



Basically intelligent: ontological and rationalistic perspectives on artificial intelligence in the context of basic income discourse

Básicamente inteligente: perspectivas ontológicas y racionalistas sobre la inteligencia artificial en el contexto del discurso de la renta básica

Shawn Christopher Vigil

Charles University, Czech Republic

2quillswriting@gmail.com

<http://orcid.org/0009-0003-8918-9654>

DOI <https://doi.org/10.48204/2805-1815.8472>

INFORMACIÓN DEL ARTÍCULO	ABSTRACT/RESUMEN
Recibido el: 28/5/2025 Aceptado el: 31/8/2025 Keywords: Applied ethics, basic income, artificial intelligence, political philosophy, ontology Palabras clave: Ética aplicada, renta básica, inteligencia artificial, filosofía política, ontología	Abstract: Basic income is a novel social welfare policy proposal that looks to preserve liberal-egalitarian principles by offering a cash entitlement delivered regularly to every individual in a given society without any stipulations (e.g., work or income requirements). The interest in such kinds of programs has grown larger in the context of exponential technological advancement, with anxieties about the prospect of AI displacing large portions of human labour abounding. However, while the problem of automation has been addressed in the basic income literature, very little philosophical treatment of it has been offered. The present essay aims to fill this gap by elucidating, evaluating, and articulating philosophical arguments that lie at the intersection of AI and ethics. The first argument deals with the question of ontology, viz., whether it is possible in principle for AI to perform all tasks associated with human labour. This argument is explored through a critique of Searle's well-known arguments against the computational theory of mind, together with Dreyfus's phenomenological perspective on the significance of context for sense-making. It is suggested that even if AI might not be able to authentically instantiate intelligence of a general kind, it might nevertheless be capable of adequately performing all tasks associated with human labour. The second argument deals with economic reasoning, viz.,

whether it would be rational for firms to substitute human labour for AI. It is suggested that micro- and macro-economic rationales betray each other and therefore cannot reliably discount the possibility of significant or complete displacement of human labour. Given that AI remains in principle a possible threat to socio-economic welfare via its relation to labour markets, we end by considering how basic income is uniquely situated to remedy the situation.

Resumen:

La renta básica es una novedosa propuesta de política de bienestar social que busca preservar los principios liberales e igualitarios al ofrecer un derecho a una prestación económica que se entrega regularmente a cada individuo de una sociedad determinada, sin ninguna condición (por ejemplo, requisitos de trabajo o ingresos). El interés en este tipo de programas ha crecido en el contexto del avance tecnológico exponencial, con la creciente inquietud ante la posibilidad de que la IA desplace gran parte del trabajo humano. Sin embargo, si bien el problema de la automatización se ha abordado en la literatura sobre la renta básica, se le ha ofrecido muy poco tratamiento filosófico. El presente ensayo pretende llenar este vacío elucidando, evaluando y articulando argumentos filosóficos que se encuentran en la intersección de la IA y la ética. El primer argumento aborda la cuestión de la ontología, es decir, si es posible, en principio, que la IA realice todas las tareas asociadas con el trabajo humano. Este argumento se explora mediante una crítica de los conocidos argumentos de Searle contra la teoría computacional de la mente, junto con la perspectiva fenomenológica de Dreyfus sobre la importancia del contexto para la construcción de sentido. Se sugiere que, aunque la IA no sea capaz de instanciar auténticamente inteligencia de tipo general, podría ser capaz de realizar adecuadamente todas las tareas asociadas con el trabajo humano. El segundo argumento aborda el razonamiento económico, es decir, si sería racional que las empresas sustituyeran el trabajo humano por la IA. Se sugiere que las lógicas micro y macroeconómicas se contradicen entre sí y, por lo tanto, no pueden descartar con fiabilidad la posibilidad de un desplazamiento significativo o completo del trabajo humano. Dado que la IA sigue siendo, en principio, una posible amenaza para el bienestar socioeconómico a través de su relación con los mercados laborales, concluimos considerando cómo la renta básica está en una posición privilegiada para remediar la situación.

