

Artificial Intelligence in the Knowledge Economy*

Enrique Ide & Eduard Talamàs

December 17, 2024

Abstract

The rise of Artificial Intelligence (AI) has the potential to reshape the knowledge economy by enabling problem solving at scale. This paper introduces a framework to analyze this transformation, incorporating AI into an economy where humans form hierarchical firms to use their time efficiently: Less knowledgeable individuals become “workers” solving routine problems, while more knowledgeable individuals become “solvers” assisting workers with exceptional problems. We model AI as a technology that transforms computing power into “AI agents,” which can either operate autonomously (as co-workers or solvers/co-pilots) or non-autonomously (only as co-pilots). We show that basic autonomous AI displaces humans towards specialized problem solving, leading to smaller, less productive, and less decentralized firms. In contrast, advanced autonomous AI reallocates humans to routine work, resulting in larger, more productive, and more decentralized firms. While autonomous AI primarily benefits the most knowledgeable individuals, non-autonomous AI disproportionately benefits the least knowledgeable. However, autonomous AI achieves higher overall output. These findings reconcile seemingly contradictory empirical evidence and reveal key tradeoffs involved in regulating AI autonomy.

*Department of Economics, IESE Business School, Carrer de Arnús i de Garí 3-7, 08034 Barcelona, Spain (eide@iese.edu and etalamas@iese.edu). We extend our deepest gratitude to Luis Garicano, whose invaluable suggestions have fundamentally shaped the presentation of this work. We have also benefited enormously from the comments of Ricardo Alonso, Nano Barahona, Nicolás Figueroa, Andrea Galeotti, Ben Golub, Jason Hartline, Niko Matouschek, Mike Ostrovsky, Cristóbal Otero, Sebastián Otero, Esteban Rossi-Hansberg (the Editor), four anonymous referees, and seminar participants at Bocconi, Brown, BSE Summer Forum, BU, CMU-Tepper, Cornell, Duke-Fuqua, EC’24, ESAM 2024, HEC Paris, IESE, LBS, MIT-Sloan, NBER Org. Econ. (Spring 2024), Notre Dame, Nottingham, NU Kellogg-MEDS, NU Kellogg-Strategy, OpenAI, Penn State, PUC-Chile, Queen’s, SED BCN 2024, UAB, UAndes-Chile, UIB, UPenn, UPF, and UW-Madison. We acknowledge the financial support of IESE through the High Impact Initiative-course 2024/2025. We declare we have no relevant or material financial interests that relate to the research described in this paper.

1 Introduction

Artificial Intelligence (AI) foundation models—such as GPT-4, Gemini, and Claude Sonnet—are driving a new wave of automation, enabling machines to undertake sophisticated knowledge work such as coding, research, and problem-solving.¹ While the potential of this technology to reshape the landscape of work is undeniable, its precise implications remain the subject of growing controversy (Meserole, 2018; Brynjolfsson, 2022; Johnson and Acemoglu, 2023; Autor, 2024; Acemoglu, 2024). This controversy arises from two factors. First, it is unclear whether lessons from previous waves of automation—primarily aimed at repetitive tasks—can inform our understanding of AI’s effects (Muro et al., 2019; Agrawal et al., 2019). Second, because AI is still in its infancy and individuals and firms are experimenting with its use (McElheran et al., 2023; Humlum and Vestergaard, 2024), current empirical evidence cannot yet capture its full equilibrium effects.

In this paper, we provide a new framework for studying the equilibrium effects of AI on the future of work. The novelty of our approach lies in focusing on a specific type of automation to understand AI’s distinctive impact: the automation of knowledge work. This is an important part of the economy where production is constrained primarily by time and knowledge, making firms central to organizing and optimizing the use of available knowledge (Garicano and Rossi-Hansberg, 2015).² From this perspective, AI stands apart from previous automation technologies in its ability to alleviate both time and knowledge constraints, potentially leading to a fundamental reorganization of production.

More precisely, we embed AI in a canonical model of a knowledge economy: The knowledge hierarchies first introduced by Garicano (2000). Using this framework, we describe the equilibrium effects of AI on organizational and labor outcomes, including (i) occupational choice, (ii) firm productivity, size, and decentralization, and (iii) the distribution of labor income. Our analysis highlights how the impact of AI crucially depends on its problem-solving capabilities and its degree of autonomy.

Our starting point—the pre-AI economy—is the baseline model of Garicano and Rossi-Hansberg (2004), Antràs, Garicano, and Rossi-Hansberg (2006), and Fuchs, Garicano, and Rayo (2015). Labor and knowledge are the sole inputs in production. Humans are endowed with one unit of time and are heterogeneous in terms of knowledge. Individuals use their time to pursue production opportunities but encounter problems of varying difficulty during the production process. Output is produced when an individual can successfully solve the problem she confronts, which occurs when her knowledge exceeds the problem’s difficulty. If a human cannot solve a problem on her own, she may seek help from another human. Help, however, is costly in terms of time.

¹Foundation models are AI systems characterized by their ability to handle a wide variety of tasks through training on large-scale datasets using self-supervised learning techniques (Bommasani et al., 2022; Council of Economic Advisers, 2024). Although generative AI systems are the most prominent examples of foundation models today, the concept encompasses non-generative tasks as well, such as classification and prediction.

²Knowledge-intensive industries, such as legal and business services, healthcare, and finance, accounted for approximately 38% of U.S. GDP and 36% of Japan’s GDP in 2017 (National Science Board, 2018).

In the competitive equilibrium of the pre-AI economy, humans either pursue production opportunities on their own (becoming “independent producers”) or join hierarchical firms to optimize the use of their time and knowledge. These firms exhibit *management by exception*: Less knowledgeable individuals become “workers” who try to solve problems first, while more knowledgeable individuals become “solvers” specializing in assisting workers with the exceptional problems they are unable to solve.³

Our innovation is to incorporate AI into this otherwise canonical setting. Our approach is based on three key properties of modern AI, documented in Section 2. First, AI is inherently scalable: unlike human knowledge, whose use is constrained by the time of the individual possessing it, AI-encoded knowledge can be leveraged across all available computing power (“compute”). Second, the rise of flexible, general-purpose foundation models—adaptable to diverse tasks with minimal additional training—enables firms to rely on a shared AI model, eliminating the need to develop costly domain-specific models. Third, AI can harness compute to automate knowledge work.

Guided by these principles, we model AI as a technology that is available to all firms and that transforms compute into “AI agents,” all endowed with the same exogenously fixed level of AI-encoded knowledge. As a benchmark, we first consider the case where these agents can do exactly the same as humans. This AI is “autonomous” as it can function both as a “co-worker” (capable of pursuing production opportunities) and as a “co-pilot” (providing advice). A key contribution of our analysis is to contrast this benchmark with the case of “non-autonomous” AI, where—for technological or regulatory reasons—AI agents are limited to operating solely as co-pilots.

In the post-AI economy, firms decide not only their organizational structure but also whether to automate part of their production process. To do so, they must rent one unit of compute per human replaced. The total amount of compute in the economy is exogenous, while its price is endogenously determined in equilibrium. We also assume that compute is abundant relative to human time, making human time the binding constraint in human-AI interactions. This assumption reflects the exponential growth in computational capacity over the past two centuries (Nordhaus, 2007) and the fact that current AI models process information and generate outputs 10 to 100 times faster than humans (Amodei, 2024).

Regarding autonomous AI, we show that its impact critically depends on whether it is “basic” or “advanced.” The introduction of basic autonomous AI—which creates AI agents with the knowledge of human pre-AI workers—induces the most knowledgeable pre-AI solvers to focus on supporting AI-driven production. This reallocation forces human workers to seek assistance from less knowledgeable individuals, pushing the most knowledgeable pre-AI workers to take specialized problem-

³There is both anecdotal and systematic empirical evidence showing the emergence of such “knowledge hierarchies” (see, e.g., Garicano and Hubbard, 2012; Caliendo and Rossi-Hansberg, 2012; Caliendo et al., 2015, 2020). For instance, Alfred Sloan (1924), a former head of General Motors (GM), once wrote: “We do not do much routine work with details. They never get up to us. I work fairly hard, but on exceptions.”

solving roles. As a result, firms become smaller, less productive, and less decentralized. In this scenario, the least knowledgeable individuals are harmed by the technology—they face competition from AI in production tasks and must rely on less knowledgeable individuals for support—whereas the most knowledgeable individuals benefit by using the technology to leverage their knowledge at a low cost.

In contrast, the introduction of an advanced autonomous AI—which creates AI agents with the knowledge of human pre-AI solvers—induces the least knowledgeable humans to seek assistance from AI instead of humans. Consequently, the least knowledgeable pre-AI solvers lose their problem-solving roles and are pushed into routine production work, leading to larger, more productive, and more decentralized firms. In this scenario, both the least and the most knowledgeable individuals benefit from the technology. The least knowledgeable gain access to a cost-effective tool for solving complex problems, while the most knowledgeable continue to thrive by leveraging the advanced AI for routine tasks.

Recent anecdotal evidence from professional services industries illustrates these reorganizations. For example, [Beane and Anthony \(2024\)](#) document that junior analysts in investment banking are increasingly distanced from senior partners as AI takes over tasks that once required junior involvement. Similarly, a 2022 survey of M&A lawyers found that AI, by automating routine tasks such as document review, has allowed junior lawyers to shift their focus to more complex responsibilities, including client engagement and legal analysis ([Litera, 2022](#)).

The impact of AI on labor outcomes also critically depends on its autonomy. As noted earlier, autonomous AI can function both as a co-worker and as a solver or co-pilot, while non-autonomous AI is limited to acting exclusively as a co-pilot. Our main finding is that non-autonomous AI tends to benefit the least knowledgeable individuals the most, in contrast to autonomous AI, which primarily benefits the most knowledgeable individuals. However, overall output is higher with autonomous than non-autonomous AI.

Intuitively, a non-autonomous AI primarily benefits the least knowledgeable individuals because it can only be used as a solver. This removes the competition between the AI and these individuals for production work while allowing them to solve relatively difficult problems without human assistance. Additionally, the lower opportunity cost of non-autonomous AI—stemming from its inability to independently handle production tasks—makes it more affordable. This affordability allows the least knowledgeable individuals to capture a larger share of AI-generated benefits. For the most knowledgeable individuals, the effect is reversed. Non-autonomous AI is less advantageous than autonomous AI because it lacks the capacity to handle routine work and competes with these individuals in advisory roles.

Our findings reveal a trade-off between output and inequality that policymakers encounter when regulating the autonomy of AI. They also provide a framework to reconcile two seemingly contradictory findings in the emerging empirical literature on AI's impact on work. On the one hand,

experimental evidence increasingly suggests that AI disproportionately benefits the least knowledgeable individuals and reduces performance inequality (e.g., [Dell’Acqua et al., 2023](#); [Noy and Zhang, 2023](#); [Brynjolfsson et al., 2023](#); [Peng et al., 2023](#); [Wiles et al., 2023](#); [Caplin et al., 2024](#)). On the other hand, non-experimental evidence indicates that firms generally view AI as a technology that complements high-skilled knowledge workers while substituting for their low-skilled counterparts ([Berger et al., 2024](#); [Alekseeva et al., 2024](#)). Although these findings may seem contradictory, they can be rationalized by our framework as we explain in Section 7: They align with the distinct impacts of non-autonomous AI and basic autonomous AI, respectively.

Related Literature

This paper advances two streams of literature. First, it contributes to the literature on knowledge hierarchies by examining the impact of AI—a technology capable of automating knowledge work at scale. Second, it enriches the literature on automation by integrating AI’s unique characteristics and exploring its organizational effects, providing new insights into how AI reshapes labor outcomes and organizational structures.

The literature on knowledge hierarchies originates with [Garicano \(2000\)](#), who introduces the model and identifies conditions under which knowledge hierarchies are optimal when agents are homogeneous.⁴ [Garicano and Rossi-Hansberg \(2004, 2006\)](#) extend this model to include heterogeneous agents, exploring how the endogenous organization of knowledge work affects labor income inequality. [Fuchs et al. \(2015\)](#) further develop this literature by characterizing the equilibrium contractual arrangements when there is asymmetric information about knowledge.⁵ Building on this foundation, we develop a framework to examine how recent advances in AI are transforming knowledge work.

In the context of knowledge hierarchies, the most closely related paper is [Antràs et al. \(2006\)](#). They study the effects of offshoring by comparing the equilibrium of a closed economy with one in which firms can form international teams. Our paper differs from theirs in two key respects. First, while offshoring gives firms access to a population with varying knowledge levels, AI gives firms access to a technology that can automate knowledge work at scale. As discussed in Section 7, this leads to qualitatively different outcomes. Second, we introduce AI-specific dimensions—such as AI’s knowledge, autonomy, and computational resources—to explore how AI’s impact is shaped by these dimensions. These considerations, absent from [Antràs et al. \(2006\)](#), are crucial for understanding AI’s

⁴The literature on knowledge hierarchies is part of the larger literature on organizational economics. Some recent theoretical contributions to this literature include [Dessein and Santos \(2006\)](#), [Cremer et al. \(2007\)](#), [Alonso et al. \(2008\)](#), and [Dessein et al. \(2016\)](#). See [Gibbons and Roberts \(2013\)](#) for a survey.

⁵Other important contributions to the literature on knowledge hierarchies include [Garicano and Hubbard \(2007, 2009\)](#), [Garicano and Rossi-Hansberg \(2012\)](#), [Bloom et al. \(2014\)](#), [Garicano and Hubbard \(2018\)](#), [Caicedo et al. \(2019\)](#), [Gumpert et al. \(2022\)](#), [Adhvaryu et al. \(2023\)](#), and [Carmona and Laohakunakorn \(2024\)](#).

unique effects.⁶

Our paper also contributes to the literature on automation, which uses task-based models to study the effects of automation on labor outcomes, inequality, and economic growth. The first important recent contribution to this literature is due to Zeira (1998), who shows how automation can lead to a decline in the labor share as the economy develops. Acemoglu and Restrepo (2018) contend that, by depressing wages, automation also encourages the creation of new tasks in which labor has a comparative advantage. Further developing these ideas, Acemoglu and Restrepo (2022) and Acemoglu et al. (2024) study how different types of technological advances—such as automation and labor-augmenting technologies—affect wages, inequality, and productivity.

In a complementary line of research, Autor et al. (2003) and Acemoglu and Autor (2011) argue that routine tasks, which are easier to automate than other types of tasks, are typically handled by those workers in the middle of the skill and wage distribution. Hence, the advent of automation can explain the emergence of employment and wage polarization. Acemoglu and Loebbing (2024) add to this point by showing that automating middle-skill tasks becomes more profitable when the cost of capital decreases.⁷

Our contribution to this literature lies in focusing on a specific type of automation to understand AI’s distinctive impact: the automation of knowledge work. This is an important part of the economy that has been partly shielded from previous automation waves. Unlike other areas, the production of knowledge work is constrained primarily by time and knowledge. Firms thus play a crucial role in organizing the efficient use and communication of their members’ knowledge (Garicano and Rossi-Hansberg, 2015).

From this perspective, AI stands apart from previous automation technologies due to its potential to alleviate time and knowledge constraints, enabling organizations to better leverage the available knowledge. To capture this, we embed AI in a model that explicitly considers the endogenous organization of knowledge work. This approach allows us to capture AI’s unique effects on labor and organizational outcomes, accounting for the knowledge of the existing workforce, communication technologies, AI capabilities, and the availability of computational resources.⁸

For example, one key advantage of our framework is its ability to examine how AI autonomy influences labor outcomes. Our analysis shows that non-autonomous AI tends to benefit the least knowledgeable individuals the most, while autonomous AI primarily advantages the most knowl-

⁶Another related contribution is Pancs’s (2018, ch. 5) textbook example, which shows how superintelligent robots—matching the abilities of the most capable humans—eliminate labor income inequality. In contrast, we demonstrate that AI’s effects on labor income distribution are nuanced, depending critically on AI’s proficiency and autonomy.

⁷Other important contributions include Aghion et al. (2017), Moll et al. (2022), Azar et al. (2023), Korinek and Suh (2024), Acemoglu (2024), and Jones (Forthcoming).

⁸Economic theory appears particularly relevant in this context because the economy is in a state of transition. As Veldkamp and Chung (2024, p. 464) note: “Empirical work takes past trends or covariances and extrapolates to form predictions. But in transitions, the future often looks systematically different from the past.”

edgeable individuals. To the best of our knowledge, these results have no parallel in the existing literature on automation and AI.

Moreover, our model also addresses a key gap in existing research: the role of firms as central actors in adopting and modulating the effects of new automation technologies. While task-based models of automation provide valuable insights into labor outcomes, they often overlook how firms influence the adoption and impact of these technologies. By explicitly incorporating the endogenous organization of knowledge work, our approach forces us to introduce firms into the analysis. Hence, our model provides insights into how AI leads to the reorganization of production and its effects on productivity, size, and organizational structures. Furthermore, it reveals how changes in firm composition create new pathways through which AI impacts labor outcomes.

For instance, we show that in the case of a basic autonomous AI, the most knowledgeable human solvers switch toward using AI for routine production work. As a result, a lower-quality pool of solvers is left to assist those who continue as workers in the post-AI world, reducing their productivity. As discussed in Section 7, this outcome is consistent with anecdotal evidence from sectors such as investment banking and legal services.

Roadmap

The rest of the paper is organized as follows. Section 2 identifies three key properties of modern AI. These properties shape the assumptions of our baseline model, introduced in Section 3, which focuses on autonomous AI. Section 4 characterizes the equilibrium with autonomous AI, while Section 5 examines the impact of this technology relative to the pre-AI equilibrium. Section 6 explores the implications of non-autonomous AI. Finally, Section 7 concludes by discussing additional implications of our analysis and linking our findings to emerging empirical evidence about AI's effects on the labor market.

2 Three Key Properties of Modern AI

In this section, we document three key properties of modern AI that are essential for understanding both the enthusiasm and concerns surrounding its impact on labor and organizational outcomes: (i) scalability, (ii) the advent of foundation models, and (iii) the potential of foundation models to automate knowledge work. These three properties guide the assumptions of our baseline model in Section 3.

1. The Crucial Difference Between Human and Artificial Intelligence is Scalability.—The primary constraint on the use of human knowledge is human time. As [Garicano and Rossi-Hansberg \(2015\)](#) explain, while knowledge itself can theoretically be used repeatedly as an input in production, its application in practice is restricted by the time and availability of the individual who possesses it.

Artificial intelligence, by contrast, operates under an entirely different paradigm. Once knowledge is encoded in an AI model, it can be deployed across all available computational resources without such limitations. This reflects the nonrival nature of digital information, which has a negligible marginal cost of reproduction (Brynjolfsson and McAfee, 2016; Goldfarb and Tucker, 2019). As Waisberg and Hudek (2021, p. 177) put it:

You can ask a question at a seminar or webinar, but you can only get the data as provided by the host or speaker, and a given speaker can only be in one seminar at a time... [In contrast, AI] scales; no longer are you limited by the time of a single domain expert. If their knowledge and behavior are encoded into an AI, you can have as many copies of that AI working at the same time as you want.

2. *The Advent of Foundation Models.*— Foundation models represent a transformative shift in AI. These are models characterized by their ability to handle a wide variety of tasks through training on large-scale datasets using self-supervised learning techniques (Bommasani et al., 2022; Council of Economic Advisers, 2024). Models like GPT-4, Gemini, and Claude Sonnet exemplify this approach, serving as flexible, general-purpose platforms that can be adapted to diverse downstream tasks with minimal additional training. This adaptability enables firms to leverage shared AI models, removing the necessity for the costly development of proprietary, domain-specific systems.⁹ As a result, a small number of powerful AI models can function across multiple domains, driving significant economic and industry-wide impact.¹⁰

Foundation models have demonstrated their versatility across various industries, from customer service automation to legal research and financial analysis. For instance, Klarna, a fintech company, automated over 700 customer service roles using a version of OpenAI’s models. In the legal field, Harvey leverages OpenAI’s technology to automate research and document analysis. Similarly, Claude Sonnet powers applications at Lonely Planet for itinerary creation and at Bridgewater Associates for tasks like research, coding, and data visualization.¹¹

Moreover, evidence suggests that general-purpose models often outperform specialized AI sys-

⁹For instance, Andreessen Horowitz, a venture capital firm, surveyed 100 Fortune 500 companies across industries such as technology, telecom, banking, healthcare, and energy. The survey found that 94% of these firms were customizing existing foundation models, while only 6% were creating their own specialized models (Wang and Xu, 2024).

¹⁰Using the language of computer scientists, foundation models exhibit *emergence*—the ability to develop unexpected capabilities that enable them to perform tasks beyond their explicit training. Combined with scalability, emergence then leads to *homogenization*, defined as the growing reliance on a few powerful models across multiple domains (Bommasani et al., 2022).

¹¹For more details on these applications—and to explore additional examples—visit OpenAI’s customer stories (<https://openai.com/news/stories/>), Anthropic’s customer page (<https://www.anthropic.com/customers>), Amazon AI Customer Stories (<https://aws.amazon.com/ai/generative-ai/customers>), and Google Cloud’s Gemini AI blog (<https://cloud.google.com/blog/products/ai-machine-learning/bringing-gemini-to-organizations-everywhere>) (accessed October 7, 2024).

tems. For instance, ChatGPT—a fine-tuned version of GPT-4—has demonstrated superior performance to BloombergGPT—a Large Language Model (LLM) trained specifically on financial data—in various finance-related tasks (Li et al., 2023). Similarly, Nori et al. (2023) demonstrate that GPT-4 surpasses specialist models across medical benchmarks and argue that this advantage extends to other domains, including accounting, law, nursing, and clinical psychology. This shows that even in fields that traditionally relied on highly specialized models, general-purpose foundation models offer better performance, increased flexibility, and lower costs.

3. *AI Foundation models are Driving a New Automation Wave.*— The adaptability, efficiency, and scalability of foundation models are fueling a new wave of automation: the automation of knowledge work. This distinguishes AI foundation models from previous digitalization waves, such as the advent of computers, enterprise software, and the rise of the internet and mobile technologies. As Sarah Tavel, a general partner at the venture capital firm Benchmark, explains, this shift is fundamentally altering how businesses deliver value (Stebbins, 2024b):

What AI enables is actually a very different unit of work that you sell, which is doing the work. And so you are almost a software company that looks like a services business that is able to sell the full work product—the outcome—as opposed to selling software that an employee has to learn to use and then gets a productivity boost from.

This shift raises the crucial question of whether current AI models are sophisticated enough for such responsibilities. Sarah Tavel remains optimistic (Stebbins, 2024b):

There are certainly use cases that can be automated right now, and we see, we see a lot of them... beyond that, there is no question that as the foundation models improve, we are going to see the ability to take on more and more complex tasks.

This optimism is grounded in the fact that much of today’s economy is driven by knowledge work, an area where AI excels compared to industries reliant on manual labor. For example, Bommasani et al. (2022) report that more than 13% of U.S. occupations have a *primary task* that is related to *writing*, based on data from the U.S. Department of Labor’s O*NET database. Collectively, these roles account for a total annual wage bill exceeding USD 675 billion.

Furthermore, projections suggest that AI’s impact will extend far beyond language-related tasks. Goldman Sachs (2023) estimates that two-thirds of U.S. occupations are at least partially exposed to automation by AI, with about a quarter of all current work tasks potentially automatable when weighted by employment share. Similarly, Eloundou et al. (2024) project that 46% of jobs could have more than half of their tasks impacted by AI foundation models in the near future.

3 The Baseline Model: Autonomous AI

The three key properties of modern AI outlined in the previous section motivate us to develop a model with the following three features. First, AI systems must be scalable, meaning that the use of knowledge embedded in the system is constrained only by computational resources. Second, the rise of flexible, general-purpose foundation models—adaptable to diverse tasks with minimal additional training—enables firms to rely on a shared AI model, eliminating the need to develop costly in-house models. Third, AI harnesses computational power to automate knowledge work.

The structure of this section is as follows: Section 3.1 introduces the model, Section 3.2 discusses several of its assumptions, and Section 3.3 describes the pre-AI equilibrium.

3.1 The Model

The Pre-AI Economy.— There is a unit mass of humans, each endowed with one unit of time and exogenous knowledge $z \in [0, 1]$. The distribution of knowledge in the population is given by a cumulative distribution function G with density g . The only assumption we impose on this distribution is that g is continuous and strictly positive for all z . The knowledge of each individual is perfectly observable.

There is a large measure of identical competitive firms. Production occurs inside firms, which are the residual claimants of all output. Labor and knowledge are the sole inputs in production. Firms have no fixed costs and, for simplicity, two layers at most.

Single-layer firms hire a single human to produce. This “independent producer” devotes her full unit of time to pursuing a single production opportunity. Each production opportunity is linked to a problem whose difficulty x is ex-ante unknown and distributed uniformly on $[0, 1]$, independently across problems. If the knowledge of the human engaging in production exceeds the problem’s difficulty, she solves the problem and produces one unit of output. Otherwise, no output is produced. The assumption that $x \sim U[0, 1]$ is simply a normalization, as the distribution of knowledge in the population is arbitrary. Under this normalization, an individual’s knowledge z is interpreted as the fraction of problems she can solve on her own.

Two-layer firms hire one “solver” and multiple “workers,” where all workers have the same knowledge. This restriction is without loss because—as we show below—the equilibrium matching arrangement between workers and solvers exhibits strict positive assortative matching.

As in single-layer firms, each worker in a two-layer firm devotes her time to a single production opportunity. The difference is that if the worker cannot solve the problem on her own, she can ask the solver for help. If the solver’s knowledge exceeds the problem’s difficulty, she communicates the solution to the corresponding worker, who then produces a unit of output. Otherwise, no production takes place. However, communication is costly in that whenever a worker asks the solver for

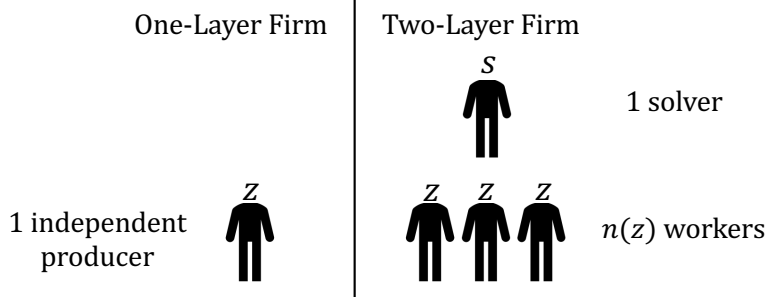


Figure 1: The Two Possible Firm Configurations in the Pre-AI World

help, that exchange consumes $h \in (0, 1)$ units of the solver's time. Hence, a two-layer organization optimally hires exactly $n(z)$ workers of knowledge z to fully exploit its solver's time, where:

$$n(z) \equiv \frac{1}{h \times (1 - z)}$$

Figure 1 depicts the two possible firm configurations of the pre-AI world, where the letter attached to each human corresponds to her knowledge.

Autonomous AI.—We model AI as a technology that transforms compute into “AI agents,” all with the same exogenously fixed level of knowledge, $z_{AI} \in [0, 1]$.¹² As a benchmark, we first consider the case where these agents can do exactly the same as humans. This means that the AI agents can function both as a “co-worker” (capable of pursuing production opportunities) and as a “co-pilot” (providing advice). In Section 6, we explore the case where the AI agents are limited to providing problem-solving guidance, rendering them non-autonomous. We begin by analyzing the autonomous case because it generates the most anxiety and debate, fueled by technology firms and venture capitalists who both believe in—and actively promote—the advancement of AI in this direction (see Section 2).

We assume that all firms have access to AI. Thus, in contrast to the pre-AI economy, firms decide not only their organizational structure but also whether to use this technology. Firms that adopt AI differ from those that do not only by automating a portion of their production process. To do so, they must rent one unit of compute per human replaced. This implies that compute is measured in units equivalent to human time. The amount of compute in the economy, denoted by μ , is exogenous. We also assume there are more production opportunities than the human time and compute necessary to pursue them all, i.e., production opportunities are abundant relative to the available resources.

Figure 2 illustrates the five possible post-AI firm configurations. In addition to single- and two-layer firms that hire only humans (the only possible pre-AI firms), there are three additional possible configurations: Single-layer automated firms (which use AI as an independent producer), bottom-automated firms (which use AI exclusively as a worker), and top-automated firms (which use AI

¹²We assume that $z_{AI} < 1$ because the equilibrium has a discontinuity at $z_{AI} = 1$. In the Online Appendix, we consider the case where $z_{AI} = 1$.

