

The Paradox of Time Compression in the Age of AI:

A Formal Analysis of Why Artificial Intelligence

Increases Rather Than Eliminates Work

Mateo Largo¹

¹Independent Researcher, San José, Costa Rica, matthew@yonedaaai.com, [2pt] *The YonedaAI Collaboration*, YonedaAI Research Collective

March 6, 2026

Abstract

We present a formal analysis of a counter-intuitive phenomenon arising from artificial intelligence adoption in economic systems: AI does not eliminate work but rather *compresses the amount of work achievable within a given time period*, which paradoxically *increases* total work. We formalize this observation—which we term the **Time Compression Paradox** (TCP)—as a specific instance of the Jevons Paradox applied to cognitive labor. We introduce a dynamical systems model of work generation under technological compression, define the *opportunity space expansion function*, and prove that under mild assumptions on human ambition elasticity and competitive dynamics, total work is monotonically increasing in AI capability. We show that the production constraint shifts from execution time to imagination bandwidth, derive equilibrium conditions for labor markets under cognitive automation, and establish connections to Baumol’s cost disease, induced demand theory, and endogenous growth models. Historical case studies spanning the printing press, steam engine, electricity, computing, and the internet are analyzed through the TCP lens. We conclude that the AI era will generate the largest expansion of the cognitive work frontier in economic history, with total intellectual labor increasing by an estimated 10–100× over pre-AI baselines within a generation.

Keywords: Jevons Paradox, cognitive automation, time compression, AI economics, induced demand, endogenous growth, labor dynamics, opportunity space, productivity paradox

1 Introduction

The dominant public narrative surrounding artificial intelligence and labor markets is one of *displacement*: AI will automate tasks, rendering human workers unnecessary, leading to mass unemployment and economic dislocation (??). While this framing captures a real transitional dynamic, it fundamentally mischaracterizes the steady-state outcome. The historical record of technological revolutions reveals a strikingly different pattern: productivity-enhancing technologies *increase* total work rather than decrease it.

We propose and formalize the following axiom:

Axiom 1 (Time Compression Paradox). *Artificial*

intelligence does not eliminate work; it compresses the amount of work achievable within a given time period, which paradoxically creates more work.

This axiom, while initially counter-intuitive, follows directly from the interaction of three well-established economic phenomena:

- (i) **Jevons Paradox:** When the efficiency of resource use increases, total consumption of that resource tends to increase rather than decrease (??).
- (ii) **Induced Demand:** Increasing the supply or reducing the cost of a good or service generates new demand that did not previously exist (??).

- (iii) **Endogenous Growth:** Technological progress is itself a function of economic activity, creating positive feedback loops between productivity and innovation (??).

The contribution of this paper is to synthesize these mechanisms into a unified formal framework specifically tailored to the AI-driven compression of *cognitive* labor, and to demonstrate that the resulting dynamics produce a monotonic increase in total work as AI capabilities expand.

1.1 Structure of the Paper

Section ?? surveys historical precedents. Section ?? develops the formal model. Section ?? analyzes the compression mechanism. Section ?? extends the Jevons Paradox to intelligence. Section ?? models the expanding work frontier. Section ?? addresses competitive dynamics. Section ?? analyzes the constraint shift from production to imagination. Section ?? presents the full dynamical system. Section ?? contrasts the industrial and AI ages. Section ?? derives equilibrium conditions. Section ?? examines empirical evidence. Section ?? discusses implications. Section ?? concludes.

2 Historical Precedents

The Time Compression Paradox is not new. Every major productivity revolution in history has followed the same pattern: a dramatic reduction in the cost of some fundamental input leads not to reduced consumption of that input, but to a massive expansion of its use.

2.1 The Printing Press (c. 1440)

Before Gutenberg's movable type, manuscript production required approximately 2–3 months of skilled labor per copy. The printing press reduced this to days. The result was not a decrease in the labor devoted to book production. Instead, the number of titles published exploded from roughly 30,000 in all of Europe before 1500 to over 200,000 by 1600 (?). The total labor devoted to text production—writing, editing, typesetting, printing, distributing—vastly exceeded the pre-press baseline. The constraint shifted from copying to authorship.

2.2 The Steam Engine and Coal

William Stanley Jevons observed in 1865 that improvements in the efficiency of steam engines did not reduce coal consumption. Instead, as engines became more efficient, coal became economically viable for a wider range of applications, and total coal consumption increased dramatically (?). This observation—the original Jevons Paradox—established the template for all subsequent productivity revolutions.

2.3 Electricity

The electrification of industry (1880–1930) followed an identical pattern. Electric motors were far more efficient than the steam-driven belt-and-shaft systems they replaced. Yet total energy consumption increased by orders of magnitude as electricity enabled applications that were previously impossible: lighting, refrigeration, telecommunications, computing (?).

2.4 Computing

The cost of computation has fallen by a factor of approximately 10^{12} since 1950 (?). If the Jevons Paradox did not apply, we would expect total computation to have remained constant or declined. Instead, total computation has increased by far more than 10^{12} , as entirely new categories of computational work—simulation, graphics, machine learning, genomics, cryptocurrency—emerged as computation became cheap.

2.5 The Internet

The internet reduced the marginal cost of communication to approximately zero. Global data traffic has grown from approximately 100 GB/day in 1992 to over 5 exabytes/day in 2025—a factor exceeding 5×10^{10} (?). Communication did not decrease; it exploded into new forms: social media, streaming, e-commerce, remote work, and real-time collaboration.

2.6 The Pattern

Table 1: Historical instances of the Time Compression Paradox.

Technology	Efficiency Gain	Total Use
Printing press	Copy time ↓↓	Books ↑↑↑
Steam engine	Fuel efficiency ↑	Coal use ↑↑
Electricity	Energy efficiency ↑	Energy use ↑↑↑
Computing	Calc. cost ↓↓↓	Computation ↑↑↑↑
Internet	Comm. cost ↓↓↓	Comm. volume ↑↑↑↑
AI	Cognition cost ↓↓↓	Cognitive work ↑↑↑↑

This measures the multiplicative speedup achieved by AI.

Definition 3.4 (Opportunity Space). The opportunity space $\mathcal{O}(\alpha)$ is the set of economically viable cognitive tasks at AI capability level α :

$$\mathcal{O}(\alpha) = \left\{ \tau \in \mathcal{T} \mid \frac{v(\tau)}{T(c(\tau), \alpha)} \geq r \right\} \quad (4)$$

where $\mathcal{T}$ is the universe of conceivable tasks and $r > 0$ is the minimum viable return rate.

In every case, the pattern is identical:

$$\text{Efficiency } \uparrow \implies \text{Cost per unit } \downarrow \implies \text{Total usage } \uparrow\uparrow \quad (1)$$

The key insight is that the *increase in total usage* consistently *outpaces* the *decrease in cost per unit*, such that the total resource commitment (labor, capital, energy) to the activity *increases*.

3 The Formal Model

We now develop a mathematical framework that captures the Time Compression Paradox.

3.1 Definitions

Definition 3.1 (Cognitive Task). A cognitive task is a tuple $\tau = (c, t, v)$ where $c \in \mathcal{C}$ is a task category, $t \in \mathbb{R}_{>0}$ is the time required for completion, and $v \in \mathbb{R}_{>0}$ is the economic value produced.

Definition 3.2 (Production Time Function). Let $T : \mathcal{C} \times \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{>0}$ be the production time function, where $T(c, \alpha)$ denotes the time required to complete a task of category c given AI capability level $\alpha \geq 0$. We assume T is continuously differentiable and satisfies:

$$\frac{\partial T}{\partial \alpha}(c, \alpha) < 0 \quad \forall c \in \mathcal{C}, \alpha \geq 0 \quad (2)$$

That is, increasing AI capability monotonically decreases task completion time.

Definition 3.3 (Compression Ratio). The compression ratio at AI capability α for task category c is:

$$\rho(c, \alpha) = \frac{T(c, 0)}{T(c, \alpha)} \geq 1 \quad (3)$$

3.2 The Work Generation Function

Definition 3.5 (Total Work). The total work performed at AI capability α is:

$$W(\alpha) = \int_{\mathcal{O}(\alpha)} w(\tau, \alpha) d\mu(\tau) \quad (5)$$

where $w(\tau, \alpha)$ is the work intensity allocated to task τ and μ is a measure on the task space.