Introduction

It is not as though we have not encountered this narrative before: A benevolent (or at least well- intentioned) intelligence bringing forth its progeny, presumably using itself as the schematic for its design. And this creation is destined to cultivate the world so as to transform it into a paradise that wants for nothing. In myths we have seen this, but nothing quite like it in reality – until now. The first three waves of industrialization rapidly and radically transformed not just the material world and society, but also our understanding of and relationship with them and, perhaps more importantly still, ourselves. The so-called

fourth wave in which we are currently enveloped stands to be just as or perhaps even more rapid and transformative. And for the first time in history, the human mind is meeting in the world what it has only previously met in imagination: An intelligence not dissimilar to its own, and even something more.

Artificial intelligence (AI) confronts humanity with deep ontological and practical questions. In regard to the former, it challenges notions of intelligence, consciousness, and humanity as such; in regard to the latter, it forces us to reckon with the possibility that any beings sophisticated enough to perform most, or all human functions will render us redundant, placing us in a precarious socio- economic situation. Basic income discourse occupies itself with questions of the second kind, although, of course, it is underpinned by questions of the first kind. In any register we must ask, how can we ensure that if and when labour becomes exceedingly scarce or disappears altogether, welfare does not vanish along with it? Advocates argue that basic income is the only policy which can provide the security needed in the context of such a society. Others deny the eventuality wholesale, finding nothing especially novel in the most recent wave of technological development (LSE, 2025). If it is true that AI cannot and will not have the radically displacing effects we imagine, then it becomes a non-issue; the argument is irrelevant in any discussion of social welfare, and we can safely leave off with fantastical projections of a post-work society and concentrate on more familiar and realistic arguments. If, however, it would in principle be possible for AI to perform all tasks relevant to human labour – from the most primitive to the most intellectually demanding – then we must seriously consider what safeguards we should have on standby in the case of our eventual complete substitution. The question then becomes, can and will AI threaten human labour such that we are left in a desperate situation that only a basic income can remedy?

Technological innovation is nothing new, nor are its effects on markets and economies. Those who are unconcerned about AI often appeal to history: Economist Heidi Schierholz, for example, observes that when new technologies are introduced, there is indeed temporary displacement in certain sectors, but they are counterbalanced by developments in others, resulting in relative stasis at a minimum and economic growth at best (Vox, 2017). Contrary to the predictions of dystopian alarmists, AI rather seems

poised to facilitate general welfare in the form of increased innovation and productivity. The Future of Jobs Report 2025, published by the World Economic Forum (2025), casts AI as a major influence in employment trends in its projection of a net growth of 78 million jobs. And amid fears that as high as 47% of jobs could face technological replacement, more modest calculations can put that number as low as 9%, far less cause for any serious concern (Frey & Osborne, 2017, p.114; Arntz & Zierahn, 2016). Technological unemployment (that is, unemployment instigated by technological progress) is therefore not near the existential threat that it is sometimes sensationalized to be, as human labour will likely continue to be complemented rather than substituted by automation.

These kinds of observations and arguments provide little comfort for the less optimistically- minded. We can grant that historical trends reveal predictable patterns and nevertheless retain the suspicion that something unprecedented is couched in this new frontier (Ford, 2015). After all, tasks that were replaced in earlier eras were largely mechanical and routine, whereas the capabilities of newer technologies are becoming increasingly sophisticated, their potential seemingly limitless. It is one thing for an automaton to perform the isolated task of assembling specific raw materials at a station in an assembly line, and quite another for it to diagnose skin conditions, evaluate legal documents, produce art, provide advice on personal affairs, write computer code and academic essays, or balance financial accounts. With such promise and uncertainty, we might temper our historically-informed confidence that things will carry on as they always have. Furthermore, the kind of work people will be compelled to pursue as a consequence of technological replacement and unemployment might not be sufficiently accommodating. For example, perhaps workers forced out of their industries simply do not have the interest or talents necessary to adapt to any newly developed sectors; parallel to the previous example of the automaton, a manual labourer could probably as easily chop timber as weld metal, but it would be a perhaps too demanding and even unreasonable expectation that he or she leave such work altogether and learn to code instead. And where would wayward labourers go if the newly developed sectors became unsustainably saturated? These and related possibilities raise further concerns that pernicious features of current economic systems (e.g., inequality) could be exacerbated

with technological progress. The ‘godfather of AI’ and Nobel laureate, Geoffrey Hinton, powerfully characterizes the situation thusly:

