



Munich Personal RePEc Archive

Total Output of the Future

Kurniady, Alvin

4 September 2025

Online at <https://mpra.ub.uni-muenchen.de/126078/>
MPRA Paper No. 126078, posted 10 Oct 2025 01:31 UTC

Total Output of the Future

Author: Alvin Kurniady

Author's Current Affiliation: None (The Author is an Independent Researcher)

Author's Past Affiliations: University of Southern California, Los Angeles, U.S.A (Alumni, B.S);
Pepperdine University, Malibu, U.S.A (Alumni, M.B.A)

Corresponding Author's Email: alvinkurniady@gmail.com

This manuscript has been submitted to the *Journal of Economic Growth* (Springer, 2025). The same version is circulated as a working paper on MPRA for early dissemination and feedback. This paper has not been peer-reviewed.

Abstract This paper develops a theory of growth with self-learning AI. I decompose “technologies” into a non-self-learning component $T(t)^\lambda$ and a recursive self-learning term $(S(t) \cdot D(t))^{\theta(t)}$, where $\theta(t) = \theta_0 + \rho \cdot \log(S(t) \cdot D(t))$ links capability gains to deployed self-learning technologies S and data D . I present two complementary production functions. Version 1 highlights distributional channels by separating AI-complementary vs. AI-substitutable labor and human capital. Version 2 is measurement-oriented, mapping the self-learning stock to AI-specific physical capital, labor forces, and human capital, thereby operationalizing S . The model yields sharp regime conditions: with small/approximately constant $\theta(t)$, the economy exhibits a balanced growth path (BGP); when $\theta(t)$ becomes large enough to push effective returns above one, growth accelerates. A log-space recursion implies a quadratic bound for $\log((SD)^{\theta(t)})$, establishing no finite-time singularity. The framework produces testable predictions—notably the need for both linear and quadratic terms in $\log(SD)$ in empirical specifications—and clarifies bottlenecks: insufficient AI-specific capital or low-quality data can hold down $\theta(t)$ and prevent acceleration even with advanced systems. Policy implications follow directly: scale compute and energy, raise H_{AI} , L_{AI} , and improve data governance/quality. The contribution is conceptual and theory-only, positioning the mechanism for subsequent empirical work while providing a tractable structure for cross-country comparisons in an economy increasingly driven by recursive, autonomous innovation.

Keywords Self-learning artificial intelligence · Economic growth theory · Recursive learning dynamics · AI-specific capital and labor · Balanced growth path vs. accelerating growth · Cross-country growth models

JEL Codes O41

1 Introduction

Artificial intelligence with self-learning capability is altering the fundamental mechanics of growth. Unlike conventional technologies, self-learning systems improve through recursive feedback—learning from deployment scale and data—so productivity can rise without proportional additions of labor or human capital. Existing macro frameworks either fold AI indistinguishably into a generic “technology” term or treat it as static capital, thereby missing the distinct dynamics of recursive learning. This paper proposes a technology block that explicitly separates non-self-learning technologies $T(t)^\lambda$ from a self-learning component $(S(t) \cdot D(t))^{\theta(t)}$, where the recursive component $\theta(t) = \theta_0 + \rho \cdot \log(S(t) \cdot D(t))$ increases with deployed self-learning AI $S(t)$ and usable data $D(t)$. I embed this block in two production-function versions: Version 1 focuses on distributional channels by distinguishing AI-complementary and AI-substitutable labor and human capital; Version 2 is measurement-oriented, mapping the stock of self-learning AI to AI-specific physical capital, labor forces, and human capital. Together they deliver conditions under which an economy remains on a balanced growth path (BGP) or transitions to sustained acceleration, provide a clean “no singularity” result, and generate testable predictions (notably a quadratic term in $\log(SD)$ arising from the recursive exponent).

This paper is purely theoretical. The contribution is conceptual: the model shows how recursive self-learning generates nonlinear growth dynamics, distributional consequences, and bottlenecks driven by AI-specific capital and data availability. I do not attempt a direct empirical test here; instead, I emphasize testable predictions that can guide future empirical work. By maintaining a theory-only scope, the paper positions itself alongside established growth-theoretic contributions that first introduced new mechanisms formally before empirical work followed (e.g., Romer 1990; Jones 1995). The value of the framework lies in clarifying the conditions under which recursive self-learning AI transforms growth dynamics, not in claiming empirical verification at this stage.

Substantively, the framework yields four headline insights. First, the regime switch—BGP versus accelerating growth—is governed by $\theta(t)$: when $\theta(t)$ remains small/approximately constant, both versions admit a BGP; when deployment and data raise $\theta(t)$ enough to push effective returns above one, growth accelerates. Second, technology disaggregation matters: modeling technology as $A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)}$ (where $A_0 > 0$ is a time-invariant normalization; results are invariant to the normalization $A_0 = 1$) reveals data–deployment complementarities and recursive

amplification that are hidden when AI is treated as static capital. Third, bottlenecks in AI-specific capital (compute, energy, data-center capacity) and in data quality can hold down $\theta(t)$, preventing acceleration even when sophisticated self-learning systems exist. Fourth, the model generates distributional predictions—gains for AI-complementary skills and erosion for AI-substitutable skills—and a roadmap for policy levers (scaling K_{AI} , improving data, building H_{AI} , L_{AI}). These themes organize the analysis that follows.

2 Literature Review

Current existing growth theories either have not taken into account AI's ability to self-learn or do not have adequate model or equation to accurately calculate total output (Y) that is useful for cross-countries comparisons, in the world where self-learning AI become increasingly important. The significant older models, including Solow (1956), Mankiw, Romer and Weil (1992), Lucas (1988), Romer (1990) and Jones (1995), clearly do not take into account AI's ability to self-learn. Table 1 summarizes the mathematical representations of total output in five foundational growth models: Solow (1956), Mankiw, Romer and Weil (1992), Lucas (1988), Romer (1990), and Jones (1995).

Table 1 Total Output Functions

Model	Y (Total Output)
Solow (1956)	$Y = K^\alpha (AL)^{1-\alpha}$
MRW (1992)	$Y = K^\alpha H^\beta (AL)^{1-\alpha-\beta}$
Lucas (1988)	$Y = AK^\alpha (uHL)^{1-\alpha}$
Romer (1990)	$Y = (\int x_i^\alpha di) L_Y^{1-\alpha}$
Jones (1995)	$Y = K^\alpha (AL_Y)^{1-\alpha}$

Where:

Y: Total Output (in all models)

K: Physical capital (in all models)

L: Labor (in all models)

A: Technology (in all models)

α : Capital's share in total output (in Solow, Lucas and Jones)

H: Human capital (in all models)

α, β : Output elasticities of capital and human capital, respectively (in MRW)
 u : Fraction of human capital allocated to production (while the rest goes to learning) (in Lucas)
 x_i : Amount of intermediate input i used (in Romer)
 N : Number of available intermediate inputs (in Romer)
 L_Y : Labor allocated to final goods production (in Romer and Jones)
 α : Output elasticity of intermediate inputs (in Romer)

All of the models discussed above do not distinguish between technologies that possess autonomous learning capabilities and those that do not. This distinction is crucial because self-learning technologies exhibit fundamentally different behaviors with significant implications for total output. Unlike conventional technologies that lack this capacity, self-learning AI can improve its capabilities and, consequently, its productivity without requiring additional human capital or labor input. This is a core feature of self-learning AI technologies that is not captured by existing frameworks. As such, all of above existing models are inadequate, especially in cross-countries comparisons, for forecasting total output in an economy where the contributions of self-learning AI are increasingly central to economic growth. Therefore, to accurately model output of the future, it is essential to disaggregate these two types of technology and account for their distinct effects.

Even some of the more recent papers still have not adopted AI's self-learning capabilities into their models. Any model that does not take into account AI's ability to self-learn is inadequate to calculate total output of a future (especially when cross-countries analysis is performed) because, in the future, self-learning AI are deployed and used by general public. For example, in Acemoglu's working paper *The Simple Macroeconomics of AI* (2024), AI's self-learning capabilities is not included in the model (Acemoglu, 2024). In this paper, Acemoglu treats AI as usual traditional capital, which do not self-learn (Acemoglu, 2024). Another recent paper that does not take into account AI's ability to self-learn is the paper by Jacobo-Romero, Carvalho and Freitas. In this paper, total output (Y) is defined as the following: $Y = (L + AX)^\alpha \cdot H^{1-\alpha}$ (Note: Y is total output; L is low-skilled labor; H is high-skilled labor; X is automation input; A is automation substitution factor; α is output elasticity of $(L + AX)$) (Jacobo-Romero, Carvalho, & Freitas, 2022). This paper does not model AI's self-learning capabilities. Furthermore, this paper treats AI as a static automation tool that substitutes for labor in certain tasks, without modeling learning, adaptation or recursive improvement. Another more recent paper that does not take into account

AI's ability to self-learn is the paper by Wang, Sarker, Alam and Sumon. In this paper, AI is treated like ordinary machines, which are static and do not have self-learning capabilities (Wang et al., 2021). This paper's model allows AI to increase growth, affect wages and shift labor demand but, in this paper, AI cannot self-learn.

There are recent papers that have adopted AI's self-learning capabilities into their model. One of the most notable papers that have adopted AI's self-learning capabilities into their model is the paper by Trammell and Korinek. In this paper, Trammel and Korinek introduce the following equations (Trammell & Korinek, 2023):

$$A'_t = A_t^\Phi [(C_t K_t)^\rho + (D_t S_t L_t)^\rho]^{\lambda/\rho}$$

A'_t : Rate of change of productivity (technological progress)

A_t^Φ : Recursive feedback from current knowledge stock (this is the self-learning)

A_t : Technology level

C_t : Capital-augmenting tech

D_t : Labor-augmenting tech

K_t : Capital used in R&D

L_t : Total labor supply

S_t : Share of labor allocated to R&D

λ : Returns to scale in research

ρ : Substitution between capital and labor in research

Above equation allows AI to recursively accelerate its own development.

$$Y_t = A_t (1 - S_t) L_t$$

Y_t depends on A_t and, therefore, it is clear that total output is significantly affected by AI's ability to self-learn. Despite the incorporation of AI's ability to self-learn in the model, Trammel and Korinek's model has a significant weakness that needs to be addressed. The weakness is that, in this model, self-learning AI are not separated from non-self-learning AI. In the real world, there are self-learning AI and non-self-learning AI; these two types of AI "behave" differently, and the difference has major implication on the total output. That is why Trammel and Korinek's model still needs to be improved.