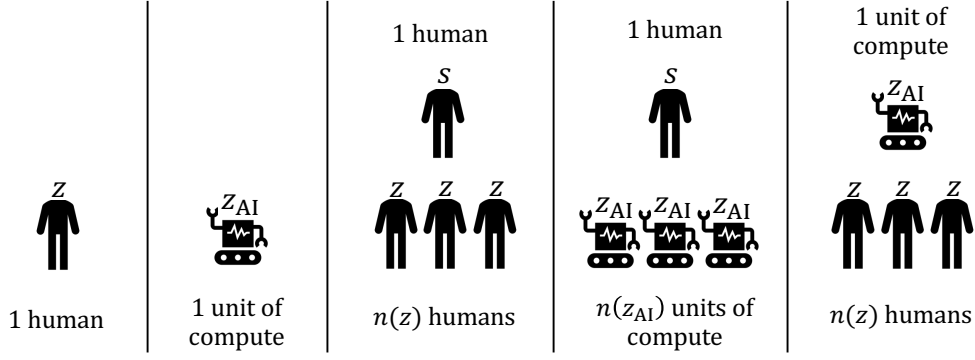


Figure 2: The Five Possible Firm Configurations in the Post-AI World

exclusively as a solver). Note that a firm will never use AI in both layers of the organization because an AI solver knows the solution to the same set of problems as an AI worker.

An important implication of considering the endogenous organization of knowledge work is that it blurs the boundary between AI as an “automation technology” and as a “co-pilot” or “problem-solving assistant.” This distinction is often framed around whether AI replaces or augments human effort (see e.g., [Brynjolfsson, 2022](#)). However, as our model makes clear, an AI co-pilot can be viewed as a technology that automates the work of human solvers.

Our modeling approach is motivated by the three key properties of modern AI described in Section 2. First, AI is scalable, meaning that once knowledge z_{AI} is encoded into the technology, such knowledge can be utilized simultaneously across all available computational resources. Second, all firms have equal access to AI through a shared model. Third, AI functions as an automation technology, enabling firms to replace humans in the different roles within the knowledge economy.

Wages, Prices, and Profits.—Let $w(z)$ be the wage of a human with knowledge z and denote by r the rental rate of one unit of compute. All agents in this economy are risk neutral and maximize their income. We normalize the value of each unit of output to one.

The problem of a firm is to decide (i) whether to use humans or AI in production (and the knowledge of the humans hired, when appropriate), (ii) whether to operate as a one-layer or two-layer organization, and (iii) in the case of two-layer organizations, whether to use a human or AI as a solver (and the knowledge of the human hired, when appropriate).

The profits of a single-layer organization are as follows:

$$\Pi_1 = \begin{cases} z - w(z) & \text{if the firm hires a human with knowledge } z \\ z_{AI} - r & \text{if the firm uses AI} \end{cases}$$

In other words, the profit of a single-layer nonautomated firm is the expected output z of its independent producer net of her wage $w(z)$. The profit of a single-layer automated firm is the expected output of AI as an independent producer z_{AI} minus the cost of renting one unit of compute r .

The profit of a two-layer organization, in turn, depends on whether it uses AI as a solver (i.e., automates the top layer, a “*tA*” firm), uses AI as a worker (i.e., automates the bottom layer, a “*bA*” firm), or it does not use AI (an “*nA*” firm):

$$\begin{aligned}\Pi_2^t(z) &= n(z)[z_{AI} - w(z)] - r && (\text{where } z \leq z_{AI}) \\ \Pi_2^b(s) &= n(z_{AI})[s - r] - w(s) && (\text{where } z_{AI} \leq s) \\ \Pi_2^n(s, z) &= n(z)[s - w(z)] - w(s) && (\text{where } z \leq s)\end{aligned}$$

where z and s denote the knowledge of a human worker and a human solver, respectively, and we use the fact that no two-layer firm hires a solver who is less knowledgeable than its workers. In all three cases, the profit of a firm is its expected output minus the cost of the resources it uses. For instance, in the case of a *tA* firm that hires workers with knowledge z , its total expected output is $n(z)z_{AI}$, while the cost of resources is $n(z)w(z) + r$.

Competitive Equilibrium.— A competitive equilibrium comprises the following objects: (1) an allocation of compute across its various potential uses; (2) a partition of the human population into distinct occupations; (3) a description of the matching arrangements between different types of humans driven by the hiring decisions of *nA* firms; and (4) a wage schedule, $w(z)$, and a rental rate for compute, r .

With this in mind, let μ_i , μ_w , and μ_s be the amount of compute rented for independent production, production in two-layer firms, and the supervision of humans, respectively. We denote by I the set of humans hired as independent producers, and by W_p and W_a the set of human workers assisted by human solvers and AI, respectively.¹³ Similarly, we denote by S_p the set of humans who assist other humans, and by S_a the set of humans who assist the production work of AI. Since a *tA* firm that hires human workers with knowledge $z \in W_a$ rents one unit of compute, we have that $\mu_s = \int_{W_a} h(1 - z)dG(z)$. Similarly, since a *bA* firm that hires a human solver with knowledge $s \in S_a$ rents $n(z_{AI})$ units of compute, then $\mu_w = n(z_{AI}) \int_{S_a} dG(z)$.

Finally, let $m : W_p \rightarrow S_p$ represent the function describing the pointwise matching arrangement generated by the hiring decisions of *nA* firms. Specifically, $m(z)$ denotes the knowledge of the solver assisting workers with knowledge z .¹⁴ As mentioned earlier, the equilibrium exhibits strict positive assortative matching, so m is monotone increasing. Additionally, it must satisfy the following resource constraint:

$$(1) \quad \int_Y h(1 - u)dG(u) = \int_{m(Y)} dG(u), \text{ for all } Y \subseteq W_p$$

¹³We use the subscript “*p*” (for people) instead of “*h*” (for humans), to avoid any confusion with the helping cost h .

¹⁴As discussed in Fuchs et al. (2015), $m(z)$ is an approximation of a firm that hires a small interval of human workers $(z - \varepsilon_1, z + \varepsilon_1)$ and a small interval of human solvers $(m(z) - \varepsilon_2, m(z) + \varepsilon_2)$ with the requirement that the mass of solvers in the firm $\int_{m(z) - \varepsilon_2}^{m(z) + \varepsilon_2} dG(u)$ is equal to h times the mass of problems left unsolved by the workers in the firm $\int_{z - \varepsilon_1}^{z + \varepsilon_1} (1 - u)dG(u)$. The latter requirement captures that *nA* firms optimally hire $n(z)$ workers of knowledge z .

where, with slight abuse of notation, $m(Y)$ represents the set of solvers assigned to the set of workers $Y \subseteq W_p$. This condition ensures that the total time required to consult on problems left unsolved by the workers in Y equals the total time available from the solvers matched to these workers, $m(Y)$. In essence, it reflects that nA firms optimally hire $n(z)$ workers with knowledge level z to fully utilize the time of their solvers.

Definition (Competitive Equilibrium). An equilibrium consists of nonnegative amounts (μ_i, μ_w, μ_s) , sets (W_p, W_a, I, S_p, S_a) , a strictly increasing matching function $m : W_p \rightarrow S_p$ satisfying (1), a wage schedule $w : [0, 1] \rightarrow \mathbb{R}_{\geq 0}$ and a rental rate of compute $r \in \mathbb{R}_{\geq 0}$, such that:

1. Firms optimally choose their structure (while earning zero profits).
2. Markets clear: (i) $\mu_i + \mu_w + \mu_s = \mu$, and (ii) the union of the sets (W_p, W_a, I, S_p, S_a) is $[0, 1]$ and the intersection of any two of these sets has measure zero.

Note that the human workers in W_p are endogenously matched with the human solvers in S_p according to the pointwise matching function m . In contrast, all the humans in W_a are matched with an AI solver (which has knowledge z_{AI}), while all the humans in S_a are matched with AI workers (whose knowledge is z_{AI}).

Compute is “Abundant” Relative to Human Time.— We focus on the case in which compute is sufficiently abundant so that the binding constraint in human-AI interactions is human time, not compute. In other words, there is more compute available than the one demanded by tA and bA firms, so the leftover compute must be rented by single-layer automated firms.

We consider this the most relevant case due to the exponential growth in computational capacity over the past two centuries (Nordhaus, 2007) and the fact that current AI models process information and generate outputs 10 to 100 times faster than humans (Amodei, 2024). However, to explore the effects of this assumption, we also examine the case of limited compute in the Online Appendix.

The following is a sufficient condition for compute to be abundant relative to time:¹⁵

$$(2) \quad \int_0^{z_{AI}} n(z)^{-1} dG(z) + n(z_{AI})(1 - G(z_{AI})) < \mu$$

To understand this condition, note that a firm will never hire a solver who is less knowledgeable than its workers. Consequently, the compute allocation that maximizes human-AI interactions subject to the constraint mentioned above involves (i) matching every human who is less knowledgeable than AI with an AI solver and (ii) matching every human who is more knowledgeable than AI with $n(z_{AI})$ units of compute using AI for production work. Since (i) uses $\mu_w = \int_0^{z_{AI}} h(1 - z) dG(z)$ units of compute, while (ii) uses $\mu_s = n(z_{AI})(1 - G(z_{AI}))$ units of compute, if condition (2) holds, then there are not enough humans to interact with all the available compute.

¹⁵Note that for any distribution G and helping cost $h \in (0, 1)$, there exists a finite μ that satisfies this condition for all $z_{AI} \in [0, 1]$. The reason for this is that the left-hand side of (2) is continuous in $z_{AI} \in (0, 1)$ and is bounded as $z_{AI} \rightarrow 0$ and $z_{AI} \rightarrow 1$ (it converges to $1/h$ and $g(1)/h + \int_0^1 n(z)^{-1} dG(z)$, respectively).

Some Notation.— For future reference, we define $W \equiv W_a \cup W_p$ and $S \equiv S_a \cup S_p$ as the overall set of human workers and solvers of the economy, respectively. We also denote by $e : S_p \rightarrow W_p$ the inverse of the matching function m . That is, e is the “employee matching function” denoting the knowledge of the human worker matched with a human solver with knowledge $s \in S_p$. This function always exists given that, as argued above, the equilibrium matching function is strictly increasing.

Finally, for any arbitrary set B , we denote by $\text{int}B$ and by $\text{cl}B$ the interior and closure of B , respectively. We also use $B \preceq B'$ to indicate that the set $B \subseteq [0, 1]$ “lies below” the set $B' \subseteq [0, 1]$. Formally, $B \preceq B'$ if $\sup B \leq \inf B'$. For example, $W_a \preceq W_p$ means that the best worker assisted by AI is weakly less knowledgeable than the worst worker assisted by a human.

3.2 Discussion of the Model

Before proceeding with the analysis, we discuss how to map AI in the model to AI in the real world. We also justify why we focus on a single AI algorithm and two-layer firms, describe how to translate compute in the model into real-world measures, and acknowledge the limitations of our modeling approach.

Interpretation.— As previously noted, autonomous AI can operate both as a co-worker, capable of pursuing production opportunities, and as a co-pilot, providing assistance.¹⁶ A real-world example of an AI co-worker is Klarna, which automated over 700 customer service agents using OpenAI’s models (see Section 2). Another is Agentforce by Salesforce, a tool that not only identifies business leads using company data but also contacts prospects and schedules meetings on behalf of salespeople (Rosenbush, 2024). Similarly, All Day TA, a startup founded by Professors Kevin Bryan and Joshua Gans from the University of Toronto, leverages foundation models to fully automate the tasks of teaching assistants for university courses.

In contrast, Brynjolfsson et al. (2023) provide an example of an AI co-pilot in the customer support industry: their study highlights the productivity effects of equipping human agents with a real-time, AI-based conversational assistant to manage customer inquiries. Another example is GitHub Copilot, which aids developers by suggesting code snippets, autocompleting functions, and recommending improvements to code quality. Similarly, BrainTrust AIR assists HR professionals by reviewing candidates, grading applications, and generating detailed scorecards to streamline recruitment decisions.

Regarding AI autonomy, for simplicity we focus on the two extreme cases. First, we consider a scenario where AI agents are fully autonomous, capable of performing tasks equivalent to those of humans with knowledge z_{AI} . Then, in Section 6, we examine the opposite case, where AI agents

¹⁶The model is silent about where production opportunities come from. The role of organizations is to decide who to assign to routine work, and who to assign to specialized problem solving. Hence, workers and independent producers—human or AI—should not be thought of as being responsible for generating production opportunities; they are only responsible for pursuing them.

are entirely non-autonomous and cannot complete any tasks independently. Most AI applications, however, operate between these extremes, with autonomy levels varying by sector. For example, tools like Harvey automate many routine legal tasks, such as document review and contract analysis. Yet, they remain unable to engage directly with clients or argue cases in court.

Because our framework models the broader knowledge economy, we therefore interpret our results as approximations of AI’s impact based on its average level of autonomy across the different relevant sectors. Accordingly, the practical implications of our analysis will likely fall between the predictions of these two benchmark cases. To determine which extreme better reflects current and future realities, the critical question is: How many tasks traditionally performed by humans can AI systems like Harvey handle autonomously, and how will this capability evolve over time?

Why Only One AI?— AI’s scalability and the fact that all firms have access to AI not only enable the widespread use of the technology but also shape the competitive landscape. In the Online Appendix, we show that if multiple AI technologies are available and all of them require similar computational resources to create AI agents, only the most advanced AI will be used in equilibrium. This follows because compute naturally flows toward the highest-performing model, creating a Bertrand-like competition for compute among AI models.

This competition for compute underpins our assumption of focusing on a single AI technology. It is also consistent with the challenges of competing in the foundation model space, as articulated by Aravind Srinivas, founder of PerplexityAI (Stebbing, 2024a):

[It] is a losing game... every time you end up finishing a large training run, you burn a lot of money, you have a great model, and then you watch it be destroyed by the next update... So where are you recovering all that money back? (...) Nobody wants to use the API [i.e., the model] if somebody else is offering a better model at a cheaper price.

Computer scientists and venture capitalists refer to this phenomenon as the *commoditization of foundation models* (Navani, 2023; Stebbings, 2024a,b).

Nevertheless, as we also show in the Online Appendix, if less advanced AI technologies require substantially less compute than their more advanced counterparts, they could coexist alongside them. This extension is consistent with the development of “smaller” foundation models like GPT-4o mini and Claude 3 Haiku. Furthermore, it highlights that human heterogeneity is fundamentally different from AI heterogeneity: While humans with varying knowledge can be employed in equilibrium despite requiring the same “computational resources” (i.e., time) to handle problems, AI heterogeneity arises only when the AI models differ substantially in both knowledge and resource requirements.¹⁷

¹⁷Furthermore, it also implies that while the introduction of a single autonomous AI technology corresponds to the introduction of an exogenous mass μ of AI agents with knowledge z_{AI} , the introduction of two or more distinct autonomous

Why Only Two Layers?— To keep the analysis transparent and provide simple, testable predictions, we focus on the simplest hierarchical organizations—those with at most two layers. As discussed in Section 7, this simple environment aligns well with emerging experimental and observational evidence on AI’s impact on knowledge work. It is also sufficiently rich to demonstrate the reorganizations AI can induce and how these depend on AI’s autonomy and problem-solving capabilities.

Allowing for more complex organizational structures could undoubtedly reveal additional subtleties in how AI reorganizes knowledge work. However, our goal is to strike a balance between simplicity and analytical richness. Introducing a more elaborate model risks obscuring the key mechanisms and insights we aim to emphasize. Furthermore, as we show in the Online Appendix, many insights we uncover—such as the conditions under which the least and most knowledgeable individuals benefit from AI—naturally generalize to settings with an arbitrary number of layers. Nevertheless, a more comprehensive examination of AI’s effects on more complex organizational structures is a promising area for future research.

No Unemployment.— In defining our competitive equilibrium, we excluded the possibility of unemployment. This is without loss of generality in our setting, as we assume that total available resources—human time and compute—are scarce relative to production opportunities (even though compute is abundant relative to time). This scarcity ensures that, in equilibrium, it is always worthwhile for every human to be employed in some capacity, even if AI is more knowledgeable than a fraction of the human population.

In the Online Appendix, we explore an extension where compute is so abundant that it outstrips production opportunities. In that case, we show that AI leads to technological unemployment: All humans who are less knowledgeable than AI become unemployed. Organizations, however, still display a hierarchical structure in that the most knowledgeable humans specialize in tackling the problems that AI cannot solve.

On Compute.— Our main goal is to analyze how AI affects human labor outcomes. For this reason, we take the AI technology as given. Thus, in our framework, “compute” refers to the resources required during the inference phase—the stage in which an AI system generates outputs based on new input data after it has already been trained. The availability of compute is limited by semiconductor hardware and energy.

To compare AI output with human labor, we express compute in terms of human time. In practice, compute is measured in *floating-point operations per second* (FLOPs), which reflect the number of arithmetic calculations the system performs per second. This is the fundamental unit of measurement in computing.

However, in the context of LLMs, *tokens* serve as a more useful unit of measurement for our notion

AI technologies, each with different resource requirements, does not correspond to the introduction of multiple exogenous masses of AI agents with varying knowledge levels.

of compute. This is because LLM providers like OpenAI and Anthropic often use a pay-as-you-go pricing model based on the number of input and output tokens processed.¹⁸ A token is a unit of text that can range from an entire word to part of a word or even a single character, depending on the text’s complexity. On average, one token corresponds to about four characters in English.¹⁹

Using tokens as a measure of compute allows us to anchor our comparison with real-world usage. For example, evaluations by Model Evaluation and Threat Research (METR)—a non-profit organization that assesses risks posed by advanced AI systems—estimate that for relatively complex tasks, such as fixing bugs in an object-relational mapping library, one hour of human labor by a STEM graduate with 3+ years of technical experience corresponds to the processing of approximately 191,000 tokens (METR, 2024). However, LLMs completed only around 30% of the tasks that humans in the study accomplished.

Limitations.— Our primary goal is to analyze how the ability of recent AI models to automate knowledge work affects human labor outcomes. While we believe our framework—covering both autonomous and non-autonomous AI models—provides a reasonable first approximation, it has certain limitations.

First, beyond automating knowledge work, AI performs other tasks, such as voice recognition and automatic translation, which could reduce communication costs for firms. Additionally, AI-driven tools can facilitate the accumulation of human capital through personalized, adaptive tutoring—an aspect excluded from our analysis. While these applications are important, technological advances that reduce communication costs or make knowledge acquisition cheaper are not new. Similar developments occurred during the ICT revolution of the 1980s and 1990s, and their implications are well-understood (see, e.g., Garicano and Rossi-Hansberg, 2006; Antràs et al., 2008; Bloom et al., 2014; Garicano and Rossi-Hansberg, 2015). In Section 7, we contrast some implications of AI as an automation technology with those of technologies that reduce communication or knowledge acquisition costs.

Second, in light of the rise of foundation models, we assume negligible customization costs and equal access to AI across firms. This serves as a useful first approximation, since customizing foundation models typically costs several orders of magnitude less than developing AI systems from scratch (Appenzeller et al., 2023; Xia et al., 2024), with costs likely to decline further.²⁰ However, cer-

¹⁸For instance, as of October 2024, OpenAI charged USD 2.50 per 1 million input tokens and USD 10 per 1 million output tokens for their flagship foundation model, GPT-4o. Source: <https://openai.com/api/pricing/> (accessed October 17, 2024).

¹⁹Regarding the conversion of inference tokens into FLOPs, a commonly used approximation for transformer-based LLMs is $\text{FLOPs} \approx 2 \times N \times \text{tokens}$, where N represents the model’s number of parameters (Appenzeller et al., 2023; Kaplan et al., 2020). While the parameter count in most well-known LLMs is proprietary, GPT-3 is publicly known to have 175 billion parameters (Brown et al., 2020).

²⁰Developing more effective and cost-efficient techniques for customizing foundation models is an especially active area in the computer science literature. Techniques such as RAG, LoRA, QLoRA, and PETE, or variations thereof, are being proposed at an accelerated pace. For an overview of the most commonly used methods, see Shah and Patel (2023).

tain applications still require domain-specific AI models, as demonstrated by programs like Google’s AlphaFold. Developing a multi-sector model of the knowledge economy—where some sectors rely on shared foundation models while others require specialized models—presents a promising avenue for future research. Similarly, AI systems often exhibit varying capabilities across industries. For example, [Goldman Sachs \(2023\)](#) estimates that Generative AI’s automation potential in healthcare support is roughly half that in office and administrative support. A multi-sector framework that accounts for these differences would offer valuable insights into how AI might drive inter-industry reallocation.

Third, we treat the AI technology as a given and model the stock of compute as an endowment. This allows us to bypass the complexities of modeling the AI training technology and the industrial organization of the AI sector—such as foundation model developers, application developers, and compute providers. It also keeps the model static. While our model touches on some of these issues—see the earlier discussion on the commoditization of foundation models—a detailed analysis of the AI sector and the drivers of AI and compute development are promising areas for future research.

3.3 Benchmark: The Pre-AI Equilibrium

We begin by presenting a partial characterization of the equilibrium without AI (for the full characterization, see Appendix A). This is also the equilibrium when the amount of compute μ is zero and was originally described by [Fuchs et al. \(2015\)](#).²¹ Note that in this case $W_a = S_a = \emptyset$, so $W = W_p$ and $S = S_p$.

Proposition 1. *In the absence of AI, there is a unique equilibrium. The equilibrium is efficient (i.e., it maximizes total output) and has the following features:*

- *Occupational stratification:* $W \preceq I \preceq S$.
- *Positive assortative matching:* The function $m : W \rightarrow S$ is strictly increasing.
- $W \neq \emptyset$ and $S \neq \emptyset$. However, $I \neq \emptyset$ if and only if $h > h_0 \in (0, 1)$.

Moreover, the wage function w is continuous, strictly increasing, and convex (strictly so when $z \in W \cup S$), and is given by:

- $w(z) = m(z) - w(m(z))/n(z)$ for all $z \in W$.
- $w(z) = z$ for all $z \in I$.
- $w(z) = C + \int_{\inf S}^z n(e(u))du > z$ for all $z \in S$, where $C > \inf S$ when $h < h_0$ (and $C = \inf S$ otherwise).

In particular, $w(z) > z$ for all $z \notin \text{cl}I$ (so $w(z) > z$ for all $z \in [0, 1]$ when $h < h_0$).

Proof. See Appendix A. □

²¹Note that an economy with $\mu = 0$ is different than that with $\mu > 0$ but $z_{AI} = 0$. This is because, even if AI cannot solve any problems, it can still draw them, thus enlarging the production possibility frontier of the economy.

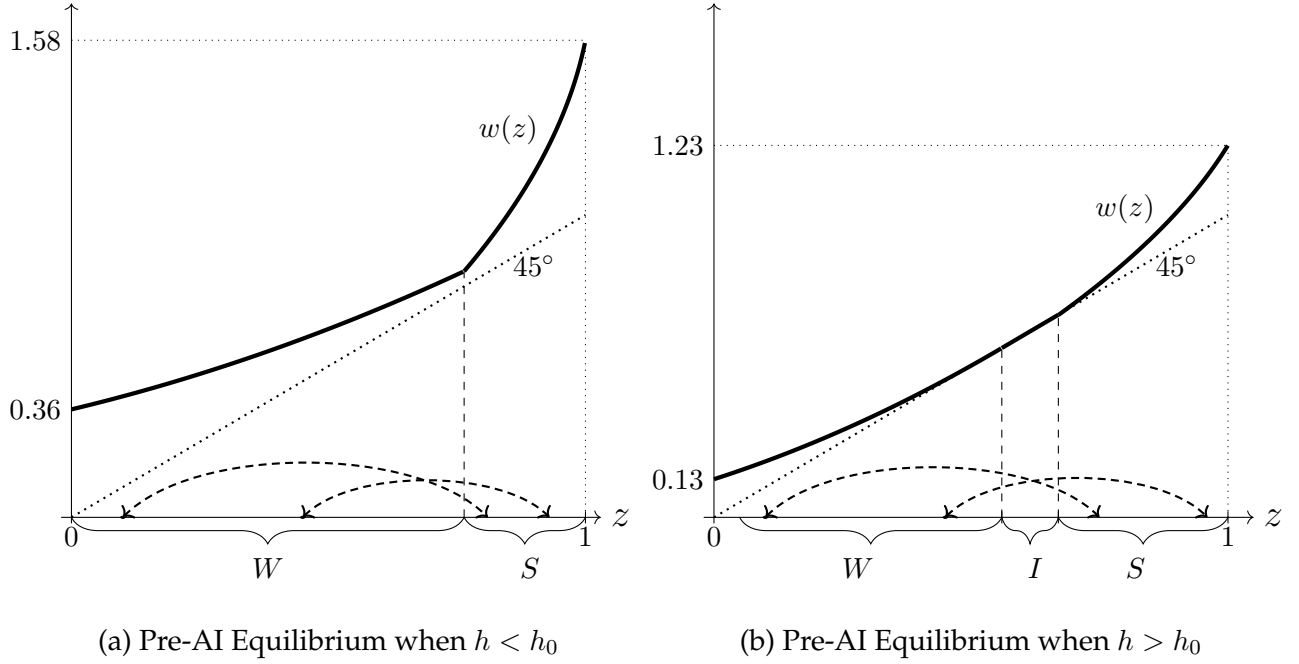


Figure 3: Illustration of the Pre-AI Equilibrium.

Notes. Distribution of knowledge: $G(z) = z$. Parameter values: For panel (a), $h = 1/2 (< h_0 = 3/4)$, while for panel (b), $h = 0.8125 (> h_0 = 3/4)$. The thick line depicts the equilibrium wage function. The dashed arrows illustrate the matching between workers and solvers. See the Online Appendix for a detailed characterization of the pre-AI equilibrium when human knowledge is uniformly distributed.

The equilibrium without AI—which we illustrate in Figure 3—has several salient features. First, it is efficient, as the First Welfare Theorem holds in this economy (with perfect competition, complete information, and only pecuniary externalities). Moreover, the wage of an individual with knowledge z corresponds to her marginal product, defined as the increase in output from introducing one additional human with this knowledge into the economy.

Second, it exhibits occupational stratification: Workers are less knowledgeable than independent producers, who are, in turn, less knowledgeable than solvers. Intuitively, more knowledgeable agents have a comparative advantage in specialized problem solving, as this allows them to leverage their knowledge by applying it to more than one problem. Hence, $(W \cup I) \preceq S$. Similarly, less knowledgeable agents have a comparative advantage in assisted rather than independent production work, as they are less likely to succeed on their own. Consequently, $W \preceq I$.

Third, there is strict positive assortative matching: Conditional on W and S , more knowledgeable workers in W match with more knowledgeable solvers in S . The reason is that worker and solver knowledge are complements: For a given team of workers, a more knowledgeable solver increases expected output, while for any given solver, more knowledgeable workers increase team size.

Fourth, the set of workers and the set of solvers are always nonempty, but the set of independent producers is empty when the communication cost h is below a threshold $h_0 \in (0, 1)$. Intuitively,

lower values of h make two-layer organizations more attractive than single-layer ones, leading all humans to be employed by two-layer organizations.

Fifth, workers and solvers earn strictly more than their expected output as independent producers (except in the case of the most knowledgeable worker and the least knowledgeable solver when $h \geq h_0$). The reason for this is that the marginal value of knowledge is strictly lower than 1 for workers (as their knowledge is used to free up solver time),²² exactly equal to 1 for independent producers (as their expected output equals their knowledge), and strictly greater than 1 for solvers (as they can leverage their knowledge by applying it to more than one problem). Hence, if solvers were to earn their expected output as independent producers, then all two-layer firms would try to hire the most knowledgeable agents as solvers. Similarly, if workers were to earn their expected output as independent producers, then all two-layer firms would try to hire the least knowledgeable agents as workers.

Finally, the equilibrium wage function w is continuous, strictly increasing, and convex (strictly so when $z \in W \cup S$). The wages for the different occupations are obtained as follows. For independent producers, $w(z) = z$ is an immediate implication of the zero-profit condition of single-layer firms. For workers and solvers, consider the problem of a two-layer organization that recruited $n(z)$ workers with knowledge $z \in W$ and is deciding which solver $s \in S$ to hire:

$$\max_{s \in S} \left\{ n(z)[s - w(z)] - w(s) \right\}$$

The corresponding first-order condition evaluated at $s = m(z)$ implies that $w'(m(z)) = n(z)$, or, equivalently, $w'(z) = n(e(z))$ for any $z \in S$. Thus $w(z) = C + \int_{\inf S}^z n(e(u))du$ for any $z \in S$, where the constant C is chosen so that the wage function is continuous, as the latter condition is necessary for market clearing. The wages of workers are then determined by the zero-profit condition of two-layer organizations: $w(z) = m(z) - w(m(z))/n(z)$.