We are talking about having a huge increase in productivity, so there is going to be more goods and services for everybody, so everybody ought to be better off. But actually, it is going to be the other way around, and it is because we live in a capitalist society. And so, what is going to happen is this huge increase in productivity is going to make much more money for the big companies and the rich, and it is going to increase the gap between the rich and the people who are going to lose their jobs... If the profits just go to the rich, that is just going to make society worse. (Nobel Prize, 2024)

A radical shift in policy – and even in institutional structures – could well be in order. It is true that we have encountered technological innovation before, but perhaps nothing quite like this. And even if we are able to adapt to some extent, we might not be able to adapt as we have in the past.

Of course, our predictions are going to vary with our assumptions and methodological choices.

If, according to our preferred methods and observations, we determine that an eventuality is highly unlikely, we might justifiably judge that allocating resources in anticipation thereof would be inefficient and a fortiori unethical, insofar as those resources could have been invested elsewhere and manifestly increased welfare. Some eventualities might, however, be of grave enough consequence that, if we cannot disqualify them outright, we ought to nevertheless have a contingency plan in the ready. To truly allay our concerns, the more effective strategy would be to find principled reasons why automation could not possibly result in the state of affairs that dystopian alarmists imagine. Two arguments readily present themselves: The first is an ontological claim to the effect that AI simply cannot perform some of the important tasks associated with human labour, and the second is an economic claim to the effect that even if such technology could be achieved, it would be irrational to implement it in such a way that significantly displaces human labour.

Understanding Ontology

In order to evaluate the first claim, we must first determine which tasks, if any, associated with human labour could not possibly be done in principle by AI. To date, it is already manifestly evident that many mechanical tasks can be automated, and the list of more intellectually demanding, call them ‘cognitive tasks’, grows year in and year out. All of these tasks fall under the category of what we call ‘weak AI’, which is the kind of artificial intelligence capable of performing very well-defined tasks with at least some degree of human oversight. This is contrasted with the notion of ‘strong AI’, otherwise called ‘Artificial General Intelligence’ (AGI), which is the kind of artificial intelligence that would be virtually identical to human intelligence. The question becomes, which tasks, if any, associated with human labour require intelligence of this second kind?

Now, when we think about what labour entails, we can deconstruct any given occupation into sets of tasks and skills, where the former are understood as that to be done and the latter as those competencies needed to perform tasks. Take, for example, caretaking: Caretakers must be able to: maintain records, which requires literacy skills (both traditional and digital); assist with domestic chores like cleaning, shopping, or facilitating health regimens, which require physical skills and sometimes special technical skills (like operating automobiles or other instruments relevant to specific industries); communicate with dependents, which requires soft skills (and which, in addition to linguistic skills, also require emotional intelligence, empathy, sound judgment, etc.). Often, caretakers can assume even more demanding roles, such as being moral educators, confidants, or companions.

Thus, to be a competent caretaker is to be able to engage in a variety of tasks using a diverse set of skills in creative ways. As this example clearly illustrates, labour is an incredibly complex phenomenon. But how much of it necessarily evades the potential of AI? That is, how much of this cannot be done in the absence of intelligence of a general kind?

Granted that at this admittedly nascent stage of technological development an entirely integrated machine has not been realized, it takes no great effort of imagination to conceive of a multifunctional automaton. All the ingredients are already there: The caring professions have already begun deploying robot assistants for everything from

executing precise surgical procedures, to running errands, managing records, maintaining clean environments, and even providing emotional support (Falcone, 2024; Yazar, 2025). Prima facie, the challenge appears to be merely the technical one of putting everything together. Nevertheless, even if all the requisite skills could be consolidated into a single machine, we might find it wanting in important ways. To be sure, such a machine might be able to perform mechanical and cognitive tasks – sometimes even better than its human counterparts – but it would not be able to do them in the way a human does. While this might not be an issue for some tasks, and for some it is indubitably an advantage, for the most important, uniquely human activities, it might be an insurmountable shortcoming.