Another paper that has incorporated AI's ability to self-learn is authored by Julia Puaschunder. In her paper, the total output (Y) is defined as followed (Puaschunder, 2022):

$$Y(t) = [A(t)K(t)]^\alpha \cdot [A(t)L(t)]^\beta [A(t)I(t)]^{1-\alpha-\beta}$$

$Y(t)$: Total output at time t

$A(t)$: Technology level

$K(t)$: Capital

$L(t)$: Labor

$I(t)$: Information stock (derived from AI's self-learning capabilities, big data, Internet access, etc.)

α, β : Output elasticities of capital and labor

Clearly, this model takes into account AI's ability to self-learn by multiplying A (technology level) by I (information stock). The weakness of this model is that technologies that self-learn are not separated from technologies that do not self-learn. These two types of technologies should be separated and treated differently in the model because they "behave" differently, and the difference has major implication on the total output (Y).

Another more recent significant paper that has taken into account AI's ability to self-learn is the paper by Besiroglu, Emery-Xu and Thompson. In this paper, Y (total output) is defined as followed: $Y(t) = [(1 - \alpha_k)K(t)]^\alpha \cdot [A(t)(1 - \alpha_l)L(t)]^{1-\alpha}$ (Besiroglu, Emery-Xu, & Thompson, 2023).

$Y(t)$: Total output at time t

$K(t)$: Total capital stock

$L(t)$: Total labor force

α_k : Share of capital allocated to R&D

α_l : Share of labor allocated to R&D

α : Output elasticity of capital

$A(t)$: Technology

This paper also introduces the following equation: $A'_t = B \cdot (\alpha_k K(t))^\beta \cdot (\alpha_l L(t))^\gamma \cdot A(t)^\theta$ (Besiroglu et al., 2023)

A'_t : Rate of technological progress

B : Constant multiplier/productivity shift

$\alpha_k K(t)$: Capital allocated to R&D

$\alpha_l L(t)$: Labor allocated to R&D

β : Elasticity of idea production with respect to R&D capital

γ : Elasticity of idea production with respect to R&D labor

θ : Recursive self-improvement parameter

The AI's self-learning capabilities are captured by the $A(t)^\theta$ term. When $\theta > 0$, past knowledge increases future innovations, which means self-learning exists. Just like Trammel and Korinek paper and Puaschunder paper, Besiroglu et al. paper does not separate self-learning AI from non-self-learning AI.

Another more recent paper that has taken into account AI's ability to self-learn is the paper by Farach, Cambon and Spataro. In this paper, total output (Y) is a function of capital (K), human labor (L) and digital labor (D); $Y(t) = F(K, L, D)$ (Farach, Cambon, & Spataro, 2025). Digital labor is defined as the cognitive work performed by AI systems, such as chatbots, code assistants and diagnostic agents (Farach et al., 2025). Only some digital labors, such as certain diagnostic agents, self-learn; not all digital labors self-learn. Even though this paper's model has taken into account AI's ability to self-learn, it still has meaningful weaknesses. This model does not separate digital labors that self-learn from digital labors that do not self-learn. Furthermore, in this model, capital (K) is not separated into AI-specific capital and general capital (non-AI-specific capital). AI-specific capital includes GPUs, power infrastructures, data centers, etc. General capital is everything else that is not used to build and run AI. For policymakers, this separation is extremely important. In a model with such separation, especially with cross-countries data, policy makers can understand what happens to an economy if AI-specific capital is large or small.

Another more recent paper that has taken into account AI's ability to self-learn is the paper by Aghion, Jones and Jones. In this paper, the authors introduce the following equation (Aghion, Jones, & Jones, 2017): $Y_t = A_t \cdot (\beta_t^{1-\rho} K_t^\rho + (1 - \beta_t)^{1-\rho} L_t^\rho)^{1/\rho}$

Y_t : Total output

A_t : Total Factor Productivity (TFP)

β_t : Fraction of tasks that have been automated

K_t : Capital input

L_t : Labor input

ρ : Substitution parameter between capital-automated and labor-performed tasks

Furthermore, the authors introduce the following equation (Aghion et al., 2017):

$$\dot{A}_t = A_t^\Phi [(B_t K_t)^\rho + (C_t S_t)^\rho]^{1/\rho}$$

$\dot{A}_t$: Rate of change of TFP

A_t : TFP

Φ : Recursive exponent

B_t : Capital-augmenting efficiency

K_t : Capital input used in R&D

C_t : Labor-augmenting efficiency

S_t : Effective labor allocated to research

ρ : Elasticity of substitution parameter

The term A_t^ϕ captures AI's self-learning capabilities. Despite the sophistication of this paper's model, the model has a major weakness. Like some models that were already mentioned above, this model does not separate self-learning AI from non-self-learning AI.

3 Model Assumptions

Self-learning AI already exist today, and the examples are AlphaZero, MuZero and AutoML systems. AlphaZero can self-learn to master complex games through self-play without any prior human knowledge, except the game rules (Silver et al., 2017). MuZero not only can self-learn to master complex games but also it can do it without even being given the game rules (Schrittwieser et al., 2020). Furthermore, MuZero can self-learn how the world works in certain fields, such as physics (Schrittwieser et al., 2020). ChatGPT, Gemini, Perplexity, Claude and many other AI, on the other hand, are often mistakenly believed to be self-learning AI but the reality is that they are not self-learning AI. ChatGPT, Gemini, Perplexity and Claude do not update their knowledge after deployment; for these AI, knowledge update requires explicit action by engineers or researchers. These statements are confirmed by the famous Apple paper, "The Illusion of Thinking", which shows that Claude 3.7 Sonnet, DeepSeek-R1, DeepSeek-V3 and o3-mini (by OpenAI) simulate thinking via pre-programmed patterns but do not adapt or improve in the way a self-learning AI would (Shojaee et al., 2025). Self-learning AI are not currently being used by general public. Self-learning AI are currently being used by only specific institutions (such as high-ranking universities, including Stanford University and UC Berkeley), labs and companies (such as Google, Amazon and Tesla) that have the expertise to use them safely. However, the fact that self-learning AI already exist means that it is only a matter of time for mass deployment to the general public. In the future, the uses of self-learning AI will be widespread. When Internet first existed in the late 1960s, only researchers and scientists used it. Now, the uses of Internet are very common and widespread; the same will happen for AI. Once a new technology, which can increase productivities and profitability very significantly, already exists, it is only a matter of time before there are widespread

adoptions and uses of that technology. The amount of gains from improving and spreading such technology is just simply too great for any party to stop the developments and the spreading of such technology. Companies always seek greater profitability, and they have the lobbyists, the money and the resources to ensure the developments and the mass deployments of such technology are not meaningfully disrupted. Moreover, there is the following attitudes among companies and lawmakers: “If we do not do it, then somebody else (or another country) will do it”. Many scientists support the idea that Artificial General Intelligence (AGI) is inevitable. While the timing of the first AGI varies depending on which expert is being asked, the majority of the scientific community believe in the inevitability of AGI. Whether AGI will come into existence or not, self-learning AI already exist, and it will only become more capable and more powerful. What is definitely inevitable is the widespread uses of self-learning AI in any field, including physics, biology, finance, etc. The last two statements are very important assumptions of this paper. Every existing paper or model has not sufficiently incorporated AI’s ability to self-learn; each of existing papers or models has meaningful flaw/s; some of the flaws have been discussed in the Literature Review section above. This paper is aimed to provide production functions, which have no meaningful flaw, of today and the future, in which self-learning AI are widely used by general public.

4 The Proposed Equations (New Theory)

4.1 Proposed Production Functions

My proposed equation is built from MRW (Mankiw-Romer-Weil) equation, which is the following (Mankiw et al., 1992): $Y(t) = K(t)^\alpha H(t)^\beta [A(t)L(t)]^{1-\alpha-\beta}$

$Y(t)$: Total output at time t

$K(t)$: Physical capital stock at time t

$H(t)$: Human capital stock at time t

$A(t)$: Level of technology at time t

$L(t)$: Total labor force at time t

α : Output elasticity of physical capital

β : Output elasticity of human capital

Self-learning AI is not represented properly in the MRW model. I intend to improve the MRW model.

In this paper, I introduce two main equations. **My first proposed equation (version 1 of my model) is as followed:**

$$Y(t) = A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)} \cdot K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2} \cdot L_C(t)^{\beta_1} \cdot L_S(t)^{\beta_2} \cdot [H_{C0} \cdot e^{h_{ct}} \cdot (1 + \mu \log(SD)) + H_{S0} \cdot e^{h_{st}} \cdot (1 - \nu \log(SD))]^\delta$$

$Y(t)$: Total output

A_0 : Initial TFP

$T(t)$: Stock of non-self-learning technologies, which includes non-self-learning AI and non-AI technologies

λ : Output elasticity of non-self-learning technologies

$S(t)$: Stock of deployed, self-learning AI systems

$D(t)$: Volume and quality of data available for learning

$\theta(t)$: Recursive learning exponent

$K_{AI}(t)$: Capital invested specifically for AI systems, such as GPUs and data centers

α_1 : Output elasticity of AI-specific physical capital

$K_G(t)$: General capital and these capital exclude capital for AI systems

α_2 : Output elasticity of general physical capital

$L_C(t)$: AI-complementary labor

β_1 : Output elasticity of AI-complementary labor

$L_S(t)$: AI-substitutable labor

β_2 : Output elasticity of AI-substitutable labor

H_{C0} : Initial stock of AI-complementary human capital

H_{S0} : Initial stock of AI-substitutable human capital

h_c : Baseline growth of H_C that comes from traditional learning, such as university degree, without self-learning AI

h_s : Baseline growth of H_S that comes from traditional learning, such as university degree, without self-learning AI

μ : Elasticity that determines how much smarter AI-complementary workers get as self-learning AI scale grows

ν : Erosion rate, which is elasticity of how badly self-learning AI substitutes AI-substitutable human capital

δ : Output elasticity of human capital

Note: $0 < [\nu \log(SD)] < 1$. Each exponent in the equation above > 0 . Always set: $A_0 = 1$. A_0 is a pure scale factor that can be absorbed by units T, S, D.