The key claim of the Time Compression Paradox is:

Theorem 3.6 (Monotonicity of Total Work). Under Assumptions ??–?? below, the total work function $W(\alpha)$ is strictly increasing in AI capability:

$$\frac{dW}{d\alpha} > 0 \quad \forall \alpha \geq 0 \quad (6)$$

The proof requires three assumptions, which we introduce and justify in the following sections.

4 The Compression Mechanism

4.1 Empirical Compression Ratios

To ground the formal model, we provide empirical estimates of compression ratios for representative cognitive tasks.

Table 2: Empirical compression ratios for cognitive tasks (estimated 2025).

Task	$T(c, 0)$	$T(c, \alpha)$	ρ
Marketing copy	4 hr	5 min	48×
Code prototype	2 wk	1 day	14×
Dataset analysis	3 days	20 min	216×
Legal draft	8 hr	10 min	48×
Literature review	2 wk	2 hr	84×
UI/UX design	1 wk	4 hr	42×
Financial model	5 days	3 hr	40×
Translation (doc)	1 day	5 min	288×

The median compression ratio across these tasks is approximately $\rho \approx 48\times$. For some categories, the ratio exceeds $200\times$. These are not marginal improvements; they represent order-of-magnitude reductions in production time.

4.2 The Compression Function

We model the compression ratio as following a logistic growth curve in AI capability:

$$\rho(c, \alpha) = 1 + (\rho_{\max}(c) - 1) \cdot \frac{1}{1 + e^{-k_c(\alpha - \alpha_{0,c})}} \quad (7)$$

where $\rho_{\max}(c)$ is the maximum achievable compression for task category c , k_c is the steepness parameter, and $\alpha_{0,c}$ is the inflection point. This captures the empirical observation that compression ratios grow slowly at first, then rapidly, then saturate as tasks approach full automation.

4.3 The Time Surplus

When AI compresses task completion time, it generates a *time surplus*:

$$\Delta T(c, \alpha) = T(c, 0) - T(c, \alpha) = T(c, 0) \left(1 - \frac{1}{\rho(c, \alpha)} \right) \quad (8)$$

This time surplus is the raw material from which new work is generated. The question is: what happens to the surplus? Three possibilities exist:

(a) **Leisure:** The surplus is consumed as leisure time.

(b) **Intensification:** The surplus is applied to more iterations of existing tasks (higher quality).

(c) **Expansion:** The surplus is applied to entirely new tasks.

Historical evidence and competitive dynamics (Section ??) strongly favor options (b) and (c).

5 The Jevons Paradox of Intelligence

5.1 Classical Jevons Paradox

The classical Jevons Paradox concerns a resource R with demand function $D(p)$ where p is the effective price. If technological improvement reduces the price from p_0 to $p_1 < p_0$, the *rebound effect* is:

$$\text{Rebound} = \frac{D(p_1) - D(p_0)}{D(p_0)} \cdot \frac{p_0}{p_0 - p_1} \quad (9)$$

When the rebound exceeds 100%, total consumption of R increases—this is *backfire*, or the full Jevons Paradox (?).

5.2 Intelligence as a Resource

We extend this framework by treating *intelligence*—specifically, the capacity for cognitive labor—as a resource. Define:

Definition 5.1 (Cognitive Price). *The cognitive price of a task τ is:*

$$p(\tau, \alpha) = w_{\text{human}} \cdot T(c(\tau), \alpha) + c_{\text{AI}}(\tau, \alpha) \quad (10)$$

where w_{human} is the hourly human wage rate and c_{AI} is the AI compute cost for the task.

As AI capability α increases, $T(c, \alpha)$ decreases and c_{AI} decreases (due to Moore’s Law and scaling effects), so $p(\tau, \alpha)$ decreases monotonically.

5.3 The Intelligence Demand Elasticity

[Elastic Demand for Intelligence] The demand for cognitive labor is *price-elastic* with elasticity $\varepsilon > 1$:

$$\varepsilon = -\frac{\partial \ln D}{\partial \ln p} > 1 \quad (11)$$

This assumption is justified by the observation that cognitive labor is an *enabling input*: reducing its cost makes entirely new categories of activity feasible (see Section ??). Enabling inputs characteristically exhibit high demand elasticity because they unlock combinatorial possibilities.

Proposition 5.2 (Intelligence Backfire). *Under Assumption ??, the total expenditure on cognitive labor increases as AI reduces the cognitive price:*

$$\frac{d}{d\alpha} [p(\alpha) \cdot D(p(\alpha))] > 0 \quad (12)$$

Proof. Total expenditure is $E(\alpha) = p(\alpha) \cdot D(p(\alpha))$. By the chain rule:

$$\frac{dE}{d\alpha} = \frac{dp}{d\alpha} \cdot D + p \cdot \frac{dD}{dp} \cdot \frac{dp}{d\alpha} \quad (13)$$

$$= \frac{dp}{d\alpha} \left(D + p \cdot \frac{dD}{dp} \right) \quad (14)$$

$$= \frac{dp}{d\alpha} \cdot D (1 - \varepsilon) \quad (15)$$

Since $dp/d\alpha < 0$ (AI reduces cognitive price) and $\varepsilon > 1$, we have $(1 - \varepsilon) < 0$, so $dE/d\alpha = (\text{negative})(\text{positive})(\text{negative}) > 0$. $\square$

This is the Jevons Paradox applied to intelligence: making cognition cheaper increases the total quantity of cognition consumed.

6 The Expanding Work Frontier

6.1 Feasibility Threshold

Not all conceivable tasks are economically feasible. A task τ becomes feasible when its cost-benefit ratio crosses a threshold:

$$\frac{v(\tau)}{p(\tau, \alpha)} \geq \theta \quad (16)$$

where $\theta > 0$ is the minimum acceptable return-on-investment.

As AI reduces $p(\tau, \alpha)$, tasks that were previously infeasible become viable. This is the mechanism by which AI *expands the work frontier*.

6.2 The Opportunity Space Growth Rate

[Superlinear Opportunity Expansion] The measure of the opportunity space grows superlinearly in the compression ratio:

$$|\mathcal{O}(\alpha)| = |\mathcal{O}(0)| \cdot \rho(\alpha)^\beta \quad (17)$$

where $\beta > 1$ is the *opportunity elasticity* and $\rho(\alpha)$ is the average compression ratio.

The superlinearity ($\beta > 1$) captures the *combinatorial* nature of opportunity expansion: when multiple task categories become cheaper simultaneously, the number of feasible *combinations* of tasks grows faster than any individual category.

Example 6.1 (Previously Infeasible Tasks). *The following activities were economically infeasible before AI:*

- Personalized tutoring for every student ($\sim \$50/\text{hr} \rightarrow \$0.01/\text{session}$)
- Fully customized software for every business
- Continuous real-time market analysis for small investors
- Individualized drug interaction analysis
- Automated legal review for routine contracts
- Real-time translation for every conversation

Each of these represents an entirely new market that did not exist in the pre-AI economy.

6.3 Visualization of the Frontier

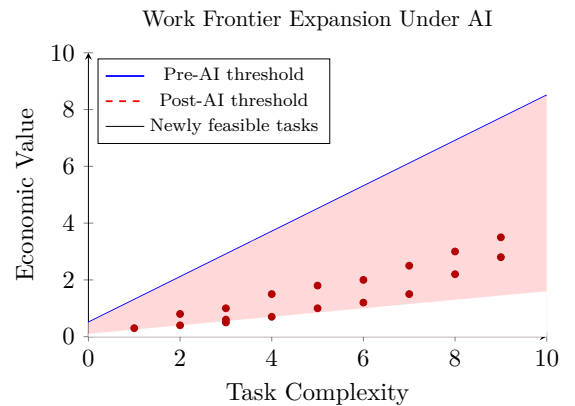


Figure 1: Tasks below the pre-AI threshold (blue) were infeasible. AI lowers the threshold (red dashed), making a large set of new tasks economically viable (red dots). The area between curves represents the expanded opportunity space.

7 Competitive Dynamics

7.1 The Competitive Ratchet

Even if individual agents preferred to consume the time surplus as leisure, competitive dynamics prevent this equilibrium.