The thing that is ostensibly missing, and a fortiori cannot possibly be instantiated in a machine, is authentic understanding. The famous ‘Chinese Room’ thought experiment formulated by Searle (1980) challenges the computational theory of mind upon which early AI was predicated and seeks to advance the thesis that syntax alone is not sufficient for semantics, or to put it otherwise, that we cannot move from purely formal symbols and operations to meanings. As the argument goes, suppose a monolingual English speaker is isolated in a closed room, equipped with nothing more than a set of materials which instructs how to manipulate symbols so as to produce coherent sentences in Chinese.

Messages in Chinese are anonymously delivered to the individual from outside the room through a small slot, and the individual follows the instruction materials, producing coherent responses and sending them back. To those on the outside, it appears that their interlocutor is a competent Chinese speaker, but in fact he is not; he is simply taking input, manipulating symbols by following a set of instructions, and then generating output. If this is indeed analogous to how computers operate, then they only have the appearance of understanding, rather than authentic understanding. For they no more understand the inputs and outputs than the hypothetical individual in the room understands Chinese.

And since human beings do have authentic understanding, the computational model must be false or otherwise incomplete. Human faculties consist of more than mere computation.

This thought experiment has generated vigorous debate that carries on to this day. Some maintain that while the individual producing the responses according to the script might not have understanding, the system as a whole nevertheless can be said to (cf. Copeland, 2002); others contend that if understanding cannot be attributed to the individual or the system, then it cannot be attributed to human agents either (or else if it can be attributed to one, it should likewise be attributed to the other) (cf. Dennett, 2013). For our purposes, the main question is, what is understanding's role in labour? Are there any tasks that an automaton could not adequately perform without understanding in some deep sense?

Before we answer such questions, we must first elucidate what exactly is meant by understanding. As an intuitive starting point, we might simply claim that it is reasonable to attribute understanding to an agent as long as it demonstrates behaviours associated therewith: If one is given a command and carries out the task appropriately; if one is asked a question and produces a plausible response; if one can pose a relevant question on a topic; etc., then it would seem that for all intents and purposes, and as far as we can possibly know, the agent understands. This is essentially the idea behind the Turing Test in all its iterations. If an automaton can interact as well as any human agent in the environment, then what exactly marks the insuperable difference? If it is something radically subjective (something that 'it is like' to be the thing in question, or a 'beetle' in a box to which only one has access), we might doubt our ability to determine the presence of understanding from without at all, whether in a human or a machine (Nagel, 1974; Wittgenstein, 2009). And if understanding amounts to overt behavioural demonstrations, then what before seemed to be merely appearance becomes less so.

Therefore, we would need a rather different conception of understanding in order to maintain a categorical difference between human and artificial intelligence. Dreyfus (1979), channelling Heidegger and Merleau-Ponty, provides just such an account: What a supposedly intelligent machine lacks is what we might call situatedness. Human beings necessarily find themselves always already in contexts, through which, and only through which, the world makes sense. When we understand things, it is not a matter of assembling impersonal, disparate bits of data, analysing them according to prefigured rules, and then operating accordingly, but seeing things as they are to us, encountering

them in ways that are conditioned by idiosyncratic histories and interpretations (which inform and work upon each other), and making uncertain but hopeful choices:

Our sense of the situation we are in determines how we interpret things, what significance we place on the facts, and even what counts as facts for us at any given time. But our sense of the situation we are in is not just our belief in a set of facts, nor is it a product of independent facts or context-free features of our environment... We never get into a situation from outside any situation whatsoever, nor do we do so by means of context-free data. (Dreyfus, 1989, pp. 43-44)

Locating oneself in a personal and global narrative, evaluating and feeling certain ways about the characters and events that populate it, choosing to attend certain of them more or less or rather than others – all of which mutually determine the ways the experiences unfold and continue ever unfolding in a sprawling hermeneutic circle – this is the situation of the human being. And this, presumably, is precisely what the machine lacks. A fully integrated AI might be able to recount all the existing philosophical and scientific scholarship ever recorded, but could it take an interest in any of it? Could it find itself at stake in anything? Could it pose itself an original question that would prompt innovation and change the way it relates to what it pretends to know and shape what it might want to know in the exploration of unfamiliar terrain? The questions that it makes sense for us to ask and the projects to which we choose to dedicate ourselves emerge out of a background of meaning that is not of a purely formal nature. And in the absence of this situatedness, an agent remains suspended in a vacuum, as it were, paralyzed from the lack of sense needed to inspire it to move in a particular direction (and in that particular direction rather than another). “A glaze cleansed of everything past does not see things as they truly are; it sees precisely – nothing” (Reid, 2019, p. 46). To understand is to have a sense of situations, to contemplate and be engaged with this complicated and interconnected world as it discloses itself to us through time.