For simplicity, in what follows, we restrict attention to $h < h_0$. This implies that there are no independent producers in the pre-AI equilibrium. In a previous version of this paper (Ide and Talamàs, 2024a), we show that virtually all of our results—except one, discussed in Section 5.4—extend to the case where $h \geq h_0$.

4 The Equilibrium with Autonomous AI

Our objective is to understand the effects of autonomous AI by comparing the pre- and post-AI equilibrium. In this section, we provide a partial characterization of the post-AI equilibrium, focusing

²²A marginal increase in the knowledge of a worker with knowledge z liberates $h < 1$ units of her solver's time. This allows her firm to hire $1/(1 - z)$ extra workers, with an expected net output gain of $(m(z) - w(z))/(1 - z) < 1$ (as the wage $w(z)$ of that worker is strictly greater than z in the interior of W).

on the essential information needed for the comparison (which is carried out in Section 5). A complete characterization of this equilibrium can be found in Appendix B.

For future reference, we index the post-AI equilibrium using the superscript “*” (note that the pre-AI equilibrium has no superscript). Furthermore, recall that $W^* \equiv W_a^* \cup W_p^*$ is the overall set of human workers and that $S^* \equiv S_a^* \cup S_p^*$ is the overall set of human solvers.

Proposition 2. *In the presence of an autonomous AI, there is a unique equilibrium. The equilibrium is efficient and has the following features:*

- *Occupational stratification:* $W^* \preceq I^* \preceq S^*$.
- *No worker is better than AI; no solver is worse than AI:* $W^* \preceq \{z_{AI}\} \preceq S^*$.
- *Positive assortative matching:* $m^* : W_p^* \rightarrow S_p^*$ is strictly increasing and $W_a^* \preceq W_p^*$ and $S_p^* \preceq S_a^*$.

Furthermore, AI is always used for independent production, and whether it is also used as a worker or as a solver depends on its knowledge level **relative to the pre-AI equilibrium**.

- *If $z_{AI} \in W$, then AI is necessarily used as a worker (and possibly also as a solver).*
- *If $z_{AI} \in S$, then AI is necessarily used as a solver (and possibly also as a worker).*

Finally, the rental rate of compute r^* is equal to z_{AI} , and the wage function w^* is continuous, strictly increasing, and convex (strictly so when $z \in W_p^* \cup S_p^*$), and is given by:

- $w^*(z) = z_{AI}(1 - 1/n(z)) > z$ for all $z \in W_a^*$.
- $w^*(z) = m^*(z) - w^*(m^*(z))/n(z)$ for all $z \in W_p^*$.
- $w^*(z) = z$ for all $z \in I^*$.
- $w^*(z) = C^* + \int_{\inf S_p^*}^z n(e^*(u))du$ for all $z \in S_p^*$, where $C^* = \inf S_p^*$.
- $w^*(z) = n(z_{AI})(z - z_{AI}) > z$ for all $z \in S_a^*$.

In particular, $w^*(\sup W^*) = \sup W_p^*$, $w^*(z_{AI}) = z_{AI}$, and $w^*(\inf S^*) = \inf S_p^*$.

Proof. See Appendix B. □

Figure 4 illustrates the post-AI equilibrium in two scenarios. Panel (a) shows a “basic” AI used as a worker and independent producer but not as a solver. Panel (b) presents a more “advanced” AI used in all three roles. In both cases, AI is the “best worker” in the economy, so it is assisted by the most knowledgeable human solvers. In panel (b), AI is also the “worst solver,” so it assists the production work of the least knowledgeable human workers. Regardless of the case, humans with knowledge level z_{AI} earn their output as independent producers, being perfect substitutes for a unit compute using AI (priced at $r^* = z_{AI}$).

Proposition 2 has three parts. The first part states the basic properties of the equilibrium. The second part describes how firms use humans and AI as a function of AI’s knowledge. Finally, the third part characterizes the equilibrium prices.

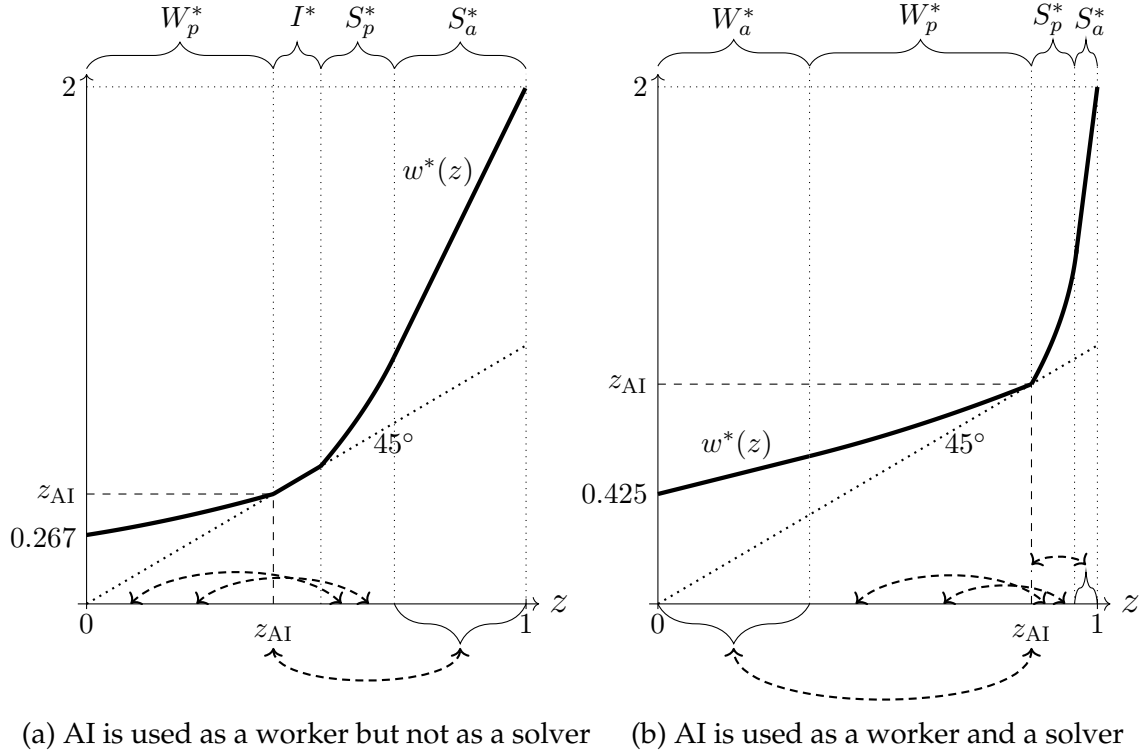


Figure 4: Illustration of the Post-AI Equilibrium for Two Different Values of z_{AI}

Notes. Distribution of knowledge: $G(z) = z$. Parameter values: Both panels have $h = 1/2$. For panel (a), $z_{AI} = 0.425$, while for panel (b), $z_{AI} = 0.85$. The thick line depicts the equilibrium wage function. The dashed arrows illustrate the matching between workers and solvers, both humans and AI. Human workers in W_p^* are endogenously matched with the human solvers in S_p^* according to m^* . All the humans in W_a^* are matched with an AI solver, while all humans in S_a^* are matched with AI workers. See the Online Appendix for a detailed characterization of the post-AI equilibrium when human knowledge is uniformly distributed.

Let us consider the first part. Like the pre-AI equilibrium, the post-AI equilibrium is efficient and exhibits occupational stratification and positive assortative matching. This is because the First Welfare Theorem continues to hold in this economy, and AI does not change the fact that (i) more knowledgeable agents have a comparative advantage in specialized problem solving, (ii) less knowledgeable agents have a comparative advantage in assisted rather than independent production work, and (iii) there are complementarities between worker and solver knowledge.

The remaining properties of the equilibrium follow from occupational stratification and positive assortative matching, in addition to the fact that AI can be used at scale. Indeed, because AI can be leveraged across all units of compute, and compute is abundant relative to time, a fraction of the available compute must be used for independent production. Hence, occupational stratification implies that no worker can be better than AI and no solver can be worse than AI (i.e., $W^* \preceq \{z_{AI}\} \preceq S^*$). Moreover, by positive assortative matching, if AI is used as a worker, then it is supervised by the most knowledgeable human solvers (i.e., $S_p^* \preceq S_a^*$), and if AI is used as a solver, then it assists the

least knowledgeable human workers (i.e., $W_a^* \preceq W_p^*$).

The second part of the proposition—how firms use humans and AI as a function of z_{AI} —is driven by the differences in the comparative advantages of agents at different knowledge levels. If AI has the knowledge of a pre-AI worker, then it must be efficient to allocate some compute to assisted production work. This is because a set of humans who are more knowledgeable than AI were receiving assistance in the pre-AI equilibrium, and AI is less likely than these humans to succeed on its own. Hence, when $z_{AI} \in W$, AI is necessarily used for assisted and independent production work.

Similarly, if AI has the knowledge of a pre-AI solver, it must be efficient to allocate some AI agents to assist humans. This is because a set of humans who are less knowledgeable than AI provided assistance in the pre-AI equilibrium, and AI—being more knowledgeable than them—has a comparative advantage in providing assistance. Thus, when $z_{AI} \in S$, AI must be used for independent production and specialized problem solving. Determining whether AI is simultaneously used in all three roles when $z_{AI} \in W$ or $z_{AI} \in S$ is subtle and not fundamental for our main results. For this reason, we relegate such details to Appendix B.

Finally, the third part of the proposition addresses equilibrium prices. The result that $r^* = z_{AI}$ follows because some compute must be used for independent production, and single-layer automated firms must break even.²³ The wages of the humans employed by tA and bA firms—those in W_a^* and S_a^* —are determined by the zero-profit conditions of these firms, along with the fact that $r^* = z_{AI}$:

$$\begin{aligned} n(z)[z_{AI} - w^*(z)] - r^* &= 0 \implies w^*(z) = z_{AI}(1 - 1/n(z)), \text{ for any } z \in W_a^* \\ n(z_{AI})[z - r^*] - w^*(z) &= 0 \implies w^*(z) = n(z_{AI})(z - z_{AI}), \text{ for any } z \in S_a^* \end{aligned}$$

The wages of the human workers and solvers hired by nA firms—those in W_p^* and S_p^* —are derived following a similar logic as in the pre-AI equilibrium. The only difference is that the best workers and worst solvers earn their expected output as independent producers (while they earn strictly more than that in the pre-AI equilibrium). The reason for this is that the equilibrium wage function must be continuous, and there is always some independent production in the post-AI equilibrium (done by AI and possibly also humans).

5 The Impact of Autonomous AI

We now analyze the impact of an autonomous AI by comparing the pre- and post-AI equilibrium. In essence, we compare the pre-AI equilibrium of Figure 3 with the post-AI equilibrium of Figure 4, as depicted in Figure 5. This comparison reveals substantial changes in (i) occupational choices, (ii) the

²³Note that z_{AI} is also the marginal product of compute: the increase in output from introducing one additional unit of compute into the economy. Indeed, since compute is abundant relative to time, any additional unit will necessarily be rented by a single-layer automated firm, meaning the resulting increase in output is z_{AI} .

matching between workers and solvers, and (iii) the distribution of labor income. These changes also affect the distribution of firm size, productivity, and decentralization.

For the analysis, we use the following terminology:

- $z_{AI} \in \text{int}W$: “AI has the knowledge of a pre-AI worker.”
- $z_{AI} \in \text{int}S$: “AI has the knowledge of a pre-AI solver.”
- $z_{AI} \in W \cap S$: “AI has the knife-edge knowledge of both a pre-AI worker and a pre-AI solver.”

For brevity, we focus on the cases $z_{AI} \in \text{int}W$ and $z_{AI} \in \text{int}S$. The results for the knife-edge case $z_{AI} \in W \cap S$ are in the Online Appendix.

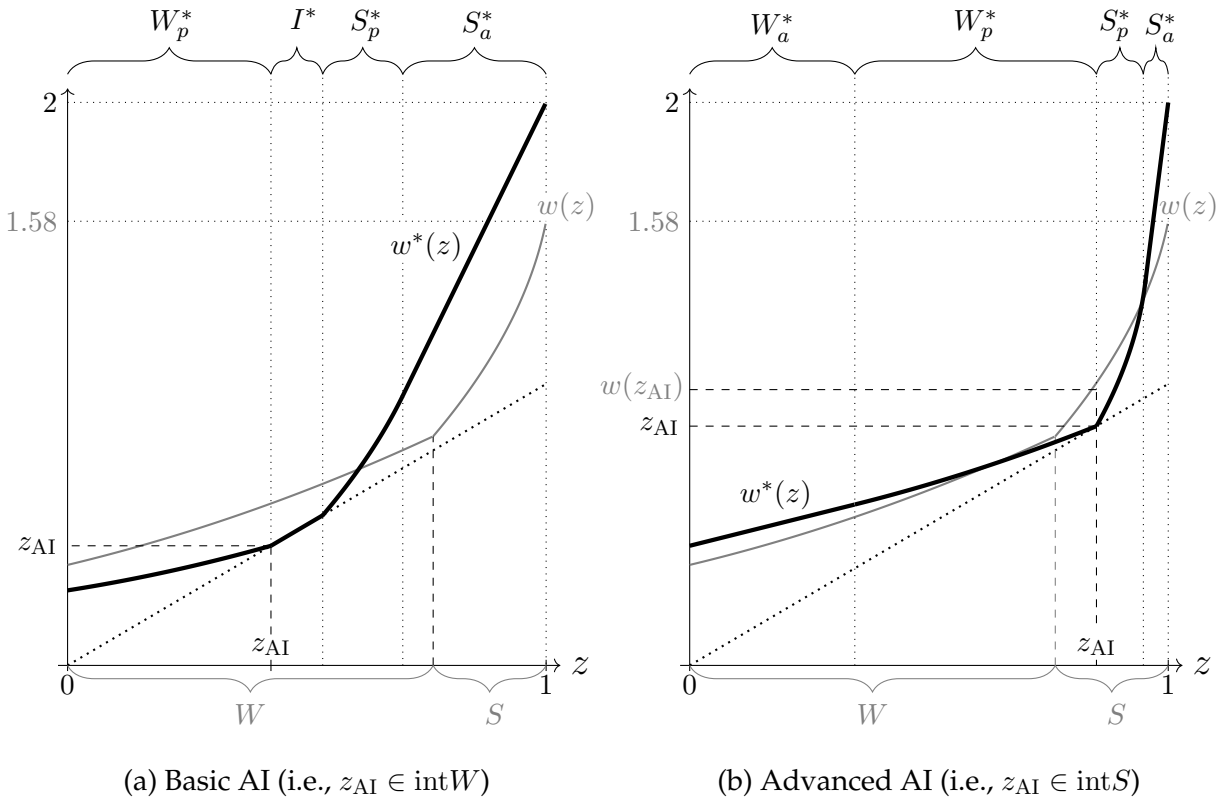


Figure 5: Comparison of the Pre- and Post-AI Equilibrium

Notes. Distribution of knowledge: $G(z) = z$. Parameter values: For both panels, $h = 1/2$ ($< h_0 = 3/4$). For panel (a), $z_{AI} = 0.425$, while for panel (b), $z_{AI} = 0.85$. In the post-AI equilibrium depicted in panel (a), AI is used as a worker and independent producer. In panel (b), AI is used in all three possible roles. See the Online Appendix for a detailed characterization of the pre-AI and post-AI equilibrium when human knowledge is uniformly distributed.

5.1 Occupational Displacement

We begin by analyzing the effects of AI on occupational choice.

Proposition 3.

- If $z_{AI} \in \text{int}W$, AI displaces humans from routine to specialized problem solving, i.e., $W^* \subset W$ and $S^* \supset S$.
- If $z_{AI} \in \text{int}S$, AI displaces humans from specialized to routine problem solving, i.e., $W^* \supset W$ and $S^* \subset S$.

Proof. See Appendix C.1. □

Note that the set of humans doing production work in either one- or two-layer organizations is the complement of the set of human solvers. Hence, when $z_{AI} \in \text{int}W$, AI decreases not only the number of humans doing assisted production work (i.e., workers) but also the number of humans doing any type of production work (i.e., workers and independent producers). Similarly, when $z_{AI} \in \text{int}S$, AI increases not only the number of humans doing assisted production work but also the number of humans doing any type of production work.

According to Proposition 3, the impact of AI on occupational choice is determined by the knowledge of AI *relative to the pre-AI equilibrium*: When AI has the knowledge of a pre-AI worker, it displaces humans from routine production work to specialized problem solving. In contrast, when AI has the knowledge of a pre-AI solver, the displacement goes in the opposite direction.

Intuitively, when AI has the knowledge of a pre-AI worker, it serves as a relatively inexpensive technology for routine production work, thus reducing workers' wages and increasing the attractiveness of creating two-layer firms. The result is a surge in the demand for solvers to match with the less expensive, more abundant "workers" (both humans and AI), which induces the most knowledgeable routine workers of the pre-AI equilibrium to switch to specialized problem solving.

Similarly, when AI has the knowledge of a pre-AI solver, it serves as a relatively inexpensive technology for specialized problem solving. This leads to the creation of additional two-layer organizations, increasing the demand for workers to match with the less expensive, more abundant "solvers" (both humans and AI). As a result, the least knowledgeable pre-AI solvers switch to routine production work.

5.2 Distribution of Firm Size, Productivity, and Span of Control

This section analyzes the effects of AI on the distribution of firm size, productivity, and span of control. We focus on two-layer organizations and begin with the following preliminary result:

Corollary 1. *AI necessarily increases the number (i.e., the measure) of two-layer firms.*

Proof. See Appendix C.2. □

We take *firm size* as its output, and *firm productivity* as its output divided by the units of time or compute used in production (this is a measure of “average” productivity). Hence, a firm’s productivity is equal to the knowledge of its solver. We define a firm’s *span of control* as the amount of resources under its solver’s supervision. Hence, the span of control of a firm that hires human workers with knowledge z is $n(z)$, while that of a firm that uses AI for production is $n(z_{\text{AI}})$. Note that span of control is a measure of decentralization because a solver can supervise more resources only if its workers solve more problems on their own.²⁴

Proposition 4.

- If $z_{\text{AI}} \in \text{int}W$, then AI decreases the maximum span of control, the minimum firm productivity, and both the minimum and maximum firm size.
- If $z_{\text{AI}} \in \text{int}S$, then AI increases the maximum span of control, the minimum firm productivity, and both the minimum and maximum firm size.

Proof. See Appendix C.3. □

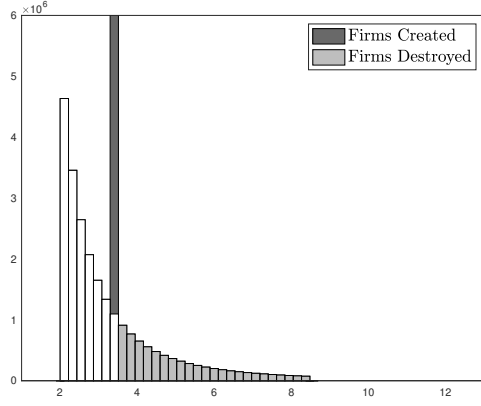
Proposition 4 is illustrated in Figure 6, which depicts the pre- and post-AI distributions of firm decentralization, productivity, and size for the two cases illustrated in Figure 5 (i.e., $z_{\text{AI}} = 0.425$ and $z_{\text{AI}} = 0.85$). The dark-shaded and light-shaded bars correspond to the firms created and destroyed by AI’s introduction, respectively. Hence, the pre-AI distribution is equal to the white bars plus the light-shaded bars, while the post-AI distribution is equal to the white bars plus the dark-shaded bars.

As the figure shows, when $z_{\text{AI}} \in \text{int}W$, AI destroys the largest and most decentralized firms and leads to the creation of firms that are smaller and less productive than the smallest and least productive pre-AI firms. Similarly, when $z_{\text{AI}} \in \text{int}S$, AI destroys the smallest and least productive firms and leads to the creation of firms that are larger and more decentralized than the largest and most decentralized pre-AI firms.

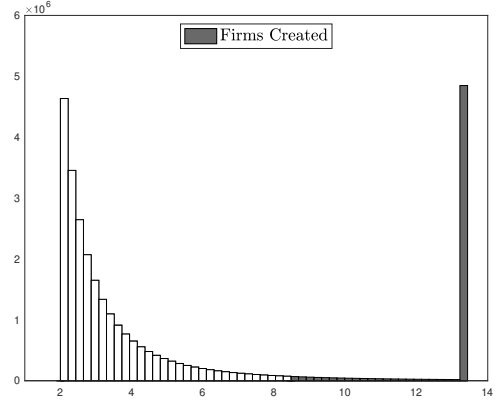
For intuition, let us first consider productivity and span of control. Recall that when $z_{\text{AI}} \in \text{int}W$, AI induces the most knowledgeable pre-AI workers to switch to specialized problem-solving (Proposition 3). In terms of productivity, this implies that AI leads to the creation of less productive firms, as the newly created class of solvers is less knowledgeable than the pre-AI solvers. In terms of span of control, this displacement of humans across occupations implies that AI destroys the most decentralized firms, as the most knowledgeable pre-AI workers become solvers post-AI. The intuition for the case $z_{\text{AI}} \in \text{int}S$ is analogous.

Now, consider the effects of AI on the distribution of firm size. Note that positive assortative matching implies that more productive firms are also more decentralized. Moreover, the least knowledgeable workers always have knowledge $z = 0$, and the most knowledgeable solvers always have

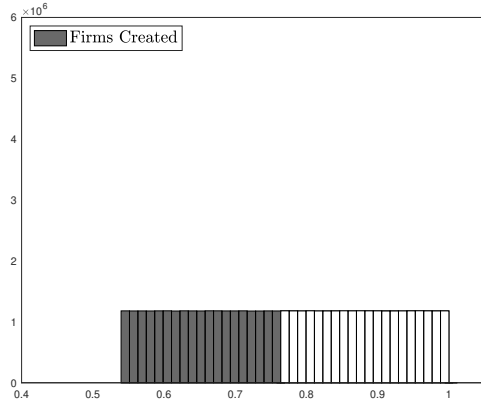
²⁴We focus on this measure (rather than the number of *humans* under the supervision of a given solver) because the spirit of our baseline model is that a unit of compute using AI is indistinguishable from a human with knowledge z_{AI} .



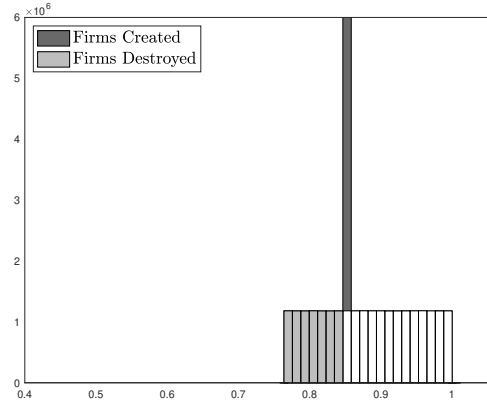
(a) Span of Control - $z_{AI} \in \text{int}W$



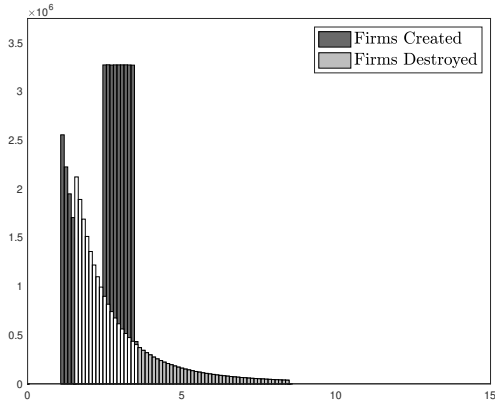
(b) Span of Control - $z_{AI} \in \text{int}S$



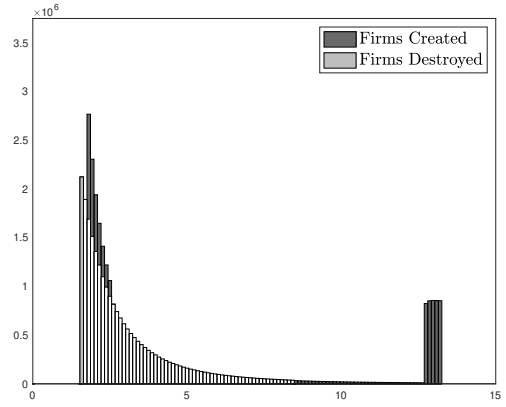
(c) Productivity - $z_{AI} \in \text{int}W$



(d) Productivity - $z_{AI} \in \text{int}S$



(e) Size - $z_{AI} \in \text{int}W$



(f) Size - $z_{AI} \in \text{int}S$

Figure 6: The Effects of AI on the Distribution of Firms' Span of Control, Productivity, and Size

Notes. This histogram is based on a human population of $N = 100 \times 10^6$ individuals. Distribution of knowledge: $G(z) = z$. For both panels, $h = 1/2$. Moreover, for panel (a), $z_{AI} = 0.425$, while for panel (b), $z_{AI} = 0.85$. Pre-AI, there are 23.6×10^6 firms. Post-AI, the number of firms increases to 46.7×10^6 in panel (a) and to 29.3×10^6 in panel (b). See the Online Appendix for the exact expressions of the pre- and post-AI distributions of firm size, productivity, and span of control.

knowledge $z = 1$. Hence, the size of the smallest firm is $n(0)$ times the minimum firm productivity, while the size of the largest firm is 1 times the maximum span of control.

Consequently, when $z_{AI} \in \text{int}W$, AI reduces the minimum and maximum firm size, as it decreases the maximum span of control and the minimum firm productivity. Similarly, when $z_{AI} \in \text{int}S$, AI increases the minimum and maximum firm size, as it increases the maximum span of control and the minimum firm productivity.

We end this subsection by showing that when AI has the knowledge of a pre-AI solver, and this knowledge is sufficiently high, AI also leads to the creation of superstar firms that have “scale without mass.” This term was coined by Brynjolfsson et al. (2008) and Autor et al. (2020) to refer to those firms that attain large market shares with relatively few employees.

More precisely, we say that AI creates superstar firms with scale but no mass if the following two conditions are satisfied: (i) the largest post-AI firms are larger than the largest pre-AI firms, and (ii) the largest post-AI firms are bottom-automated (that is, they use a single human solver to supervise production work by AI).

Corollary 2. *There exists a $\zeta \in \text{int}S$ such that for all $z_{AI} \geq \zeta$, AI leads to the creation of superstar firms with scale but no mass.*

Proof. See Appendix C.4. □

Intuitively, AI increases the maximum firm size only when $z_{AI} \in \text{int}S$. In such a case, we know that AI is necessarily used as a solver and possibly also as a worker (Proposition 2). The condition that $z_{AI} \geq \zeta$ is necessary and sufficient for AI to be used in both roles, in which case the largest post-AI firms are bA firms (as AI is the best worker in the economy). These superstar firms with scale but no mass can, in fact, be seen in panel (f) of Figure 6.

5.3 The Effects of AI on Workers and Solvers who Are Not Occupationally Displaced

We now analyze the impact of AI on the productivity of pre-AI workers who remain workers post-AI (i.e., workers who are not occupationally displaced) and the span of control of pre-AI solvers who remain solvers post-AI (i.e., solvers who are not occupationally displaced).

In line with our definition of firm productivity, we define a worker’s (average) productivity as her expected output per unit of time. Thus, it is equal to the knowledge of the solver with whom she is matched. Similarly, a solver’s span of control is equal to the resources under her supervision. Hence, her span of control is increasing in the knowledge of the workers with whom she is matched.

As the next proposition shows, AI affects not only those humans who directly interact with it but also those who do not interact with it. The reason for this is that when a subset of humans switches from matching with other humans to matching with AI, it induces a complete reorganization of all

matches in the economy. For the proposition that follows, let us recall that $e(z)$ is the employee function of the pre-AI equilibrium.

Proposition 5.