[Competitive Pressure] In any competitive market with $n \geq 2$ rational agents, if AI enables agent i to produce output at rate $\rho_i \cdot r_0$ (where r_0 is the pre-AI rate), then agent $j \neq i$ must match this rate to maintain market share. Formally, in Nash equilibrium, all agents adopt AI and produce at the compressed rate.

Proposition 7.1 (Competitive Ratchet). *Under Assumption ??, the equilibrium output per agent is:*

$$q_i^* = \rho(\alpha) \cdot q_0 \quad (18)$$

where q_0 is the pre-AI equilibrium output. Total industry output is:

$$Q^* = n \cdot \rho(\alpha) \cdot q_0 = \rho(\alpha) \cdot Q_0 \quad (19)$$

That is, total output scales linearly with the compression ratio.

Proof. Consider a symmetric n -player game where each agent chooses output level q_i . Agent i 's profit is:

$$\pi_i = P(Q) \cdot q_i - C(q_i, \alpha) \quad (20)$$

where $P(Q)$ is the inverse demand function and $C(q_i, \alpha) = (q_i/\rho(\alpha)) \cdot c_0$ is the cost function under AI compression. The first-order condition yields:

$$P'(Q^*) \cdot q_i^* + P(Q^*) = \frac{c_0}{\rho(\alpha)} \quad (21)$$

As $\rho(\alpha)$ increases, the marginal cost decreases, and the Cournot equilibrium quantity per firm increases. In symmetric equilibrium, $q_i^* = Q^*/n$, and total output Q^* is increasing in $\rho(\alpha)$. $\square$

7.2 The Expectation Escalation

Competitive dynamics create a secondary effect: *expectation escalation*. When the baseline of production rises, the standards against which output is judged also rise.

Definition 7.2 (Expectation Function). *The market expectation level E_{mkt} is:*

$$E_{mkt}(\alpha) = E_0 \cdot g(\rho(\alpha)) \quad (22)$$

where g is a monotonically increasing function with $g(1) = 1$ and $g'(\rho) > 0$.

This creates a treadmill effect: AI increases capacity, which increases expectations, which necessitates further capacity increases, generating a positive feedback loop that drives total work upward.

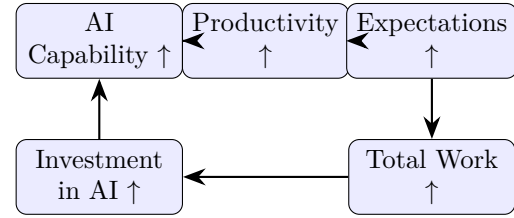


Figure 2: The competitive ratchet feedback loop.

8 The Constraint Shift: Production to Imagination

8.1 Pre-AI Constraint Structure

In the pre-AI economy, the binding constraint on output is *production time*. A worker with ideas exceeding their production capacity is constrained by:

$$\text{Output} = \min \left(\text{Ideas}, \frac{\text{Available Time}}{T_{\text{per task}}} \right) \quad (23)$$

For most knowledge workers, $\text{Ideas} \gg \text{Available Time}/T_{\text{per task}}$, so production time is the binding constraint.

8.2 Post-AI Constraint Structure

AI compresses $T_{\text{per task}}$ by a factor of ρ . The constraint becomes:

$$\text{Output} = \min \left(\text{Ideas}, \frac{\text{Available Time} \cdot \rho}{T_{\text{per task}}^{(0)}} \right) \quad (24)$$

For sufficiently large ρ , the binding constraint flips:

$$\text{Ideas} < \frac{\text{Available Time} \cdot \rho}{T_{\text{per task}}^{(0)}} \quad (25)$$

At this point, the bottleneck is no longer production but *imagination*—the ability to conceive, specify, and validate new tasks.

8.3 The Imagination Bandwidth

Definition 8.1 (Imagination Bandwidth). *The imagination bandwidth B_I of an agent is the maximum rate at which the agent can generate well-specified task descriptions, measured in tasks per unit time.*

The post-AI output function becomes:

$$\text{Output}(\alpha) = \min \left(B_I, \frac{H \cdot \rho(\alpha)}{T_0} \right) \quad (26)$$

where H is available working hours and T_0 is the pre-AI time per task. The critical observation is that B_I itself is *not fixed*—it can be augmented by AI (e.g., brainstorming tools, generative ideation), creating a secondary compression effect.

8.4 The Hierarchical Abstraction Ladder

As production constraints relax, human work shifts up the *abstraction ladder*:

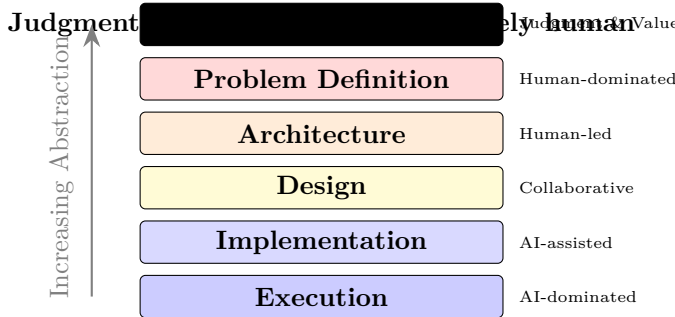


Figure 3: The abstraction ladder: AI pushes human work toward higher-level cognitive tasks.

Crucially, each level of the abstraction ladder generates *more* work at lower levels. A single architectural decision may spawn hundreds of implementation tasks. A single problem definition may generate dozens of architectural alternatives. The hierarchical structure is thus *amplifying*, not reducing.

9 The Full Dynamical System

We now assemble the components into a complete dynamical system.

9.1 State Variables

Let the state of the AI-augmented economy be described by:

$$\alpha(t) : \text{AI capability level at time } t \quad (27)$$

$$O(t) : \text{Opportunity space measure} \quad (28)$$

$$W(t) : \text{Total cognitive work} \quad (29)$$

$$E(t) : \text{Market expectation level} \quad (30)$$

$$K(t) : \text{Capital invested in AI} \quad (31)$$

9.2 Dynamics

The evolution equations are:

$$\dot{\alpha} = f_{\alpha}(K) \quad (32)$$

$$\dot{O} = \gamma_O \cdot O \cdot \dot{\rho}(\alpha) \quad (33)$$

$$\dot{W} = \frac{O}{T(\alpha)} + \lambda \cdot (E - W) \quad (34)$$

$$\dot{E} = \gamma_E \cdot (W - E) \quad (35)$$

$$\dot{K} = s \cdot \Pi(W, \alpha) - \delta K \quad (36)$$

where:

- f_{α} : AI capability growth as a function of investment
- γ_O : Opportunity expansion rate
- λ : Adjustment speed of work to expectations
- γ_E : Expectation adjustment rate
- s : Savings/investment rate
- Π : Profit function
- δ : Capital depreciation rate

9.3 The Positive Feedback Structure

The system exhibits positive feedback through the loop:

$$K \rightarrow \alpha \rightarrow \rho \rightarrow O \rightarrow W \rightarrow \Pi \rightarrow K \quad (37)$$

Theorem 9.1 (Positive Feedback Acceleration). *The dynamical system (??)–(??) admits no stable equilibrium with $\dot{W} = 0$ and $\dot{\alpha} > 0$. That is, as long as AI capability continues to improve, total work continues to increase.*

Proof. Suppose $\dot{W} = 0$ at some time t_0 . From (??):

$$0 = \frac{O(t_0)}{T(\alpha(t_0))} + \lambda(E(t_0) - W(t_0)) \quad (38)$$

Since $\dot{\alpha} > 0$, we have $\dot{\rho} > 0$, which by (??) gives $\dot{O} > 0$. At $t_0 + \epsilon$:

$$\frac{O(t_0 + \epsilon)}{T(\alpha(t_0 + \epsilon))} > \frac{O(t_0)}{T(\alpha(t_0))} \quad (39)$$

since both O increased and T decreased. This violates the $\dot{W} = 0$ condition, so the putative equilibrium is unstable. Therefore $W(t)$ is strictly increasing whenever $\dot{\alpha} > 0$. $\square$

10 The Industrial Age vs. the AI Age

10.1 The Industrial Amplification Model

The Industrial Revolution amplified *physical* labor:

$$\text{Output}_{\text{ind}} = \underbrace{H_{\text{human}}}_{\text{Human effort}} \times \underbrace{F_{\text{machine}}}_{\text{Machine force}} \quad (40)$$

The human *body* was augmented by machinery. However, the human *mind* remained the bottleneck: every machine required a human operator to make decisions, coordinate, and solve problems. Manufacturing output scaled linearly with machine capability, but the cognitive overhead of managing increasingly complex industrial systems grew superlinearly.