If my proposed model is compared with the MRW model, then the comparisons are as followed:

The Technologies Part

$A(t)^{1-\alpha-\beta}$ in the MRW model is replaced by the following in my proposed model: $A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)}$. *Unlike in MRW where A (technology) and L (total labor force) cannot be separated because, in MRW, both A and L are exogenous, in my proposed production function, technology and total labor force can be separated because each of them is endogenous.*

Contribution (to total output) from self-learning AI is not properly represented in the MRW model whereas, in my proposed model, contribution from self-learning AI (to total output) is properly represented by the following part: $[S(t) \cdot D(t)]^{\theta(t)}$. Contribution (to total output) from non-self-learning AI and all of non-AI technologies are represented by $T(t)^\lambda$ in my proposed model. A_0 is just the initial TFP that is always set to 1. In my model, self-learning AI are clearly separated from non-self-learning AI, which is categorized the same as every non-AI technology. $S(t)$ is stock of deployed, self-learning AI systems. Volume and quality of data available for AI's self-learning are represented by $D(t)$.

I model the effective self-learning signal in production as the product $S(t) D(t)$. The multiplicative form captures strong complementary between deployment scale and data. If either deployed self-learning capacity or useable data is near zero, then effective self-learning is near zero. This mirrors evidence and theory that treat data as a nonrival input that raises returns when combined with algorithms and scale (Jones & Tonetti, 2020), and that data accumulation interacts with firm scale to amplify performance (Farboodi et al., 2019). Formal models of data externalities also show that one unit of data's value rises with the presence of other data and learning users (Ichihashi, 2021). Economically, S belongs in the signal because only deployed systems generate usage and feedback loops that create and exploit data (Acemoglu & Restrepo, 2018; Farboodi et al., 2019). D belongs because the volume and quality of usable data condition how much the deployed systems can learn from their environment (Jones & Tonetti, 2020; Ichihashi, 2021). The multiplicative SD term is thus not ad hoc; it encodes the complementarity that theory predicts, and the scaling evidence corroborates.

How fast the self-learning AI improve depends on $\theta(t)$, the recursive learning exponent. $\theta(t)$ is calculated as followed: $\theta(t) = \theta_0 + \rho \cdot \log(S(t) \cdot D(t))$. θ_0 is baseline learning capacity and ρ is the sensitivity of recursive improvement, which measures how strongly scale leads to smarter AI. The log is used to create diminishing returns from scale, which is common in AI capability scaling laws. This formulation mirrors observed empirical regularities in AI scaling: each doubling of compute or training data raises effective capability by roughly a constant increment in log-space (Hestness et al., 2017; Kaplan et al., 2020; Hoffmann et al., 2022; Bahri et al., 2024).

An important implication of this specification, $\theta(t) = \theta_0 + \rho \cdot \log(S(t) \cdot D(t))$, is that the recursive block generates a quadratic term in $\log(SD)$. To see this, note that the self-learning component enters as $[S(t) \cdot D(t)]^{\theta(t)}$. Taking logs yields $\theta(t) \log(SD)$. Because $\theta(t)$ itself is defined as $\theta_0 + \rho \cdot \log(SD)$, substitution gives $\theta(t) \log(SD) = \theta_0 \log(SD) + \rho [\log(SD)]^2$. Thus, any regression or empirical specification derived from the model must include both the linear and quadratic terms in $\log(SD)$. The presence of this quadratic is not arbitrary but rather the unique empirical signature of the recursive exponent. It captures the curvature implied by recursive learning: output rises with $\log(SD)$ but at a rate that itself grows with the scale of deployed systems and data.

In the MRW model, $A(t)$ is defined as followed: $A(t) = A_0 \cdot e^{gt}$ (Mankiw et al., 1992). Therefore, the growth rate of $A(t)$ in the MRW model is g . Clearly, there is no recursive or self-learning component in MRW model. In my model, the “technologies part” is: $A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)}$. A_0 does not change. The growth rate contribution of $T(t)^\lambda$ is the following:

$$\frac{d}{dt} \log (T(t)^\lambda) = \lambda \cdot \frac{d}{dt} \log T(t) = \lambda \cdot \frac{T'(t)}{T(t)} = \lambda \cdot g_T$$

The growth rate contribution of $[S(t) \cdot D(t)]^{\theta(t)}$ is: $\rho(g_S + g_D) \cdot \log(SD) + \theta(t) \cdot (g_S + g_D)$. Below is the calculation.

Let's define the following: $SL(t) = [S(t) \cdot D(t)]^{\theta(t)}$; note: SL stands for Self-Learning.

$$g_{SL(t)} = \frac{SL'(t)}{SL(t)}$$

$$\log SL(t) = \theta(t) \cdot \log[S(t) \cdot D(t)]$$

$$\frac{d}{dt} \log SL(t) = \dot{\theta}(t) \cdot \log (SD) + \theta(t) \cdot \frac{d}{dt} \log(SD)$$

Since:

$$\frac{d}{dt} \log(SD) = \frac{d}{dt} \log S + \frac{d}{dt} \log D = \frac{\dot{S}}{S} + \frac{\dot{D}}{D} = g_S + g_D$$

$$\dot{\theta}(t) = \rho \cdot \frac{d}{dt} \log(SD) = \rho \cdot (g_S + g_D)$$

$$g_{SL(t)} = \rho(g_S + g_D) \cdot \log(SD) + \theta(t) \cdot (g_S + g_D)$$

S itself is a function of AI-specific human capital, AI-specific labor and AI-specific physical capital. $S = f(H_{AI}, L_{AI}, K_{AI})$. H_{AI} is AI-specific human capital; L_{AI} is AI-specific labor; K_{AI} is AI-specific physical capital. Total human capital = $H_{AI} + H_{NAI}$ (Note: H_{NAI} is non-AI human capital). Total labor forces = $L_{AI} + L_{NAI}$ (Note: L_{NAI} is non-AI labor). Total physical capital = $K_{AI} + K_G$.

Therefore, my second production function is as followed:

$$Y(t) = A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)} \cdot K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2} \cdot L_{AI}(t)^{\varnothing_1} \cdot L_{NAI}(t)^{\varnothing_2} \cdot H_{AI}(t)^{\partial_1} \cdot H_{NAI}(t)^{\partial_2}$$

I label this equation as version 2 of my model. Each exponent in this equation > 0 . Always set: $A_0 = 1$. A_0 is a pure scale factor that can be absorbed by units T, S, D. Just like the version 1 of my production function, in the version 2, technology can be separated from total labor force because each of them is endogenous.

This part “ $Y(t) = A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)} \cdot K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2}$ ” of the version 2 equation is exactly the same as the same part of the version 1 equation. $\varnothing_1$ is output elasticity of AI-specific labor forces; $\varnothing_2$ is output elasticity of non-AI labor forces; ∂_1 is output elasticity of AI-specific human capital; ∂_2 is output elasticity of non-AI human capital. Since T consists of non-AI technologies and non-self-learning AI, then: $T = T_{NAI} + T_{AI}$ (Note: T_{NAI} is non-AI technologies; T_{AI} is non-self-learning AI). Thus, T_{AI} is also a function of AI-specific human capital, AI-specific labor and AI-specific physical capital. $T_{AI} = g(H_{AI}, L_{AI}, K_{AI})$.

As stated earlier, $\theta(t) = \theta_0 + \rho \cdot \log(S(t) \cdot D(t))$. In both versions of my production functions, $\theta(t)$ is the only exponent whose value is determined by its base, which is S and D. This fact is what makes it recursive, a distinct feature of self-learning AI. The MRW model does not have an exponent whose value is determined by its base and, therefore, the MRW does not have a recursive component; this means that the MRW model is insufficient in an economy with deployed self-learning AI. With all fairness, the MRW paper was published in 1992, a year in which self-learning AI had not yet existed. The technologies part is exactly the same in both of versions of my model. In both my production functions, S(t) is defined as followed: $S(t) = s_0 \cdot K_{AI}(t)^\tau \cdot H_{AI}(t)^\varpi \cdot L_{AI}(t)^\Gamma$. s_0 : Scale parameter, converting units of inputs into the units of S(t) and capturing baseline technology efficiency in deployed self-learning AI; note: $s_0 > 0$.

τ, ϖ, Γ : Output elasticities of AI-specific physical capital, AI-specific human capital and AI-specific labor, respectively, for self-learning AI; note: $\tau, \varpi, \Gamma > 0$.

The rate of change of $S(t)$ is: $S(t) \cdot [\tau (\dot{K}_{AI}(t)/K_{AI}(t)) + \varpi (\dot{H}_{AI}(t)/H_{AI}(t)) + \Gamma (\dot{L}_{AI}(t)/L_{AI}(t))]$.

The growth rate of $S(t)$ is: $g_S = \tau (\dot{K}_{AI}(t)/K_{AI}(t)) + \varpi (\dot{H}_{AI}(t)/H_{AI}(t)) + \Gamma (\dot{L}_{AI}(t)/L_{AI}(t))$.