- If $z_{AI} \in \text{int}W$, then:
 - The productivity of $z \in W^* \subset W$ is strictly lower post-AI than pre-AI.
 - The span of control of $z \in S \subset S^*$ is strictly larger post-AI than pre-AI if $e(z) < z_{AI}$, and strictly smaller post-AI than pre-AI if $e(z) > z_{AI}$.
- If $z_{AI} \in \text{int}S$, then:
 - The productivity of $z \in W \subset W^*$ is strictly higher post-AI than pre-AI if $z < e(z_{AI})$, and strictly lower post-AI than pre-AI if $z > e(z_{AI})$.
 - The span of control of $z \in S^* \subset S$ is strictly larger post-AI than pre-AI.

Proof. See Appendix C.5. □

Proposition 5 is illustrated in Figure 7 when $z_{AI} \in \text{int}W$ and AI is used as a worker and as an independent producer (but not as a solver). As the figure shows, AI worsens the matches of all pre-AI workers who remain workers (i.e., those humans in W_p^*), reducing their productivity. Moreover, AI improves the match of the least knowledgeable solvers who are solvers both pre- and post-AI—increasing their span of control—but worsens the match of the most knowledgeable solvers who are solvers both pre- and post-AI.

For intuition, suppose first that AI has the knowledge of a pre-AI worker. In this case, (i) AI is the best worker in the post-AI economy and is therefore assisted by the best solvers (Proposition 2), and (ii) the knowledge of the worst solvers decreases with AI's introduction (Proposition 3). Both effects worsen the pool of solvers available for those pre-AI workers who remain workers post-AI, thus reducing their productivity.

Regarding the solvers who are not occupationally displaced, the match of the best solvers worsens with AI's introduction because post-AI, they assist production by AI, while pre-AI, they assist humans who are more knowledgeable than AI. However, the match of the worst solvers improves with AI's introduction because post-AI, the least knowledgeable workers are assisted by the newly appointed solvers (humans and possibly AI), who are less knowledgeable than the worst pre-AI solvers.

Similarly, when AI has the knowledge of a pre-AI solver, (i) AI is the worst solver in the post-AI economy and therefore assists the least knowledgeable workers (Proposition 2), and (ii) the knowledge of the best workers increases with AI's introduction (Proposition 3). Both effects improve the pool of workers available for pre-AI solvers who remain solvers post-AI, thus increasing their span of control.

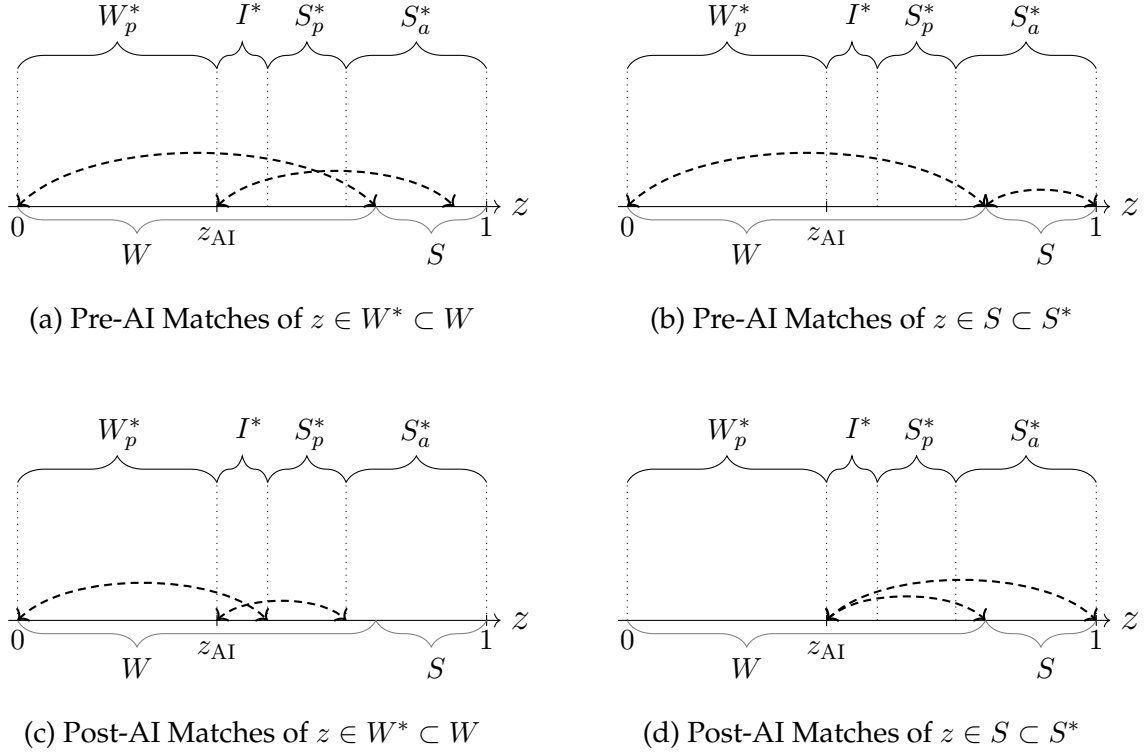


Figure 7: An Illustration of Proposition 5 - $z_{AI} \in \text{int}W$

Notes. Distribution of knowledge: $G(z) = z$. Parameter values: For all panels, $h = 1/2$ and $z_{AI} = 0.425$. In the post-AI equilibrium, AI is used as a worker and independent producer. The dashed arrows illustrate the matching between workers and solvers. See the Online Appendix for a detailed characterization of the pre-AI and post-AI equilibrium when human knowledge is uniformly distributed.

Regarding the workers who are not occupationally displaced, the match of the worst workers improves with AI's introduction because post-AI, they are assisted by AI, while pre-AI, they are assisted by humans who are less knowledgeable than AI. However, the match of the best workers worsens with AI's introduction because post-AI, the best solvers assist the production work of the newly appointed workers (humans and possibly AI) who are more knowledgeable than the best pre-AI workers.

Together, Propositions 4 and 5 show that even when AI destroys the most decentralized firms (i.e., when $z_{AI} \in \text{int}W$), a fraction of the non-occupationally displaced solvers are reallocated to more decentralized firms. Similarly, even when AI destroys the least productive firms (i.e., when $z_{AI} \in \text{int}S$), a fraction of the non-occupationally displaced workers are reallocated to less productive firms.

5.4 Labor Income

Finally, we analyze the effects of AI on total labor income and its distribution. We begin with the following result:

Lemma 1. *Total output and total labor income increase with the introduction of AI.*

The proof of this result is intuitive and relatively straightforward, so we provide it here as part of the main text. Moreover, this proof is instructive because it shows that not all gains from AI accrue to labor.

Proof. The fact that AI increases total output follows because AI expands the production possibility frontier, and the First Welfare Theorem holds in this setting. The fact that AI increases total labor income, in turn, follows from two observations. First, if all compute is assigned to independent production, then AI does not affect total labor income (as it does not interact with humans in the workplace). Second, capital income is equal to μz_{AI} regardless of how compute is used. Hence:

$$\text{Total output post-AI} = \text{Total labor income post-AI} + \mu z_{AI}$$

Consequently, given that the post-AI equilibrium is efficient, unique, and does not allocate all compute to independent production, the following must hold:

$$\text{Total output post-AI} > \text{Total labor income pre-AI} + \mu z_{AI}$$

Hence, total labor income must be strictly greater post-AI than pre-AI. $\square$

We now turn to analyzing the impact of AI on the distribution of labor income. Given wages are a reflection of an individual's marginal product, understanding whose wage increases or decreases with the introduction of AI amounts to understanding which humans are complemented by the technology (in the sense that AI increases their marginal product) and which humans are substituted by the technology (in the sense that AI decreases their marginal product). Note that an agent's marginal product is not necessarily equal to her productivity (as defined in Section 5.3) because an agent's introduction affects the output of other agents by changing worker-solver matches (which creates pecuniary externalities).

Disentangling the distributional effects of AI is not trivial due to the existence of two potentially countervailing forces. On the one hand, AI changes the composition of firms and, therefore, the quality of matches (as shown by Proposition 5). On the other hand, by mimicking humans with knowledge z_{AI} , AI changes the relative scarcities of different knowledge levels, affecting how each firm's output is divided between workers and solvers.

We first show that *if* a human with knowledge $z < z_{AI}$ wins from AI's introduction, then all those humans with knowledge $z' < z$ must also be winners. Similarly, *if* a human with knowledge $z > z_{AI}$

wins from AI's introduction, then all those humans with knowledge $z' > z$ must also be winners. Given that humans with knowledge z_{AI} are always worse off after AI's introduction, this implies that the winners from AI are necessarily located at the extremes of the knowledge distribution:

Lemma 2. Define $\Delta(z) \equiv w^*(z) - w(z)$. Then $\Delta(z_{AI}) < 0$ and:

- If $\Delta(z) > 0$ for some $z \in [0, z_{AI}]$, then $\Delta(z') > 0$ for all $z' \in [0, z]$.
- If $\Delta(z) > 0$ for some $z \in [z_{AI}, 1]$, then $\Delta(z') > 0$ for all $z' \in [z, 1]$.

Proof. See Appendix C.6. □

Next, we characterize when there are winners below or above z_{AI} :

Proposition 6.

- (i) There exists a set $[0, z] \subseteq [0, z_{AI}]$ of winners if and only if $z_{AI} > \bar{z}_{AI}$, where $\bar{z}_{AI} \in \text{int}W$.
- (ii) There always exists a set $[z, 1] \subseteq (z_{AI}, 1]$ of winners.

Proof. See Appendix C.7. □

Figure 5, at the beginning of this section, provides an illustration of Proposition 6. As shown in panel (a) of the figure, only the most knowledgeable humans win when z_{AI} is relatively low (more precisely, all $z > 0.48$ are winners). In contrast, as shown in panel (b), both the least and the most knowledgeable humans win when z_{AI} is relatively high (more precisely, all $z < 0.63$ and all $z > 0.95$ are winners in this case).

We now provide intuition for this proposition. We do this using Figure 8, which builds on Figure 5(a) and illustrates the pre- and post-AI scenarios when $z_{AI} \in \text{int}W$. Based on Lemma 2, to verify the existence of winners below or above z_{AI} , it is sufficient to analyze the impact of AI on the wages of the least and most knowledgeable individuals. Accordingly, Figure 8(a) shows how AI affects the match of individuals with $z = 0$ and the wages of their match before and after AI, while Figure 8(b) provides the same analysis for individuals with $z = 1$.

Consider Figure 8(a) first. As the figure shows, AI's introduction has two opposing effects on the wages of the least knowledgeable individuals. On the one hand, AI reduces their wages by lowering their productivity, as they are now matched with a less knowledgeable solver (a negative “match” effect). On the other hand, AI increases their wages by allowing them to appropriate a larger share of the firm's output (a positive “share” effect). This occurs because the least knowledgeable solvers (with whom they are matched) now earn exactly their expected output as independent producers—their wages fall precisely on the 45° line—whereas pre-AI, they earned strictly more.

Thus, when $z_{AI} \in \text{int}W$, AI benefits the least knowledgeable individuals if the positive share effect outweighs the negative match effect. This happens when z_{AI} is sufficiently high (i.e., when $z_{AI} > \bar{z}_{AI}$,

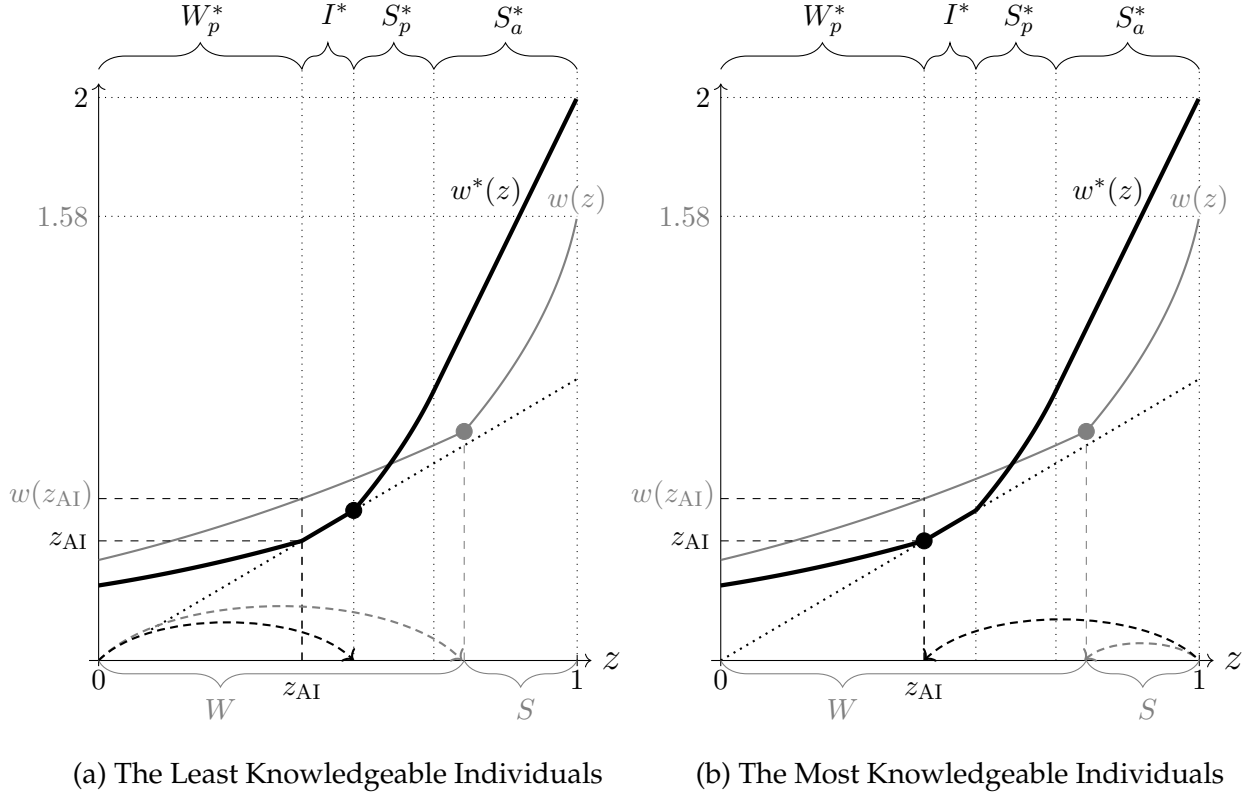


Figure 8: Decomposing the Effects of AI on Wages

Notes. Distribution of knowledge: $G(z) = z$. Parameter values: $h = 1/2 (< h_0 = 3/4)$ and $z_{AI} = 0.425$, so AI is used as a worker and independent producer. Panel (a) shows how AI affects the match of $z = 0$ and the wages of their match before and after AI, while panel (b) provides the same analysis for $z = 1$. See the Online Appendix for a detailed characterization of the pre-AI and post-AI equilibrium when human knowledge is uniformly distributed.

where $\bar{z}_{AI} \in \text{int}W$). Conversely, when $z_{AI} \in \text{int}S$, the least knowledgeable individuals always benefit from AI's introduction because both the match and share effects are positive: AI always improves the match of the least knowledgeable individuals in this case (see Proposition 5).

Now consider Figure 8(b). As in the case of the least knowledgeable individuals, AI generates both a negative match effect and a positive share effect for the most knowledgeable individuals. However, in this case, the positive share effect always dominates, explaining why there is always a set of winners above z_{AI} .

The key driver behind this result is that $h < h_0$, meaning the size of the team each solver assists is large. Consequently, even a small decrease in the wages of the most knowledgeable workers (with whom the most knowledgeable individuals are matched) is amplified by team size, helping the positive share effect outweigh any negative match effect. In a previous version of this paper (Ide and Talamàs, 2024a), we show that AI can sometimes harm the most knowledgeable individuals when $h \geq h_0$.

Finally, it is important to note that the finding that the most knowledgeable individuals always

benefit from AI’s introduction not only relies on $h < h_0$ but also on the fact that $z_{AI} \in [0, 1)$, meaning there is no “superintelligent AI.” As Proposition 6 make clear, when $h < h_0$ and $z_{AI} \in [0, 1)$, the most knowledgeable humans necessarily benefit from AI’s introduction, even as z_{AI} approaches 1. However, as shown in the Online Appendix, the most knowledgeable humans necessarily lose from the technology’s introduction when $z_{AI} = 1$.

Intuitively, this difference arises because when $z_{AI} < 1$, the humans with $z \in (z_{AI}, 1]$ can still use AI to leverage their knowledge. In contrast, when $z_{AI} = 1$, this is no longer possible, as AI can solve the same range of problems as the most knowledgeable humans.

6 Non-Autonomous AI

In our baseline model, AI is a technology that transforms units of compute into AI agents that can do exactly the same as humans with knowledge z_{AI} . This type of AI is “autonomous,” as it creates agents that not only provide solutions to problems but can also attempt to solve them (i.e., they can actively pursue production opportunities). We started with this case because it generates the most anxiety and debate, fueled by technology firms and venture capitalists who both believe in—and actively promote—the advancement of AI in this direction (see Section 2).

However, for technical or regulatory reasons, many existing AI systems still lack the ability to operate autonomously. In this section, we analyze the impact of non-autonomous AI. While autonomous AI creates agents that can function as both co-workers and solvers/co-pilots—capable of performing production work and giving advice depending on the context—non-autonomous AI creates agents that can only act as a co-pilot.

Our main finding is that a non-autonomous AI tends to benefit the least knowledgeable individuals the most, contrasting with autonomous AI, where the greatest beneficiaries are the most knowledgeable individuals. However, overall output is higher with autonomous than non-autonomous AI.

Formally, the model is identical to the baseline, except that AI cannot generate or pursue production opportunities. Thus, single-layer automated firms and bottom-automated firms are unable to produce output, so a human will never be hired to act as a solver in a bottom-automated firm (i.e., $S_a = \emptyset$ and $\mu_w = 0$). In this context, we interpret the amount of compute rented for independent production, μ_i , as the compute that is left idle (i.e., not productive).

To ensure a meaningful comparison of equilibrium outcomes between autonomous and non-autonomous AI, we continue to limit firms to a maximum of two layers. Additionally, when comparing these AI systems, we assume that the available compute is identical and abundant relative to time in both scenarios. In the Online Appendix, we extend this analysis to different compute levels. For future reference, we use the superscript “ $\star$ ” to denote equilibrium outcomes with non-autonomous

AI, distinguishing them from those with autonomous AI, which are marked by “*.”

6.1 The Impact of Non-Autonomous AI

Since compute is abundant relative to time but AI lacks autonomy, the equilibrium rental rate of compute is zero, $r^* = 0$. The reason is that some compute must be rented for independent production work, which cannot generate output without AI autonomy. However, in contrast to the case of an autonomous AI, the wage for an individual with knowledge z_{AI} is no longer equal to the equilibrium rental rate of compute ($w^*(z_{AI}) \neq r^* = 0$). This difference occurs because an AI agent is no longer a perfect substitute for a human with knowledge z_{AI} .

The following proposition completes the characterization of the post-AI equilibrium. For this proposition, recall that $w(z)$ denotes the wage schedule of the pre-AI equilibrium.

Proposition 7. *In the presence of a non-autonomous AI, there is a unique equilibrium. The equilibrium is efficient, maximizes labor income, and the rental rate of compute is $r^* = 0$. Moreover,*

- *If $z_{AI} \leq w(0)$, then AI is not used in equilibrium, i.e., $W_a^* = \emptyset$. Hence, human occupations and wages coincide with the pre-AI ones.*
- *If $z_{AI} > w(0)$, then only the least knowledgeable individuals use AI as a solver, i.e., $W_a^* \subseteq (W_p^* \cup I^* \cup S_p^*)$, where all these sets are not empty except possibly I^* .*

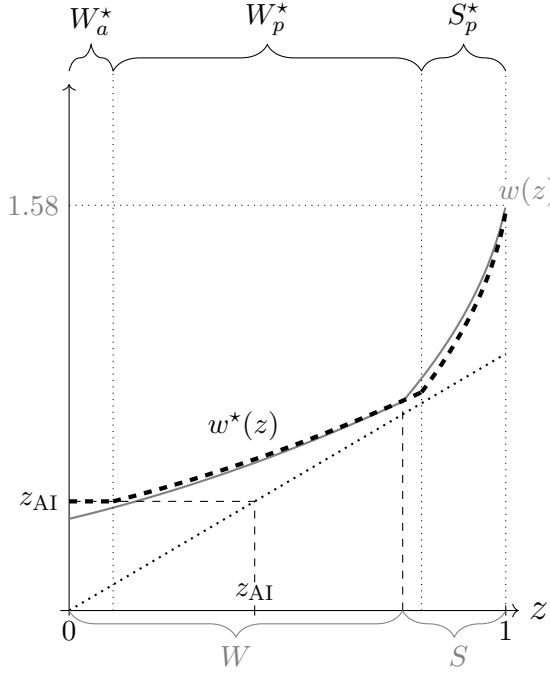
In any case,

1. *Overall output is strictly higher with autonomous than non-autonomous AI.*
2. *Non-autonomous AI creates losers:*
 $\exists z \in (0, 1]$ such $w^*(z) \leq w(z)$ (with strict inequality if $z_{AI} > w(0)$).
3. *The least knowledgeable benefit more from non-autonomous AI than from no AI or autonomous AI:*
 $\exists \epsilon > 0$ such that, for all $z \in [0, \epsilon)$, $w^*(z) \geq \max\{w(z), w^*(z)\}$ (with strict inequality if $z_{AI} > w(0)$).
4. *The most knowledgeable benefit more from autonomous AI than from non-autonomous AI:*
 $\exists \epsilon > 0$ such that, for all $z \in (1 - \epsilon, 1]$, $w^*(z) \leq w^*(z)$ (with strict inequality if $z \neq 1$).

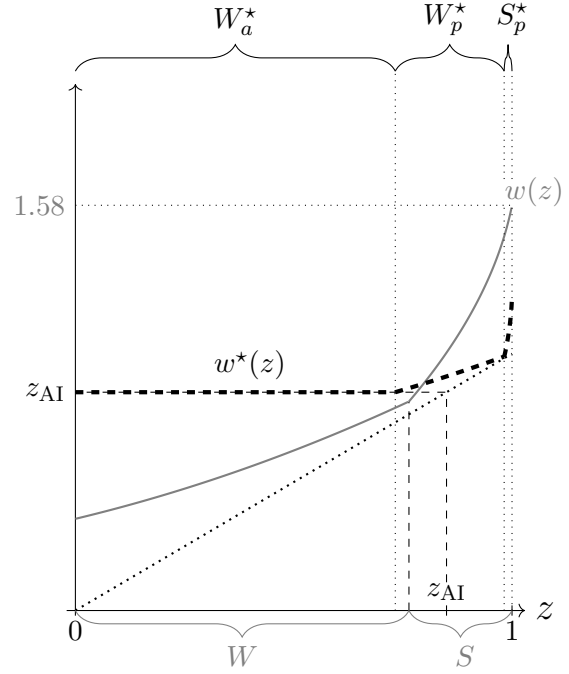
Proof. See the Online Appendix. □

Figure 9 provides an illustration of Proposition 7. Panels (a) and (b) show the effects of a non-autonomous AI compared to the pre-AI scenario, distinguishing between Basic AI and Advanced AI. Panels (c) and (d) compare the equilibrium wages for autonomous and non-autonomous AI under the same parameter values as panels (a) and (b).

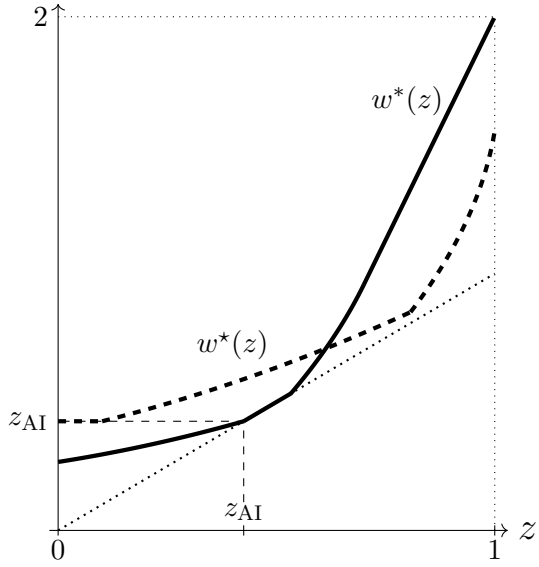
According to this proposition, a non-autonomous AI is used in equilibrium only if its knowledge level is sufficiently advanced (i.e., if $z_{AI} > w(0)$). In this case, individuals with knowledge below a



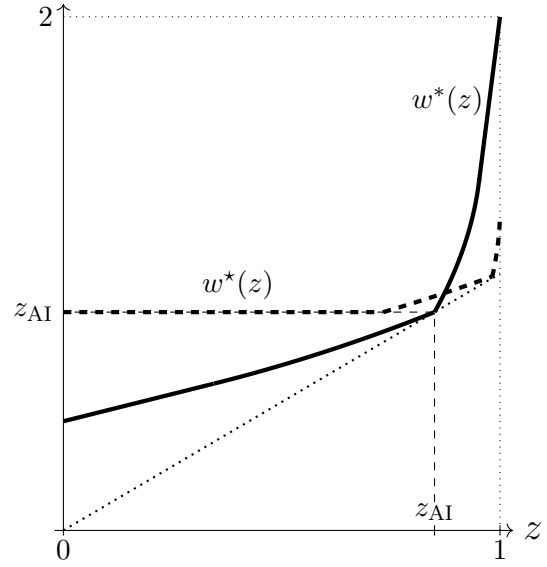
(a) Basic AI (i.e., $z_{\text{AI}} \in \text{int} W$)



(b) Advanced AI (i.e., $z_{\text{AI}} \in \text{int} S$)



(c) Autonomy vs. No Autonomy
(Basic AI)



(d) Autonomy vs. No Autonomy
(Advanced AI)

Figure 9: The Effects of a Non-Autonomous AI

Notes. Distribution of knowledge: $G(z) = z$. Parameter values: For all panels, $h = 1/2$. Moreover, for panel (a), $z_{\text{AI}} = 0.425$, while for panel (b), $z_{\text{AI}} = 0.85$. The thick gray line represents the pre-AI equilibrium wage function w . The gray line depicts the equilibrium wage function of the pre-AI equilibrium, w . The thick black dashed line depicts the equilibrium wage function with a non-autonomous AI, w^* . The thick black line represents the equilibrium wage function with an autonomous AI, w^* . In panels (a) and (b), the most knowledgeable individuals lose relative to the pre-AI equilibrium from the introduction of a non-autonomous AI. This result, however, is not general; it is possible to construct examples in which the most knowledgeable individuals also win from a non-autonomous AI.

certain threshold work in automated firms—using AI as a solver or co-pilot—while those above the threshold work in non-automated firms. Moreover, total output with non-autonomous AI is strictly lower than with autonomous AI because non-autonomy imposes a binding constraint on compute usage, leaving some compute idle.

In terms of wages, the introduction of non-autonomous AI creates both winners and losers in equilibrium, similar to autonomous AI. However, the distribution of gains and losses changes dramatically. Non-autonomous AI always benefits the least knowledgeable individuals relative to the pre-AI equilibrium. Moreover, these benefits exceed those of an autonomous AI. In contrast, non-autonomous AI may lead to gains or losses compared to the pre-AI scenario for the most knowledgeable individuals. In any case, their wages with non-autonomous AI are always lower than with autonomous AI.

Intuitively, a non-autonomous AI always benefits the least knowledgeable individuals because it is used exclusively as a solver. This removes the competition between the AI and these individuals for production work while allowing them to solve problems of greater difficulty without human assistance. This advantage is enhanced by the fact that non-autonomous AI has a lower opportunity cost than autonomous AI, as it cannot independently perform production tasks. This affordability enables the least knowledgeable individuals to appropriate a larger share of the benefits generated by AI.

For the most knowledgeable individuals, the effect is reversed. Non-autonomous AI is less advantageous than autonomous AI because it lacks the capacity to handle routine work and may compete with these individuals in advisory roles. However, these factors do not necessarily leave the most knowledgeable individuals worse off relative to the pre-AI equilibrium. By raising wages for the least knowledgeable, a non-autonomous AI may reduce the wages of the least knowledgeable pre-AI solvers, enabling the most knowledgeable individuals to hire these solvers—who remain relatively knowledgeable within the population—as workers at relatively low wages. This, under certain parameters, can increase the wages of the most knowledgeable individuals compared to the pre-AI equilibrium.