But how did we find ourselves in this situation, and is it really impossible for a machine to be situated? After all, there was a time before homo sapiens, and a time before collective and individual stories began to be written. Likewise, AI has a history following the progression from mechanization to digitization, and it could also presumably act as if it had a history. In this respect, the parameters appear similar for biological and artificial

intelligences. A human agent has and acts as if she has a history; an automaton has and can act as if it has a history. If there is a meaningful difference, we must probe the qualification implied in the modal verb 'can'. Whereas a human agent in fact has a biography which she also acts as if she has, any biography that an AI would be programmed to act as if it had would amount to a fabrication. The human agent in fact had a mother, experienced the joy of success and the disappointment of failure, nursed the wounds of a broken heart, impacted others, forged enriching friendships. And all these experiences inform and influence the way that she acts in the world that she encounters in every new situation. Meanwhile, the machine might be made to act as if it had similar experiences while in fact having none. Probing such a machine, it could no doubt recount a persuasively rich narrative of historical development, childhood memories, thwarted intentions, and future hopes. And all of it would be as artificial as the intelligence itself. Something about the unreality of these experiences might dispose us to reject the ontological claim that AI can genuinely find itself situated and have understanding in the same deep sense that a human agent is and has. However, we might consider analogous cases in human agents in which we might be hesitant to discount unreal experiences. For example, the 'alters' of those suffering dissociative identity disorder (DID) or persons in the throes of dissociative fugues do not in fact have the biographies they recount and feel as though they actually do; in these conditions, they display all the other complex faculties and capacities of 'real' persons – they believe certain things to have happened to them, have impressions of states of affairs, and evaluate and feel ways about things as consequences thereof, i.e., they behave as situated and understanding agents. The unreality of their experiences would not permit us to treat them as if they were not worth acknowledging, or worse still, imply that those identities do not remain morally considerable beings. If we want to deny them legitimacy on the basis of things actually having been the case, then we would need to demonstrate how that, and only that, is determinant of authentic situatedness and understanding, rather than the confluence of everything else associated therewith absent actual experience. On the other hand, if acting as if one was situated and understands is sufficient, then the categorical difference between human and machine again begins to fissure.

What matters for this discussion is whether it would be enough for an AI acting as if it was a situated, understanding being to adequately perform activities for which we think situatedness and understanding are necessary. We turn back to the example of the caring professions. Of the tasks with which caretakers are charged, perhaps the most challenging are those related to interpersonal interactions (i.e., those requiring soft skills). Breaking bad news to loved ones, offering emotional support or moral tuition, earning the trust of others, etc. are all highly delicate matters that we might regard as quintessentially human (i.e., tasks associated with intelligence of the second kind). In order to successfully navigate these sorts of situations, one must have a sense of them in the robust way heretofore elaborated. Empathy is comforting because we know that she who empathizes understands our experience, and not just in a descriptive way. We are receptive to advice because she who offers it is someone we trust, who has relevant expertise not just from erudition but also through lived experience. If we were to receive the same kind of support from an AI that has merely been programmed to behave as if it had lived experience to which it could appeal in order to offer empathy and advice, we might not be so receptive and comforted. It might strike us as fraudulent. But then, we also routinely connect with fictions. Are the lessons we take from *Antigone* and *Hamlet* less substantive because their ontological status is contentious? Or are they situated, understanding beings only as long as the covers of their tomes remain open? And if they are not, would their appearance as such for the duration of their stories not still have a lasting effect? If we take them to be ‘unreal’ or only real for the duration of their stories, and what we learn from them nevertheless impresses itself upon us lastingly, then why should we discredit an AI whose situatedness and understanding amounted to a fiction?