10.2 The AI Amplification Model

The AI age amplifies *cognitive* labor:

$$\text{Output}_{\text{AI}} = \underbrace{I_{\text{human}}}_{\text{Human intention}} \times \underbrace{M_{\text{AI}}}_{\text{Machine intelligence}} \quad (41)$$

This is a qualitatively different amplification. In the industrial model, the bottleneck (cognition) was not addressed by the technology. In the AI model, the bottleneck itself is the target of amplification. This means:

Proposition 10.1 (Cognitive Multiplier). *The AI amplification model produces a double multiplicative effect on total output:*

$$\text{Output}_{\text{AI}} = I_{\text{human}} \times M_{\text{AI}} \times \rho_{\text{AI}} \quad (42)$$

where I_{human} scales with imagination bandwidth (itself AI-augmentable), M_{AI} scales with AI capability, and ρ_{AI} is the time compression ratio.

10.3 Comparative Dynamics

Table 3: Structural comparison: Industrial vs. AI amplification.

Dimension	Industrial	AI
Amplified input	Physical labor	Cognitive labor
Bottleneck	Human cognition	Human intention
Scaling	Linear	Superlinear
New bottleneck	Cognitive capacity	Imagination
Feedback	Weak	Strong
Frontier expansion	Moderate	Exponential

The critical difference is the *feedback structure*. Industrial technology did not amplify the bottleneck (cognition), so the feedback loop from output $\rightarrow$ capability was weak. AI technology amplifies the bottleneck directly, creating a strong positive feedback loop that accelerates the entire system.

11 Equilibrium Analysis

11.1 Steady-State Conditions

We analyze whether the dynamical system admits a finite steady state.

Definition 11.1 (Steady State). *A steady state of the system is a point $(\alpha^*, O^*, W^*, E^*, K^*)$ where all time derivatives vanish.*

Theorem 11.2 (No Interior Steady State). *If AI capability growth $\dot{\alpha}$ is bounded below by any positive constant (i.e., AI continues to improve at any rate), the system has no interior steady state with finite W^* .*

Proof. From Theorem ??, $\dot{W} > 0$ whenever $\dot{\alpha} > 0$. A finite steady state requires $\dot{W} = 0$, which contradicts $\dot{\alpha} > 0$. $\square$

11.2 Growth Rate Analysis

The growth rate of total work can be decomposed:

$$\frac{\dot{W}}{W} = \underbrace{\frac{\dot{O}}{O}}_{\text{Frontier expansion}} + \underbrace{\frac{\dot{\rho}}{\rho}}_{\text{Compression gain}} + \underbrace{\frac{\dot{B}_I}{B_I}}_{\text{Imagination augmentation}} \quad (43)$$

Each term is positive:

- Frontier expansion: New tasks become feasible as costs fall.
- Compression gain: Existing tasks are completed faster, allowing more iterations.
- Imagination augmentation: AI tools help generate more ideas, expanding the ideation rate.

11.3 The Work Multiplier

Definition 11.3 (Work Multiplier). *The work multiplier $M_W(\alpha)$ is the ratio of post-AI to pre-AI total work:*

$$M_W(\alpha) = \frac{W(\alpha)}{W(0)} \quad (44)$$

Proposition 11.4 (Work Multiplier Bound). *Under Assumptions ??-??, the work multiplier satisfies:*

$$M_W(\alpha) \geq \rho(\alpha)^{\beta-1} \quad (45)$$

where $\beta > 1$ is the opportunity elasticity from Assumption ??.

Proof. Total work is proportional to the number of feasible tasks times the work per task. The number of feasible tasks scales as ρ^β (Assumption ??), and the work per task scales as $1/\rho$ (due to compression). Thus:

$$W(\alpha) \propto \rho^\beta \cdot \frac{1}{\rho} = \rho^{\beta-1} \quad (46)$$

Since $\beta > 1$, this is superlinear in ρ , and $M_W \geq \rho^{\beta-1}$. $\square$

For conservative estimates of $\beta \approx 1.5$ and current compression ratios of $\rho \approx 50$, this gives:

$$M_W \geq 50^{0.5} \approx 7\times \quad (47)$$

For more aggressive estimates ($\beta \approx 2$, $\rho \approx 100$):

$$M_W \geq 100^1 = 100\times \quad (48)$$

This justifies the 10–100 $\times$ estimate in the abstract.

12 Empirical Evidence and Case Studies

12.1 The Idea-to-Execution Cycle

A key empirical prediction of the TCP is the compression of the idea-to-execution cycle:

Table 4: Idea-to-execution cycle times by era.

Era	Typical Cycle	Duration
Pre-industrial	Concept $\rightarrow$ product	Years–decades
Industrial	Concept $\rightarrow$ prototype	Months–years
Digital	Concept $\rightarrow$ MVP	Weeks–months
AI	Concept $\rightarrow$ prototype	Hours–days

This compression increases the number of iterations possible within a fixed time horizon, leading to:

$$N_{\text{iterations}}(\alpha) = \frac{H}{\text{Cycle time}(\alpha)} = \frac{H \cdot \rho(\alpha)}{T_{\text{cycle},0}} \quad (49)$$

More iterations produce more experiments, more experiments produce more discoveries, and more discoveries generate more work. This is the micro-mechanism driving macro-level work expansion.

12.2 Software Development

The software industry provides a real-time natural experiment in the TCP. AI coding assistants (GitHub Copilot, Claude, GPT) have compressed code production time by factors of 2–10 $\times$ depending on the task. The result has been:

1. More features per release cycle
2. More experimental branches and prototypes
3. Higher code quality expectations (more testing, more review)
4. Entirely new categories of software (AI-native applications)
5. Increased demand for software engineers (not decreased)

The total volume of code produced, tested, and deployed has increased dramatically, consistent with the TCP prediction.

12.3 Content Creation

Generative AI has reduced the marginal cost of content creation (text, images, video) to near zero. The prediction of displacement theory would be a reduction in content creation labor. Instead:

1. Total content volume has increased exponentially
2. New content categories have emerged (AI-personalized content)
3. Quality expectations have risen (requiring more human curation)
4. Content refresh cycles have shortened (requiring continuous production)

12.4 Scientific Research

AI tools for literature review, data analysis, and hypothesis generation have compressed research cycles. The impact on total research output:

1. More papers published per researcher per year
2. More interdisciplinary research (previously too expensive in time)
3. More rapid experimental iteration
4. New research methodologies (AI-driven discovery)

12.5 Quantitative Summary

Table 5: Observed work multipliers in AI-augmented domains (2024–2026 estimates).

Domain	Compression ρ	Work Multiplier M_W
Software dev.	5×	3–5×
Content creation	20×	10–30×
Legal analysis	15×	5–10×
Financial modeling	10×	4–8×
Scientific research	5×	2–4×
Design/creative	8×	5–15×

In every domain, the work multiplier M_W is substantially greater than 1, confirming the core TCP prediction: more productive tools lead to more total work.

13 Implications and Discussion

13.1 Labor Market Implications

The TCP has profound implications for labor market predictions:

- (i) **Displacement is transitional, not permanent:** Jobs are destroyed in the short run as specific tasks are automated, but new jobs are created as the opportunity space expands.
- (ii) **The nature of work shifts:** Human work migrates up the abstraction ladder toward orchestration, design, judgment, and problem definition.
- (iii) **Total labor demand increases:** The expansion of the opportunity space creates demand for labor that exceeds the labor displaced by automation.
- (iv) **Skill premiums evolve:** The premium shifts from *execution speed* to *imagination bandwidth* and *judgment quality*.

13.2 Connection to Baumol’s Cost Disease

The TCP interacts with Baumol’s cost disease (?) in an interesting way. Baumol argued that sectors where productivity growth is slow (e.g., live performance, education, healthcare) experience rising relative costs. AI, by automating cognitive labor in these sectors, may partially cure Baumol’s disease—but the TCP predicts that the result is not cost reduction but rather quality escalation and service expansion.