As noted earlier, we have limited firms to a maximum of two layers to ensure a meaningful comparison of equilibrium outcomes between autonomous and non-autonomous AI. This restriction provides an additional benefit: it implies that problems within the organization can be escalated only once, effectively introducing a cost of seeking help beyond the helping cost h (which is borne solely by the recipient of the question). Notably, this implies that even with free compute, individuals with slightly less knowledge than AI prefer to seek help from another human rather than from AI, as the probability of AI providing helpful assistance is low (see Figure 9).²⁵

²⁵The equilibrium is continuous with respect to the availability of compute, meaning all our results hold when the price of compute is strictly positive but sufficiently close to zero. In the Online Appendix, we also examine the case where compute is significantly less abundant. The results of this section extend to that scenario as well, albeit with certain qualifications.

We believe this is a reasonable outcome, as a similar situation would arise if problems could be escalated multiple times but the agents seeking assistance bore part of the time cost associated with requesting help. Without these constraints—such as with an arbitrary and endogenous number of layers and no time cost for the asking party—all individuals less knowledgeable than AI would always rely on the system for help. In this scenario, introducing AI effectively shifts the knowledge of all agents with $z < z_{AI}$ towards z_{AI} . Nevertheless, the key result of this section—that the least knowledgeable individuals prefer non-autonomous AI while the most knowledgeable favor autonomous AI—remains valid even under this alternative case.²⁶

For completeness, we conclude this section by examining the effects of non-autonomous AI on (i) human occupational choices, (ii) the distribution of firm size, productivity, and decentralization, and (iii) the productivity and span of control of workers and solvers who are not displaced, compared to the pre-AI equilibrium.

Proposition 8. *When $z_{AI} > w(0)$, the effects of a non-autonomous AI relative to the pre-AI equilibrium are the following:*

- *Occupational Displacement:*
 - *AI always displaces humans from specialized problem-solving to routine work, i.e., $W^* \supset W$ and $S^* \subset S$.*
- *Distribution of Firm Productivity, Size, and Span of Control:*
 - *When $z_{AI} \in \text{int}W$, AI decreases the minimum firm productivity and the minimum firm size, while increasing the maximum span of control and the maximum firm size.*
 - *When $z_{AI} \in \text{int}S$, AI increases the maximum span of control, the minimum firm productivity, and both the minimum and maximum firm size.*
 - *AI never leads to the creation of superstar firms that have scale but no mass.*
- *Productivity of Non-Occupationally Displaced Workers:*
 - *When $z_{AI} \in \text{int}W$, the productivity of all workers is lower post-AI than pre-AI.*
 - *When $z_{AI} \in \text{int}S$, the productivity of the least knowledgeable workers is higher post-AI than pre-AI. The productivity of the most knowledgeable workers is lower post-AI than pre-AI.*
- *Span of Control of the Non-Occupationally Displaced Solvers.*
 - *The span of control of all solvers is larger post-AI than pre-AI.*

Proof. See the Online Appendix. □

²⁶With more than two layers, an advanced non-autonomous AI could also amplify the expertise of the most knowledgeable individuals by relieving them of the need to assist less knowledgeable individuals (who, post-AI, turn to AI for assistance before seeking help from other humans). However, the key insight remains unchanged: the most knowledgeable individuals still prefer autonomous AI over non-autonomous AI because the former offers the additional advantage of being able to perform routine work for them.

7 Discussion and Final Remarks

In this paper, we propose a new framework for analyzing the impact of recent advances in AI on organizational and labor outcomes. Our approach is grounded in the premise that AI can automate knowledge work—an increasingly critical part of the economy where time and knowledge are the scarce inputs in production. Viewed from this perspective, AI stands apart from previous automation technologies due to its potential to alleviate time and knowledge constraints, enabling organizations to better leverage their members’ knowledge.

To conclude, we discuss some additional implications of our analysis and link our findings to emerging evidence and stylized facts about AI’s effects on the labor market.

Co-pilots vs. Co-workers.—Our analysis uncovers distinct labor implications when comparing AI co-pilots—which provide problem-solving assistance—to AI co-workers—which independently execute production tasks. Compared to the pre-AI scenario, AI co-pilots primarily benefit less knowledgeable individuals, reducing inequality in earnings and performance. In contrast, AI co-workers offer the greatest advantages to highly knowledgeable individuals, amplifying the value of expertise. Co-pilots emerge when AI is either (i) not autonomous or (ii) autonomous but sufficiently advanced, whereas co-workers can arise for any knowledge level but only in the case in which AI is autonomous.

Figure 10 illustrates these differences by comparing the pre-AI and post-AI equilibria under two scenarios: Autonomous AI (Panel (a)) and non-autonomous AI (Panel (b)) at $z_{AI} = 0.55$. At this knowledge level, autonomous AI functions exclusively as a co-worker, while non-autonomous AI operates solely as a co-pilot.

Intuitively, AI co-pilots “level the playing field” because they provide the least knowledgeable individuals with a low-cost tool to solve problems of greater difficulty. This enables a broader range of workers to address complex problems, aligning with the insights of Autor (2024).²⁷ In contrast, AI co-workers amplify expertise by allowing highly knowledgeable individuals to leverage their knowledge cheaply and at scale.

These findings are particularly relevant in light of emerging experimental evidence on AI’s labor effects. Existing studies consistently show that AI disproportionately benefits the least knowledgeable individuals and reduces performance inequality (e.g., Dell’Acqua et al., 2023; Noy and Zhang, 2023; Brynjolfsson et al., 2023; Peng et al., 2023; Wiles et al., 2023; Caplin et al., 2024).²⁸ However, these

²⁷In Autor’s (2024, p. 2) view, “[AI] can enable a larger set of workers equipped with necessary foundational training to perform higher-stakes decision-making tasks currently arrogated to elite experts.”

²⁸For instance, Brynjolfsson et al. (2023) reports that deploying an AI-based conversational assistant significantly improved productivity for less-skilled agents in a Fortune 500 customer support center, with minimal impact on experienced ones. Similarly, Caplin et al. (2024) reported that AI assistance was more beneficial for lower-ability individuals than for higher-ability ones when tasked with estimating whether people in photographs appeared to be over 21 years old.

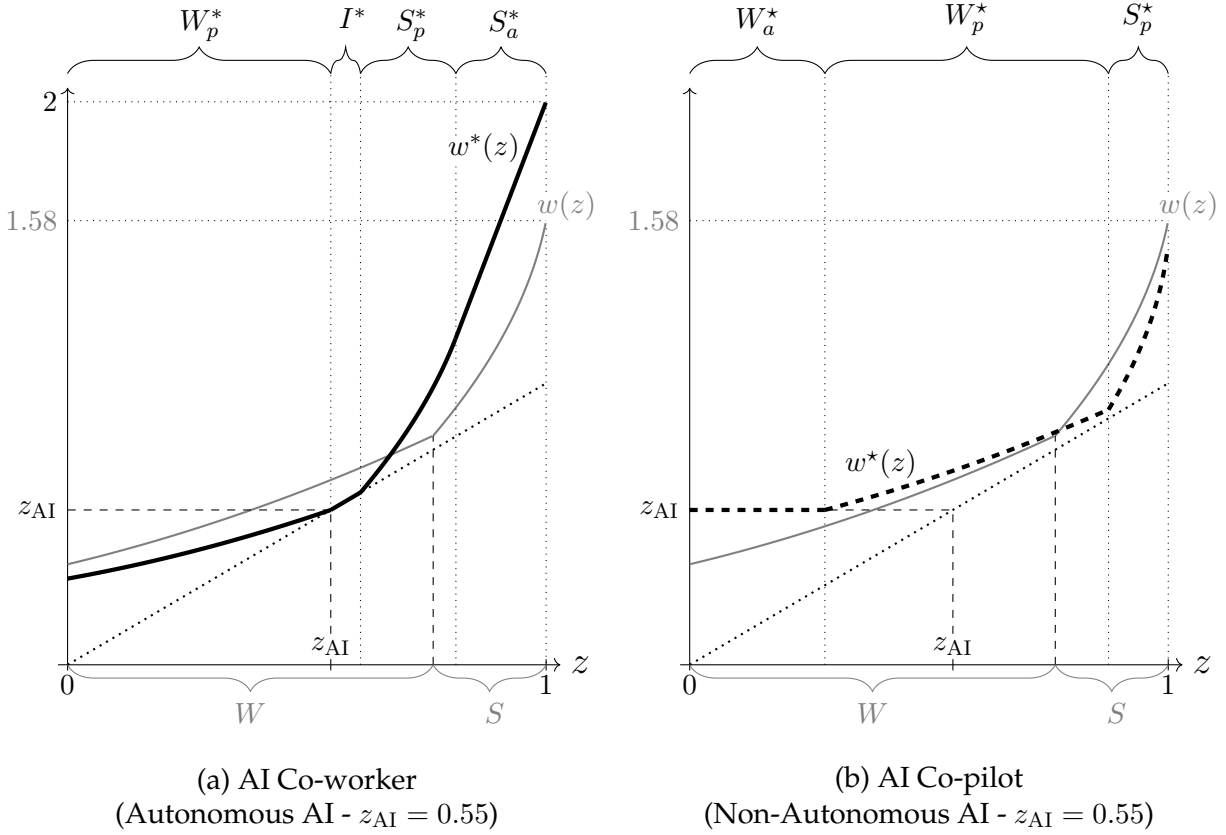


Figure 10: Co-workers vs. Co-pilots

Notes. Distribution of knowledge: $G(z) = z$. Parameter values: For both panels, $h = 1/2$ ($< h_0 = 3/4$) and $z_{AI} = 0.55$. The thick gray line represents the pre-AI equilibrium wage function, w . The thick black dashed line depicts the equilibrium wage function with a non-autonomous AI, w^* . The thick black line represents the equilibrium wage function with an autonomous AI, w^* . When $z_{AI} = 0.55$, autonomous AI functions exclusively as a co-worker, whereas non-autonomous AI operates solely as a co-pilot.

studies focus exclusively on AI co-pilots. While our framework supports these conclusions, it also highlights that they may not extend to autonomous AI designed to perform independent knowledge work, such as conducting research, drafting documents, or performing data analysis.

In fact, non-experimental evidence suggests that firms perceive current iterations of generative AI more as a basic co-worker than as a co-pilot. For instance, using high-frequency job posting data from LinkUp, [Berger et al. \(2024\)](#) document a marked decline in job postings for white-collar roles requiring low levels of knowledge, such as low-level office jobs, following the introduction of ChatGPT in November 2022. At the same time, job postings for high-knowledge roles, such as executive positions, increased. This trend indicates that firms view generative AI as a technology that complements high-skilled knowledge workers while substituting for their low-skilled counterparts. Similarly, [Alekseeva et al. \(2024\)](#) provide evidence that AI adoption changes managerial roles by heightening the demand for complex problem-solving skills while diminishing the demand for routine administrative skills.

We believe that further empirical studies designed to disentangle and estimate the distinct labor market implications of AI co-workers and AI co-pilots represent an important avenue for future research.

Reorganizations in Professional Services.— Our analysis highlights how AI reshapes firm structures by reorganizing knowledge work. For example, basic autonomous AI induces the most knowledgeable pre-AI solvers to focus on supporting AI-driven production. This reallocation forces human workers to seek assistance from less knowledgeable individuals, pushing the most knowledgeable pre-AI workers to take specialized problem-solving roles (see Section 5).

There is anecdotal evidence of these reorganizations in professional services industries. For instance, in 2022, Litera, a legal technology provider, surveyed over 200 Mergers and Acquisitions (M&A) lawyers from firms with more than 80 lawyers in the United States, United Kingdom, and Canada. Many respondents expressed concern that AI tools, particularly those used for routine tasks like document review, are limiting learning opportunities for junior lawyers who traditionally performed these tasks. At the same time, most agreed that by automating such tasks, AI allows junior lawyers to focus on more complex work, such as client engagement and legal analysis (Litera, 2022).

Similarly, Beane and Anthony (2024) document how junior analysts in investment banking are increasingly isolated from senior partners, as these partners now rely on AI to handle tasks previously requiring junior involvement. This shift has also raised concerns about skill development among less experienced workers. Beane (2024) underscores the issue, stating:

The way we’re handling AI is keeping young workers from learning skills... [This is because AI] allows experts to do more, independently, so they don’t need younger, less-experienced workers to help them out anymore—so those novices are left without mentors to teach them the skills they need to do their job.

Ide and Talamàs (2024b) study the effects of AI on the duration and efficiency of apprenticeships and relational knowledge transfers using the framework of Garicano and Rayo (2017).

AI vs. Traditional Robots, Offshoring and Immigration.— In this paper, we model AI as a technology that transforms compute into AI agents capable of performing knowledge work, either autonomously or non-autonomously. AI, therefore, is akin to a labor shock—it is an automation technology. This paper maps the distinctive characteristics of this shock to labor and organizational outcomes.

The automation shock induced by AI differs fundamentally from that caused by traditional robots. Traditional robots excel at executing codifiable tasks at scale and are, therefore, best described as a “putty-clay” technology (Martinez, 2021):²⁹ Once deployed, their functionality remains fixed, with

²⁹Putty-clay models assume that capital and technology are flexible only prior to investment decisions, whereas putty-putty models allow for flexibility both before and after investments are made (Atkeson and Kehoe, 1999).

limited capacity for adaptation or improvement.³⁰

In contrast, AI foundation models and compute exhibit “putty-putty” characteristics: Compute can be rented on demand through cloud computing providers and costlessly redeployed across various AI foundation models.³¹ These models, in turn, are highly flexible, adaptable, and designed for self-improvement. Consequently, while traditional robots exist in vintages with differing capabilities, compute consistently gravitates toward the most advanced and high-performing AI models, as discussed in Section 3.2.

Similarly, the labor shock induced by AI diverges sharply from classical shocks such as offshoring or immigration. Non-autonomous AI, for instance, has no clear analog in these frameworks. Moreover, even for autonomous AI, the introduction of the technology diverges significantly from offshoring or immigration. The key distinction lies in AI’s scalability. Once knowledge is encoded into an AI system, it can be deployed across all units of compute, which are widely available. This scalability amounts to the introduction of a large population of homogeneous agents.

In contrast, large human populations are inherently heterogeneous. Consequently, offshoring or immigration typically involve either the introduction of a relatively small population of homogeneous agents, or the introduction of a large population of heterogeneous agents. As we formally show in the Online Appendix, this difference leads to qualitatively different outcomes.

AI vs. ICTs.— A significant implication of our model is that not all firms adopt AI to *automate knowledge work*. This characteristic makes AI, as an automation technology, fundamentally different from earlier waves of digitalization—such as the introduction of computers or email—which saw near-universal adoption across firms. Partial adoption is a common characteristic of automation technologies. Both empirical observations and standard models of automation (e.g., [Acemoglu and Loebbing, 2024](#)) suggest that it is often not cost-effective to automate all tasks that can potentially be automated.

For autonomous AI, partial adoption occurs because deploying AI involves a strictly positive opportunity cost: the foregone output that an autonomous AI agent could generate as an independent producer. As a result, adopting autonomous AI may not be economically viable for all firms, consistent with the evidence presented by [Svanberg et al. \(2024\)](#). For non-autonomous AI, partial adoption arises even when the opportunity cost is zero. This is because AI’s advice is useless in firms with

³⁰For instance, a robotic arm in a car assembly line might be programmed and equipped to install windshields with remarkable precision. However, if the same robot needed to switch to assembling dashboards, it would often require costly hardware reconfiguration, reprogramming, and potentially even new components. Relatedly, when a better robot design becomes available, it is costly to update the existing robots to the new design.

³¹As noted in Section 3.2, generative AI model providers typically use a pay-as-you-go pricing model based on compute volume. Furthermore, the ease with which these models can be swapped is amplified by platforms like Amazon Bedrock, which integrate multiple foundation models within the same cloud environment. Indeed, according to a recent survey by Andreessen Horowitz, “Most enterprises are designing their applications so that switching between models requires little more than an API change. Some companies are even pre-testing prompts so the change happens literally at the flick of a switch” (see [Wang and Xu, 2024](#)).

very knowledgeable human workers.

However, as discussed in Section 3.2, the utility of AI extends beyond automating knowledge work. Thus, firms may adopt AI technologies for purposes unrelated to automation, such as reducing communication costs or lowering the cost of acquiring knowledge. As noted in the same section, the effects of such technologies have been extensively studied (see, e.g., [Garicano and Rossi-Hansberg, 2006](#); [Antràs et al., 2008](#); [Bloom et al., 2014](#); [Garicano and Rossi-Hansberg, 2015](#)). These effects are fundamentally different from the ones of AI as an automation technology.

For instance, better communication technologies—modeled as a reduction in the communication cost h —always displace humans from specialized problem-solving towards routine work as the very best solvers can now supervise larger teams. In contrast, the displacement generated by AI as an automation technology depends on its autonomy and its problem-solving capabilities (see Propositions 3 and 8).

AI in Advanced vs. Developing Economies.— An important implication of our analysis is that the effects of AI depend on its capabilities relative to human pre-AI occupations. Since these pre-AI occupations are endogenous—shaped by the knowledge of the human population and the level of communication technologies—the same AI technology can produce different outcomes in advanced and developing economies.

Indeed, in economies with superior communication technologies or a more knowledgeable population, the knowledge threshold required to become a solver in the pre-AI equilibrium is higher (see the Online Appendix for details). As a result, the same autonomous AI technology might be considered “basic” in an advanced economy—displacing humans from routine tasks to specialized problem-solving—while being regarded as “advanced” in a developing economy, thereby displacing humans in the opposite direction.^{32,33}

Autonomy vs. No Autonomy: Policy Implications.— The rise of AI and its potential to transform labor markets has captured policymakers’ attention, as evidenced by the recent report of [Council of Economic Advisers \(2024\)](#) and the [EU Artificial Intelligence Act \(2024\)](#). Our findings reveal key tradeoffs in regulating AI autonomy. While autonomous AI generates higher overall output than non-autonomous AI, it also exacerbates labor income inequality. This occurs because autonomous AI primarily benefits the most knowledgeable individuals, whereas non-autonomous AI disproportionately benefits the least knowledgeable, as shown in Proposition 7. Using this framework to quantify the tradeoff between output and inequality associated with autonomous and non-autonomous AI is an important direction for future research.

³²The reverse situation, however, is not possible: An autonomous AI that is advanced for an advanced economy must necessarily also be advanced for a developing economy.

³³This argument implicitly assumes the absence of cross-national teams. [Ide and Talamàs \(2024c\)](#) study the impact of AI on offshoring and the international division of knowledge work using the framework developed in this paper.

APPENDIX

A The Pre-AI Equilibrium: Complete Characterization

In this Appendix, we provide the full characterization of the equilibrium without AI. As noted in the main text, this equilibrium was first described in general by [Fuchs et al. \(2015\)](#). Proposition 1 follows directly from Lemmas A.1 and A.2 and Corollary A.1.

To start, note that the First Welfare Theorem holds in this setting. Thus, the competitive equilibrium is efficient. The following lemma shows that in the pre-AI world, there is a unique efficient allocation:

Lemma A.1. *In the absence of AI, there is a unique surplus maximizing allocation:*

- When $h \geq h_0 \in (0, 1)$, then $W = [0, \underline{z}]$, $I = (\underline{z}, \bar{z})$, and $S = [\bar{z}, 1]$. Moreover, workers and solvers are matched according to the strictly increasing function $m(z; \bar{z})$ given by $\int_{\bar{z}}^{m(z; \bar{z})} dG(u) = \int_0^z h(1-u)dG(u)$. Finally, the cutoffs $0 < \underline{z} < \bar{z} < 1$ satisfy:

$$(3) \quad \frac{1}{h} - \bar{z} = \int_{\bar{z}}^1 n(e(u; \bar{z}))du, \text{ and } m(\underline{z}; \bar{z}) = 1, \text{ where } e(z; \bar{z}) = m^{-1}(z; \bar{z})$$

- When $h < h_0$, then $W = [0, \hat{z}]$, $I = \emptyset$, and $S = [\hat{z}, 1]$. Moreover, workers and solvers are matched according to the strictly increasing function $m(z; \hat{z})$ given by $\int_{\hat{z}}^{m(z; \hat{z})} dG(u) = \int_0^z h(1-u)dG(u)$. Finally, the cutoff $\hat{z} \in (0, 1)$ satisfies $m(0; \hat{z}) = \hat{z}$, $m(\hat{z}; \hat{z}) = 1$, and:

$$(4) \quad \frac{1}{h} - \hat{z} > \int_{\hat{z}}^1 n(e(z; \hat{z}))dz, \text{ where } e(z; \hat{z}) = m^{-1}(z; \hat{z})$$

Proof. For the proof see [Fuchs et al. \(2015, Lemma 2\)](#). □

Note that the efficient (and, therefore, equilibrium) allocation satisfies occupational stratification and strict positive assortative matching. Moreover, $W \neq \emptyset$ and $S \neq \emptyset$ but $I \neq \emptyset$ if and only if $h > h_0 \in (0, 1)$. The next step is characterizing the wages that support the allocation of Lemma A.1 as a competitive equilibrium:

Lemma A.2. *In the absence of AI, the equilibrium wage function is given by:*

- When $h \geq h_0$:

$$w(z) = \begin{cases} m(z; \bar{z}) - w(m(z; \bar{z}))/n(z) & \text{if } z \in W \\ z & \text{if } z \in I \\ \bar{z} + \int_{\bar{z}}^z n(e(u; \bar{z}))du & \text{if } z \in S \end{cases}$$

- When $h < h_0$:

$$w(z) = \begin{cases} m(z; \hat{z}) - w(m(z; \hat{z}))/n(z) & \text{if } z \in W \\ \hat{z} + \frac{1}{1+n(\hat{z})} \left\{ \frac{1}{h} - \hat{z} - \int_{\hat{z}}^1 n(e(u; \hat{z}))du \right\} + \int_{\hat{z}}^z n(e(u; \hat{z}))du & \text{if } z \in S \end{cases}$$

Proof. See the Online Supplement of [Fuchs et al. \(2015, specifically pp. 1–4\)](#). $\square$

We end by showing that w is continuous, strictly increasing, and convex (strictly so when $z \in W \cup S$) and that $w(z) > z$ for all $z \notin \text{cl}I$.

Corollary A.1. *Irrespective of whether $h \geq h_0$, the equilibrium wage function $w(z)$ is continuous, strictly increasing, and (weakly) convex. Moreover, it satisfies:*

- $w(z) > z$ for $z \in W$ (except possibly at $z = \sup W$), and $w'(z) \in (0, 1)$, and $w''(z) > 0$ for $z \in \text{int}W$.
- $w(z) = z$ for all $z \in I$.
- $w(z) > z$ for $z \in S$ (except possibly at $z = \inf S$), and $w'(z) > 1$, and $w''(z) > 0$ for $z \in \text{int}S$.

Proof. Consider first $h \geq h_0$. To show continuity, it is sufficient to verify that $w(z)$ is continuous at $z = \underline{z}$ and $z = \bar{z}$. Continuity at $z = \underline{z}$ follows from the fact that $\lim_{z \uparrow \underline{z}} w(z) = 1 - w(1)h(1 - \underline{z}) = \underline{z} = \lim_{z \downarrow \underline{z}} w(z)$, where we are using the fact that $m(\underline{z}; \bar{z}) = 1$ and that $w(1) = \int_{\bar{z}}^1 n(e(u; \bar{z})) du + \bar{z} = \frac{1}{h}$ (due to condition (3)). Continuity at $z = \bar{z}$ follows from Lemma A.2 because it implies that $\lim_{z \downarrow \bar{z}} w(z) = w(\bar{z}) = \bar{z}$.

We now show the remaining properties of $w(z)$. Note that if $z \in \text{int}S$, then $w'(z) = n(e(z; \bar{z})) > 1$, which also implies that $w''(z) > 0$ because both $n(z)$ and $e(z; \bar{z})$ are strictly increasing in their arguments. Given that $\lim_{z \downarrow \bar{z}} w(z) = \bar{z}$, the previous results then imply that $w(z) > z$ for all $z \in (\bar{z}, 1]$. Similarly, if $z \in \text{int}W$, then $w'(z) = hw(m(z; \bar{z})) > 0$, which immediately implies that $w''(z) > 0$, since both $w(z)$ and $m(z; \bar{z})$ are strictly increasing in their arguments. Given that $\lim_{z \uparrow \underline{z}} w(z) = \underline{z}$, the previous results then imply that $w(z) > z$ for all $z \in [0, \underline{z})$. Finally, $w'(z) \in (0, 1)$ follows from the fact that:

$$w'(z) = hw(m(z; \bar{z})) = h \left[\frac{m(z; \bar{z}) - w(z)}{h(1 - z)} \right] < 1$$

where the second-to-last equality comes from the firms' zero-profit condition, and the last inequality because $m(z; \bar{z}) \leq 1$ and $w(z) > z$ when $z \in \text{int}W$.

Now consider $h \leq h_0$. To show continuity, it is sufficient to verify that $w(z)$ is continuous at $z = \hat{z}$. Given that $m(\hat{z}; \hat{z}) = 1$, from Lemma A.2, we have that:

$$\lim_{z \uparrow \hat{z}} w(z) = 1 - \frac{1}{n(\hat{z})}w(1) \quad \text{and} \quad \lim_{z \downarrow \hat{z}} w(z) = \frac{1}{1+n(\hat{z})} \left\{ n(\hat{z}) - \int_{\hat{z}}^1 n(e(u; \hat{z})) du \right\}$$

Note then that $w(1) = \lim_{z \downarrow \hat{z}} w(z) + \int_{\hat{z}}^1 n(e(u; \hat{z})) du$. Combining the latter with the expression for $\lim_{z \uparrow \hat{z}} w(z)$ above and rearranging terms yields $\lim_{z \uparrow \hat{z}} w(z) = \lim_{z \downarrow \hat{z}} w(z)$. The proof that $\lim_{z \downarrow \hat{z}} w(z) > \hat{z}$, in turn, can be found in [Fuchs et al. \(2015, Online Supplement, p. 4\)](#). Finally, the proofs for the remaining properties of $w(z)$ (i.e., their monotonicity and convexity, among others) follow the exact same logic as in the case when $h > h_0$. $\square$

B Equilibrium with Autonomous AI: Complete Characterization

In this Appendix, we provide a complete characterization of the equilibrium with autonomous AI. As noted in the main text, we focus on $h < h_0$. The proof of Proposition 2 is a direct implication of Lemmas B.1 and B.2 below.

We begin the characterization with the following set of results:

Lemma B.1. *Any equilibrium with AI has the following features:*

- *The equilibrium is efficient.*
- *Some compute must be allocated to independent production: $\mu_i^* > 0$.*
- *The price of compute is equal to AI's knowledge: $r^* = z_{\text{AI}}$.*
- *Occupational stratification: $W^* \preceq I^* \preceq S^*$.*
- *No worker is better than AI; no solver is worse than AI: $W^* \preceq \{z_{\text{AI}}\} \preceq S^*$.*
- *Positive assortative matching: $m^* : W_p^* \rightarrow S_p^*$ is strictly increasing and $W_a^* \preceq W_p^*$ and $S_p^* \preceq S_a^*$.*

Proof. • *The Equilibrium is Efficient.*— This result follows because the First Welfare Theorem holds in our setting (with perfect competition, complete information, and only pecuniary externalities).

• *Some compute must be allocated to independent production.*— This result follows because compute is abundant relative to human time. Hence, there are not enough humans to interact with AI inside two-layer organizations.

• *The price of compute is equal to AI's knowledge.*— This follows because the single-layer firms using AI must obtain zero profits.

• *Occupational stratification.*— Notice that the First Welfare Theorem holds in our setting. Hence, a competitive equilibrium must be efficient. Occupational stratification then follows because any surplus maximizing allocation must satisfy it. The proof of this last result is analogous to the proof of Lemma 1 in Fuchs et al. (2015).

• *No worker is better than AI; no solver is worse than AI.*— This result follows from occupational stratification and the fact that some compute must necessarily be used for independent production.