None of this is to say that there is no difference between something actually having been the case and something only imagined having been the case. If a lover acted as though she had betrayed her beloved, and the beloved, through her living as though that was the case, believed himself to have been betrayed, even though the betrayal never actually occurred, then we might say that they are both living under an illusion, even while pragmatically the betrayal is as real as if it had actually happened. We would consider them obstinate at a minimum if after having learned that the betrayal was an illusion they

insisted on carrying on believing it. The point, however, is not whether or not they are mistaken, but only that they have a sense of a situation, whatever it happens to be.

Herewith the ontological gap might not yet be closed. It is not clear that acting as if one is a situated, understanding being is equivalent to being a situated, understanding being. However, it might be enough to convincingly display situatedness and understanding in order to perform those tasks in which they are required. AI can already perform mechanical and cognitive tasks, and it seems possible in principle that it could even perform the most human of tasks. As long as this remains an open possibility, we cannot rest assured that human labour is unquestionably secured.

Economic (Ir)rationality

We can now evaluate the second argument, viz., that even if it is possible to develop AI such that it could perform all tasks relevant to human labour, it would be irrational to implement it to such an extent that it would significantly displace it. According to standard economic models, rational agents act so as to maximize utility. This principle offers a fairly straightforward protocol for firms considering automation: A task should be automated if doing so would result in lower input costs compared to human labour (since lowering input costs would effectively translate to higher profits, i.e., more utility). It follows, therefore, that if automating a significant portion of tasks would be more costly than hiring human labour, then it would be irrational to invest in automation. But would this case ever obtain, and if so, should we expect it to endure?

The premise upon which this argument depends is the empirical one, viz., that either or both the initial investment in or maintenance of labour-saving technology in the form of AI would in fact amount to greater costs for a firm than human labour. Estimates of what it would cost to achieve and deploy AGI (or something sufficiently comparable) at scale are notoriously dubious and speculative at best. But in any case, in order to advance the argument on these grounds, we would need to reject the premise that economic optimists offered earlier, i.e., the historical claim that there is nothing unprecedented in the trends. Of course, new technologies are initially quite expensive, but as they develop, their costs tend to go down. Automobiles and personal computers, once luxuries reserved for the incredibly wealthy, are now ubiquitous. If we are convinced that past trends will

continue, then it would seem we should have no reason to suppose that the projected costs of AI will pose a serious obstacle, at least not in perpetuity. As AI becomes more integrated into our economy and society, we should expect costs to become less prohibitive, as has always been the case with novel commodities. To reject this would be to open the possibility that something economically unprecedented is couched in this new frontier, which the economic optimist wants to simultaneously reject, rendering the position inconsistent. It is also worth noting that some of the most ambitious AI enthusiasts apparently have no concern about costs anyway. Sam Altman, CEO of OpenAI, one of the leading organizations in this space, expresses the sentiment, “Whether we burn \$500 million a year or \$5 billion—or \$50 billion a year—I don’t care, I genuinely don’t... As long as we can figure out a way to pay the bills, we’re making AGI...” (Hetzner, 2024).

To carry the argument through to its logical conclusion, let us suppose that the technology is both possible and economically expedient. If a firm is to remain rational, then it would be compelled to substitute its human workforce with fully capable AI, since failing to do so would result in lower utility. But in order to determine whether such a policy would actually be rational, we have to consider what the wider effects of significantly displacing human labour would be. The labourers that face technological unemployment are at the same time the consumers to whom firms intend to sell their goods and services; and as long as their finances are a function largely (and in some cases exclusively) of earned income, it is only too obvious that rendering the workforce inert would be to extinguish the purchasing power of a large portion, if not all, of the consumer base. Consequently, firms would have no one to whom to sell their goods and services, their profits would disappear, and both individual and aggregate utility would plummet. Therefore, significantly displacing the human workforce would be irrational.

Have we any reason to suppose that private firms would carry out this sort of holistic calculation? Again, if history is any indication, we might be wary; simply witness the behaviours of monopolistic robber barons and the very need for anti-trust law. If self-defeating self-interest is not an inherent feature of certain economic systems, we would be forgiven if we suspected that it is nevertheless a feature of the psychology of unscrupulous entrepreneurs who appear to have an insatiable appetite for acquisition. As long as rogue economies and governments, along with the actors that remain all too eager

to exploit them, cannot be reliably disqualified even by strong normative principles, we need some kind of systemic mechanisms to safeguard against them.