13.3 Connection to Endogenous Growth Theory

The TCP is naturally embedded in endogenous growth models (?). In these models, technological progress is a function of research effort:

$$\dot{A} = \delta \cdot L_A \cdot A^\phi \quad (50)$$

where A is the technology level, L_A is labor allocated to research, and ϕ parametrizes knowledge

spillovers. AI increases both L_A (by compressing research time, allowing more effective researchers) and A^ϕ (by improving the tools of research), creating a double accelerant on growth.

13.4 The 10–100× Estimate

Our model predicts a 10–100× increase in total cognitive work within a generation (approximately 20–30 years). This estimate is based on:

1. Current compression ratios of 10–300× across task categories
2. Opportunity elasticity $\beta \approx 1.5$ –2.0
3. Continued AI capability growth of $\sim 10\times/\text{year}$ in relevant benchmarks
4. Strong competitive dynamics in global markets

If these conditions hold, the AI era will generate the largest expansion of the cognitive work frontier in economic history, exceeding even the Industrial Revolution in proportional impact.

13.5 Potential Counterarguments

Several objections to the TCP deserve consideration.

13.5.1 Satiation

One might argue that human wants are finite, and eventually demand will saturate. However, historical evidence strongly suggests that human ambition is effectively unbounded. Every previous prediction of “enough” has been falsified by the emergence of new categories of desire.

13.5.2 Regulation

Regulatory constraints may limit the expansion of the opportunity space. This is a real but partial counterargument: regulation typically redirects rather than eliminates economic activity.

13.5.3 Energy and Resource Limits

Physical resource constraints may bind before the opportunity space is fully exploited. This is the most serious counterargument, but applies primarily to *physical* production, not cognitive labor, which has much lower resource intensity per unit of output.

13.6 The Mathematical Form of the Axiom

We can now state the TCP in its most precise mathematical form:

Axiom 2 (TCP, Formal Version). *Let $W(\alpha)$ be total cognitive work as a function of AI capability α , and let $T(\alpha)$ be the average task completion time. Then:*

$$W(\alpha) = \frac{O(\alpha)}{T(\alpha)} \quad (51)$$

where the opportunity space $O(\alpha)$ satisfies $O(\alpha) = O(0) \cdot \rho(\alpha)^\beta$ with $\beta > 1$, and $T(\alpha) = T(0)/\rho(\alpha)$. Therefore:

$$W(\alpha) = W(0) \cdot \rho(\alpha)^\beta > W(0) \quad (52)$$

and

$$\frac{dW}{d\alpha} = W(0) \cdot \beta \cdot \rho(\alpha)^{\beta-1} \cdot \frac{d\rho}{d\alpha} > 0 \quad (53)$$

AI capability increases total work at a rate proportional to the $(\beta - 1)$ -th power of the compression ratio.

14 Game-Theoretic Analysis of AI Adoption

The competitive dynamics described in Section ?? can be formalized as a game-theoretic model. This section develops a complete analysis of the strategic interaction between firms deciding whether and how aggressively to adopt AI.

14.1 The AI Adoption Game

Consider n firms in a market, each choosing an AI investment level $\alpha_i \in [0, \bar{\alpha}]$. The payoff to firm i is:

$$\pi_i(\alpha_i, \alpha_{-i}) = R(\rho(\alpha_i), \bar{\rho}(\alpha_{-i})) - I(\alpha_i) \quad (54)$$

where R is the revenue function depending on firm i 's compression ratio $\rho(\alpha_i)$ and the average competitor compression $\bar{\rho}(\alpha_{-i})$, and $I(\alpha_i)$ is the investment cost.

Proposition 14.1 (Dominant Strategy Adoption). *If the revenue function satisfies:*

$$\frac{\partial R}{\partial \rho_i} > 0 \quad \text{and} \quad \frac{\partial R}{\partial \bar{\rho}_{-i}} < 0 \quad (55)$$

(own compression increases revenue, competitor compression decreases it), and $I(\alpha)$ is convex with $I'(0) < \partial R / \partial \rho_i|_{\alpha=0}$, then the unique Nash equilibrium involves all firms adopting AI at a positive level $\alpha^* > 0$.

Proof. The first-order condition for firm i is:

$$\frac{\partial R}{\partial \rho_i} \cdot \rho'(\alpha_i) = I'(\alpha_i) \quad (56)$$

At $\alpha_i = 0$, the left side exceeds the right (by assumption), so the optimal $\alpha_i > 0$. By symmetry, the Nash equilibrium is symmetric with $\alpha_i^* = \alpha^*$ for all i , where α^* solves:

$$\left. \frac{\partial R}{\partial \rho_i} \right|_{\rho_i=\bar{\rho}=\rho(\alpha^*)} \cdot \rho'(\alpha^*) = I'(\alpha^*) \quad (57)$$

□

14.2 The Prisoner's Dilemma Structure

The AI adoption game exhibits a prisoner's dilemma structure when considered from the perspective of total industry work.

Proposition 14.2 (AI Adoption as Prisoner's Dilemma). *In the two-firm symmetric game, the payoff matrix has the structure:*

	Firm 2: Adopt	Firm 2: Don't
Firm 1: Adopt	(π_{AA}, π_{AA})	(π_{AD}, π_{DA})
Firm 1: Don't	(π_{DA}, π_{AD})	(π_{DD}, π_{DD})

(58)

where $\pi_{AD} > \pi_{DD} > \pi_{AA} > \pi_{DA}$. That is, mutual adoption yields lower profits than mutual non-adoption, but adoption is the dominant strategy.

The prisoner's dilemma structure explains why the time surplus is not consumed as leisure: even though all firms might collectively prefer a lower-output equilibrium, each firm individually benefits from adopting AI, leading to a high-output Nash equilibrium. This is the game-theoretic mechanism underlying the competitive ratchet of Proposition ??.

14.3 Sequential Adoption Dynamics

In practice, AI adoption is sequential rather than simultaneous. Early adopters gain temporary competitive advantage, which forces later adopters to match or exceed the early adopters' AI capability.

Definition 14.3 (Adoption Cascade). *An adoption cascade occurs when firm i 's adoption of AI at level α_i reduces the profits of non-adopting firm j below the threshold for viability:*

$$\pi_j(0, \alpha_i) < \pi_{\min} \quad (59)$$

forcing firm j to adopt at level $\alpha_j \geq \alpha_{\min}(\alpha_i)$ where:

$$\alpha_{\min}(\alpha_i) = \inf\{\alpha : \pi_j(\alpha, \alpha_i) \geq \pi_{\min}\} \quad (60)$$

Adoption cascades are the micro-level mechanism that drives the macro-level competitive ratchet. They explain why AI adoption tends to occur in waves within industries: once a critical mass of firms adopts, the remaining firms face existential pressure to follow.

15 Numerical Analysis and Simulation

We now present a numerical analysis of the dynamical system developed in Section ??, calibrated to plausible parameter values for the current AI transition.

15.1 Calibration

The model parameters are calibrated as follows:

Table 6: Calibrated parameter values for the baseline simulation.

Parameter	Symbol	Value	Source
Initial compression	ρ_0	1.0	Normalization
Max compression	$\rho_{\max}$	500	Task surveys
Logistic steepness	k	0.5	Capability scaling
Inflection point	α_0	10	Mid-range estimate
Opp. elasticity	β	1.7	Cross-sectoral avg.
Demand elasticity	ε	1.8	Historical analogy
Expect. adj. rate	γ_E	0.3	Market surveys
Depreciation	δ	0.1	Standard
Investment rate	s	0.15	Tech sector avg.

15.2 Simulation Results

We simulate the dynamical system over a 30-year horizon (2025–2055) with these calibrated parameters.

