returns is more precise than either side of the current debate expects.

The elasticity of innovation output with respect to deliberate staffing slack is substantial but sharply non-constant: a median of 0.49 at a five per cent slack ratio, 0.27 at ten per cent, and indistinguishable from zero by fifteen. The profit-maximising slack ratio sits at a median of 8.9 per cent of paid professional time under present technology and 7.4 per cent under an AI-augmented technology, against an innovation-maximising ratio of 18.5 per cent. Lean utilisation targets at five per cent are not a defensible trade-off but a dominated policy: they surrender roughly thirteen per cent of innovation output to capture about one per cent of firm value. That they nevertheless look attractive is a consequence of the geometry — the value curve is flat where the innovation curve is steep — and not of any error in the reasoning of the managers who adopt them.



Three findings should change how the decision is framed. Deliberate slack is lumpy: in one parameter configuration in eight the optimum is zero, and token allocations below the contiguity threshold are dominated by both zero and a serious commitment. The size of the AI productivity gain is nearly irrelevant to the optimal slack ratio, which is governed instead by the innovation value share, the fragmentation of professional time, and the strength of the dedicated research function. And the private optimum lies well below the social optimum: at the conservative spillover multipliers documented in the R&D literature, roughly one per cent of firm value buys about eight per cent more innovation output, which is the strongest case available for treating designated innovation time as a legitimate object of public support rather than a private indulgence.

The tension named in the title is real, but it is smaller and more tractable than the debate suggests. Firms are not choosing between efficiency and innovation. Most are simply operating on the flat part of one curve and the steep part of another, without the instrumentation to know it.

REFERENCES

1. R. M. Cyert and J. G. March, *A Behavioral Theory of the Firm*. Englewood Cliffs, NJ: Prentice-Hall, 1963.
2. H. Leibenstein, "Allocative efficiency vs. 'X-efficiency'," *American Economic Review*, vol. 56, no. 3, pp. 392–415, 1966.
3. M. C. Jensen, "Agency costs of free cash flow, corporate finance, and takeovers," *American Economic Review*, vol. 76, no. 2, pp. 323–329, 1986.
4. A. N. Freitas, F. R. Serra, L. Guerrazzi, V. Scazzioti, and I. C. Scafuto, "Evolution of research on resource slack and future directions," *Management Review Quarterly*, 2025, doi: 10.1007/s11301-025-00540-6.
5. N. I. Depino-Besada, X. H. Vázquez, A. Sartal, and L. López-Manuel, "An integrative review and research agenda on organizational slack: addressing assumptions, conflicted theoretical proposals and conclusions," *Management Review Quarterly*, vol. 76, no. 2, pp. 1201–1249, 2026, doi: 10.1007/s11301-025-00502-y.
6. U.S. Census Bureau, "Large firms with at least 20 employees biggest AI users," *Business Trends and Outlook Survey*, May 2026. [Online]. Available: <https://www.census.gov/library/stories/2026/05/ai-use-businesses.html>
7. J. S. Allen, "Monitoring AI adoption in the U.S. economy," *FEDS Notes*, Board of Governors of the Federal Reserve System, Apr. 3, 2026.
8. World Intellectual Property Organization, *Global Innovation Index 2025*. Geneva: WIPO, 2025.
9. E. Penrose, *The Theory of the Growth of the Firm*. Oxford, U.K.: Oxford Univ. Press, 1959.
10. L. J. Bourgeois, "On the measurement of organizational slack," *Academy of Management Review*, vol. 6, no. 1, pp. 29–39, 1981.
11. N. Nohria and R. Gulati, "Is slack good or bad for innovation?" *Academy of Management Journal*, vol. 39, no. 5, pp. 1245–1264, 1996.
12. K. Mellahi and A. Wilkinson, "A study of the association between level of slack reduction following downsizing and innovation output," *Journal of Management Studies*, vol. 47, no. 3, pp. 483–508, 2010, doi: 10.1111/j.1467-6486.2009.00872.x.
13. T. M. Amabile, C. N. Hadley, and S. J. Kramer, "Creativity under the gun," *Harvard Business Review*, vol. 80, no. 8, pp. 52–61, 2002.
14. A. Agrawal, C. Catalini, A. Goldfarb, and H. Luo, "Slack time and innovation," *Organization Science*, vol. 29, no. 6, pp. 1056–1073, 2018, doi: 10.1287/orsc.2018.1215.
15. M. Rashidi, M. A. Seyed Naghavi, B. Rezaeemanesh, and R. Vaezai, "The framework of human resource slack," *Management and Labour Studies*, vol. 50, no. 2, pp. 236–256, 2025, doi: 10.1177/0258042X241277264.
16. T. Chen, H. Park, and T. Rajwani, "Diverse human resource slack and firm innovation: evidence from politically connected firms," *International Business Review*, vol. 33, no. 2, art. 102244, 2024, doi: 10.1016/j.ibusrev.2023.102244.
17. W. M. Cohen and D. A. Levinthal, "Absorptive capacity: a new perspective on learning and innovation," *Administrative Science Quarterly*, vol. 35, no. 1, pp. 128–152, 1990.
18. E. Brynjolfsson, D. Li, and L. R. Raymond, "Generative AI at work," *Quarterly Journal of Economics*, vol. 140, no. 2, pp. 889–942, 2025.
19. N. S. K. Sastry, "An impact of departmental structure on task achievement," *EPRA International Journal of Economics, Business and Management Studies*, vol. 11, no. 9, 2024.
20. N. S. K. Sastry, "Revitalizing leadership to combat workforce burnout," *EPRA International Journal of Economics, Business and Management Studies*, vol. 11, no. 9, 2024.
21. N. S. K. Sastry and Manjula Mallya M., "Revolutionizing HR hiring with AI analytics: enhancing efficiency and attracting top talent," *EPRA International Journal of Research and Development*, vol. 10, no. 9, pp. 125–131, 2025.
22. N. S. K. Sastry and A. Selvarasu, "A study of impact on performance appraisal on employee's engagement," *International Journal of Managerial Studies and Research*, vol. 2, no. 11, pp. 10–22, 2014.
23. N. S. K. Sastry and Manjula Mallya M., "Strategic HRM for economic resilience: skill diversification, labor market dynamics and technology disruption," *EPRA International Journal of Research and Development*, vol. 10, no. 11, pp. 408–414, 2025.
24. N. S. K. Sastry, "Understanding the challenges of women micro-entrepreneurs in Karnataka," *International Journal of Managerial Studies and Research*, vol. 13, no. 8, pp. 12–16, 2025, doi: 10.20431/2349-0349.1308002.
25. P. K. Vineeth Kumar, J. J. Jijesh, H. Radhakrishna, P. G. Kulkarni, Manjula Mallya M., and N. S. K. Sastry, "Assessing the effects of enterprise resource planning on the higher education sector in India," *E3S Web of Conferences*, vol. 692, art. 03012, 2026, doi: 10.1051/e3sconf/202669203012.