Basic Income as a Remedy to Technological Unemployment

Insofar as significant technological unemployment remains a realistic possibility, how should we respond? It would seem that existing or potential labour-centric apparatuses (e.g., unemployment insurance or a Negative Income Tax) would likely not be able to meet the demands of an increasingly post-work society, since the labour upon which they are predicated is precisely what would be missing. Naturally, then, we would need some program that is not dependent upon labour to function, and the candidate that appears to be uniquely suited for such a situation is basic income. This is because unlike other programs, basic income is unconditional and universal (that is, available to all citizens of a polity without any stipulations for its receipt), and in a society in which work becomes incredibly scarce or disappears altogether, it would not be viable to place work or wealth conditions on the receipt of resources.

We might still question both of these features, however. Granted that economic contributions to social welfare programs might become obsolete in a post-work society, we might nevertheless ponder noneconomic features thereof, e.g., solidarity. The ideas of desert and reciprocity that run deep in the psychology of social beings might be assumed to persist even in a society in which the economy does not depend upon human labour for production and remuneration, challenging the notion of unconditionality in any prospective basic income program. Susskind anticipates that in such a society, the question would become one of contributive justice, rather than distributive justice (Susskind, 2020). That is, while material provisions will be secured, people will still have the intuition that everyone ought to contribute to society in some way in order to have a share in its products. The failure to satisfy these intuitions could undermine solidarity, leading to social fragmentation. A basic income that is unconditional and universal might adequately respond to problems of distribution by securing the material needs of members of society, but it might be inadequate to respond to problems of contribution. Susskind, therefore, proposes a 'conditional basic income' (CBI), according to which people would be required to make noneconomic contributions to society in order to

receive benefits. “If some people are not able to contribute through the work that they do, then they will be required to do something else for the community instead; if they cannot make an economic contribution, they will be asked to make a noneconomic one in its place” (Susskind, 2020, p. 192). The kinds of noneconomic contributions could be variable, e.g., community education projects or creation of cultural works. The intuition that people should reciprocate and be deserving must be satisfied in order to preserve social solidarity, and this can be done by requiring noneconomic contributions as a condition for the receipt of benefits.

But imposing such a condition seems to imply that should one not fulfil it, then she is essentially condemned to suffer. In a society such as has been imagined, where economic contributions are no longer functional, noneconomic ones exhaust all else that could conceivably be demanded.

Susskind's proposal is, therefore, not categorically different than existing social welfare programs that compete with an unconditional and universal basic income. Existing conditional programs stipulate that one should be able and willing to work or else be legitimately indisposed (from age or disability); those who are able but unwilling are not entitled to any benefits. Presumably, the same allowances would remain in a post-work society; we would not expect anything from the elderly or lame, but we would still be demanding of the able but unwilling. In a labour-centric society, they are compelled to make economic contributions or suffer. In a post-work society, they are compelled to make noneconomic contributions or suffer. Though the affix changes, the result remains the same. In making distribution contingent upon contribution, the problems of both never abate. Imposing noneconomic contributions when no other contributions could be made would be to satisfy a psychological imperative rather than a moral one, to prioritize an intense but contingent preference over material and ethical necessity.

This is not to say, of course, that the psychological cannot be at once moral. Dissatisfaction, the thwarting of the will, the absence or erosion of special relationships, etc. all inspire moral action and therefore cannot be dismissed as illusory or otherwise insignificant. The point, however, is that the belief that one should not be entitled to the products of society unless she has contributed thereto cannot take priority when the capacity to contribute has been severed. The satisfaction one feels at the principle of

reciprocity being unnecessarily fulfilled pales in comparison to the privation to which one necessarily condemns others in order to feel it.

We might also have reservations about commodifying creative and humanitarian, i.e., noneconomic, activities. Engaging in community or cultural work is often considered a virtue in itself and subjecting it to economic interests could pervert its nature. For example, the over justification effect, in which intrinsic motivation is decreased or even extinguished by the introduction of extrinsic incentives, can demotivate people from engaging in these important projects. The value of this kind of work comes from the work itself and the subjective feelings we have about it. Making it a condition for economic benefits threatens the nature and performance thereof, which could lead to the very social fragmentation a CBI is meant to preserve.