• *Positive assortative matching.*— The emergence of positive assortative matching—which follows from the supermodularity of the profits of two-layer organizations—is proven in Eeckhout and Kircher (2018, Proposition 1, p. 94) in a more general setting that nests our setting. Positive assortative matching then implies that the matching function is strictly increasing and that $W_a^* \preceq W_p^*$ and $S_p^* \preceq S_a^*$ (since no worker is better than AI and no solver is worse than AI). $\square$

The next corollary is a direct implication of Lemma B.1:

Corollary B.1. *An equilibrium allocation must take one of the following four potential configurations:*

• **Type 1 configuration:**

$$W_a^* = \emptyset, W_p^* = [0, z_{AI}], I^* = (z_{AI}, \bar{z}_1^*), S_p^* = [\bar{z}_1^*, \bar{z}_1^*], S_a^* = [\bar{z}_1^*, 1], \text{ where } z_{AI} < \bar{z}_1^* \leq \bar{z}_1^* \leq 1$$

$$\text{So } \mu_w^* = \int_{\bar{z}_1^*}^1 n(z_{AI}) dG(z), \mu_s^* = 0, \mu_i^* = \mu - \mu_w^*$$

• **Type 2 configuration:**

$$W_a^* = [0, \bar{z}_2^*], W_p^* = [\bar{z}_2^*, z_{AI}], I^* \subseteq \{z_{AI}\}, S_p^* = [z_{AI}, \bar{z}_2^*], S_a^* = [\bar{z}_2^*, 1], \text{ where } 0 \leq \bar{z}_2^* \leq z_{AI} \leq \bar{z}_2^* \leq 1$$

$$\text{So } \mu_w^* = \int_{\bar{z}_2^*}^1 n(z_{AI}) dG(z), \mu_s^* = \int_0^{\bar{z}_2^*} n(z)^{-1} dG(z), \mu_i^* = \mu - \mu_w^* - \mu_s^*$$

• **Type 3 configuration:**

$$W_a^* = [0, \bar{z}_3^*], W_p^* = [\bar{z}_3^*, \bar{z}_3^*], I^* = (\bar{z}_3^*, z_{AI}), S_p^* = [z_{AI}, 1], S_a^* = \emptyset, \text{ where } 0 < \bar{z}_3^* \leq \bar{z}_3^* < z_{AI}$$

$$\text{So } \mu_w^* = 0, \mu_s^* = \int_0^{\bar{z}_3^*} n(z)^{-1} dG(z), \mu_i^* = \mu - \mu_s^*$$

• **Type 4 configuration:**

$$W_a^* = \emptyset, W_p^* = [0, \bar{z}_4^*], I^* = (\bar{z}_4^*, \bar{z}_4^*) \ni z_{AI}, S_p^* = [\bar{z}_4^*, 1], S_a^* = \emptyset, \text{ where } 0 \leq \bar{z}_4^* < \bar{z}_4^* \leq 1$$

$$\text{So } \mu_w^* = 0, \mu_s^* = 0, \mu_i^* = \mu$$

Proof. As mentioned above, the proof of this corollary is a direct implication of Lemma B.1. Note that in a Type 2 configuration, I^* can be either $\{z_{AI}\}$ or $\emptyset$ because the human with knowledge z_{AI} is indifferent between any of the three roles. However, this is irrelevant for all practical purposes because I^* has measure zero. $\square$

Intuitively, in a Type 1 configuration, AI is used as a worker and independent producer. In a Type 2 configuration, AI is used in all three possible roles (i.e., as a worker, an independent producer, and a solver). In a Type 3 configuration, AI is used as a solver and independent producer, while in a Type 4 configuration, AI is used exclusively as an independent producer.

Now, recall that W and S are the sets of human workers and solvers in the pre-AI equilibrium. For $z_{AI} \in W$, define the function $m_w : [0, z_{AI}] \rightarrow [z_{AI}, 1]$ by $\int_{z_{AI}}^{m_w(z; z_{AI})} dG(u) = \int_0^z h(1-u) dG(u)$ and note that $z_{AI} \in W$ implies that $m_w(z_{AI}; z_{AI}) \leq 1$. Let then $e_w(z; z_{AI}) \equiv m_w^{-1}(z; z_{AI})$ and define:

$$\Gamma_w(x) \equiv n(x)(m_w(x; x) - x) - x - \int_x^{m_w(x; x)} n(e_w(u; x)) du$$

Similarly, for $z_{AI} \in S$, define the function $e_s : [z_{AI}, 1] \rightarrow [0, z_{AI}]$ by $\int_z^1 dG(u) = \int_{e_s(z; z_{AI})}^{z_{AI}} h(1-u) dG(u)$, note that $z_{AI} \in S$ implies that $e_s(z_{AI}; z_{AI}) \geq 0$, and define:

$$\Gamma_s(x) \equiv \frac{1}{h} - x - \int_x^1 n(e_s(u; x)) du$$

Consider the following partition of the knowledge space (note that $W \cup S = [0, 1]$ when $h < h_0$): $\mathcal{R}_1 \equiv \{z \in W : \Gamma_w(z) \leq 0\}$, $\mathcal{R}_2 \equiv \{z \in W : \Gamma_w(z) > 0\} \cup \{z \in S : \Gamma_s(z) > 0\}$, and $\mathcal{R}_3 \equiv \{z \in S : \Gamma_s(z) \leq 0\}$. The following lemma—which is the main result of this appendix—characterizes in detail the post-AI equilibrium:

Lemma B.2. *In the presence of AI, there is a unique competitive equilibrium. It is given as follows:*

- If $z_{AI} \in \mathcal{R}_1$, then the equilibrium allocation is Type 1. The equilibrium cutoffs $\underline{z}_1^*$ and $\bar{z}_1^*$ satisfy:

$$\bar{z}_1^* = m_1^*(z_{AI}; \underline{z}_1^*) \quad \text{and} \quad n(z_{AI})(m_1^*(z_{AI}; \underline{z}_1^*) - z_{AI}) = \underline{z}_1^* + \int_{\underline{z}_1^*}^{m_1^*(z_{AI}; \underline{z}_1^*)} n(e_1^*(z; \underline{z}_1^*)) dz$$

where $m_1^* : [0, z_{AI}] \rightarrow [\underline{z}_1^*, \bar{z}_1^*]$ is given by $\int_{\underline{z}_1^*}^{m_1^*(z; \underline{z}_1^*)} dG(u) = \int_0^z h(1-u)dG(u)$ and $e_1^*(z; \cdot) = (m_1^*)^{-1}(z; \cdot)$

- If $z_{AI} \in \mathcal{R}_2$, then the equilibrium allocation is Type 2. The equilibrium cutoffs $\underline{z}_2^*$ and $\bar{z}_2^*$ satisfy:

$$\bar{z}_2^* = m_2^*(z_{AI}; \underline{z}_2^*) \quad \text{and} \quad n(z_{AI})(m_2^*(z_{AI}; \underline{z}_2^*) - z_{AI}) = z_{AI} + \int_{z_{AI}}^{m_2^*(z_{AI}; \underline{z}_2^*)} n(e_2^*(z; \underline{z}_2^*)) dz$$

where $m_2^* : [\underline{z}_2^*, z_{AI}] \rightarrow [z_{AI}, \bar{z}_2^*]$ is given by $\int_{z_{AI}}^{m_2^*(z; \underline{z}_2^*)} dG(u) = \int_{\underline{z}_2^*}^z h(1-u)dG(u)$ and $e_2^*(z; \cdot) = (m_2^*)^{-1}(z; \cdot)$

- If $z_{AI} \in \mathcal{R}_3$, then the equilibrium allocation is Type 3. The equilibrium cutoffs $\underline{z}_3^*$ and $\bar{z}_3^*$ satisfy:

$$\bar{z}_3^* = e_3^*(1; \underline{z}_3^*) \quad \text{and} \quad \frac{1}{h} = z_{AI} + \int_{z_{AI}}^1 n(e_3^*(z; \underline{z}_3^*)) dz$$

where $m_3^* : [\underline{z}_3^*, \bar{z}_3^*] \rightarrow [z_{AI}, 1]$ given by $\int_{z_{AI}}^{m_3^*(z; \underline{z}_3^*)} dG(u) = \int_{\underline{z}_3^*}^z h(1-u)dG(u)$ and $e_3^*(z; \cdot) = (m_3^*)^{-1}(z; \cdot)$

The equilibrium matching function is given by $m^*(z) = m_j^*(z; \underline{z}_j^*)$ if $z_{AI} \in \mathcal{R}_j$, while the equilibrium wage w^* is continuous, strictly increasing, and (weakly) convex, and satisfies:

- $w^*(z) = z_{AI}(1 - 1/n(z)) > z$ for all $z \in W_a^*$.
- $w^*(z) = m^*(z) - w^*(m^*(z))/n(z)$ for all $z \in W_p^*$.
- $w^*(z) = z$ for all $z \in I^*$.
- $w^*(z) = \inf S_p^* + \int_{\inf S_p^*}^z n(e^*(u))du$ for all $z \in S_p^*$.
- $w^*(z) = n(z_{AI})(z - z_{AI}) > z$ for all $z \in S_a^*$.

For the interested reader, in the Online Appendix, we also provide a version of Lemma B.2 for the case in which human knowledge is uniformly distributed. In such a case, the equilibrium expressions (i.e., the wage function w^* , the equilibrium cutoffs, and the partition $\mathcal{R}_j$ for $j = 1, 2, 3$) can be obtained in (almost) closed form.

According to Lemma B.1, AI is always used for independent production. Lemma B.2 adds to this point by stating that if $z_{AI} \in W$, then AI is also used as a worker and possibly as a solver, while if $z_{AI} \in S$, then AI is also used as a solver and possibly as a worker. Indeed, if $z_{AI} \in W$, then the post-AI equilibrium is either Type 1 (in which AI is a worker but not a solver) or Type 2 (in which AI is both a worker and solver). Similarly, if $z_{AI} \in S$, then the post-AI equilibrium is either Type 3 (in which AI is a solver but not a worker) or Type 2.

Before formally proving this lemma, we informally derive the equilibrium in one of the regions to provide insight into its construction:

Informal Construction of the Equilibrium.— Suppose that $z_{AI} \in \mathcal{R}_2$. By Corollary B.1, we know that such an equilibrium must lead to the following partition of the human population:

$$W_a^* = [0, \underline{z}_2^*], \quad W_p^* = [\underline{z}_2^*, z_{AI}], \quad I^* = \emptyset, \quad S_p^* = [z_{AI}, \bar{z}_2^*], \quad S_a^* = [\bar{z}_2^*, 1], \quad \text{where } 0 \leq \underline{z}_2^* \leq z_{AI} \leq \bar{z}_2^* \leq 1$$

As mentioned in the main text, given that the equilibrium price of compute is $r^* = z_{AI}$, the zero-profit condition of a tA firm pins down the wage $w^*(z) = z_{AI}(1 - 1/n(z))$ of a human worker with knowledge $z \in W_a^*$. Similarly, the zero-profit condition of a bA firm determines the wage $w^*(z) = n(z_{AI})(z - z_{AI})$ of a human solver with knowledge $z \in S_a^*$.

Now consider those firms that do not use AI, i.e., the nA firms. First, let $m_2^*(z)$ be the equilibrium matching function in this case. As mentioned in Section 3.1, this function must be strictly increasing and satisfy $\int_{z_{AI}}^{m_2^*(z)} dG(u) = \int_{z_2^*}^z h(1-u)dG(u)$ for all $z \in [z_2^*, z_{AI}]$ (which is simply constraint (1) in this specific case). Moreover, given that $\sup W_p^* = z_{AI}$ and $\sup S_p^* = \bar{z}_2^*$, it must also be that $m_2^*(z_{AI}) = \bar{z}_2^*$.

Note then that for any given $z \in W_p$, there exists a unique increasing function $m_2^*(z)$ that satisfies both constraints. It is given by the solution to the differential equation $m_2^{*'}(z) = h(1-z)g(z)/g(m_2^*(z))$ with border condition $m_2^*(z_{AI}) = \bar{z}_2^*$ (which comes from differentiating both sides of the resource constraint). We denote such a unique function by $m_2^*(z; \bar{z}_2^*)$ (as it depends on $\bar{z}_2^*$ through its domain).

Consider then the problem of an nA firm that recruited $n(z)$ workers of type $z \in W_p^*$ and is deciding which solver $z \in S_p^*$ to hire: $\max_{s \in S_p^*} \Pi_2^{nA}(s, z) = n(z)[s - w(z)] - w(s)$. The corresponding first-order condition evaluated at $s = m_2^*(z; \bar{z}_2^*)$ implies that $w'(z) = n(e_2^*(z; \bar{z}_2^*))$ for any $z \in S_p^*$. Thus, $w^*(z) = C^* + \int_{z_{AI}}^z n(e_2^*(u; \bar{z}_2^*))du$ for any $z \in S_p^*$. The wages of the workers hired by such firms are then determined by the zero-profit condition of nA firms: $w^*(z) = m_2^*(z; \bar{z}_2^*) - w^*(m_2^*(z; \bar{z}_2^*)/n(z))$ for any $z \in W_p^*$.

The final step is determining the constant C^* , the cutoff $\bar{z}_2^*$, and arguing that no firms have incentives to deviate. To do this, note that the least knowledgeable human solver has the same knowledge as AI, i.e., $\inf S_p^* = z_{AI}$. Hence, her wage must be equal to the price of one unit of compute, so $C^* = z_{AI}$. Moreover, the most knowledgeable solver hired by an nA firm has the same knowledge as the least knowledgeable solver of a bA firm. As a result, these two individuals must receive the same wage, i.e., $\lim_{z \uparrow \bar{z}_2^*} w^*(z) = \lim_{z \downarrow \bar{z}_2^*} w^*(z)$. Since $\bar{z}_2^* = m_2^*(z_{AI}; \bar{z}_2^*)$, we obtain the following:

$$(5) \quad n(z_{AI})(m_2^*(z_{AI}; \bar{z}_2^*) - z_{AI}) = z_{AI} + \int_{z_{AI}}^{m_2^*(z_{AI}; \bar{z}_2^*)} n(e_2^*(z; \bar{z}_2^*))dz$$

which is the equilibrium condition in the statement of Lemma B.2. It is then possible to prove that there is a unique cutoff $\bar{z}_2^*$ that satisfies this condition and that such a cutoff is contained in $(0, z_{AI}]$ if and only if $z_{AI} \in \mathcal{R}_2$. This explains why this equilibrium can only arise in such a region of the parameter space. The fact that $C^* = z_{AI}$ and that $\bar{z}_2^*$ satisfies (5) then implies that $w^*(z)$ is also continuous at the juncture between W_p^* and S_p^* :

$$\begin{aligned} \lim_{z \downarrow \bar{z}_2^*} w^*(z) &= m_2^*(\bar{z}_2^*; \bar{z}_2^*) - \frac{w^*(m_2^*(\bar{z}_2^*; \bar{z}_2^*))}{n(\bar{z}_2^*)} = z_{AI} \left(1 - \frac{1}{n(\bar{z}_2^*)}\right) = \lim_{z \uparrow \bar{z}_2^*} w^*(z) \\ \lim_{z \uparrow z_{AI}} w^*(z) &= m_2^*(z_{AI}; \bar{z}_2^*) - \frac{w^*(m_2^*(z_{AI}; \bar{z}_2^*))}{n(z_{AI})} = z_{AI} = \lim_{z \downarrow z_{AI}} w^*(z) \end{aligned}$$

This is sufficient to guarantee that $w^*(z)$ is continuous in all its domain, i.e., for all $z \in [0, 1]$.

From here, arguing that no firm has incentives to deviate is straightforward. Indeed, note that the wage function is continuous, strictly increasing, and weakly convex. This implies that if a firm does

not have incentives to deviate “locally,” then it does not have incentives to deviate globally either. That bA firms do not want to deviate locally, i.e., hire a different human solver in S_p^* , follows because such deviation also leads to no profits. Similar reasoning also explains why tA and nA firms do not have incentives to deviate locally either. $\square$

We now formally prove the lemma. This is done in two steps. In the first step, we show that the outcomes described in the lemma are indeed an equilibrium by verifying that (i) markets clear, and (ii) firms are maximizing their profits while obtaining zero profits. In the second step, we prove that the equilibrium is unique.

Proof of Step 1. We show this part only for $z_{AI} \in \mathcal{R}_1$, as the other two cases are analogous. We begin by verifying market clearing in the market for compute. By Corollary B.1, it is immediate that $\mu_i^* + \mu_w^* + \mu_s^* = \mu$. Moreover, the total time required to consult on the problems left unsolved by AI is equal to the total time available of solvers in S_a^* , i.e., $h(1 - z_{AI})\mu_w^* = \int_{\underline{z}_1^*}^1 dG(z)$.

We now turn to verifying labor market clearing. First, it must be that the total time required to consult on the problems left unsolved by the human workers in the interval $[0, z] \subseteq W_p^*$ is equal to the total time available of human solvers in the interval $[\underline{z}_1^*, m_1^*(z; \underline{z}_1^*)] \subseteq S_p^*$. This resource constraint is satisfied as $m_1^*(z; \underline{z}_1^*)$ is given by $\int_{\underline{z}_1^*}^{m_1^*(z; \underline{z}_1^*)} dG(u) = \int_0^z h(1 - u)dG(u)$ and $\bar{z}_1^* = m_1^*(z_{AI}; \underline{z}_1^*)$.

Second, it must be that the union of the sets $(W_p^*, W_a^*, I^*, S_p^*, S_a^*)$ is $[0, 1]$, and the intersection of any two of these sets has measure zero. By Corollary B.1, this occurs if and only if $z_{AI} < \underline{z}_1^* \leq \bar{z}_1^* \leq 1$. Verifying that that $\underline{z}_1^* \leq \bar{z}_1^*$ is straightforward: It follows because $m_1^*(z; \underline{z}_1^*)$ is strictly increasing in z plus the fact that $\underline{z}_1^* = m_1^*(0; \underline{z}_1^*)$ and $\bar{z}_1^* = m_1^*(z_{AI}; \underline{z}_1^*)$.

Showing that $z_{AI} < \underline{z}_1^*$ requires more work. Note that $\underline{z}_1^*$ is given by the solution $\Gamma_1(\underline{z}_1^*; z_{AI}) = 0$, where $\Gamma_1(x; z_{AI}) \equiv n(z_{AI})(m_1^*(z_{AI}; x, z_{AI}) - z_{AI}) - x - \int_x^{m_1^*(z_{AI}; x, z_{AI})} n(e_1^*(z; x, z_{AI}))dz$ (here we are making explicit that $m_1^*(\cdot)$ and $e_1^*(\cdot)$ also depend indirectly on z_{AI} through the boundary of the set $W_p^* = [0, z_{AI}]$ to avoid any type of confusion³⁴). It is then not difficult to prove that $\Gamma_1(x; z_{AI})$ is strictly increasing in x , so $z_{AI} < \underline{z}_1^*$ if and only if $\Gamma_1(0; z_{AI}) < 0$. Furthermore, note that $m_1^*(z; 0, z_{AI})$ satisfies $\int_0^{m_1^*(z; 0, z_{AI})} dG(u) = \int_0^z h(1 - u)dG(u)$ for $z \in [0, z_{AI}]$, which implies that $m_1^*(z; 0, z_{AI}) = m_w(z; z_{AI})$ (and, therefore, that $e_1^*(z; 0, z_{AI}) = e_w(z; z_{AI})$). Hence, $\Gamma_1(0; z_{AI}) = \Gamma_w(z_{AI}) < 0$ as $z_{AI} \in \mathcal{R}_1$.

Finally, we show that $\bar{z}_1^* \leq 1$. To prove it, we show that $\underline{z}_1^* < \hat{z}$ and then use this result to conclude that $\bar{z}_1^* \leq 1$. As a first step, note that $m_1^*(z; \hat{z}, \hat{z})$ satisfies $\int_{\hat{z}}^{m_1^*(z; \hat{z}, \hat{z})} dG(u) = \int_0^z h(1 - u)dG(u)$ for $z \in [0, \hat{z}]$, so $m_1^*(z; \hat{z}, \hat{z}) = m(z; \hat{z})$ where $m(z; \hat{z})$ is the pre-AI matching function. Recall then that $\underline{z}_1^*$ is the unique solution $\Gamma_1(\underline{z}_1^*; z_{AI}) = 0$, where $\Gamma_1(x; z_{AI})$ is strictly increasing in x . It is not difficult to prove that $\Gamma_1(x; z_{AI})$ is also strictly decreasing in z_{AI} for any given x . We then claim that $\Gamma_1(\hat{z}; z_{AI}) > 0$, which immediately implies that $\underline{z}_1^* < \hat{z}$. Indeed, note that $\Gamma_1(\hat{z}; z_{AI}) \geq \Gamma_1(\hat{z}; \hat{z}) = \frac{1}{h} - \hat{z} - \int_{\hat{z}}^1 n(e(z; \hat{z}))dz$, where the first inequality follows because $\Gamma_1(x; z_{AI})$ is strictly decreasing in

³⁴In the statement of Lemma B.2, we simply wrote $m_1^*(z; \underline{z}_1^*)$ instead of $m_1^*(z; \underline{z}_1^*, z_{AI})$ to avoid cluttering notation.

z_{AI} and $z_{AI} \leq \hat{z}$ (as $z_{AI} \in W$) and the last equality follows because $e_1^*(z; \hat{z}, \hat{z}) = e(z; \hat{z})$ for all $z \in [\hat{z}, 1]$. However, by Lemma A.1, we know that $\frac{1}{h} - \hat{z} - \int_{\hat{z}}^1 n(e(z; \hat{z}))dz > 0$ when $h < h_0$, so $\Gamma_1(\hat{z}; z_{AI}) > 0$.

Having proven that $\underline{z}_1^* < \hat{z}$, we now show that $\bar{z}_1^* \leq 1$. By construction, $\bar{z}_1^*$ satisfies $\int_{\bar{z}_1^*}^{\bar{z}_1^*} dG(u) = \int_0^{z_{AI}} h(1-u)dG(u)$. Note, moreover, that $z \in W$ implies that $\int_0^{z_{AI}} h(1-u)dG(u) \leq \int_0^{\hat{z}} h(1-u)dG(u) = \int_{\hat{z}}^1 dG(u)$ (as $z_{AI} \leq \hat{z}$ and $\int_0^{\hat{z}} h(1-u)dG(u) = \int_{\hat{z}}^1 dG(u)$). Hence, $\int_{\bar{z}_1^*}^{\bar{z}_1^*} dG(u) \leq \int_{\hat{z}}^1 dG(u)$. Given that $\underline{z}_1^* < \hat{z}$, it must be that $\bar{z}_1^* < 1$.

Having verified market clearing, we now show that, in the candidate equilibrium, firms maximize their profits while obtaining zero profits. Given how the wages are constructed, it is clear that firms are optimizing their profits “locally” while obtaining zero profits. Thus, we only need to consider “global deviations.” As discussed above, to discard such global deviations, it is sufficient to show that w^* is continuous, strictly increasing, and weakly convex. This is what we prove next.

To show continuity, it suffices to verify that w^* is continuous at the junctures of (i) S_p^* and S_a^* , (ii) I^* and S_p^* , and (iii) W_p^* and I^* . For (i), note that $\lim_{z \uparrow \bar{z}_1^*} w^*(z) = \bar{z}_1^* + \int_{\bar{z}_1^*}^{\bar{z}_1^*} n(e_1^*(z; \underline{z}_1^*))dz$ and $\lim_{z \downarrow \bar{z}_1^*} w^*(z) = n(z_{AI})(\bar{z}_1^* - z_{AI})$. Hence, from the conditions determining $\underline{z}_1^*$ and $\bar{z}_1^*$, we obtain that $\lim_{z \uparrow \bar{z}_1^*} w^*(z) = \lim_{z \downarrow \bar{z}_1^*} w^*(z)$. For (ii) note that by construction, $\lim_{z \uparrow \underline{z}_1^*} w^*(z) = \underline{z}_1^* = \lim_{z \downarrow \underline{z}_1^*} w^*(z)$. Finally, for (iii) note that $\lim_{z \uparrow z_{AI}} w^*(z) = \bar{z}_1^* - w^*(\bar{z}_1^*)/n(z_{AI}) = z_{AI} = \lim_{z \downarrow z_{AI}} w^*(z)$ given that $w^*(\bar{z}_1^*) = n(z_{AI})(\bar{z}_1^* - z_{AI})$.

With continuity at hand, proving that w^* is strictly increasing and weakly convex is straightforward: The logic is analogous to the proof of Corollary A.1 (which shows that the pre-AI wage function satisfies these two properties). Consequently, in the case $z_{AI} \in \mathcal{R}_1$, the outcome described in the statement is indeed a competitive equilibrium. $\square$

Proof of Step 2. We show this part only for $z_{AI} \in \mathcal{R}_1$, as the other two cases follow the same logic. In particular, we show that if $z_{AI} \in \mathcal{R}_1$, then there cannot be any other type of equilibrium.

Suppose first by contradiction that there is a Type 3 equilibrium. By Corollary B.1, we know that such an equilibrium must lead to the following partition of the human population:

$$W_a^* = [0, \underline{z}_3^*], W_p^* = [\underline{z}_3^*, \bar{z}_3^*], I^* = (\bar{z}_3^*, z_{AI}), S_p^* = [z_{AI}, 1], S_a^* = \emptyset, \text{ where } 0 < \underline{z}_3^* \leq \bar{z}_3^* < z_{AI}$$

By Lemma B.3, we have that if $z_{AI} \in \mathcal{R}_1$, then $z_{AI} < \hat{z}$. This implies the following:

$$\int_{\underline{z}_3^*}^{\bar{z}_3^*} h(1-u)dG(u) < \int_0^{z_{AI}} h(1-u)dG(u) < \int_0^{\hat{z}} h(1-u)dG(u) = \int_{\hat{z}}^1 dG(u) < \int_{z_{AI}}^1 dG(u)$$

which violates the resource constraint that the total time required to consult on the problems left unsolved by the human workers in the interval W_p^* must equal to the total time available of human solvers in the interval S_p^* , i.e., $\int_{\underline{z}_3^*}^{\bar{z}_3^*} h(1-u)dG(u) = \int_{z_{AI}}^1 dG(u)$. Hence, a Type 3 configuration cannot arise when $z_{AI} \in \mathcal{R}_1 \subset W$.

Now suppose for contradiction that there is a Type 2 equilibrium. As explained above (see “Informal Construction of the Equilibrium”), for this to be an equilibrium, there must exist a cutoff

$\underline{z}_2^* \in [0, z_{AI}]$ with the property that $\Gamma_2(\underline{z}_2^*; z_{AI}) = 0$, where $\Gamma_2(x; z_{AI}) \equiv n(z_{AI})(m_2^*(z_{AI}; x) - z_{AI}) - z_{AI} - \int_{z_{AI}}^{m_2^*(z_{AI}; x)} n(e_2^*(z; x))dz$. It is then not difficult to prove that $\Gamma_2(x; z_{AI})$ is strictly decreasing in x . Hence, there exists at most one $\underline{z}_2^*$ that satisfies $\Gamma_2(\underline{z}_2^*; z_{AI}) = 0$, and for $\underline{z}_2^* \geq 0$, it must be that $\Gamma_2(0; z_{AI}) \geq 0$, where $m_2^*(z; 0)$ is given by $\int_{z_{AI}}^{m_2^*(z; 0)} dG(u) = \int_0^z h(1-u)dG(u)$ for $z \in [0, z_{AI}]$. However, from this last condition we have that $m_2^*(z; 0) = m_w(z; z_{AI})$ (as $m_2^*(z; 0)$ and $m_w(z; z_{AI})$ satisfy the same condition), so $\Gamma_2(0; z_{AI}) = \Gamma_w(z_{AI})$. Consequently, for $\underline{z}_2^* \geq 0$, we need that $\Gamma_w(z_{AI}) \geq 0$, which contradicts the fact that $z_{AI} \in \mathcal{R}_1$.

Finally, suppose for contradiction that there is a Type 4 equilibrium. By Corollary B.1, we know that such an equilibrium must lead to the following partition of the human population:

$$W_p^* = \emptyset, W_p^* = [0, \underline{z}_4^*], I^* = (\underline{z}_4^*, \bar{z}_4^*) \ni z_{AI}, S_p^* = [\bar{z}_4^*, 1], S_a^* = \emptyset$$

Moreover, following similar reasoning as that developed above for a Type 2 equilibrium, for this to be an equilibrium, (i) it must be that $\bar{z}_4^* = m_4^*(0; \underline{z}_4^*)$, where $m_4^*(0; \underline{z}_4^*)$ satisfies $m_4^*(\underline{z}_4^*; \underline{z}_4^*) = 1$ and $\int_{m_4^*(0; \underline{z}_4^*)}^{m_4^*(z; \underline{z}_4^*)} dG(u) = \int_{\underline{z}_4^*}^z h(1-u)dG(u)$ for $z \in [0, \underline{z}_4^*]$, and (ii) there must exist a cutoff $\underline{z}_4^* < m_4^*(0; \underline{z}_4^*)$ such that $1/h - m_4^*(0; \underline{z}_4^*) - \int_{m_4^*(0; \underline{z}_4^*)}^1 n(e_4^*(z; \underline{z}_4^*))dz = 0$. However, by Lemma A.1, we know that the unique solution to these equilibrium conditions is $\underline{z}_4^* = \underline{z}$ and $\bar{z}_4^* = m_4^*(0; \underline{z}_4^*) = \bar{z}$, where $\underline{z}$ and $\bar{z}$ are the equilibrium cutoffs of the pre-AI equilibrium when $h > h_0$. This, however, implies that this configuration is an equilibrium only if $z_{AI} \in I^* = (\underline{z}, \bar{z})$, which contradicts the assumption that $z_{AI} \in \mathcal{R}_1$. $\square$

We end this appendix with some properties of the partition $(\mathcal{R}_1, \mathcal{R}_2, \mathcal{R}_3)$ that will be useful later. In particular, the next lemma states that when z_{AI} is sufficiently low, the equilibrium is Type 1, so AI is used as a worker and as an independent producer (but not as a solver) in that case. Similarly, when $z_{AI} = \hat{z}$ or $z_{AI} \rightarrow 1$ (where $\hat{z}$ is the knowledge cutoff to become a solver in the pre-AI equilibrium), the equilibrium is Type 2, so AI is used in all three possible roles.