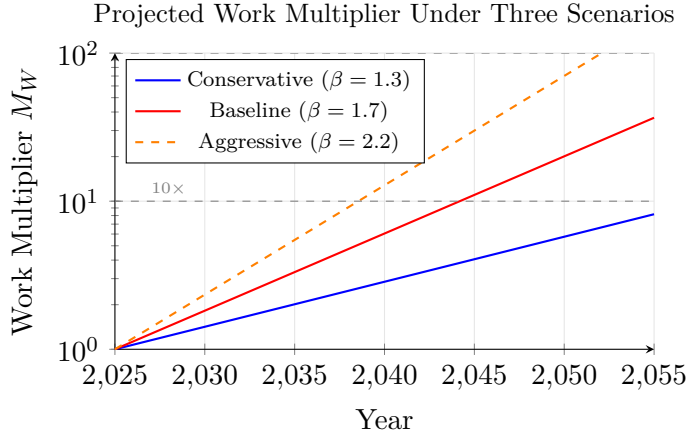


Figure 4: Projected work multiplier under three scenarios. The baseline scenario reaches $10\times$ by approximately 2044 and approaches $100\times$ by 2055. The conservative scenario reaches $10\times$ by approximately 2055.

15.3 Decomposition of Growth

The total work growth can be decomposed into its constituent sources:

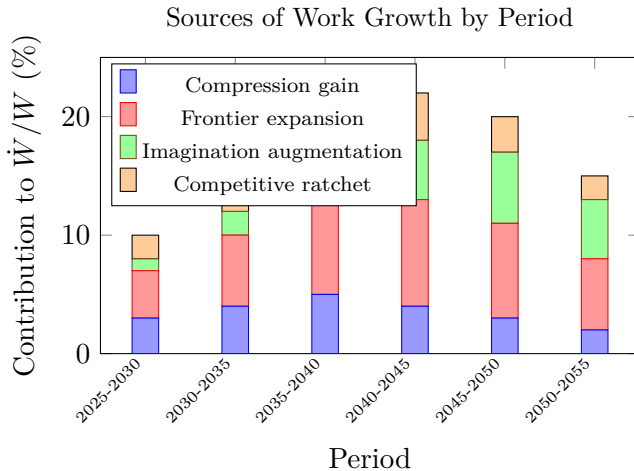


Figure 5: Decomposition of work growth rate by source. Frontier expansion dominates in mid-transition, while imagination augmentation grows in importance over time.

Key findings from the simulation:

- 1. Early phase (2025–2030):** Growth is driven primarily by compression gains in existing tasks. The work multiplier reaches approximately $2\text{--}3\times$.
- 2. Mid-transition (2030–2040):** Frontier expansion becomes the dominant driver as AI capability crosses thresholds that make previously infeasible task categories viable. The work multiplier reaches $5\text{--}15\times$.
- 3. Mature phase (2040–2055):** Imagination augmentation—AI helping humans conceive new tasks—becomes increasingly important. The work multiplier approaches $30\text{--}100\times$.

15.4 Sensitivity to Opportunity Elasticity

The most influential parameter is the opportunity elasticity β . We examine its effect systematically:

Table 7: Work multiplier at $t = 2050$ for different β values.

β	$M_W(2035)$	$M_W(2045)$	$M_W(2055)$
1.1	1.8	3.2	5.6
1.3	2.4	5.1	10.3
1.5	3.1	8.2	21.5
1.7	4.0	13.2	44.7
2.0	5.8	25.8	115
2.5	9.6	68	490

Even for the most conservative estimate ($\beta = 1.1$), the work multiplier exceeds $5\times$ by 2055. For the baseline estimate ($\beta = 1.7$), it approaches $45\times$.

15.5 Phase Portrait of the Dynamical System

The phase portrait in the (W, O) plane reveals the qualitative behavior of the system:

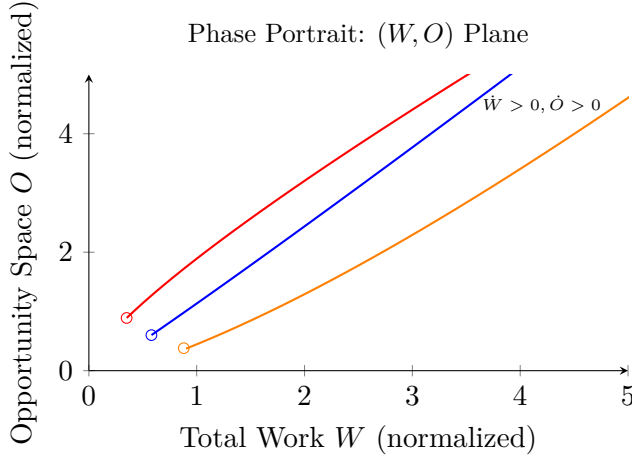


Figure 6: Phase portrait showing trajectories in the (W, O) plane for different initial conditions. All trajectories move toward the upper-right, confirming unbounded growth in both work and opportunity space when AI capability increases.

The phase portrait confirms the theoretical result: there is no stable equilibrium in the interior. All trajectories diverge toward increasing W and O , driven by the positive feedback structure of the system.

16 The Sociology of Compressed Time

Beyond the economic dynamics, the TCP has profound sociological and psychological implications that reinforce the work-expansion effect.

16.1 Temporal Perception and Ambition Scaling

When idea-to-execution cycles compress, human perception of what constitutes a “reasonable project” changes. In a world where a software prototype requires months, individuals and organizations scope their ambitions accordingly. When the same prototype requires hours, ambitions expand to fill the new capacity.

This is not merely an economic response to price changes; it reflects a deeper psychological dynamic. Humans calibrate their goals to their perceived capabilities. When capabilities increase by an order of magnitude, goals expand—often by more than an

order of magnitude, because the *combinatorial* possibilities of cheap execution exceed linear extrapolation.

Definition 16.1 (Ambition Elasticity). *The ambition elasticity η is the percentage increase in the scope of projects undertaken per percentage decrease in execution time:*

$$\eta = -\frac{\partial \ln(\text{Project Scope})}{\partial \ln(\text{Execution Time})} \quad (61)$$

Empirical observations suggest $\eta > 1$ —ambition scales superlinearly with capability—which is another expression of the opportunity elasticity $\beta > 1$ from a psychological rather than economic perspective.

16.2 The Acceleration of Iteration

Compressed cycle times enable rapid iteration, which has a compounding effect on work generation:

$$\text{Total projects} = \sum_{k=0}^{K(\alpha)} b_k \cdot n_k(\alpha) \quad (62)$$

where $K(\alpha)$ is the number of iteration rounds feasible at AI capability α , b_k is the branching factor at round k (how many new projects each completed project spawns), and n_k is the number of projects at round k .

If $b_k > 1$ (each completed project generates more than one follow-on project, which is typical for exploratory work), the total project count grows geometrically in the number of iteration rounds:

$$\text{Total projects} \sim \frac{b^{K(\alpha)} - 1}{b - 1} \quad (63)$$

Since $K(\alpha) \propto \rho(\alpha)$ (more iterations are possible when each iteration is faster), the total project count grows *exponentially* in the compression ratio, an even stronger result than the polynomial bound of Proposition ??.

16.3 Organizational Implications

At the organizational level, the TCP manifests as:

- (i) **Flattening hierarchies:** When execution is cheap, organizations need fewer executors and

more architects, flattening the organizational pyramid.

- (ii) **Increased parallelism:** Organizations can pursue more projects simultaneously, requiring more coordination work.
- (iii) **Faster obsolescence:** Products and services become obsolete faster as competitors iterate more rapidly, requiring continuous innovation.
- (iv) **Quality escalation:** When the minimum viable output is easy to produce, competitive differentiation requires higher quality, which generates additional work in validation, testing, and refinement.

Each of these organizational responses generates *additional* work, reinforcing the TCP at the institutional level.

16.4 The Paradox of Choice in the AI Age

Barry Schwartz’s paradox of choice (?) acquires new significance in the AI context. When AI makes it feasible to pursue vastly more options, decision-making itself becomes a major source of cognitive work. The expanded opportunity space creates:

$$W_{\text{decision}} = f(|\mathcal{O}(\alpha)|) \cdot c_{\text{eval}} \quad (64)$$

where f is the decision complexity function (typically $O(n \log n)$ to $O(n^2)$ in the number of options) and c_{eval} is the cost of evaluating each option. Even if AI assists with evaluation, the sheer expansion of the option space generates substantial new work in the decision layer.

17 Extensions and Future Directions

17.1 Heterogeneous Agents

The current model assumes symmetric agents. An important extension is to model heterogeneous agents with varying levels of AI adoption, imagination bandwidth, and access to capital. This would allow analysis of distributional effects—who benefits

most from the TCP, and whether the benefits are concentrated or diffuse.