Conclusion

A fully technological society in which human labour is completely replaced by AI remains something of a utopian fantasy. Should it ever be the case that the problems of scarcity and distribution were comprehensively resolved by technological innovation, then our economic systems would be radically transformed, as we might not even need to concern ourselves with the threat of privation at all, nor the conventions of exchange necessitated thereby. It is difficult to imagine non-trivial reasons to exclude anyone in a world cultivated by our artificial progeny that wants for nothing. In a near-full technological society, in which resources fall into the hands of the few still gainfully active members of society, redistribution of society's products would have to be a function of something other than the labour to which the disenfranchised many no longer have recourse. Though these circumstances seem like distant and uncertain possibilities, it would be as irresponsible to dismiss them on such grounds as it would be to pass the ecological buck to future generations. The foregoing has attempted to render it plausible that even if AI does not reach the level of true AGI – that is, intelligence on the level of that of humanity, with authentic understanding and a sense of its place in the context of a world and its complex histories – it could well be possible in principle that all relevant tasks great and small associated with human labour can be satisfactorily performed thereby. There may be something unprecedented in these new frontiers. Our rationality can betray us. And if

things cannot be resolutely supposed to carry on the same as they always have, then we must be prepared as we have never been.

Bibliographic references

Arntz, M., Gregory, T., & Zierahn U. (2016). The risk of automation for jobs in OECD countries: A comparative analysis. *OECD Social, Employment and Migration Working Papers*, 189, 1-35.

Copeland, B. J. (2002). The Chinese room from a logical point of View. In J. Preston & J.M. Bishop (Eds.), *Views into the Chinese room: New essays on Searle and artificial intelligence* (pp. 104- 122). Oxford University Press

Dennett, D.C. (2013). *Intuition pumps and other tools for thinking*. W.W. Norton & Company.

Falcone, S. (2024, March 21). 6 AI robots that are changing healthcare in 2024. Nurse.org. <https://nurse.org/articles/nurse-robots>

Ford, M. (2015). *Rise of the robots: Technology and the threat of a jobless future*. Basic Books.

Frey, C.D., & Osbourne M.A. (2017). The future of employment: How susceptible are jobs to computerization. *Technological Forecasting and Social Change*, 114, 254-280.

Hetzner, C. (2024, May 3). OpenAI's Sam Altman doesn't care how much AGI will cost: Even if he spends \$50 billion a year, some breakthroughs for mankind are priceless. Fortune.com. <https://fortune.com/2024/05/03/openai-sam-altman-microsoft-agi-artificial-generalintelligence-costs/>

LSE. (2025, March 11). Is AI really taking our jobs? The future of work explained | LSE research [Video]. YouTube. <https://www.youtube.com/watch?v=QqfunA6aSS4>

Nagel, T. (1974). What is it like to be a bat? *The Philosophical Review*, 83(4), (1974), 435-450.

Nobel Prize. (2024, December 18). Nobel minds 2024 [Video]. YouTube. https://www.youtube.com/watch?v=1tELIYbO_U8

Reid, J. (2019). *Heidegger's moral ontology*. Cambridge University Press.

Searle, J.R. (1980). Minds, brains, and programs. *The Behavioral and Brain Sciences*, 3, 417-457.

Susskind, D. (2020). *A world without work: Technology, automation, and how we should respond*. Metropolitan Books.

- Vox. (2017, November 13). The big debate about the future of work, explained [Video]. Youtube. <https://www.youtube.com/watch?v=TUmyygCMMGA>
- Wittgenstein, L. (2009). *Philosophical Investigations*: 4th Edition. (P.M.S. Hacker & J. Schulte, Eds.). (G.E.M. Anscombe, P.M.S. Hacker, & J. Schulte, Trans.). Blackwell Publishing, Ltd. (1953).
- World Economic Forum. (2025, January 7). The future of jobs report 2025. World Economic Forum. <https://www.weforum.org/publications/the-future-of-jobs-report-2025/digest/>
- Yazar, K. (2025, March 18). AI Nurse Robots that are Changing Healthcare. *TechTarget*. <https://www.techtarget.com/whatis/feature/AI-nurserobots-that-are-changing-healthcare>