Lemma B.3. (i) $z_{AI} \in \mathcal{R}_1$ if $z_{AI} \in [0, \epsilon)$ with $\epsilon \downarrow 0$, (ii) $z_{AI} = \hat{z} \in \mathcal{R}_2$ (where $\{\hat{z}\} = W \cap S$), and (iii) $z_{AI} \in \mathcal{R}_2$ if $z_{AI} \in [1 - \epsilon, 1)$ with $\epsilon \downarrow 0$.

Proof. First we show that $z_{AI} \in \mathcal{R}_1$ if $z_{AI} \in [0, \epsilon)$ with $\epsilon \downarrow 0$. Note that $\Gamma_w(0) = 0$ (as $m_w(z; 0) = 0$), and that $\Gamma'_w(0) = -1$, as:

$$\Gamma'_w(x) = \frac{1}{h} - 1 - \frac{1}{h(1-x)} + \frac{m_w(x; x) - x}{h(1-x)^2} + h \int_x^{m_w(x; x)} \frac{n(e_w(z; x))^3 g(x)}{g(e_w(z; x))} dz$$

Hence, if $z_{AI} \in [0, \epsilon)$ with $\epsilon \downarrow 0$, then $z_{AI} \in W$ and $\Gamma_w(z_{AI}) \leq 0$, so $z_{AI} \in \mathcal{R}_1$.

Next, we prove that $\hat{z} \in \mathcal{R}_2$ by showing that $\Gamma_s(\hat{z}) = \Gamma_w(\hat{z}) > 0$. Note that $m_w(z; \hat{z}) = m(z; \hat{z})$, where $m(z; \hat{z})$ is the matching function of the pre-AI equilibrium. This implies that $m_w(\hat{z}; \hat{z}) = 1$, so $\Gamma_w(\hat{z}) = (1/h) - \hat{z} - \int_{\hat{z}}^1 n(e(z; \hat{z}))dz > 0$, where the last inequality follows from Lemma A.1. Similarly, note that $e_s(z; \hat{z}) = e(z; \hat{z})$, where $e(z; \hat{z}) = m^{-1}(z; \hat{z})$. Thus, $\Gamma_s(\hat{z}) = (1/h) - \hat{z} - \int_{\hat{z}}^1 n(e(z; \hat{z}))dz$, so $\Gamma_s(\hat{z}) = \Gamma_w(\hat{z}) > 0$.

Finally, we show that $z_{AI} \in \mathcal{R}_2$ if $z_{AI} \in [1 - \epsilon, 1)$ with $\epsilon \downarrow 0$. First, note that if $z_{AI} \in [1 - \epsilon, 1)$ with $\epsilon \downarrow 0$, then $z_{AI} \in S$. Second, it is not difficult to prove that $\Gamma_s(1) = 1/h - 1 > 0$, so by continuity, $\Gamma_s(z_{AI}) > 0$ for all $z_{AI} \in [1 - \epsilon, 1)$ with $\epsilon \downarrow 0$. Hence, if $z_{AI} \in [1 - \epsilon, 1)$ with $\epsilon \downarrow 0$, then $z_{AI} \in S$ and $\Gamma_s(z_{AI}) > 0$, so $z_{AI} \in \mathcal{R}_2$. $\square$

C Proofs Omitted from Section 5

C.1 Proof of Proposition 3

• $z_{AI} \in \text{int}W$.— By Lemma B.2, the AI equilibrium is either Type 1 or Type 2. If it is Type 2, then $\sup W^* = \inf S^* = z_{AI}$, so $W^* \subset W$ and $S^* \supset S$, since $z_{AI} < \hat{z}$ when $z_{AI} \in \text{int}W$. If it is Type 1, then $\sup W^* = z_{AI}$ and $\inf S^* = \underline{z}_1^*$. That $\sup W^* = z_{AI}$ implies that $W^* \subset W$ because $z_{AI} < \hat{z}$. That $\inf S^* = \underline{z}_1^*$ implies that $S^* \supset S$ given that $\underline{z}_1^* < \hat{z}$ (as shown in the proof of Lemma B.2). $\square$

• $z_{AI} \in \text{int}S$.— By Lemma B.2, the AI equilibrium is either Type 2 or Type 3. Suppose first that it is Type 2. Then $\sup W^* = \inf S^* = z_{AI}$, so $W^* \supset W$ and $S^* \subset S$ given that $z_{AI} > \hat{z}$ when $z_{AI} \in \text{int}S$.

Now, suppose the equilibrium is Type 3. Then, $\inf S^* = z_{AI} > \hat{z}$, immediately implying that $S^* \subset S$. To prove that $W^* \supset W$, we show that $\sup W^* = \bar{z}_3^* > \hat{z}$. Indeed, recall that $\bar{z}_3^*$ is given by:³⁵

$$\bar{z}_3^* = e_3^*(1; \underline{z}_3^*, z_{AI}) \quad \text{and} \quad \frac{1}{h} = z_{AI} + \int_{z_{AI}}^1 n(e_3^*(z; \underline{z}_3^*, z_{AI})) dz$$

where $\int_{z_{AI}}^{m_3^*(z; \underline{z}_3^*, z_{AI})} dG(u) = \int_{\underline{z}_3^*}^z h(1-u)dG(u)$ for $z \in [\underline{z}_3^*, \bar{z}_3^*]$

Using the fact that $m_3^*(\bar{z}_3^*; \underline{z}_3^*, z_{AI}) = 1$, the equilibrium conditions that determine $\underline{z}_3^*$ and $\bar{z}_3^*$ can be written as follows:³⁶

$$\underline{z}_3^* = \tilde{e}_3^*(z_{AI}; \bar{z}_3^*, z_{AI}) \quad \text{and} \quad \frac{1}{h} = z_{AI} + \int_{z_{AI}}^1 n(\tilde{e}_3^*(z; \bar{z}_3^*, z_{AI})) dz$$

where $G(z) = 1 - \int_{\tilde{e}_3^*(z; \bar{z}_3^*, z_{AI})}^{\bar{z}_3^*} h(1-u)dG(u)$ for $z \in [z_{AI}, 1]$

Define $\Gamma_3(x; z_{AI}) \equiv \frac{1}{h} - z_{AI} - \int_{z_{AI}}^1 n(\tilde{e}_3^*(z; x, z_{AI})) dz$. It is not difficult to prove that $\Gamma_3(x; z_{AI})$ is strictly decreasing in x and strictly increasing in z_{AI} . Moreover, $\bar{z}_3^*$ is given by the unique solution to $\Gamma_3(\bar{z}_3^*; z_{AI}) = 0$. To prove that $\bar{z}_3^* > \hat{z}$, it suffices to show that $\Gamma_3(\hat{z}; z_{AI}) > 0$. Since $z_{AI} > \hat{z}$, we have that $\Gamma_3(\hat{z}; z_{AI}) > \Gamma_3(\hat{z}; \hat{z}) = \frac{1}{h} - \hat{z} - \int_{z_{AI}}^1 n(e(z; \hat{z})) dz > 0$, where the second-to-last inequality follows because $\tilde{e}_3^*(z; \hat{z}, \hat{z}) = e(z; \hat{z})$ for all $z \in [\hat{z}, 1]$, and the last inequality comes from the pre-AI equilibrium characterized in Lemma A.1. $\square$

³⁵To avoid any type of confusion, we are making explicit that $m_3^*(z; \underline{z}_3^*, z_{AI})$ and $e_3^*(1; \underline{z}_3^*, z_{AI})$ depend on both $\underline{z}_3^*$ and z_{AI} (in the statement of Lemma B.2, we simply wrote $m_3^*(z; \underline{z}_3^*)$ instead of $m_3^*(z; \underline{z}_3^*, z_{AI})$ to avoid cluttering notation).

³⁶Note that $\tilde{e}_3^*(z; \bar{z}_3^*, z_{AI})$ is the equilibrium employee function indexed by $\bar{z}_3^*$ instead of $\underline{z}_3^*$.

C.2 Proof of Corollary 1

Since all two-layer organizations hire a single solver, to prove this corollary, it is sufficient to show that AI increases the overall number of solvers in the economy. When $z_{\text{AI}} \in \text{int}W$, this follows because $S^* \supset S$. When $z_{\text{AI}} \in \text{int}S$, more humans become workers after AI's introduction (i.e., $W^* \supset W$). Hence, the overall number of solvers—human plus AI—must increase as each worker requires the same amount of help post-AI as that required pre-AI. $\square$

C.3 Proof of Proposition 4

Consider first span of control. The most decentralized firm of the pre-AI economy has a span of control $n(\sup W)$, while the most decentralized firm of the post-AI world has a span of control of $n(\sup W^*)$. That AI decreases the maximum span of control when $z_{\text{AI}} \in \text{int}W$, and increases it when $z_{\text{AI}} \in \text{int}S$, follows from occupational displacement, i.e. Proposition 3, and the fact that $n(z) = [h \times (1 - z)]^{-1}$ is strictly increasing in z .

Consider next productivity. The least productive firm of the pre-AI economy has productivity $\inf S$, while the least productive firm of the post-AI world has productivity $\inf S^*$. That AI decreases the minimum firm productivity when $z_{\text{AI}} \in \text{int}W$, and increases it when $z_{\text{AI}} \in \text{int}S$, follows directly from occupational displacement (i.e. Proposition 3).

Finally, consider size. As noted in the main text, the size of the smallest firm is $n(0)$ times the minimum firm productivity, while the size of the largest firm is 1 times the maximum span of control. Consequently, when $z_{\text{AI}} \in \text{int}W$, AI reduces the minimum and maximum firm size, as it decreases the maximum span of control and the minimum firm productivity. Likewise, when $z_{\text{AI}} \in \text{int}S$, AI increases the minimum and maximum firm size, as it increases the maximum span of control and the minimum firm productivity. $\square$

C.4 Proof of Corollary 2

First, note that when $z_{\text{AI}} \in \text{int}S$, the largest post-AI firms are larger than the largest pre-AI firms.

Second, we know by Lemma B.3 that $z_{\text{AI}} \in \mathcal{R}_2$ if $z_{\text{AI}} \in [1 - \epsilon, 1) \subset S$ with $\epsilon \downarrow 0$. Thus, by Lemma B.2, firms use AI in all three possible roles when z_{AI} is sufficiently close to 1. Since no worker is better than AI and the equilibrium exhibits positive assortative matching (Lemma B.1), we then have that AI is the best worker in the economy and is therefore supervised by the most knowledgeable humans. Hence, when z_{AI} is sufficiently close to 1, the largest firms post-AI are bottom-automated firms.

Combining these two results, we have that when z_{AI} is sufficiently close to 1, AI leads to the creation of superstar firms with scale but no mass. $\square$

C.5 Proof of Proposition 5

As noted in the main text, a worker's productivity increases if and only if her solver match improves. Similarly, a given solver's span of control increases if and only if her workers' knowledge increases.

Moreover, for the proof that follows, it is important that we make explicit that the pre-AI matching and employee functions depend on $\hat{z}$ (i.e., the knowledge threshold that separates workers and solvers in the pre-AI equilibrium). For this reason, we write $m(z; \hat{z})$ and $e(z; \hat{z})$ instead of $m(z)$ and $e(z)$, respectively.

• $z_{AI} \in \text{int}W$.— First, we show that every $z \in W^*$ has a worse solver post-AI than pre-AI. Note that if $z \in W_a^*$, then such a worker is matched with AI in the AI equilibrium. However, if this is the case, then $m(z; \hat{z}) \geq \hat{z} > z_{AI}$, as $z_{AI} \in \text{int}W$.

Proving that every $z \in W_p^*$ also has a worse solver is more involved. Since $z_{AI} \in \text{int}W$, the AI equilibrium is either Type 1 or Type 2. Suppose first that it is Type 1. Then, the matching functions pre- and post-AI are given by:

$$\begin{aligned} \int_{\underline{z}_1^*}^{m_1^*(z; \underline{z}_1^*)} dG(u) &= \int_0^z h(1-u) dG(u) \text{ for } z \in W_p^* = [0, z_{AI}] \\ \int_{\hat{z}}^{m(z; \hat{z})} dG(u) &= \int_0^z h(1-u) dG(u) \text{ for } z \in W = [0, \hat{z}] \end{aligned}$$

Thus, if $z \in W_p^* \cap W = W_p^*$, then $\int_{\underline{z}_1^*}^{m_1^*(z; \underline{z}_1^*)} dG(u) = \int_{\hat{z}}^{m(z; \hat{z})} dG(u)$, so $m_1^*(z; \underline{z}_1^*) < m(z; \hat{z})$ as $\underline{z}_1^* < \hat{z}$.

Suppose instead that the AI equilibrium is Type 2. Then, the matching functions pre- and post-AI are given by:

$$\begin{aligned} \int_{z_{AI}}^{m_2^*(z; \underline{z}_2^*)} dG(u) &= \int_{\underline{z}_2^*}^z h(1-u) dG(u) \text{ for } z \in W_p^* = [\underline{z}_2^*, z_{AI}] \\ \int_{\hat{z}}^{m(z; \hat{z})} dG(u) &= \int_0^z h(1-u) dG(u) \text{ for } z \in W = [0, \hat{z}] \end{aligned}$$

Consequently, if $z \in W_p^* \cap W = W_p^*$, then $\int_{\hat{z}}^{m(z; \hat{z})} dG(u) - \int_{z_{AI}}^{m_2^*(z; \underline{z}_2^*)} dG(u) = \int_0^{\underline{z}_2^*} h(1-u) dG(u) > 0$, which implies that $m_2^*(z; \underline{z}_2^*) < m(z; \hat{z})$ as $z_{AI} < \hat{z}$.

We now turn to solvers, i.e., those $z \in S \subset S^*$. We first claim that if $e(z; \hat{z}) = z_{AI}$, then $z \in S_a^* \cap S$, which immediately implies that if $e(z; \hat{z}) = z_{AI}$, then z has the same span of control pre- and post-AI. The proof is via the contrapositive. Suppose that $z \notin S_a^* \cap S$ (but that z is a solver). Then $z \in S_p^* \cap S$. However, if this is the case, then $z_{AI} \geq e_j^*(z; \underline{z}_j^*) > e(z; \hat{z})$, where the first inequality follows because AI is the best worker, and the second inequality follows because $e_j^*(z'; \underline{z}_j^*) > e(z'; \hat{z})$ for all $z' \in S_p^* \cap S$ if $m_j^*(z''; \underline{z}_j^*) < m(z''; \hat{z})$ for all $z'' \in W_p^* \cap W = W_p^*$ (which we already showed is true for $j = 1, 2$). Hence, $e(z; \hat{z}) \neq z_{AI}$.

The previous claim implies that if $e(z; \hat{z}) < z_{AI}$, then z 's span of control increases with AI, while if $e(z; \hat{z}) > z_{AI}$, then z 's span of control decreases with AI. Indeed, if $e(z; \hat{z}) < z_{AI}$, then either $z \in S_p^* \cap S$ or $z \in S_a^* \cap S$. In either case, z is assisting more knowledgeable workers post-AI than pre-AI: If $z \in S_p^*$, then we already know that $e_j^*(z; \underline{z}_j^*) > e(z; \hat{z})$, while if $z \in S_a^*$, then, post-AI, she is assisting AI, while

pre-AI, she was humans with knowledge $e(z; \hat{z}) < z_{AI}$. Similarly, if $e(z; \hat{z}) > z_{AI}$, then $z \in S_a^* \cap S$. The latter implies that z is assisting less knowledgeable workers post-AI than pre-AI, as post-AI, she is assisting the work of AI, while pre-AI, she was assisting humans with knowledge $e(z; \hat{z}) > z_{AI}$. $\square$

• $z_{AI} \in \text{int}S$.— First, we show that every $z \in S^* \subset S$ has a larger span of control post-AI than pre-AI. Note that if $z \in S_a^*$, then such a solver is matched with AI in the post-AI equilibrium. However, if so, then $e(z; \hat{z}) \leq \hat{z} < z_{AI}$, as $z_{AI} \in \text{int}S$.

We now show that the same holds for every $z \in S_p^*$. Since $z_{AI} \in \text{int}S$, the AI equilibrium is either Type 2 or Type 3. Suppose first that it is Type 2. Then, the employee functions pre- and post-AI are given by:

$$\begin{aligned} \int_z^{\bar{z}_2^*} dG(u) &= \int_{e_2^*(z; \bar{z}_2^*)}^{z_{AI}} h(1-u)dG(u) \text{ for } z \in S_p^* = [z_{AI}, \bar{z}_2^*] \\ \int_z^1 dG(u) &= \int_{e(z; \hat{z})}^{\hat{z}} h(1-u)dG(u) \text{ for } z \in S = [\hat{z}, 1] \end{aligned}$$

Consequently, if $z \in S_p^* \cap S = S_p^*$, then $0 < \int_{\bar{z}_2^*}^1 dG(u) = \int_{e(z; \hat{z})}^{\hat{z}} h(1-u)dG(u) - \int_{e_2^*(z; \bar{z}_2^*)}^{z_{AI}} h(1-u)dG(u)$, which implies that $e_2^*(z; \bar{z}_2^*) > e(z; \hat{z})$ (since $z_{AI} > \hat{z}$).

Suppose instead that the AI equilibrium is Type 3. Then, the employee functions pre- and post-AI are given by:

$$\begin{aligned} \int_z^1 dG(u) &= \int_{e_3^*(z; \bar{z}_3^*)}^{\bar{z}_3^*} h(1-u)dG(u) \text{ for } z \in S_p^* = [z_{AI}, 1] \\ \int_z^1 dG(u) &= \int_{e(z; \hat{z})}^{\hat{z}} h(1-u)dG(u) \text{ for } z \in S = [\hat{z}, 1] \end{aligned}$$

Consequently, for $z \in S_p^* \cap S = S_p^*$, then $\int_{\bar{z}_3^*}^1 dG(u) = \int_{e(z; \hat{z})}^{\hat{z}} h(1-u)dG(u)$. However, if this is the case, then $e(z; \hat{z}) < e_3^*(z; \bar{z}_3^*)$ since $\hat{z} < \bar{z}_3^*$.

We now turn to workers, i.e., those with $z \in W \subset W^*$. We first claim that if $z = e(z_{AI}; \hat{z})$, then $z \in W_a^* \cap W$, which immediately implies that if $z = e(z_{AI}; \hat{z})$, then z is equally productive pre- and post-AI. The proof is via the contrapositive. Suppose that $z \notin W_a^* \cap W$ (but that z is a worker). Then $z \in W_p^* \cap W$. However, if so, then $z_{AI} \leq m_j^*(z; \bar{z}_j^*) < m(z; \hat{z})$, where the first inequality follows because AI is the worst solver, and the second inequality follows because $m_j^*(z'; \bar{z}_j^*) < m(z'; \hat{z})$ for all $z' \in W_p^* \cap W$ if $e_j^*(z''; \bar{z}_j^*) > e(z''; \hat{z})$ for all $z'' \in S_p^* \cap S = S_p^*$ (which we already showed is true for $j = 2, 3$). Hence, $z \neq e(z_{AI}; \hat{z})$.

The previous claim implies that if $z < e(z_{AI}; \hat{z})$, then z is strictly more productive post-AI than pre-AI, while if $z > e(z_{AI}; \hat{z})$, then z is strictly less productive post-AI than pre-AI. Indeed, if $z < e(z_{AI}; \hat{z})$, then $z \in W_a^* \cap W$. Hence, z is matched with a better solver post-AI than pre-AI because post-AI, she is assisted by AI, while pre-AI, she was assisted by a human with knowledge $m(z; \hat{z}) < z_{AI}$. Similarly, if $z > e(z_{AI}; \hat{z})$, then $z \in W_a^* \cap W$ or $z \in W_p^* \cap W$. If $z \in W_a^* \cap W$, then z is matched with a worse solver post-AI than pre-AI because post-AI, she is assisted by AI, while pre-AI, she was assisted by a human with knowledge $m(z; \hat{z}) > z_{AI}$. Similarly, if $z \in W_p^* \cap W$, then z is also matched with a worse solver post-AI than pre-AI because we already established that $m_j^*(z; \bar{z}_j^*) < m(z; \hat{z})$ for $j = 2, 3$. $\square$

C.6 Proof of Lemma 2

For ease of exposition, we divide the proof of the lemma into three smaller claims:

Claim C.1. (i) $\Delta(z_{AI}) < 0$, (ii) $\Delta(1) > 0$, and (iii) $\Delta(0) > 0$ if $z_{AI} \in S$.

Proof. That $\Delta(z_{AI}) = z_{AI} - w(z_{AI}) < 0$ follows directly from the fact that $w(z) > z$ for all $z \in [0, 1]$ when $h < h_0$ (see Proposition 1). Consider next $\Delta(1) = w^*(1) - w(1)$. By Lemma B.2, we have that $w^*(1) = 1/h$, while by Lemma A.1 and Corollary A.1, we have that $w(1) < 1/h$. Hence, $\Delta(1) > 0$. Finally, consider $\Delta(0) = w^*(0) - w(0)$ and suppose that $z_{AI} \in S$. From Lemma B.2, we have that $w^*(0) = z_{AI}(1 - h)$ as the human with zero knowledge is assisted by AI irrespective of whether $z_{AI} \in \mathcal{R}_2$ or $z_{AI} \in \mathcal{R}_3$. Moreover, from Lemma A.1 and Corollary A.1 we have that $w(0) = \hat{z} - hw(\hat{z})$. Thus, $\Delta(0) = z_{AI}(1 - h) - (\hat{z} - hw(\hat{z})) > (1 - h)(z_{AI} - \hat{z}) \geq 0$, where the first inequality follows because $w(\hat{z}) > \hat{z}$, and the second inequality follows because $z_{AI} \geq \hat{z}$. $\square$

Claim C.2. If $\Delta(z) > 0$ for some $z \in [z_{AI}, 1)$, then $\Delta(z') > 0$ for all $z' \in [z, 1)$.

Proof. Given that $\Delta(z_{AI}) < 0$ (by Claim C.1), it suffices to show that if $\Delta(z)$ crosses zero at some $z > z_{AI}$, then it always crosses zero from below. We first consider $z_{AI} \in \text{int}W$ and then $z_{AI} \in \text{int}S$.

• $z_{AI} \in \text{int}W$.— Lemma B.2 implies that $\sup W^* = z_{AI}$, while Proposition 3 that $W^* \subset W$ and $S^* \supset S$. Hence, if $z > z_{AI}$, then z can only belong to either $I^* \cap W$, $S^* \cap W$, or $S^* \cap S$.

Now, irrespective of the presence of AI, the marginal return to knowledge is greater for solvers than for independent producers, and it is greater for independent producers than for workers. Hence, $\Delta'(z) = w^{*'}(z) - w'(z) \geq 0$ whenever z is in either $I^* \cap W$ or $S^* \cap W$. This implies that if $\Delta(z)$ crosses zero in either of these sets, then it necessarily crosses from below.

Consider then $z \in S^* \cap S$. Since $S^* = S_p^* \cup S_a^*$, here we have two cases to consider: $z \in S_p^* \cap S$ and $z \in S_a^* \cap S$. If $z \in S_p^* \cap S$, then $\Delta'(z) = n(e_j^*(z; z_j^*)) - n(e(z; \hat{z}))$, where $e_j^*(z; z_j^*)$ is the employee function in a Type $j = 1, 2$ equilibrium and $e(z; \hat{z})$ employee function in the pre-AI equilibrium. However, by Proposition 5, we know that $e_j^*(z; z_j^*) > e(z; \hat{z})$ as every $z \in S_p^* \cap S$ supervises better workers post-AI than pre-AI. Hence, in this case, $\Delta'(z) > 0$ also. Consequently, if $\Delta(z)$ crosses zero when $z \in S_p^* \cap S$, then it necessarily crosses from below.

Finally, consider the possibility that $\Delta(z)$ crosses zero at $z \in S_a^* \cap S$. In this case, $\Delta'(z) = n(z_{AI}) - n(e(z; \hat{z})) \geq 0$, so $\Delta(z)$ is no longer monotone in z in this set. Note, however, that $\Delta''(z) = -n'(e(z; \hat{z}))e'(z; \hat{z}) < 0$, so $\Delta(z)$ is concave. Moreover, if $S_a^* \neq \emptyset$, then $1 \in S_a^*$. The fact that $\Delta(z)$ is concave and that $\Delta(1) > 0$ then immediately implies that if $\Delta(z)$ crosses zero in this set, then it can only cross once and from below (otherwise, $\Delta(1) \leq 0$ contradicting the fact $\Delta(1) > 0$).

• $z_{AI} \in \text{int}S$.— From Lemma B.2, we know that $\inf S^* = z_{AI}$. Moreover, from Proposition 3, we have that $W^* \supseteq W$ and $S^* \subseteq S$. Consequently, if $z > z_{AI}$, then $z \in S^* \cap S$ necessarily. Since $S^* = S_p^* \cup S_a^*$,

here we have two cases to consider: $z \in S_p^* \cap S$ and $z \in S_a^* \cap S$.

If $z \in S_p^* \cap S$, then $\Delta'(z) = n(e_j^*(z; \underline{z}_j^*)) - n(e(z; \hat{z}))$, while if $z \in S_a^* \cap S$, then $\Delta'(z) = n(z_{AI}) - n(e(z; \hat{z}))$, where $e_j^*(z; \underline{z}_j^*)$ is the employee function in a Type $j = 2, 3$ equilibrium and $e(z; \hat{z})$ the employee function in the pre-AI equilibrium. In either case, $\Delta'(z) > 0$ since Proposition 5 states that $e_j^*(z; \underline{z}_j^*) > e(z; \hat{z})$ and $z_{AI} > e(z; \hat{z})$ when AI has the knowledge of a pre-AI solver. Consequently, $\Delta'(z) > 0$ for all $z \geq z_{AI}$, so if $\Delta(z)$ crosses zero at some $z > z_{AI}$, then it always crosses it from below. $\square$

Claim C.3. *If $\Delta(z) > 0$ for some $z \in [0, z_{AI}]$, then $\Delta(z') > 0$ for all $z' \in [0, z]$.*

Proof. Given that $\Delta(z_{AI}) < 0$, it suffices to show that if $\Delta(z)$ crosses zero at some $z < z_{AI}$, then it always crosses zero from above. We first consider $z_{AI} \in W$, and then $z_{AI} \in \text{int}S$.

• $z_{AI} \in \text{int}W$.— Lemma B.2 implies that $\sup W^* = z_{AI}$, while Proposition 3 that $W^* \subseteq W$ and $S^* \supseteq S$. Consequently, if $z < z_{AI}$, then $z \in W^* \cap W$, where $W^* = W_a^* \cup W_p^*$.

Now, if $z \in W_p^* \cap W$, then $\Delta'(z) = hw^*(m_j^*(z; \underline{z}_j^*)) - hw(m(z; \hat{z}))$, where $m_j^*(z; \underline{z}_j^*)$ is the matching function in a Type $j = 1, 2$ equilibrium, and $m(z; \hat{z})$ the matching function of the pre-AI equilibrium. However, using the firms' zero-profit condition:

$$\Delta'(z) = hw^*(m_j^*(z; \underline{z}_j^*)) - hw(m(z; \hat{z})) = \frac{m_j^*(z; \underline{z}_j^*) - m(z; \hat{z}) - \Delta(z)}{1 - z}$$

Consequently, if $\Delta(z) = 0$ at some z in this interval, say at $z = \zeta$, then $\Delta'(\zeta) = m_j^*(\zeta; \underline{z}_j^*) - m(\zeta; \hat{z}) \leq 0$, where the last inequality follows because every worker is assisted by a worse solver post-AI than pre-AI when $z_{AI} \in \text{int}W$ (as shown in Proposition 5).