In both my production functions, T_{AI} , which is non-self-learning AI, is defined as followed: $T_{AI}(t) = t_0 \cdot K_{AI}(t)^{\mathfrak{Z}} \cdot H_{AI}(t)^{\wp} \cdot L_{AI}(t)^{\varsigma}$.

t_0 : Scale parameter, converting units of inputs into the units of T_{AI} and capturing baseline technology efficiency in deployed non-self-learning AI; note: $t_0 > 0$.

$\mathfrak{Z}, \wp, \varsigma$: Output elasticities of AI-specific physical capital, AI-specific human capital and AI-specific labor, respectively, for non-self-learning AI; note: $\mathfrak{Z}, \wp, \varsigma > 0$.

The rate of change of $T_{AI}(t)$ is: $T_{AI}(t) \cdot [\mathfrak{Z} (\dot{K}_{AI}(t)/K_{AI}(t)) + \wp (\dot{H}_{AI}(t)/H_{AI}(t)) + \varsigma (\dot{L}_{AI}(t)/L_{AI}(t))]$.

The growth rate of $T_{AI}(t)$ is: $\dot{T}_{AI}(t)/T_{AI}(t) = \mathfrak{Z} (\dot{K}_{AI}(t)/K_{AI}(t)) + \wp (\dot{H}_{AI}(t)/H_{AI}(t)) + \varsigma (\dot{L}_{AI}(t)/L_{AI}(t))$.

$T_{NAI}(t) = T_{NAI}(0) \cdot e^{g_{NAI}t}$, and $g_{NAI} \geq 0$.

The growth rate of $T_{NAI}(t)$ is: $\dot{T}_{NAI}(t)/T_{NAI}(t) = g_{NAI}$.

Therefore, I have the followings:

$T(t) = T_{NAI}(t) + T_{AI}(t) = T_{NAI}(0) \cdot e^{g_{NAI}t} + t_0 \cdot K_{AI}(t)^{\mathfrak{Z}} \cdot H_{AI}(t)^{\wp} \cdot L_{AI}(t)^{\varsigma}$.

The growth rate of $T(t)$ is: $g_T = [(T_{NAI}(t)/T(t)) \cdot g_{NAI}] + [(T_{AI}(t)/T(t)) \cdot (\mathfrak{Z} (\dot{K}_{AI}(t)/K_{AI}(t)) + \wp (\dot{H}_{AI}(t)/H_{AI}(t)) + \varsigma (\dot{L}_{AI}(t)/L_{AI}(t)))]$.

I define the growth rate of $T_{AI}(t)$ as $g_{T_{AI}}(t)$, which equals to: $\mathfrak{Z} (\dot{K}_{AI}(t)/K_{AI}(t)) + \wp (\dot{H}_{AI}(t)/H_{AI}(t)) + \varsigma (\dot{L}_{AI}(t)/L_{AI}(t))$.

Therefore, the growth rate of $T(t)$ can be stated as: $g_T = [(T_{NAI}(t)/T(t)) \cdot g_{NAI}] + [(T_{AI}(t)/T(t)) \cdot g_{T_{AI}}]$.

The Physical Capital Part

$K(t)^\alpha$ in the MRW model is replaced by the following: $K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2}$ in both versions of my model. In both versions of my model, total capital = $K_{AI}(t) + K_G(t)$. $K_{AI}(t)$ is physical capital invested specifically for AI systems, such as GPUs and data centers. $K_G(t)$ is all other physical capital. Notice that, in both versions of my model, physical capital is not divided based on self-learning capabilities. This is because all AI (self-learning and non-self-learning) require many of the same physical capital, such as GPUs and data centers; furthermore, non-self-learning AI can be improved and become self-learning AI; this makes separating physical capital based on self-learning capabilities almost impossible. That is why, in both versions of my model, each physical

capital is classified based on whether it is for AI or not. The growth rate of K_{AI} is: $g_{K_{AI}} = \dot{K}_{AI} / K_{AI}$. The growth rate contribution of $K_{AI}^{\alpha_1}$ is: $\frac{d}{dt} \log K_{AI}^{\alpha_1} = \alpha_1 \cdot g_{K_{AI}}$. The growth rate of K_G is: $g_{K_G} = \dot{K}_G / K_G$. The growth rate contribution of $K_G^{\alpha_2}$ is: $\frac{d}{dt} \log K_G^{\alpha_2} = \alpha_2 \cdot g_{K_G}$. The purpose of separating physical capital into K_{AI} and K_G is to notice and understand what happen to contribution from non-self-learning technologies (which is represented by $T(t)^\lambda$), contribution from self-learning AI (which is represented by $[S(t) \cdot D(t)]^{\theta(t)}$) and, therefore, total output ($Y(t)$) when physical capital invested specifically for AI systems are increased or decreased. Cross-sectional and panel regressions across countries can be used to gain such understanding.

The Human Capital Part

For the human capital part, I will now discuss the version 1 of my model in this paragraph. $H(t)^\beta$ in MRW model is replaced by the following (in my version 1 model): $[H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD)) + H_{S0} \cdot e^{hst} \cdot (1 - v \log(SD))]^\delta$. In my version 1 model, human capital is divided into human capital that is enhanced by AI and human capital that is reduced by AI. The following part is the human capital that is enhanced by AI: $H_C = H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD))$, whereas the following part is the human capital that is reduced by AI: $H_S = H_{S0} \cdot e^{hst} \cdot (1 - v \log(SD))$. The term e^{hct} accounts for increase in AI-complementary human capital that is caused by all types of learning that do not use self-learning AI. For AI-complementary human capital, the learning that is enhanced by self-learning AI is captured by the following: $(1 + \mu \log(SD))$. μ is the elasticity that determines how much smarter AI-complementary workers get as self-learning AI scale grows. H_{C0} is just initial stock of AI-complementary human capital. The term e^{hst} accounts for increase in AI-substitutable human capital that is caused by all types of learning that do not use self-learning AI. h_s itself consists of two competing and opposite forces. The first force is from people who learn and improve their skills through reading, experiences, etc. This first force increases h_s . The second force is from non-self-learning AI that reduce people's human capital values, as these non-self-learning AI reduce human tasks or replace the humans completely. This second force reduces h_s . For AI-substitutable human capital, the reduction of the value of human capital that is caused by self-learning AI is represented by the following term: $(1 - v \log(SD))$. v is the erosion rate, which is the elasticity of how badly self-learning AI substitutes AI-substitutable human capital. H_{S0} is just the initial stock of AI-substitutable human capital. The growth rate contribution of the human capital in my model is calculated as followed:

$$H(t) = [H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD)) + H_{S0} \cdot e^{hst} \cdot (1 - \nu \log(SD))]^\delta$$

$$H(t) = Z(t)^\delta$$

$$\frac{\dot{H}(t)}{H(t)} = \delta \cdot \frac{\dot{Z}(t)}{Z(t)}$$

$$M(t) = H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD))$$

$$N(t) = H_{S0} \cdot e^{hst} \cdot (1 - \nu \log(SD))$$

$$\dot{Z}(t) = \dot{M}(t) + \dot{N}(t)$$

$$\frac{d}{dt}M(t) = H_{C0} \cdot e^{hct} [h_C (1 + \mu \log(SD)) + \mu \cdot \frac{d}{dt} \log(SD)]$$

$$\frac{d}{dt}N(t) = H_{S0} \cdot e^{hst} [h_S (1 - \nu \log(SD)) - \nu \cdot \frac{d}{dt} \log(SD)]$$

$$\dot{Z}(t) = H_{C0} \cdot e^{hct} [h_C (1 + \mu \log(SD)) + \mu \cdot \frac{d}{dt} \log(SD)] + H_{S0} \cdot e^{hst} [h_S (1 - \nu \log(SD)) - \nu \cdot \frac{d}{dt} \log(SD)]$$

$$g_{H(t)} = \delta \cdot (\{H_{C0} \cdot e^{hct} [h_C (1 + \mu \log(SD)) + \mu \cdot \frac{d}{dt} \log(SD)] + H_{S0} \cdot e^{hst} [h_S (1 - \nu \log(SD)) - \nu \cdot \frac{d}{dt} \log(SD)]\} / \{H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD)) + H_{S0} \cdot e^{hst} \cdot (1 - \nu \log(SD))\})$$

$$\text{Since } \frac{d}{dt} \log(SD) = \frac{\dot{S}}{S} + \frac{\dot{D}}{D} = g_S + g_D, \text{ then:}$$

$$\text{The growth rate contribution of human capital: } g_{H(t)} = \delta \cdot (\{H_{C0} \cdot e^{hct} [h_C (1 + \mu \log(SD)) + \mu \cdot (g_S + g_D)] + H_{S0} \cdot e^{hst} [h_S (1 - \nu \log(SD)) - \nu \cdot (g_S + g_D)]\} / \{H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD)) + H_{S0} \cdot e^{hst} \cdot (1 - \nu \log(SD))\})$$

For the human capital part, I will now discuss the version 2 of my model in this paragraph. $H(t)^\beta$ in MRW model is replaced by the following (in my version 2 model): $H_{AI}(t)^{\hat{\partial}_1} \cdot H_{NAI}(t)^{\hat{\partial}_2}$. In the version 2 of my model, human capital is categorized based on whether it is AI-specific or not; AI-specific human capital, H_{AI} , are human capital that can develop AI, whereas the rest of human capital are non-AI human capital, H_{NAI} . The growth rate of H_{AI} is: $g_{H_{AI}} = \dot{H}_{AI}/H_{AI}$. The growth rate of H_{NAI} is: $g_{H_{NAI}} = \dot{H}_{NAI}/H_{NAI}$. The growth rate contribution of $H_{AI}^{\hat{\partial}_1}$ is: $\frac{d}{dt} \log H_{AI}^{\hat{\partial}_1} = \hat{\partial}_1 \cdot g_{H_{AI}}$. The growth rate contribution of $H_{NAI}^{\hat{\partial}_2}$ is: $\frac{d}{dt} \log H_{NAI}^{\hat{\partial}_2} = \hat{\partial}_2 \cdot g_{H_{NAI}}$.