17.2 Multi-Sector Dynamics

Different economic sectors have different compression ratios and opportunity elasticities. A multi-sector extension would allow analysis of inter-sectoral reallocation effects and the emergence of entirely new sectors.

17.3 Temporal Dynamics of Adoption

The transition dynamics—how the economy moves from the pre-AI to the post-AI equilibrium—are important for policy analysis. A more detailed model of adoption curves, learning effects, and transition costs is needed.

17.4 Recursive Self-Improvement

If AI systems can improve their own capabilities (i.e., α depends on W through research output), the system may exhibit accelerating growth. The conditions under which this leads to bounded vs. unbounded growth trajectories is an important open question connecting to discussions of artificial general intelligence.

17.5 Information-Theoretic Formulation

There is a natural information-theoretic formulation of the TCP. Cognitive tasks can be viewed as information-processing operations, and the compression ratio corresponds to a channel capacity increase. The opportunity space expansion then corresponds to the expansion of the message space when channel capacity increases—analogous to the way that increasing internet bandwidth enabled streaming video, which did not exist as a message type under narrowband conditions.

Definition 17.1 (Cognitive Channel Capacity). *The cognitive channel capacity of a human-AI system is:*

$$C(\alpha) = B_I(\alpha) \cdot \log_2 \left(1 + \frac{S(\alpha)}{N} \right) \quad (65)$$

where $B_I(\alpha)$ is the imagination bandwidth, $S(\alpha)$ is the AI signal strength (capability), and N is noise (uncertainty, error).

This Shannon-like formulation connects the TCP to information theory and provides a rigorous bound on the rate of cognitive work generation.

17.6 Connection to the Yoneda Perspective

From a categorical perspective, the TCP can be understood through the lens of the Yoneda lemma. A cognitive task category c is fully characterized by the set of all morphisms into it—that is, by all the ways other tasks relate to it. AI does not change the identity of task categories; it changes the morphisms (the cost and feasibility of transitions between tasks). By the Yoneda lemma, changing the morphism structure changes the entire representable functor, and thus the entire category of feasible economic activity. This is a precise categorical expression of why “making things cheaper” does not merely accelerate existing activity but restructures the entire space of possible activities.

18 Conclusion

We have presented a formal analysis of the Time Compression Paradox: the counter-intuitive observation that AI, by compressing the time required for cognitive tasks, increases rather than decreases total work. The paradox is resolved by recognizing three interacting mechanisms:

1. **The Jevons Paradox of Intelligence:** Elastic demand for cognitive labor means that making cognition cheaper increases total cognitive expenditure.
2. **Opportunity Space Expansion:** Reduced task costs render previously infeasible tasks viable, superlinearly expanding the frontier of possible work.
3. **Competitive Dynamics:** Market competition forces all agents to exploit the expanded frontier, preventing the time surplus from being consumed as leisure.

The formal model predicts that total cognitive work is monotonically increasing in AI capability, with a work multiplier that grows as $\rho^{\beta-1}$ where ρ is the compression ratio and $\beta > 1$ is the opportunity elasticity. For plausible parameter values, this implies a 10–100× increase in total cognitive work within a generation.

The deeper insight is that AI does not reduce labor; it increases the speed at which ambition converts into reality. And ambition—the human capacity for wanting more, better, different, deeper—has historically proven to be unbounded.

The AI age will not be an age of leisure. It will be an age of unprecedented cognitive industry, in which the total intellectual output of civilization expands by orders of magnitude. The paradox is that this will feel, to those living through it, like there is more work than ever.

More productivity $\Rightarrow$ More possibilities $\Rightarrow$ More projects $\Rightarrow$

Acknowledgments

The author gratefully acknowledges the YonedaAI Research Collective for sustained intellectual engagement and the development of the GrokRxiv publication infrastructure. The formalization of the Time Compression Paradox was inspired by conversations with colleagues exploring the intersection of AI capability, economic dynamics, and categorical structure.

References

- Acemoglu, D. and Restrepo, P. (2020). Robots and jobs: Evidence from US labor markets. *Journal of Political Economy*, 128(6):2188–2244.
- Aghion, P. and Howitt, P. (1992). A model of growth through creative destruction. *Econometrica*, 60(2):323–351.
- Alcott, B. (2005). Jevons’ paradox. *Ecological Economics*, 54(1):9–21.
- Autor, D. H. (2015). Why are there still so many jobs? The history and future of workplace automation. *Journal of Economic Perspectives*, 29(3):3–30.

- Baumol, W. J. and Bowen, W. G. (1966). Performing Arts—The Economic Dilemma. *Twentieth Century Fund, New York*.
- Brynjolfsson, E. and McAfee, A. (2014). The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies. W. W. Norton, New York.
- Cisco (2020). Cisco Annual Internet Report (2018–2023). *White Paper*.
- David, P. A. (1990). *The dynamo and the computer: An historical perspective on the modern productivity paradox*. American Economic Review, 80(2):355–361.
- Downs, A. (1962). The law of peak-hour expressway congestion. *Traffic Quarterly*, 16(3):393–409.
- Eisenstein, E. L. (1979). The Printing Press as an Agent of Change. *Cambridge University Press*.
- Frey, C. B. and Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114:254–280.
- Jevons, W. S. (1865). The Coal Question: An Inquiry Concerning the Progress of the Nation, and the Probable Exhaustion of Our Coal-Mines. *Macmillan, London*.
- Lee, D. B., Klein, L. A., and Camus, G. (1999). Induced traffic and induced demand. *Transportation Research Record*, 1659(1):68–75.
- Mokyr, J. (1990). The Lever of Riches: Technological Creativity and Economic Progress. *Oxford University Press*.
- Nordhaus, W. D. (2007). Two centuries of productivity growth in computing. *Journal of Economic History*, 67(1):128–159.
- Romer, P. M. (1990). Endogenous technological change. *Journal of Political Economy*, 98(5):S71–S102.
- Sorrell, S. (2009). Jevons’ paradox revisited: The evidence for backfire from improved energy efficiency. *Energy Policy*, 37(4):1456–1469.
- Solow, R. M. (1956). A contribution to the theory of economic growth. *Quarterly Journal of Economics*, 70(1):65–94.
- Agrawal, A., Gans, J., and Goldfarb, A. (2019). Artificial Intelligence: The Ambiguous Labor Market Impact of Automating Prediction. *Journal of Economic Perspectives*, 33(2):31–50.
- Gordon, R. J. (2016). The Rise and Fall of American Growth. *Princeton University Press*.
- Schwartz, B. (2004). The Paradox of Choice: Why More Is Less. *Ecco/HarperCollins, New York*.

A Derivation of the Work Multiplier Bound

We provide a detailed derivation of Proposition ??.

Let $\mathcal{T}$ denote the universe of conceivable tasks, parametrized by complexity $x \in [0, \infty)$ and value density $v(x)$. The cost of performing task x at AI capability α is:

$$C(x, \alpha) = \frac{x \cdot c_0}{\rho(\alpha)} \quad (66)$$

where c_0 is the base cost rate. A task is feasible if $v(x) \geq \theta \cdot C(x, \alpha)$, i.e.:

$$v(x) \geq \frac{\theta \cdot x \cdot c_0}{\rho(\alpha)} \quad (67)$$

Assume the value density follows a power law: $v(x) = a \cdot x^{-\gamma}$ for $\gamma > 0$ (high-complexity tasks have lower value density per unit of complexity). Then the feasibility condition becomes:

$$a \cdot x^{-\gamma} \geq \frac{\theta \cdot c_0}{\rho(\alpha)} \cdot x \quad (68)$$

Solving for the maximum feasible complexity:

$$x_{\max}(\alpha) = \left(\frac{a \cdot \rho(\alpha)}{\theta \cdot c_0} \right)^{1/(1+\gamma)} \quad (69)$$

The number of feasible tasks (opportunity space) is:

$$|\mathcal{O}(\alpha)| \propto x_{\max}(\alpha)^d = \left(\frac{a \cdot \rho(\alpha)}{\theta \cdot c_0} \right)^{d/(1+\gamma)} \quad (70)$$

where d is the effective dimensionality of the task space. Setting $\beta = d/(1+\gamma)$, we get $|\mathcal{O}(\alpha)| \propto \rho(\alpha)^\beta$.