On the other hand, if $z \in W_a^* \cap W$, then following the same reasoning as before, we have that:

$$\Delta'(z) = hw(z_{AI}) - hw(m(z; \hat{z})) = \frac{z_{AI} - m(z; \hat{z}) - \Delta(z)}{1 - z}$$

Consequently, if $\Delta(z) = 0$ at some z in this interval, say at $z = \zeta$, then $\Delta'(\zeta) = z_{AI} - m(\zeta; \hat{z}) \leq 0$, where the inequality follows, again, from Proposition 5.

• $z_{AI} \in \text{int}S$.— In this case, Lemma B.2 implies that $\inf S^* = z_{AI}$, while Proposition 3 that $W^* \supset W$ and $S^* \subset S$. Consequently, if $z < z_{AI}$, then z can only belong to either $I^* \cap S$, $W^* \cap S$, or $W^* \cap W$.

Now, $\Delta'(z) \leq 0$ whenever z is in either $I^* \cap S$ or $W^* \cap S$, given that the marginal return to knowledge is greater for solvers than for independent producers, and it is greater for independent producers than for workers. Hence, if $\Delta(z)$ crosses zero in either of these sets, then it necessarily crosses from above.

Consider then $z \in W^* \cap W$. Since $W^* = W_p^* \cup W_a^*$, we have two cases to consider: $z \in W_p^* \cap W$ and $z \in W_a^* \cap W$. If $z \in W_p^* \cap W$, then:

$$\Delta'(z) = hw^*(m_j^*(z; \underline{z}_j^*)) - hw(m(z; \hat{z})) = \frac{m_j^*(z; \underline{z}_j^*) - m(z; \hat{z}) - \Delta(z)}{1 - z}$$

where $m_j^*(z; \underline{z}_j^*)$ is the matching function in a Type $j = 2, 3$ equilibrium, and $m(z; \hat{z})$ the matching function of the pre-AI equilibrium. Consequently, if $\Delta(z) = 0$ at some z in this interval, say at $z = \zeta$, then $\Delta'(\zeta) = m_j^*(\zeta; \underline{z}_j^*) - m(\zeta; \hat{z}) \leq 0$, where the last inequality follows because every $z \in W^* \cap W$ is assisted by a worse solver post-AI than pre-AI when $z_{AI} \in \text{int}S$.

Finally, consider the possibility that $\Delta(z)$ crosses zero at $z \in W_a^* \cap W$. In this case, $\Delta'(z) = h z_{AI} - h w(m(z; \hat{z}))$, so $\Delta''(z) = -h w'(m(z; \hat{z})) f'(z; \hat{z}) < 0$, implying that $\Delta(z)$ is concave. Moreover, if $W_a^* \neq \emptyset$, then $0 \in W_a^*$, and we know that $\Delta(0) > 0$ in this case. Consequently, if $\Delta(z)$ crosses zero in this set, then it can only cross once and from above (otherwise, $\Delta(0) \leq 0$ contradicting the fact $\Delta(0) > 0$). $\square$

C.7 Proof of Proposition 6

Part (i) (“there always exists a set $[z, 1] \subseteq (z_{AI}, 1]$ of winners”) follows directly from Lemma 2 and the fact that $\Delta(1) > 0$ (see Claim C.1). Hence, part (ii) (“there exists a set $[0, z] \subseteq [0, z_{AI}]$ of winners if and only if $z_{AI} > \bar{z}_{AI}$, where $\bar{z}_{AI} \in \text{int}W$ ”) is the only part that remains to be proven. To do so, we begin by constructing $\bar{z}_{AI}$ and then show that the statement is true.

Let $\Delta(0; z_{AI}) \equiv w^*(0; z_{AI}) - w(0)$, and define $\bar{z}_{AI}$ as the solution to $\Delta(0; \bar{z}_{AI}) = 0$. We first show that $\bar{z}_{AI}$ exists, is unique, and that $\bar{z}_{AI} \in (0, \hat{z})$. We do this by showing that $\Delta(0; z_{AI})$ crosses zero once as we move from $z_{AI} = 0$ to $z_{AI} = 1$, and that this crossing point is at a $z_{AI} < \hat{z}$. Indeed, if $z_{AI} \geq \hat{z}$, then Claim C.1 states that $\Delta(0; z_{AI}) > 0$ (as $z_{AI} \in S$ in this case). Moreover, as shown in Lemma B.3, $0 \in \mathcal{R}_1$, so when $z_{AI} = 0$, the equilibrium is always Type 1. The latter implies that $\Delta(0; 0) = \underline{z}_1^*(0)(1 - h) - w(0) = -w(0) < 0$,³⁷ where the second-to-last equality follows because $\underline{z}_1^*(0) = 0$, as can be easily be proven from the condition that determines $\underline{z}_1^*(z_{AI})$ (see the statement of Lemma B.2).

Now, when $z \in [0, \hat{z}] = W$, the equilibrium is either Type 1, in which case $w^*(0; z_{AI}) = \underline{z}_1^*(z_{AI})(1 - h)$, or Type 2, in which case $w^*(0; z_{AI}) = z_{AI}(1 - h)$. Using the equilibrium condition that determines $\underline{z}_1^*(z_{AI})$, it is not difficult to prove that (i) $\underline{z}_1^*(z_{AI})$ is strictly increasing in z_{AI} , and that (ii) $\underline{z}_1^*(z_{AI}) = z_{AI}$ whenever we switch from a Type 1 to a Type 2 equilibrium (and vice versa). Consequently, irrespective of the equilibrium type in this region, $w^*(0; z_{AI})$ is continuous and strictly increasing in z_{AI} , which implies that $\Delta(0; z_{AI})$ is also continuous and strictly increasing in z_{AI} . This result, combined with the fact that $\Delta(0; 0) < 0$ and $\Delta(0; \hat{z}) > 0$, immediately yields the desired result.

Having constructed $\bar{z}_{AI}$, we prove that there exists $z < z_{AI}$ such that $\Delta(z; z_{AI}) > 0$ if and only if $z_{AI} > \bar{z}_{AI}$. First we show that if there exists $z < z_{AI}$ such that $\Delta(z; z_{AI}) > 0$, then $z_{AI} > \bar{z}_{AI}$. To do this, we prove the contrapositive statement: If $z_{AI} \leq \bar{z}_{AI}$, then there is no such z . Indeed, as shown above, $\Delta(0; z_{AI}) \leq 0$ for all $z_{AI} \leq \bar{z}_{AI}$. Hence, Lemma 2 implies that $\Delta(z; z_{AI}) \leq 0$ for all $z < z_{AI}$.

³⁷To avoid any confusion, here we make explicit that $\underline{z}_1^*(z_{AI})$ depends on z_{AI} as seen from the condition that determines $\underline{z}_1^*(z_{AI})$ in the statement of Lemma B.2.

Finally, we prove that if $z_{AI} > \bar{z}_{AI}$, then there exists $z < z_{AI}$ such that $\Delta(z; z_{AI}) > 0$. Indeed, as shown above, $z_{AI} > \bar{z}_{AI}$ then $\Delta(0; z_{AI}) > 0$. Consequently, Lemma 2 immediately implies that there exists $\zeta \in [0, z_{AI})$ such that $\Delta(z; z_{AI}) > 0$ for $z \in [0, \zeta)$ and $\Delta(z; z_{AI}) < 0$ for $z \in (\zeta, z_{AI}]$. $\square$

References

- ACEMOGLU, D. (2024): “The Simple Macroeconomics of AI,” Working Paper.
- ACEMOGLU, D. AND D. AUTOR (2011): “Chapter 12 - Skills, Tasks and Technologies: Implications for Employment and Earnings,” Elsevier, vol. 4 of *Handbook of Labor Economics*, 1043–1171.
- ACEMOGLU, D., F. KONG, AND P. RESTREPO (2024): “Tasks At Work: Comparative Advantage, Technology and Labor Demand,” NBER Working Paper No. 32872.
- ACEMOGLU, D. AND J. LOEBBING (2024): “Automation and Polarization,” Working Paper.
- ACEMOGLU, D. AND P. RESTREPO (2018): “The Race Between Man and Machine: Implications of Technology for Growth, Factor Shares, and Employment,” *American Economic Review*, 108, 1488–1542.
- (2022): “Tasks, Automation, and the Rise in U.S. Wage Inequality,” *Econometrica*, 90, 1973–2016.
- ADHVARYU, A., V. BASSI, A. NYSHADHAM, J. A. TAMAYO, AND N. TORRES (2023): “Organizational Responses to Product Cycles,” NBER Working Paper No. 31582.
- AGHION, P., B. F. JONES, AND C. I. JONES (2017): “Artificial Intelligence and Economic Growth,” NBER Working Paper No. 23928.
- AGRAWAL, A., J. GANS, AND A. GOLDFARB (2019): “Introduction,” in *The Economics of Artificial Intelligence: An Agenda*, ed. by A. Agrawal, J. Gans, and A. Goldfarb, University of Chicago Press, 1–19.
- ALEKSEEVA, L., J. AZAR, M. GINÈ, AND S. SAMILA (2024): “AI Adoption and the Demand for Managerial Expertise,” SSRN Working Paper No. 4986379.
- ALFRED SLOAN (1924): “The Most Important Thing I Ever Learned about Management,” *System: The Magazine of Business*, 194, 140–141.
- ALONSO, R., W. DESSEIN, AND N. MATOUSCHEK (2008): “When Does Coordination Require Centralization?” *American Economic Review*, 98, 145–179.
- AMODEI, D. (2024): “Machines of Loving Grace: How AI Could Transform the World for the Better,” October, 2024. Available at <https://darioamodei.com/machines-of-loving-grace> (accessed October 14, 2024).
- ANTRÀS, P., L. GARICANO, AND E. ROSSI-HANSBERG (2006): “Offshoring in a Knowledge Economy,” *The Quarterly Journal of Economics*, 121, 31–77.
- (2008): “Organizing Offshoring: Middle Managers and Communication Costs,” in *The Organization of Firms in a Global Economy*, ed. by E. Helpman, D. Marin, and T. Verdier, Harvard University Press, 311–339.

- APPENZELLER, G., M. BORNSTEIN, AND M. CASADO (2023): “Navigating the High Costs of AI Compute,” *Andressen Horowitz*, April 27, 2023. Available at <https://a16z.com/navigating-the-high-cost-of-ai-compute/> (accessed October 16, 2024).
- ATKESON, A. AND P. J. KEHOE (1999): “Models of Energy Use: Putty-Putty versus Putty-Xlay,” *American Economic Review*, 89, 1028–1043.
- AUTOR, D. (2024): “Applying AI to Rebuild Middle Class Jobs,” NBER Working Paper No. 32140.
- AUTOR, D., D. DORN, L. F. KATZ, C. PATTERSON, AND J. VAN REENEN (2020): “The Fall of the Labor Share and the Rise of Superstar Firms,” *The Quarterly Journal of Economics*, 135, 645–709.
- AUTOR, D. H., F. LEVY, AND R. J. MURNANE (2003): “The Skill Content of Recent Technological Change: An Empirical Exploration*,” *The Quarterly Journal of Economics*, 118, 1279–1333.
- AZAR, J., M. CHUGUNOVA, K. KELLER, AND S. SAMILA (2023): “Monopsony and Automation,” Max Planck Institute for Innovation & Competition Research Paper No. 23-21.
- BEANE, M. (2024): “How AI Could Keep Young Workers From Getting the Skills They Need,” *The Wall Street Journal*, July 26, 2024. Available at <https://www.wsj.com/lifestyle/careers/ai-training-young-employees-ed08cedc> (accessed November 15, 2024).
- BEANE, M. AND C. ANTHONY (2024): “Inverted Apprenticeship: How Senior Occupational Members Develop Practical Expertise and Preserve Their Position When New Technologies Arrive,” *Organization Science*, 35, 405–431.
- BERGER, P. G., W. CAI, L. QIU, AND C. XINYI SHEN (2024): “Employer and Employee Responses to Generative AI: Early Evidence,” SSRN Working Paper No. 4874061.
- BLOOM, N., L. GARICANO, R. SADUN, AND J. VAN REENEN (2014): “The Distinct Effects of Information Technology and Communication Technology on Firm Organization,” *Management Science*, 60, 2859–2885.
- BOMMASANI, R., D. A. HUDSON, E. ADELI, R. ALTMAN, S. ARORA, S. VON ARX, M. S. BERNSTEIN, J. BOHG, A. BOSSELUT, E. BRUNSKILL, E. BRYNJOLFSSON, S. BUCH, D. CARD, R. CASTELLON, N. CHATTERJI, A. CHEN, K. CREEL, J. Q. DAVIS, D. DEMSZKY, C. DONAHUE, M. DOUMBOUYA, E. DURMUS, S. ERMON, J. ETCHEMENDY, K. ETHAYARAJH, L. FEI-FEI, C. FINN, T. GALE, L. GILLESPIE, K. GOEL, N. GOODMAN, S. GROSSMAN, N. GUHA, T. HASHIMOTO, P. HENDERSON, J. HEWITT, D. E. HO, J. HONG, K. HSU, J. HUANG, T. ICARD, S. JAIN, D. JURAFSKY, P. KALLURI, S. KARAMCHETI, G. KEELING, F. KHANI, O. KHATTAB, P. W. KOH, M. KRASS, R. KRISHNA, R. KUDITIPUDI, A. KUMAR, F. LADHAK, M. LEE, T. LEE, J. LESKOVEC, I. LEVENT, X. L. LI, X. LI, T. MA, A. MALIK, C. D. MANNING, S. MIRCHANDANI, E. MITCHELL, Z. MUNYIKWA, S. NAIR, A. NARAYAN, D. NARAYANAN, B. NEWMAN, A. NIE, J. C. NIEBLES, H. NILFOROSHAN, J. NYARKO, G. OGUT, L. ORR, I. PAPADIMITRIOU, J. S. PARK, C. PIECH, E. PORTELANCE, C. POTTS, A. RAGHUNATHAN, R. REICH, H. REN, F. RONG, Y. ROOHANI, C. RUIZ, J. RYAN, C. RÉ, D. SADIGH, S. SAGAWA, K. SANTHANAM, A. SHIH, K. SRINIVASAN, A. TAMKIN, R. TAORI, A. W. THOMAS, F. TRAMÈR, R. E. WANG, W. WANG, B. WU, J. WU, Y. WU, S. M. XIE, M. YASUNAGA, J. YOU, M. ZAHARIA, M. ZHANG, T. ZHANG, X. ZHANG, Y. ZHANG, L. ZHENG, K. ZHOU, AND P. LIANG (2022): “On the Opportunities and Risks of Foundation Models,” ArXiv preprint arXiv:2108.07258.

- BROWN, T. B., B. MANN, N. RYDER, M. SUBBIAH, J. KAPLAN, P. DHARIWAL, A. NEELAKANTAN, P. SHYAM, G. SASTRY, A. ASKELL, S. AGARWAL, A. HERBERT-VOSS, G. KRUEGER, T. HENIGHAN, R. CHILD, A. RAMESH, D. M. ZIEGLER, J. WU, C. WINTER, C. HESSE, M. CHEN, E. SIGLER, M. LITWIN, S. GRAY, B. CHESS, J. CLARK, C. BERNER, S. MCCANDLISH, A. RADFORD, I. SUTSKEVER, AND D. AMODEI (2020): "Language Models are Few-Shot Learners," ArXiv preprint arXiv:2005.14165.
- BRYNJOLFSSON, E. (2022): "The Turing Trap: The Promise & Peril of Human-Like Artificial Intelligence," *Daedalus*, 151, 272–287.
- BRYNJOLFSSON, E., D. LI, AND L. R. RAYMOND (2023): "Generative AI at Work," NBER Working Paper No. 31161.
- BRYNJOLFSSON, E. AND A. MCAFEE (2016): *The Second Machine Age*, WW Norton.
- BRYNJOLFSSON, E., A. MCAFEE, M. SORELL, AND F. ZHU (2008): "Scale Without Mass: Business Process Replication and Industry Dynamics," *Harvard Business School Technology & Operations Mgt. Unit Research Paper*.
- CAICEDO, S., J. LUCAS, ROBERT E., AND E. ROSSI-HANSBERG (2019): "Learning, Career Paths, and the Distribution of Wages," *American Economic Journal: Macroeconomics*, 11, 49–88.
- CALIENDO, L., G. MION, L. D. OPROMOLLA, AND E. ROSSI-HANSBERG (2020): "Productivity and Organization in Portuguese Firms," *Journal of Political Economy*, 128, 4211–4257.
- CALIENDO, L., F. MONTE, AND E. ROSSI-HANSBERG (2015): "The Anatomy of French Production Hierarchies," *Journal of Political Economy*, 123, 809–852.
- CALIENDO, L. AND E. ROSSI-HANSBERG (2012): "The Impact of Trade on Organization and Productivity," *The Quarterly Journal of Economics*, 127, 1393–1467.
- CAPLIN, A., D. J. DEMING, S. LI, D. J. MARTIN, P. MARX, B. WEIDMANN, AND K. J. YE (2024): "The ABC's of Who Benefits from Working with AI: Ability, Beliefs, and Calibration," NBER Working Paper No. 33021.
- CARMONA, G. AND K. LAOHAKUNAKORN (2024): "Improving the Organization of Knowledge in Production by Screening Problems," *Journal of Political Economy*, 132, 1290–1326.
- COUNCIL OF ECONOMIC ADVISERS (2024): "Chapter 7: An Economic Framework for Understanding Artificial Intelligence," in *The 2024 Economic Report of the President*, Executive Office of the President of the United States, available at <https://www.whitehouse.gov/cea/written-materials/2024/03/21/the-2024-economic-report-of-the-president/> (accessed October 8, 2024).
- CREMER, J., L. GARICANO, AND A. PRAT (2007): "Language and the Theory of the Firm," *The Quarterly Journal of Economics*, 122, 373–407.
- DELL'ACQUA, F., E. MCFOWLAND III, E. R. MOLICK, H. LIFSHITZ-ASSAF, K. KELLOGG, S. RAJENDRAN, L. KRAYER, F. CANDELON, AND K. R. LAKHANI (2023): "Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of AI on Knowledge Worker Productivity and Quality," Harvard Business School Technology & Operations Mgt. Unit Working Paper No. 24-013.

- DESSEIN, W., A. GALEOTTI, AND T. SANTOS (2016): "Rational Inattention and Organizational Focus," *American Economic Review*, 106, 1522–1536.
- DESSEIN, W. AND T. SANTOS (2006): "Adaptive Organizations," *Journal of Political Economy*, 114, 956–995.
- ECKHOUT, J. AND P. KIRCHER (2018): "Assortative Matching with Large Firms," *Econometrica*, 86, 85–132.
- ELOUNDOU, T., S. MANNING, P. MISHKIN, AND D. ROCK (2024): "GPTs are GPTs: Labor market impact potential of LLMs," *Science*, 384, 1306–1308.
- EU ARTIFICIAL INTELLIGENCE ACT (2024): "Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828," *Official Journal of the European Union*, ELI: <http://data.europa.eu/eli/reg/2024/1689/oj>.
- FUCHS, W., L. GARICANO, AND L. RAYO (2015): "Optimal Contracting and the Organization of Knowledge," *The Review of Economic Studies*, 82, 632–658.
- GARICANO, L. (2000): "Hierarchies and the Organization of Knowledge in Production," *Journal of Political Economy*, 108, 874–904.
- GARICANO, L. AND T. HUBBARD (2007): "Managerial Leverage Is Limited by the Extent of the Market: Hierarchies, Specialization, and the Utilization of Lawyers' Human Capital," *The Journal of Law & Economics*, 50, 1–43.
- GARICANO, L. AND T. N. HUBBARD (2009): "Specialization, Firms, and Markets: The Division of Labor Within and Between Law Firms," *The Journal of Law, Economics, & Organization*, 25, 339–371.
- (2012): "Learning About the Nature of Production from Equilibrium Assignment Patterns," *Journal of Economic Behavior & Organization*, 84, 136–153.
- (2018): "Earnings Inequality and Coordination Costs: Evidence from US Law Firms," *The Journal of Law, Economics, and Organization*, 34, 196–229.
- GARICANO, L. AND L. RAYO (2017): "Relational Knowledge Transfers," *American Economic Review*, 107, 2695–2730.
- GARICANO, L. AND E. ROSSI-HANSBERG (2004): "Inequality and the Organization of Knowledge," *American Economic Review Papers and Proceedings*, 94, 197–202.
- (2006): "Organization and Inequality in a Knowledge Economy," *The Quarterly Journal of Economics*, 121, 1383–1435.
- (2012): "Organizing Growth," *Journal of Economic Theory*, 147, 623–656.
- (2015): "Knowledge-Based Hierarchies: Using Organizations to Understand the Economy," *Annual Review of Economics*, 7, 1–30.
- GIBBONS, R. AND J. ROBERTS (2013): *The Handbook of Organizational Economics*, Princeton University Press.

- GOLDFARB, A. AND C. TUCKER (2019): “Digital Economics,” *Journal of Economic Literature*, 57, 3–43.
- GOLDMAN SACHS (2023): “Generative AI: Hype, or Truly Transformative?” July 5, 2023. Available at <https://www.goldmansachs.com/pdfs/insights/pages/top-of-mind/generative-ai-hype-or-truly-transformative/report.pdf> (accessed November 30, 2024).
- GUMPERT, A., H. STEIMER, AND M. ANTONI (2022): “Firm Organization with Multiple Establishments,” *The Quarterly Journal of Economics*, 137, 1091–1138.
- HUMLUM, A. AND E. VESTERGAARD (2024): “The Adoption of ChatGPT,” University of Chicago, Becker Friedman Institute for Economics Working Paper.
- IDE, E. AND E. TALAMÀS (2024a): “Artificial Intelligence in the Knowledge Economy,” ArXiv preprint arXiv:2312.05481 (version 6).
- (2024b): “Expertise, Apprenticeships, and Technological Change,” Working Paper.
- (2024c): “The Impact of AI on Global Knowledge Work,” Working Paper.
- JOHNSON, S. AND D. ACEMOGLU (2023): *Power and Progress: Our Thousand-Year Struggle Over Technology and Prosperity*, Hachette UK.
- JONES, C. I. (Forthcoming): “The A.I. Dilemma: Growth versus Existential Risk,” *American Economic Review: Insights*.
- KAPLAN, J., S. MCCANDLISH, T. HENIGHAN, T. B. BROWN, B. CHESS, R. CHILD, S. GRAY, A. RADFORD, J. WU, AND D. AMODEI (2020): “Scaling Laws for Neural Language Models,” ArXiv preprint arXiv:2001.08361.
- KORINEK, A. AND D. SUH (2024): “Scenarios for the Transition to AGI,” NBER Working Paper No. 32255.
- LI, X., S. CHAN, X. ZHU, Y. PEI, Z. MA, X. LIU, AND S. SHAH (2023): “Are ChatGPT and GPT-4 General-Purpose Solvers for Financial Text Analytics? A Study on Several Typical Tasks,” ArXiv preprint arXiv:2305.05862.
- LITERA (2022): “Litera Technology in M&A Report,” Available at <https://info.litera.com/report-2022-ma-technology-survey.html> (accessed September 5, 2024).
- MARTINEZ, J. (2021): “Putty-Clay Automation,” CEPR Working Paper No. DP16022.
- MCÉLHERAN, K., J. F. LI, E. BRYNJOLFSSON, Z. KROFF, E. DINLERSOZ, L. S. FOSTER, AND N. ZOLAS (2023): “AI Adoption in America: Who, What, and Where,” NBER Working Paper No. 31788.
- MESEROLE, C. O. (2018): “What is Machine Learning?” Tech. rep., Brookings Institution.
- METR (2024): “An Update on Our General Capability Evaluations,” August 6, 2024. Available at <https://metr.org/blog/2024-08-06-update-on-evaluations/> (accessed October 15, 2024).
- MOLL, B., L. RACHEL, AND P. RESTREPO (2022): “Uneven Growth: Automation’s Impact on Income and Wealth Inequality,” *Econometrica*, 90, 2645–2683.

- MURO, M., J. WHITON, AND R. MAXIM (2019): “What Jobs are Affected by AI?” Tech. rep., Brookings Institution.
- NATIONAL SCIENCE BOARD (2018): “Science & Engineering Indicators - 2018 Report,” Available at <https://www.nsf.gov/statistics/2018/nsb20181/report> (accessed December 4, 2024).
- NAVANI, D. A. (2023): “The Commoditization of LLMs,” *Communications of the ACM*, September 12, 2024. Available at <https://cacm.acm.org/blogcacm/the-commoditization-of-llms/> (accessed November 25, 2024).
- NORDHAUS, W. D. (2007): “Two Centuries of Productivity Growth in Computing,” *The Journal of Economic History*, 67, 128–159.
- NORI, H., Y. T. LEE, S. ZHANG, D. CARIGNAN, R. EDGAR, N. FUSI, N. KING, J. LARSON, Y. LI, W. LIU, R. LUO, S. M. MCKINNEY, R. O. NESS, H. POON, T. QIN, N. USUYAMA, C. WHITE, AND E. HORVITZ (2023): “Can Generalist Foundation Models Outcompete Special-Purpose Tuning? Case Study in Medicine,” ArXiv preprint arXiv:2311.16452.
- NOY, S. AND W. ZHANG (2023): “Experimental Evidence on the Productivity Effects of Generative Artificial Intelligence,” *Science*, 381, 187–192.
- PANCS, R. (2018): *Lectures on Microeconomics: The Big Questions Approach*, MIT Press.
- PENG, S., E. KALLIAMVAKOU, P. CIHON, AND M. DEMIRER (2023): “The Impact of AI on Developer Productivity: Evidence from GitHub Copilot,” ArXiv preprint arXiv:2302.06590.
- ROSENBUSH, S. (2024): “AI Agents Can Do More Than Answer Queries. That Raises a Few Questions,” *The Wall Street Journal*, October 16, 2024. Available at <https://www.wsj.com/articles/ai-agents-can-do-more-than-answer-queries-that-raises-a-few-questions-15009853> (accessed October 17, 2024).
- SHAH, A. AND C. PATEL (2023): “Mastering LLM Techniques: Customization,” *Nvidia Developer Technical Blog*, August 10, 2023. Available at <https://developer.nvidia.com/blog/selecting-large-language-model-customization-techniques> (accessed December 14, 2024).
- STEBBINGS, H. (2024a): “Aravind Srinivas: Will Foundation Models Commoditise?” *20VC with Harry Stebbings*, Episode #1161. Available at <https://www.thetwentyminutevc.com/aravind-srinivas> (accessed October 8, 2024).
- (2024b): “Sarah Tavel: Will Foundation Models Be Commoditised?” *20VC with Harry Stebbings*, Episode #1149. Available at <https://www.thetwentyminutevc.com/sarah-tavel> (accessed October 8, 2024).
- SVANBERG, M., W. LI, M. FLEMING, B. GOEHRING, AND N. THOMPSON (2024): “Beyond AI Exposure: Which Tasks are Cost-Effective to Automate with Computer Vision?” SSRN Working Paper No. 4700751.
- VELDKAMP, L. AND C. CHUNG (2024): “Data and the Aggregate Economy,” *Journal of Economic Literature*, 62, 458–84.
- WAISBERG, N. AND A. HUDEK (2021): *AI for Lawyers: How Artificial Intelligence is Adding Value, Amplifying Expertise, and Transforming Careers*, John Wiley & Sons.

- WANG, S. AND S. XU (2024): "16 Changes to the Way Enterprises Are Building and Buying Generative AI," *Andreessen Horowitz*, March 21, 2024. Available at <https://a16z.com/generative-ai-enterprise-2024/> (accessed December 13, 2024).
- WILES, E., Z. T. MUNYIKWA, AND J. J. HORTON (2023): "Algorithmic Writing Assistance on Jobseekers' Resumes Increases Hires," NBER Working Paper No. 30886.
- XIA, Y., J. KIM, Y. CHEN, H. YE, S. KUNDU, C. HAO, AND N. TALATI (2024): "Understanding the Performance and Estimating the Cost of LLM Fine-Tuning," ArXiv preprint arXiv:2408.04693.
- ZEIRA, J. (1998): "Workers, Machines, and Economic Growth," *The Quarterly Journal of Economics*, 113, 1091–1117.