The Labor Force Part

For the labor force part, I now discuss the version 1 of my model until I state that I am switching to discussion of the version 2. In the MRW model, the total labor force at time t is represented by $L(t)$. In version 1 of my model, the total labor force is divided into L_C (AI-complementary labor)

and L_S (AI-substitutable labor). L_C is the labor force whose human capital values are enhanced by AI. L_S is the labor force whose human capital values are reduced by AI and, as a result, some of them work less or become unemployed as AI substitute them. Total labor force at time t : $L(t) = L_C(t) + L_S(t)$. In my equation, the labor force part is the following: $L_C(t)^{\beta_1} \cdot L_S(t)^{\beta_2}$. β_1 is output elasticity of AI-complementary labor and β_2 is output elasticity of AI-substitutable labor.

Both L_C and L_S change over time. First, people in the L_S category can improve their skills and switch to the L_C category. For example, some computer programmers lose their jobs, as certain AI can write codes and do their tasks. Since these people already have background in computer science, they can learn to become AI computer programmers, which are in high demand. Once they have the skills to design, build and optimize AI systems, they switch to the L_C category. That is an example of migration from L_S to L_C , which occur when people who were in L_S improved their skills until their skills are good enough for the L_C category. The second type of migration is the migration from L_C to L_S , which occurs as AI technologies keep improving and make people's skills in the L_C category become substitutable by AI. The rate of change of L_C and L_S are defined as followed:

$$\frac{d}{dt} L_C(t) = \phi_C(t) + \Psi_L(t) - \chi(t)$$

$$\frac{d}{dt} L_S(t) = \phi_S(t) - \Psi_L(t) + \chi(t)$$

$\Psi_L(t)$: Re-skilling rate, which is the rate of people who were in L_S category and have switched to L_C category because these people have improved their skills enough for the L_C category.

$\chi(t)$: The rate of people who have switched from L_C category to L_S category because AI technologies keep improving and, as a result, these people who were in L_C category can be fully or partially replaced by AI.

$\phi_C(t)$: Inflows into L_C that come from anything other than $\Psi_L(t)$; $\phi_C(t)$ includes inflows from immigration or education that feed complementary occupations; examples: AI-engineer migrants and computer science graduates with great AI-skills.

$\phi_S(t)$: Inflows into L_S that come from anything other than $\chi(t)$; $\phi_S(t)$ are new workers whose skills are substitutable by AI; examples: high school dropouts and high school graduates, who have AI-substitutable skills and decide to work full-time, instead of pursuing higher education.

$\phi_C(t)$, $\phi_S(t)$, $\Psi_L(t)$ and $\chi(t)$ are defined as followed:

$$\phi_C(t) = \eta_C(t) \cdot \bar{L}(t)$$

$\eta_C(t)$: Baseline entry rate into L_C

$\bar{L}(t)$: Eligible working age population

$$\phi_S(t) = \pi_S(t) \cdot \bar{L}(t)$$

$\pi_S(t)$: Baseline entry rate into L_S

$\bar{L}(t)$: Eligible working age population

$$\Psi_L(t) = \vartheta_L(t) \cdot L_S(t)$$

$\vartheta_L(t)$: Share of L_S who can and do improve their skills enough to become part of L_C

$$\chi(t) = \xi \cdot L_C(t) \cdot \log(1 + [S(t) \cdot D(t)])$$

ξ : Erosion coefficient, which reflects the sensitivity of complementary roles to being reclassified as substitutable

For the labor force part, I now discuss the version 2 of my model. $L(t)^{1-\alpha-\beta}$ in the MRW model is replaced by the following (in my version 2 model): $L_{AI}(t)^{\varnothing_1} \cdot L_{NAI}(t)^{\varnothing_2}$. In the version 2, total labor force is categorized based on whether it is AI-specific or not. Unlike in the version 1, there is no migration between the 2 categories in the version 2. To have AI-specific skills, an individual must have strong computer science, computer engineering, or mathematics background. This fact prevents meaningful migration between L_{AI} and L_{NAI} . People, who do not have at least one of these backgrounds, cannot become part of L_{AI} overnight; they must learn for years before they can be part of L_{AI} . People in L_{AI} do not want to be part of L_{NAI} ; maybe there are very few people who want to, but it is safe to assume that there is no meaningful migration from L_{AI} to L_{NAI} . Therefore, in the version 2 of my model, there is no migration between L_{AI} and L_{NAI} . The growth of L_{AI} is: $g_{L_{AI}} = \dot{L}_{AI}/L_{AI}$. The growth rate of L_{NAI} is: $g_{L_{NAI}} = \dot{L}_{NAI}/L_{NAI}$. The growth rate contribution of $L_{AI}^{\varnothing_1} = \varnothing_1 \cdot g_{L_{AI}}$. The growth rate contribution of $L_{NAI}^{\varnothing_2} = \varnothing_2 \cdot g_{L_{NAI}}$.

Table 2 summarizes and compares MRW and the two versions of my model.

Table 2 Comparisons between MRW and 2 Versions of Proposed Model

	MRW	My Model (Version 1)	My Model (Version 2)
Y (Total Output)	Technologies Part · Physical Capital Part · Human Capital Part · Labor Force Part		
Technologies Part	$[A(t)]^{1-\alpha-\beta}$	$A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)}$	$A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)}$
Physical Capital Part	$K(t)^\alpha$	$K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2}$	$K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2}$
Human Capital Part	$H(t)^\beta$	$[H_{C0} \cdot e^{\text{het}} \cdot (1 + \mu \log(SD)) + H_{S0} \cdot e^{\text{hst}} \cdot (1 - \nu \log(SD))]^\delta$	$H_{AI}(t)^{\varnothing_1} \cdot H_{NAI}(t)^{\varnothing_2}$

Labor Force Part	$[L(t)]^{1-\alpha-\beta}$	$L_C(t)^{\beta_1} \cdot L_S(t)^{\beta_2}$	$L_{AI}(t)^{\varnothing_1} \cdot L_{NAI}(t)^{\varnothing_2}$
-------------------------	---------------------------	---	--

Y (total output) in version 1 is different from Y in version 2. Both the technologies part ($A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)}$) and the physical capital part ($K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2}$) are exactly the same in both versions of my model, as clearly shown in Table 2. For the labor force part, logically, if it is the same economy, then the total labor force must be the same under both version 1 and version 2. Therefore, $L_C + L_S = L_{AI} + L_{NAI}$. However, $L_C(t)^{\beta_1} \cdot L_S(t)^{\beta_2} \neq L_{AI}(t)^{\varnothing_1} \cdot L_{NAI}(t)^{\varnothing_2}$. Furthermore, $[H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD)) + H_{S0} \cdot e^{hst} \cdot (1 - \nu \log(SD))]^\delta \neq H_{AI}(t)^{\varnothing_1} \cdot H_{NAI}(t)^{\varnothing_2}$. Just like MRW model, both versions of my model are intended to be mostly applied in cross-sectional and panel regressions across countries. Therefore, it does not matter if Y in version 1 differs from Y in version 2. Version 1 should be compared with version 1, and version 2 should be compared to version 2 in cross-countries analysis.

Version 1 is more focused in demonstrating the impact of self-learning AI on human capital and the labor force; who benefit and who are hurt by deployed self-learning AI, as well as the magnitude of the gain and the pain from the deployed self-learning AI, should be better understood by version 1. **Version 2 is to calculate S, which applies to both version 1 and version 2.** I use data in version 2 to calculate S in version 1, which is the same as S in version 2.

In a world without AI, both versions of my model are just MRW model. In a world without AI, the recursive exponent $\theta(t)$ is just zero (in both versions of my production function). This means that $[S(t) \cdot D(t)]^{\theta(t)} = [S(t) \cdot D(t)]^0 = 1$. Therefore, in both versions of my model, the technologies part is just: $A_0 \cdot T(t)^\lambda$, which is the same as $[A(t)]^{1-\alpha-\beta}$ in MRW, since $A_0 \cdot T(t)$ is the same as $A(t)$ (stated in a different way), with different exponent name but essentially the same exponent value. In both versions of my model and in the world without AI, K_{AI} does not exist and, therefore, the physical capital part is just $K_G(t)^{\alpha_2}$, which is the same as $K(t)^\alpha$ in MRW model (just with different base name and different exponent name but the value of the base and the exponent are the same in the world without AI). In the world without AI and version 1 of my production function, there is no need to categorize L into L_C and L_S since AI does not exist and, therefore, there is nothing to be complemented or substituted by AI. Thus, the labor part of my version 1 is just: $L(t)^\beta$, which is the same as $L(t)^{1-\alpha-\beta}$ in MRW model (just with different exponent name but the same exponent value). In the world without AI and in my version 2 production function, L_{AI} does

not exist and, therefore, the labor part is just: $L_{NAI}(t)^{\varnothing_2}$, which is the same as $L(t)^{1-\alpha-\beta}$ (just with different base name and different exponent name but the same base value and the same exponent value). In the version 1 of my production function, the human capital part is the following: $[H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD)) + H_{S0} \cdot e^{hst} \cdot (1 - \nu \log(SD))]^\delta$. In the world without AI, μ (elasticity that determines how much smarter AI-complementary workers get as self-learning AI scale grows) is zero, and ν (elasticity of how badly self-learning AI substitutes AI-substitutable human capital) is also zero because there is no self-learning AI. This means that the human capital part is only the following: $[(H_{C0} \cdot e^{hct}) + (H_{S0} \cdot e^{hst})]^\delta$. $(H_{C0} \cdot e^{hct})$ represents H_C , and $(H_{S0} \cdot e^{hst})$ represents H_S . Since there is nothing to be complemented or substituted by AI (in the world without AI), there is no need to separate H into H_C and H_S . Therefore, the human capital part is just: $H(t)^\delta$, which is the same as $H(t)^\beta$ in MRW model (just with different exponent name but the same exponent value). In the world without AI and version 2 of my production function, the human capital part is just: $H_{NAI}(t)^{\varnothing_2}$ (since H_{AI} does not exist), which is the same as $H(t)^\beta$ in MRW model (just with different base name and different exponent name but the value of both the base and the exponent are the same). Furthermore, in the world without AI, it only makes sense to combine A with L again because, without AI, both A and L become exogenous. Therefore, in the world without AI, both my production functions are truly just MRW production function. MRW model is correct and sufficient at the time it was introduced because there is no deployed self-learning AI in 1992. However, for the present time and anytime in the future, MRW model is insufficient because deployed self-learning AI already exist and will only grow exponentially; this paper is intended to introduce a model that is appropriate for the present time and anytime in the future.