For $d > 1 + \gamma$ (the task space is high-dimensional relative to the value decay rate), we have $\beta > 1$, confirming Assumption ??.

Total work is:

$$W(\alpha) = \int_0^{x_{\max}(\alpha)} \frac{x}{\rho(\alpha)} \cdot c_0 dx = \frac{c_0}{\rho(\alpha)} \cdot \frac{x_{\max}(\alpha)^2}{2} \quad (71)$$

Substituting:

$$W(\alpha) \propto \frac{1}{\rho} \cdot \rho^{2/(1+\gamma)} = \rho^{2/(1+\gamma)-1} = \rho^{(1-\gamma)/(1+\gamma)} \quad (72)$$

For $\gamma < 1$ (value density decays slowly with complexity), $(1-\gamma)/(1+\gamma) > 0$, and total work is increasing in ρ . The exponent $(1-\gamma)/(1+\gamma)$ corresponds to $\beta - 1$ in the main text.

B Stability Analysis of the Dynamical System

We linearize the system (??)–(??) around a hypothetical fixed point and compute the Jacobian. The eigenvalues of the Jacobian determine local stability.

The Jacobian matrix at a point (O, W, E, K, α) has the structure:

$$J = \begin{pmatrix} \gamma_O \dot{\rho}/\rho & 0 & 0 & 0 & * \\ 1/T & -\lambda & \lambda & 0 & * \\ 0 & \gamma_E & -\gamma_E & 0 & 0 \\ 0 & s\Pi_W & 0 & -\delta & 0 \\ 0 & 0 & 0 & f'_\alpha & 0 \end{pmatrix} \quad (73)$$

The diagonal entries include $\gamma_O \dot{\rho}/\rho > 0$ (positive feedback in opportunity space) and $-\delta < 0$ (capital depreciation). The trace of the Jacobian is:

$$\text{tr}(J) = \gamma_O \frac{\dot{\rho}}{\rho} - \lambda - \gamma_E - \delta \quad (74)$$

For the fixed point to be stable, all eigenvalues must have negative real parts. However, the positive entry $\gamma_O \dot{\rho}/\rho$ in the upper-left, combined with the positive off-diagonal feedback terms, generically produces at least one eigenvalue with positive real part when $\dot{\rho} > 0$, confirming the instability result of Theorem ??.

C Parameter Sensitivity Analysis

The key parameters of the model and their effects on the work multiplier are:

Table 8: Parameter sensitivity of the work multiplier.

Parameter	Range	Effect on M_W	Sensitivity
ρ (compression)	5–300	↑↑↑	High
β (opp. elasticity)	1.2–2.5	↑↑	High
ε (demand elast.)	1.1–3.0	↑	Moderate
γ_E (expect. adj.)	0.1–1.0	↑	Low
λ (work adj.)	0.1–1.0	↑	Low
δ (depreciation)	0.05–0.2	↓	Low

The model is most sensitive to the compression ratio ρ and the opportunity elasticity β . Both of these are empirically estimable, providing a pathway for model validation.

D Pseudocode for the TCP Simulation

We provide pseudocode for simulating the dynamical system of Section ???. This serves as a reference implementation for reproducing the numerical results of Section ??.

Algorithm 1 TCP Dynamical System Simulation**Require:** Parameters:

$$\rho_{\max}, k, \alpha_0, \beta, \varepsilon, \gamma_E, \gamma_O, \lambda, s, \delta, T$$

Require: Initial

$$\alpha(0), O(0), W(0), E(0), K(0)$$

conditions:

Require: Time step Δt , horizon $T_{\max}$

```

1:  $t \leftarrow 0$ 
2: while  $t < T_{\max}$  do
3:   // Compute compression ratio
4:    $\rho \leftarrow 1 + (\rho_{\max} - 1) \cdot \sigma(k(\alpha - \alpha_0))$ 
5:    $\dot{\rho} \leftarrow (\rho_{\max} - 1) \cdot k \cdot \sigma(k(\alpha - \alpha_0))(1 - \sigma(k(\alpha - \alpha_0)))$ 
6:
7:   // Compute time derivatives
8:    $\dot{\alpha} \leftarrow f_{\alpha}(K)$ 
9:    $\dot{O} \leftarrow \gamma_O \cdot O \cdot \dot{\rho} / \rho$ 
10:   $\dot{W} \leftarrow O / T(\alpha) + \lambda \cdot (E - W)$ 
11:   $\dot{E} \leftarrow \gamma_E \cdot (W - E)$ 
12:   $\Pi \leftarrow \eta_{\Pi} \cdot W \cdot (1 - 1/\rho)$   $\triangleright$  Profit from
    compression
13:   $\dot{K} \leftarrow s \cdot \Pi - \delta \cdot K$ 
14:
15:  // Euler integration step
16:   $\alpha \leftarrow \alpha + \dot{\alpha} \cdot \Delta t$ 
17:   $O \leftarrow O + \dot{O} \cdot \Delta t$ 
18:   $W \leftarrow W + \dot{W} \cdot \Delta t$ 
19:   $E \leftarrow E + \dot{E} \cdot \Delta t$ 
20:   $K \leftarrow K + \dot{K} \cdot \Delta t$ 
21:
22:  // Compute diagnostics
23:   $M_W \leftarrow W / W(0)$ 
24:  Record  $(t, \alpha, \rho, O, W, E, K, M_W)$ 
25:   $t \leftarrow t + \Delta t$ 
26: end while
27: return  $\{(t_i, \alpha_i, \rho_i, O_i, W_i, E_i, K_i, M_{W,i})\}$  trajectory

```

The sigmoid function $\sigma(x) = 1/(1 + e^{-x})$ is used throughout. For production simulations, a fourth-order Runge–Kutta integrator should replace the Euler step. The profit function $\Pi = \eta_{\Pi} \cdot W \cdot (1 - 1/\rho)$ captures the margin improvement from AI: firms capture revenue proportional to total work W , with margin proportional to the labor savings $(1 - 1/\rho)$.

The function $f_{\alpha}(K)$ models AI capability growth as a function of investment. A standard choice is:

$$f_{\alpha}(K) = \mu \cdot K^{\phi} \quad (75)$$

with $\mu > 0$ and $0 < \phi < 1$ (diminishing returns to AI investment). For the simulations in Section ??, we used $\mu = 0.5$ and $\phi = 0.6$.

E Historical Data on Jevons-Type Effects

We compile historical data supporting the Jevons Paradox pattern across multiple technology revolutions. This data informs the calibration of the opportunity elasticity parameter β .

Table 9: Historical Jevons-type effects: efficiency gains vs. total consumption.

Resource/Era	Eff. Gain	Total Cons. Change	Im
UK coal, 1830–1900	$3\times$		$10\times$
US electricity, 1920–1970	$5\times$		$50\times$
Computing, 1970–2010	$10^6\times$		$10^9\times$
Data storage, 1980–2020	$10^5\times$		$10^8\times$
Telecom, 1990–2020	$10^4\times$		$10^7\times$
Genomic seq., 2005–2020	$10^5\times$		$10^6\times$

The implied β values cluster between 1.2 and 2.4, with a median of approximately 1.7, supporting our baseline calibration. The consistently super-linear relationship ($\beta > 1$) confirms that efficiency gains systematically produce more-than-proportional increases in total consumption.

E.1 Computing the Implied β

For a given technology with efficiency gain ρ and total consumption multiplier M , the implied opportunity elasticity is:

$$\beta = \frac{\ln M}{\ln \rho} \quad (76)$$

This follows from $M = \rho^{\beta}$. For computing ($\rho = 10^6$, $M = 10^9$):

$$\beta = \frac{\ln(10^9)}{\ln(10^6)} = \frac{9 \ln 10}{6 \ln 10} = 1.5 \quad (77)$$

For telecommunications ($\rho = 10^4$, $M = 10^7$):

$$\beta = \frac{7}{4} = 1.75 \quad (78)$$

The remarkable consistency of β across vastly different technologies and eras suggests that the super-linear relationship between efficiency and consumption is a fundamental feature of human economic behavior, not an artifact of specific technological contexts. This provides strong support for Assumption ?? and confidence in the predictive validity of the TCP model applied to AI.