26. N. S. K. Sastry and Manjula Mallya M., "Economic constraints and the status quo: barriers to innovation and creative problem-solving," *EPRA International Journal of Research and Development*, vol. 10, no. 9, pp. 191–197, 2025.
27. N. S. K. Sastry, "An impact study on fostering future leadership for innovation to meet global challenges," *International Journal of Managerial Studies and Research*, vol. 12, no. 2, pp. 23–26, 2024, doi: 10.20431/2349-0349.1202003.
28. Manjula Mallya M. and N. S. K. Sastry, "Building future-ready organizations: leadership strategies for young talent management," *EPRA International Journal of Research and Development*, vol. 10, no. 9, pp. 88–91, 2025.
29. N. S. K. Sastry and Manjula Mallya M., "Talent designation and social comparison: economic drivers of organizational growth – an empirical analysis," *EPRA International Journal of Economics, Business and Management Studies*, vol. 12, no. 10, 2025.
30. N. S. K. Sastry and P. Y. Ravish, "Empowering employees in the Indian workplace: the transformative power of delegation and resources," *EPRA International Journal of Research and Development*, vol. 10, no. 7, 2025.
31. N. S. K. Sastry and P. Y. Ravish, "The impact of delegated authority and resource provision on decision-making – fostering employee initiative," *EPRA International Journal of Multidisciplinary Research*, vol. 11, no. 7, 2025.
32. N. S. K. Sastry and Manjula Mallya M., "Women employee perspective strategies on Bangalore IT startups: assessing MSME frameworks, challenges and opportunities," *EPRA International Journal of Research and Development*, vol. 10, no. 7, 2025.
33. N. S. K. Sastry and T. Arunachalam, "Probabilistic safe RL: modelling and managing risk in dynamic environments," *EPRA International Journal of Research and Development*, vol. 10, no. 9, pp. 83–87, 2025.
34. N. Bloom, M. Schankerman, and J. Van Reenen, "Identifying technology spillovers and product market rivalry," *Econometrica*, vol. 81, no. 4, pp. 1347–1393, 2013.
35. B. F. Jones and L. H. Summers, "A calculation of the social returns to innovation," National Bureau of Economic Research, Cambridge, MA, Working Paper 27863, 2020.
36. Z. Griliches, "Patent statistics as economic indicators: a survey," *Journal of Economic Literature*, vol. 28, no. 4, pp. 1661–1707, 1990.
37. N. Bloom, C. I. Jones, J. Van Reenen, and M. Webb, "Are ideas getting harder to find?" *American Economic Review*, vol. 110, no. 4, pp. 1104–1144, 2020.
38. UNESCO Institute for Statistics, "Researchers in R&D per million people" and "Research and development expenditure (% of GDP)," 2024 data, processed by Our World in Data. [Online]. Available: <https://ourworldindata.org>
39. World Intellectual Property Organization, *World Intellectual Property Indicators 2025*. Geneva: WIPO, 2025.
40. J. G. March, "Exploration and exploitation in organizational learning," *Organization Science*, vol. 2, no. 1, pp. 71–87, 1991.