As far as I am aware, **there is no empirical data about the impact of deployed self-learning AI on total output and, therefore, I cannot use empirical data to support my model.** However, in the future, I expect that empirical data will confirm the usefulness of my proposed model. At present time, I can prove my model with logics. As already mentioned above, my production functions are as followed:

$$Y(t) = A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)} \cdot K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2} \cdot L_C(t)^{\beta_1} \cdot L_S(t)^{\beta_2} \cdot [H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD)) + H_{S0} \cdot e^{hst} \cdot (1 - \nu \log(SD))]^\delta \text{ (Note: this is version 1).}$$

$$Y(t) = A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)} \cdot K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2} \cdot L_{AI}(t)^{\varnothing_1} \cdot L_{NAI}(t)^{\varnothing_2} \cdot H_{AI}(t)^{\varnothing_1} \cdot H_{NAI}(t)^{\varnothing_2} \text{ (Note: this is version 2).}$$

The following few paragraphs are my proofs using logics.

First, self-learning AI already exist and have already been deployed but only have been used by a few parties, including top universities, labs and large tech companies. Since self-learning AI are expected to provide massive gain in productivities, it is only a matter of time before general public use self-learning AI directly. Since self-learning AI have recursive component, there is a need for a production function that takes into account the recursive element of self-learning AI.

Second, since both MRW model and Romer (1990) model are already widely accepted, then, as long as economists agree with my logics in the previous, this and the next paragraph, both versions of my model should also be widely accepted. I mention MRW because each of my production functions consists of technologies part, physical capital part, labor force part and human capital part, which is exactly like the MRW model. However, there is one very important difference between MRW, and both my production functions. In MRW, A and L cannot be separated because both A and L are exogenous. In both of my production functions, technologies and labor forces are separated, and this is fine because, in both my production functions, both technologies and labor forces are endogenous. This is why I mention Romer (1990) because, in Romer (1990) paper, Romer states that A is endogenous, which means that A can be modeled independently of L. Romer (1990) broke the Uzawa's restriction by making A endogenous and, Romer's model is widely accepted. Both of my production functions have the structures that are already widely accepted by economists. To adjust with the current and future economy, in which deployed self-learning AI is increasingly affecting total output, I have created two production functions that take into account self-learning AI. In both of my production functions, I just separate $A(t)^{1-\alpha-\beta}$ in the MRW model into non-self-learning technologies ($A_0 \cdot T(t)^\lambda$ in my model) and self-learning technologies ($A_0 \cdot [S(t) \cdot D(t)]^{\theta(t)}$ in my model). Furthermore, in my version 2 production function, for K, L and H in the MRW model, I separate each of them into AI-specific and non-AI. Given the exponential increase in productivity due to self-learning AI, these separations are necessary. I do not separate K based on whether a capital is for self-learning AI or not because both self-learning and non-self-learning AI use many of the same capital, such as GPUs, power and data centers; in addition, non-self-learning AI can be improved into self-learning AI. These facts make it almost impossible to separate based on self-learning capability. That is why, for K, I separate based on whether a capital is for AI or not. Similar logic applies to why I separate L and H based on whether it is AI-specific or not. To have AI-specific skills, an individual must have strong background in computer science,

computer engineering or mathematics so there is high barrier of entry to have H_{AI} and to be part of L_{AI} . People who are in L_{AI} tend to remain in L_{AI} , and people who are in L_{NAI} tend to remain in L_{NAI} . However, the barrier between engineers, who can build non-self-learning AI but not self-learning AI, and engineers, who can build self-learning AI, is not high. In fact, self-learning AI are built by engineers who previously could not build self-learning AI. AI engineers, who cannot build self-learning AI yet, can learn how to build self-learning AI much faster than non-AI workers (workers who belong in L_{NAI}). That is why, in version 2 of my model, I separate L and H based on whether it is AI-specific or not; I did not separate them based on whether it is self-learning-specific or not.

Third, since nowadays many people are concerned about what AI can do to the labor forces, there is a need of a new production function that helps in addressing this concern. The version 2's technologies and physical capital part, which is: $A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)} \cdot K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2}$, is exactly the same as version 1's. In version 1, for the human capital part and the labor force part, instead of separating based on whether they are AI-specific or non-AI-specific, I separate the human capital and the labor force based on whether they are complemented or substituted by AI. AI-complementary human capital should grow as self-learning AI grow; this fact is reflected in the following portion of the equation: $(1 + \mu \log(SD))$. AI-substitutable human capital should shrink as self-learning AI grow; this is reflected in the following portion of the equation: $(1 - \nu \log(SD))$.

4.2 Balanced Growth Path (BGP), Accelerating Growth Path or Singularity.

Self-learning AI can become smarter, more capable and more productive over time without any new human capital or any additional labor. Therefore, self-learning AI is different from any other technology that has existed; self-learning AI can turn an economy that is in balanced growth path into an economy that is in accelerating growth path. In both my production functions, whether an economy is in balanced growth path or accelerating growth path, it depends on the value of $\theta(t)$, the recursive learning exponent. In the version 1 production function, the sum of all exponents, excluding $\theta(t)$, is the following: $\lambda + \alpha_1 + \alpha_2 + \beta_1 + \beta_2 + \delta < 1$. In the version 2 production function, the sum of all exponents, excluding $\theta(t)$, is the following: $\lambda + \alpha_1 + \alpha_2 + \varnothing_1 + \varnothing_2 + \partial_1 + \partial_2 < 1$. In both my production functions, if $\theta(t)$ is small, then the sum of all exponents, including $\theta(t)$, is equal to or less than 1, which means an economy is in balanced growth path; if $\theta(t)$ is large enough, then the sum of all exponents, including $\theta(t)$, is larger than 1, which means an economy is in

accelerating growth path. If the sum of all exponents, including $\theta(t)$, equals to 1, then there is constant returns to scale (CRS); if the sum of all exponents, including $\theta(t)$, is less than 1, then there is decreasing returns to scale (DRS). If $\theta(t)$ is large enough to make an economy to be in accelerating growth path, then total output must be increasing at increasing rate; this must be caused by contributions from self-learning AI that increase at increasing rate. Aghion, Jones and Jones explicitly support that self-learning AI increase productivity at increasing rate over time (Aghion et al., 2017). Trammell and Korinek also explicitly support the idea that sufficiently advanced self-learning AI lead to increasing productivity at increasing rate (Trammell & Korinek 2023).

Singularity is not possible in both versions of my production function. By definition, singularity must satisfy the following: $\lim_{t \rightarrow t^*} Y(t) = \infty$ in finite time, which is not possible in each of my production functions. In both versions, the only block that could in principle drive an explosion is the self-learning term $SL(t) = [S(t) \cdot D(t)]^{\theta(t)}$ with $\theta(t) = \theta_0 + \rho \cdot \log(S(t) \cdot D(t))$. Let $W(t) = \log(S(t) \cdot D(t))$. This means: $\log SL(t) = \theta(t) \cdot W(t)$. Using $\dot{\theta}(t) = \rho \dot{W}(t)$ together with $\dot{W}(t) = g_S(t) + g_D(t)$, I get the following differential identity: $\frac{d}{dt} \log SL(t) = \dot{W}(t) [\theta_0 + 2\rho W(t)]$. Because $S(t) = s_0 \cdot K_{AI}(t)^\tau \cdot H_{AI}(t)^\varpi \cdot L_{AI}(t)^\Gamma$, its growth rate is a finite linear combination of the finite growth rates of K_{AI} , H_{AI} and L_{AI} ; hence, g_S is finite. $D(t)$ is treated with an ordinary finite growth rate g_D . Let $\dot{W}(t) = r(t)$. This means: $\dot{W}(t)$, which equals to: $g_S(t) + g_D(t)$, is bounded on every finite horizon by some $\bar{r} < \infty$. Integrating, $W(t) = W(0) + \int_0^t r(t) dt \leq W(0) + \bar{r}t$. Substituting this bound back gives: $\frac{d}{dt} \log SL(t) \leq \bar{r} [\theta_0 + 2\rho(W(0) + \bar{r}t)] = A + Bt$ for constants A, B . A single integration yields: $\log SL(t) \leq \log SL(0) + At + \frac{B}{2} t^2 = c_0 + c_1t + c_2t^2$. Therefore, $SL(t) \leq \exp(c_0 + c_1t + c_2t^2)$; the self-learning block can accelerate (quadratic exponent) but cannot blow up at any finite time. Since $SL(t) \leq \exp(c_0 + c_1t + c_2t^2)$ on any finite horizon, and all remaining multiplicative/sum factors in both versions are standard Cobb-Douglas or exponential terms with finite growth rates, every term is finite on $[0, T]$. Therefore, $Y(t)$ is finite for all finite t and neither version admits a finite-time singularity. This paragraph has proven that each of my production functions cannot lead to singularity.

4.3 AI's Dependency on K_{AI}

Deployed self-learning AI do not need any *additional* human capital or labor, in order to become smarter, more capable and more productive but self-learning AI will always need more of K_{AI} in the form of GPUs, data storage and memory systems, power, etc. For example, if deployed self-learning AI do not get more GPUs, then the self-learning AI will become static, which means they no longer self-learn. For deployed self-learning AI to perform optimally and continuously, growing K_{AI} is absolutely necessary. There should be clear positive correlation between deployed self-learning AI and K_{AI} . However, the strength of the correlation depends on how much progress AI engineers have made in power efficiency. The more power efficient, the lower the correlation between deployed self-learning AI and K_{AI} . Furthermore, not all K_{AI} must be added continuously for deployed self-learning AI to perform optimally and continuously. Examples of K_{AI} that do not need to be added continuously are model architecture design, training codebase and orchestration tools. These costs are one-time fixed AI-specific capital costs.

Since the value of S depends on K_{AI} , insufficient amount of K_{AI} will not create a situation where productivity increases at increasing rate. The only variable (in both my production functions) that determines whether an economy is in balanced growth path or accelerating growth path is: $\theta(t)$, which is defined as followed: $\theta(t) = \theta_0 + \rho \cdot \log(S(t) \cdot D(t))$. Since $S(t)$ is defined as followed: $S(t) = s_0 \cdot K_{AI}(t)^\tau \cdot H_{AI}(t)^\omega \cdot L_{AI}(t)^\Gamma$, this means that the value of $\theta(t)$ depends not only on $S(t)$ directly but also on K_{AI} , H_{AI} and L_{AI} indirectly. If K_{AI} becomes scarce, then S becomes smaller than what it could have been if K_{AI} is not scarce. For example, if rare earths, which is one of the necessary K_{AI} , becomes scarce, then S should be smaller than S in a scenario, in which rare earths is abundant. Since $\theta(t)$ becomes smaller when S becomes smaller and S becomes smaller when K_{AI} becomes scarce, if K_{AI} becomes small enough, then the sum of all exponents (in each of my production functions) becomes 1 or less than 1. In such case, an economy is in balanced growth path, not accelerating growth path. The value of S and, therefore, the value of $\theta(t)$, also depend on H_{AI} and L_{AI} . However, given how fast AI revolution is going at the present time, I assume that there will be no lack of H_{AI} and L_{AI} in the future. Physical capital can be scarce but human capital can always be developed and, at present time, talents for AI are growing in fast speed. At present time, there are great appetites to innovate, create and improve the best AI; with this present

situation, lack of H_{AI} or lack of L_{AI} is very unlikely in the future. Therefore, I assume that lack of K_{AI} is possible, but lack of H_{AI} or L_{AI} is very unlikely in the future.

Table 3 Consequences of Different Values of K_{AI}

Situation	Consequences in Each of My Production Functions	Consequences for the Economy
K_{AI} is large enough	Sum of all exponents in the production function is larger than 1	Economy in accelerating growth path
K_{AI} becomes too small	Sum of all exponents in the production function is 1 or less than 1	Economy in balanced growth path

At present time, deployed self-learning AI are still rare. Therefore, today, we probably still live in an economy that is in balanced growth path because both S and, therefore, $\theta(t)$ are still very small. Today, H_{AI} is still small, but it will grow exponentially in the next many years; the evidence of this claim can be seen by everyone. Today, only very few people can build self-learning AI, and nobody can build AGI yet but the current pace of the improvements in AI technologies is extremely fast; this very fast pace of improvements in AI technologies is a strong indication that H_{AI} is also increasing at very high rate. Today, U.S economy is probably still in balanced growth path because H_{AI} is still not large enough but, in the future, H_{AI} will be large enough and the only variable (in both proposed production functions) that is most likely to make us still living in an economy with balanced growth path is K_{AI} . Today, probably, the sum of all exponents (in both versions of my production function) is still 1 or less than 1. However, in the next few years, as predicted by many AI experts, deployed self-learning AI will grow at very high rate. I predict that, within the next 5 to 10 years, we already live in an economy, in which general public use self-learning AI. In this prediction, S and, therefore, $\theta(t)$ will be large enough to make us live in an economy that is in accelerating growth path. However, if for whatever reason, K_{AI} becomes insufficient, then we may still live in a balanced growth path economy in the next 5 to 10 years. K_{AI} includes many necessary elements, including power, rare earths elements, data centers, semiconductors and many other essential elements. If any of the necessary elements or a combination of the necessary elements becomes scarce, then the value of S and, therefore, $\theta(t)$ will be lower than anticipated.

4.4 The Future of AI-Specific Human Capital and Labor

$L_{AI} = L_{SLAI} + L_{NSLAI}$. L_{SLAI} is labor that can build self-learning AI, and L_{NSLAI} is labor that cannot build self-learning AI but can build non-self-learning AI. $H_{AI} = H_{SLAI} + H_{NSLAI}$. H_{SLAI} is human

capital that can build self-learning AI, and H_{NSLAI} is human capital that cannot build self-learning AI but can build non-self-learning AI. Since there are great and growing appetites to increase S as fast as possible, the followings are expected:

$$\frac{d}{dt}L_{SLAI}(t) > 0; \frac{d}{dt}H_{SLAI}(t) > 0$$

What will happen to L_{NSLAI} is unclear because there are two forces moving in the opposite direction. The first force, which reduces L_{NSLAI} , is migrations from L_{NSLAI} to L_{SLAI} as some AI-specific workers, who previously could not build self-learning AI, have improved their skills enough so that they now can build self-learning AI. The second force, which increases L_{NSLAI} , is newcomers into L_{NSLAI} ; these newcomers include some recent computer science graduates, immigrants with non-self-learning AI skills, etc. H_{NSLAI} will increase over time because non-self-learning AI can be improved to become a better version of non-self-learning AI, and the increase in the skills to do this transformation increases H_{NSLAI} .

5 More Implications of the Existence of Self-Learning AI

I now discuss the version 1 of my model until I state otherwise. My version 1 of production function is as followed:

$$Y(t) = A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{\theta(t)} \cdot K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2} \cdot L_C(t)^{\beta_1} \cdot L_S(t)^{\beta_2} \cdot [H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD))] + H_{S0} \cdot e^{hst} \cdot (1 - \nu \log(SD))^\delta$$

An increase in any variable or parameter in the equation above, with the exception of ν (which is the elasticity of how badly self-learning AI substitutes AI-substitutable human capital), increases the total output, $Y(t)$, *ceteris paribus*.

As S (deployed self-learning AI) get larger, AI-substitutable human capital (which is represented by the following: $H_{S0} \cdot e^{hst} \cdot (1 - \nu \log(SD))$) becomes smaller while AI-complementary human capital (which is represented by the following: $H_{C0} \cdot e^{hct} \cdot (1 + \mu \log(SD))$) becomes larger. Furthermore, since $\frac{d}{dt} L_S(t) = \phi_S(t) - \Psi_L(t) + \chi(t)$, and $\chi(t) = \xi \cdot L_C(t) \cdot \log(1 + [S(t) \cdot D(t)])$, if S (deployed self-learning AI) become smarter, more widespread and more productive, then L_S increase, which means that the number of workers in AI substitutable roles increases. Moreover, since $\frac{d}{dt} L_C(t) = \phi_C(t) + \Psi_L(t) - \chi(t)$, if S increase, then L_C decrease, which means the number of workers in AI complementary roles decreases. However, on average, each worker in the L_C category has become more valuable as AI-complementary human capital becomes larger and the

size of L_C decreases when S , deployed self-learning AI, become smarter, more widespread and more productive. For workers in the L_S category, when S improve, AI-substitutable human capital decreases and the size of L_S gets larger; these cause, on average, each worker in the L_S category to become less valuable. Because S (deployed self-learning AI) use data to improve themselves and the amount and the quality of the improvements depend on D (the volume and the quality of the data available for learning), S and D have the same effects on human capital and labor force. Just like S , if D increases, then AI-substitutable human capital becomes smaller while AI-complementary human capital becomes larger. Moreover, if D increases, then L_S increases and L_C decreases. D also has the same directional effect as S on the average value of workers in the L_S category and on the average value of workers in the L_C category. My model properly reflects the importance of high-quality large amount of data on the capabilities and productivities of self-learning AI, and therefore the total output. The greater the amount and the quality of the data that are used by self-learning AI, the greater the productivities of self-learning AI and therefore the total output. This claim is supported by many works, including the paper by Mohammed et al; in this paper, the authors provide a large-scale empirical study which shows that 20%-30% data degradation causes drops in performances by as much as 10%-40% (Mohammed et al., 2025).

I now discuss the version 2 of my model. The production function of my version 2 is as followed:

$$Y(t) = A_0 \cdot T(t)^\lambda \cdot [S(t) \cdot D(t)]^{0(t)} \cdot K_{AI}(t)^{\alpha_1} \cdot K_G(t)^{\alpha_2} \cdot L_{AI}(t)^{\varnothing_1} \cdot L_{NAI}(t)^{\varnothing_2} \cdot H_{AI}(t)^{\hat{\varnothing}_1} \cdot H_{NAI}(t)^{\hat{\varnothing}_2}$$

In the version 2, an increase in any variable or parameter in the equation above will increase the total output, $Y(t)$, *ceteris paribus*. As already discussed above, $S(t) = s_0 \cdot K_{AI}(t)^\tau \cdot H_{AI}(t)^\varpi \cdot L_{AI}(t)^\Gamma$. This means that version 2 of my production function can stand alone while version 1 of my production function cannot stand alone without data from the version 2. To calculate S , the value of H_{AI} and L_{AI} are needed, and these values only exist in the version 2 production function. This why I introduce 2 production functions in this paper. Without the version 2, the version 1 is incomplete. Both my production functions are very important.

My equations apply to both rich countries and poor countries but the value of each variable and each parameter in my equations can differ very greatly between rich countries and poor countries. Even among rich countries themselves, there are significant differences between the rich countries that have institutions, which produce a lot of most advanced AI, and the rich countries that do not have institutions, which produce a lot of most advanced AI. For example, there are

great differences between the United States, which produces a lot of most advanced AI, and several rich European countries (such as Germany, France and U.K), which do not produce a lot of most advanced AI and mostly only use a lot of advanced AI. Because of the continuous recursive nature of self-learning AI, the nations that fall behind in self-learning AI will very likely have great difficulties in catching up.

6 Policy Implications from the Proposed Production Functions

Both versions of my model identify three critical drivers of output growth: $S(t)$ (deployed self-learning AI), $D(t)$ (quality and volume of data for learning), and $K_{AI}(t)$ (AI-specific physical capital). Policies that increase these variables, or parameters linked to them, move the economy from balanced growth toward accelerating growth. Below are some policy recommendations.

Rich AI-Producing Countries (S , D , and K_{AI} are the highest or one of the highest)

- Increase K_{AI} via expanded AI R&D fundings, infrastructure investments, and energy capacity.
- Raise H_{AI} and L_{AI} through more aggressive advanced STEM education, better immigration policies for AI talent, and more PhD scholarships.
- Sustain D with national-scale data infrastructure and governance to ensure high-quality datasets.
- Help AI-substitutable labor forces through targeted upskilling.

Poor Countries (S , D and K_{AI} are still extremely low)

- Focus on K_G (general capital) and all human capital before investing heavily in K_{AI} .
- Use affordable AI tools to improve productivity in non-frontier sectors, indirectly raising L_C and H_C .
- Attract foreign investment to incrementally increase S and K_{AI} without diverting scarce resources from basic infrastructure.

Rich AI-Consuming Countries (S , D and K_{AI} are much lower than rich AI-producing countries but much higher than poor countries)

- Increase H_{AI} and L_{AI} by scaling AI-related education and research capacity.
- Expand K_{AI} through coordinated international investment in chips, data centers, and model development.
- Raise D through data-sharing agreements and joint ventures with leading AI producers.

7 Conclusion

This paper isolates the mechanism that makes modern AI macro-relevant: recursive learning. By disaggregating technology into a non-self-learning component $T(t)^\lambda$ and a self-learning term $(S(t) \cdot D(t))^{\theta(t)}$ with $\theta(t) = \theta_0 + \rho \cdot \log(S(t) \cdot D(t))$, the framework turns the vague idea that “AI improves itself” into a tractable object inside aggregate production. Two complementary specifications implement this object: Version 1 makes the distributional channels transparent (AI-complementary vs. AI-substitutable labor and human capital), while Version 2 renders the self-learning stock measurable from AI-specific capital, labor, and human capital.

The analysis delivers a sharp regime map. When $\theta(t)$ is small/locally constant—i.e., the recursive force is present but not dominant—both specifications admit a balanced growth path with constant factor–output ratios; when deployment and data raise $\theta(t)$ enough to push effective returns above one, the economy transitions to sustained acceleration. The recursion operates in log-space, implying that $\log((SD)^{\theta(t)})$ grows at most quadratically in time; thus, the model permits acceleration without finite-time blow-ups. These results are derived within a production structure that remains familiar—Cobb–Douglas blocks and bounded trend growth—so the interpretation of elasticities and steady-state objects stays clear.

Economically, the framework explains why self-learning AI is a powerful gap-amplifier. The term $(SD)^{\theta(t)}$ formalizes the complementarity between deployment scale and data quality, making countries that both build and deploy frontier systems drift away from pure users. It also clarifies distribution: AI-complementary skills appreciate; AI-substitutable tasks depreciate; and the availability of AI-specific capital (compute, energy, data-center capacity) and high-quality data emerges as the possible critical bottleneck. Policy levers therefore have clean targets: scale K_{AI} where power and chips bind, improve data governance/quality where information is the constraint, and expand H_{AI} , L_{AI} where talent limits $\theta(t)$.

Finally, the framework is immediately usable. Version 2 operationalizes $S(t)$ for calibration and cross-country comparison, and the technology block implies regression-ready structure: empirical specifications should include both linear and quadratic terms in $\log(SD)$ to capture the signature of recursion. In this sense, the model functions as a benchmark: simple enough to slot into standard growth accounting, precise enough to discipline debates about AI’s long-run impact, and transparent about the conditions under which economies remain balanced or accelerate.

Statements and Declarations

No funding was received to assist with the preparation of this manuscript.

The author has no competing interests to declare that are relevant to the content of this article.

References

- Acemoglu, D., & Restrepo, P. (2018). The race between man and machine: Implications of technology for growth, factor shares, and employment. *American Economic Review*, 108(6), 1488–1542. <https://doi.org/10.1257/aer.20160696>
- Acemoglu, D. (2024, April 5). *The simple macroeconomics of AI* [Unpublished manuscript]. Massachusetts Institute of Technology. Prepared for *Economic Policy*.
- Aghion, P., Jones, B. F., & Jones, C. I. (2017). *Artificial intelligence and economic growth* (NBER Working Paper No. 23928). National Bureau of Economic Research. <https://doi.org/10.3386/w23928>
- Bahri, Y., Dyer, E., Kaplan, J., Lee, J., & Evci, U. (2024). Explaining neural scaling laws. *Proceedings of the National Academy of Sciences*, 121(6), e2311878121. <https://doi.org/10.1073/pnas.2311878121>
- Besiroglu, T., Emery-Xu, N., & Thompson, N. (2023). Economic impacts of AI-augmented R&D. *arXiv*. <https://arxiv.org/abs/2212.08198v2>
- Farach, A., Cambon, A., & Spataro, J. (2025). Evolving the productivity equation: Should digital labor be considered a new factor of production? *arXiv*. <https://arxiv.org/abs/2505.09408v1>
- Farboodi, M., Mihet, R., Philippon, T., & Veldkamp, L. (2019). Big data and firm dynamics. *AEA Papers and Proceedings*, 109, 38–42. <https://doi.org/10.1257/pandp.20191001>
- Hestness, J., Narang, S., Ardalani, N., Diamos, G., Jun, H., Kianinejad, H., Patwary, M. A., Yang, Y., & Zhou, Y. (2017). Deep learning scaling is predictable, empirically. *arXiv*. <https://arxiv.org/abs/1712.00409>
- Hoffmann, J., Borgeaud, S., Mensch, A., Buchatskaya, E., Cai, T., Rutherford, E., de Las Casas, D., et al. (2022). Training compute-optimal large language models. In *Advances in Neural Information Processing Systems* (Vol. 35, pp. 30016–30030).

- https://proceedings.neurips.cc/paper_files/paper/2022/hash/c1e2faff6f588870935f114ebe04a3e5-Abstract-Conference.html
- Ichihashi, S. (2021). The economics of data externalities. *Journal of Economic Theory*, 196, 105316. <https://doi.org/10.1016/j.jet.2021.105316>
- Jacobo-Romero, M., Carvalho, D. S., & Freitas, A. (2022). Estimating productivity gains in digital automation. *arXiv*. <https://arxiv.org/abs/2210.01252v2>
- Jones, C. I. (1995). R&D-based models of economic growth. *Journal of Political Economy*, 103(4), 759–784. <https://doi.org/10.1086/261995>
- Jones, C. I., & Tonetti, C. (2020). Nonrivalry and the economics of data. *American Economic Review*, 110(9), 2819–2858. <https://doi.org/10.1257/aer.20191330>
- Kaplan, J., McCandlish, S., Henighan, T., Brown, T. B., Chess, B., Child, R., Gray, S., et al. (2020). Scaling laws for neural language models. *arXiv*. <https://arxiv.org/abs/2001.08361>
- Lucas, R. E. (1988). On the mechanics of economic development. *Journal of Monetary Economics*, 22(1), 3–42. [https://doi.org/10.1016/0304-3932\(88\)90168-7](https://doi.org/10.1016/0304-3932(88)90168-7)
- Mankiw, N. G., Romer, D., & Weil, D. N. (1992). A contribution to the empirics of economic growth. *Quarterly Journal of Economics*, 107(2), 407–437. <https://doi.org/10.2307/2118477>
- Mohammed, S., Budach, L., Feuerpfeil, M., Ihde, N., Nathansen, A., Noack, N., Patzlaff, H., Naumann, F., & Harmouch, H. (2025). The effects of data quality on machine learning performance on tabular data. *arXiv*. <https://arxiv.org/abs/2207.14529v6>
- Puaschunder, J. M. (2022). Extension of endogenous growth theory: Artificial intelligence as a self-learning entity. In *Proceedings of the Research Association for Interdisciplinary Studies (RAIS) Conference* (October 23–24, 2022). <https://doi.org/10.5281/zenodo.7372430>
- Romer, P. M. (1990). Endogenous technological change. *Journal of Political Economy*, 98(5, Part 2), S71–S102. <https://doi.org/10.1086/261725>
- Schrittwieser, J., Antonoglou, I., Hubert, T., Simonyan, K., Sifre, L., Schmitt, S., Guez, A., et al. (2020). Mastering Atari, Go, Chess and Shogi by planning with a learned model. *Nature*, 588(7839), 604–609. <https://doi.org/10.1038/s41586-020-03051-4>

- Shojaee, P., Mirzadeh, I., Alizadeh, K., Horton, M., Bengio, S., & Farajtabar, M. (2025). The illusion of thinking: Understanding the strengths and limitations of reasoning models via the lens of problem complexity. *arXiv*. <https://arxiv.org/abs/2507.12345>
- Silver, D., Schrittwieser, J., Simonyan, K., Antonoglou, I., Huang, A., Guez, A., Hubert, T., et al. (2017). Mastering the game of Go without human knowledge. *Nature*, 550(7676), 354–359. <https://doi.org/10.1038/nature24270>
- Solow, R. M. (1956). A contribution to the theory of economic growth. *Quarterly Journal of Economics*, 70(1), 65–94. <https://doi.org/10.2307/1884513>
- Trammell, P., & Korinek, A. (2023). Economic growth under transformative AI (NBER Working Paper No. 31065). National Bureau of Economic Research. <https://doi.org/10.3386/w31065>
- Wang, L., Sarker, P. K., Alam, K., & Sumon, S. (2021). Artificial intelligence and economic growth: A theoretical framework. *Scientific Annals of Economics and Business*, 68(4), 421–443. <https://doi.org/10.47743/saeb-2021-0027>