

9-30-2022

Intelligence Augmentation: Human Factors in AI and Future of Work

Souren Paul
Northern Kentucky University, souren.paul@gmail.com

Lingyao Yuan
Iowa State University, lyuan@iastate.edu

Hemant K. Jain
University of Tennessee at Chattanooga, hemant-jain@utc.edu

Lionel P. Robert Jr.
University of Michigan, lprobert@umich.edu

Jim Spohrer
ISSIP.org, spohrer@gmail.com

See next page for additional authors

Follow this and additional works at: <https://aisel.aisnet.org/thci>

Recommended Citation

Paul, S., Yuan, L., Jain, H. K., Robert, L. P., Spohrer, J., & Lifshitz-Assaf, H. (2022). Intelligence Augmentation: Human Factors in AI and Future of Work. AIS Transactions on Human-Computer Interaction, 14(3), 426-445. <https://doi.org/10.17705/1thci.00174>
DOI: 10.17705/1thci.00174

This material is brought to you by the AIS Journals at AIS Electronic Library (AISeL). It has been accepted for inclusion in AIS Transactions on Human-Computer Interaction by an authorized administrator of AIS Electronic Library (AISeL). For more information, please contact elibrary@aisnet.org.

Intelligence Augmentation: Human Factors in AI and Future of Work

Authors

Souren Paul, Lingyao Yuan, Hemant K. Jain, Lionel P. Robert Jr., Jim Spohrer, and Hila Lifshitz-Assaf



9-2022

Intelligence Augmentation: Human Factors in AI and Future of Work

Souren Paul

College of Informatics, Northern Kentucky University, souren.paul@gmail.com

Lingyao (Ivy) Yuan

Department of Information Systems and Business Analytics, Debbie & Jerry Ivy College of Business, Iowa State University, lyuan@iastate.edu

Hemant Jain

Gary W. Rollins College of Business, The University of Tennessee at Chattanooga, Hemant-jain@utc.edu

Lionel P. Robert Jr.

School of Information, Robotic Institute, University of Michigan, lprobert@umich.edu

Jim Spohrer

ISSIP.org, spohrer@gmail.com

Hila Lifshitz-Assaf

Warwick Business School, Warwick University, Hdiginnovation@gmail.com

Follow this and additional works at: <http://aisel.aisnet.org/thci/>

Recommended Citation

Paul, S., Yuan, L., Jain, H., Robert, L. P., Spohrer, J., & Lifshitz-Assaf, H. (2022). Intelligence augmentation: Human factors in AI and future of work. *AIS Transactions on Human-Computer Interaction*, 14(3), pp. 426-445.

DOI: 10.17705/1thci.00174

Available at <http://aisel.aisnet.org/thci/vol14/iss3/6>



Intelligence Augmentation: Human Factors in AI and Future of Work

Souren Paul

College of Informatics
Northern Kentucky University
souren.paul@gmail.com

Lingyao (Ivy) Yuan

Debbie & Jerry Ivy College of Business
Iowa State University
lyuan@iastate.edu

Hemant Jain

Gary W. Rollins College of Business
University of Tennessee Chattanooga
Hemant-Jain@utc.edu

Lionel P. Robert Jr.

Professor School of Information
Robotic Institute
University of Michigan, USA
lprobert@umich.edu

Jim Spohrer

ISSIP.org
spohrer@gmail.com

Hila Lifshitz-Assaf

Warwick Business School
Warwick University
Hdiginnovation@gmail.com

Abstract:

The availability of parallel and distributed processing at a reasonable cost and the diversity of data sources have contributed to advanced developments in artificial intelligence (AI). These developments in the AI computing environment are not concomitant with changes in the social, legal, and political environment. While considering deploying AI, the deployment context and the end goal of human intelligence augmentation for that specific context have surfaced as significant factors for professionals, organizations, and society. In this research commentary, we highlight some important socio-technical aspects associated with recent growth in AI systems. We elaborate on the intricacies of human-machine interaction that form the foundation of augmented intelligence. We also highlight the ethical considerations that relate to these interactions and explain how augmented intelligence can play a key role in shaping the future of human work.

Keywords: Intelligence Augmentation, Intelligence, Artificial Intelligence, Human-AI Interaction, AI Ethics, Future of Work.

Fiona Fui-Hoon Nah was the accepting senior editor for this paper.

1 Introduction

The artificial intelligence (AI) field emerged in the 1950s with the expectation of emulating and automating various aspects of human intelligence. Expert systems and initial forms of intelligent systems received considerable attention in their initial years (Nah et al., 1999; Nah & Benbasat, 2004). The euphoria somehow became more tempered with the realization that, except for automating routine tasks such as credit approval and fraud detection, AI could not tap into many other areas of human intelligence. Eventually, the “AI winter” began in the mid-1990s (Jain et al., 2021). With recent success in machine learning, new hopes and expectations have resurfaced (Faraj et al., 2018; Kittur et al., 2019). The concept of machine learning was proposed many decades ago. It enables computers to learn on their own and uses data and algorithms to emulate human learning processes. One type of machine learning methods is supervised learning, where machines learn to map labeled outputs to inputs. These methods have been applied in various ways, from classifying emails to detecting fraudulent transactions. These applications are based on function approximation in which a task is embodied in a function and the learning process focuses on improving the function’s accuracy. A major area of supervised learning is deep learning, which refers to a neural network of three or more layers. Applications of deep learning include intelligent digital assistants and self-driving cars. Another type of machine learning is unsupervised learning in which unlabeled data is analyzed and clustered to identify hidden patterns. No human intervention is necessary for these learning methods. Common applications include social network analysis and market segmentation analysis. A third type is reinforcement learning in which intelligent agents learn in an interactive manner, using feedback from their own actions. Reinforcement learning are commonly applied in computer games and robotics.

AI applications are used to improve business efficiency (e.g., recommendation systems, chatbots), assist in daily lives (e.g., voice command-based digital assistants), and automate complex operations (e.g., autopilot in vehicles) (Brynjolfsson & Mitchell, 2017; Furman & Seamans, 2019; Kellogg et al., 2020). In recent years, we have seen tremendous growth in machine learning, which the changing computing environment has facilitated (Jordan & Mitchell, 2015). Both the availability of parallel and distributed processing at a reasonable cost and the diversity of data sources have contributed to newer developments in machine learning. However, changes in the computing environment have not been accompanied by changes in the social, legal, and political environment. Some have raised concerns about the privacy of data that are used to train machine learning systems, about trust and biases in machine learning algorithms, and about using AI to fabricate fake news and deceive the masses (Wang et al., 2018). In this research commentary, we highlight some important socio-technical aspects that are associated with the recent growth in AI systems.

The recent growth in AI has also raised a debate about augmentation versus automation by machine intelligence (Raisch & Krakowski, 2021; Lebovitz et al., 2022). Scholars describe human-AI augmentation as an expansion of expertise or knowledge where humans and machines “combine their complementary strengths” to “multiply their capabilities” (Raisch & Krakowski, 2021, p. 6). Through this expansion, human-AI augmentation is expected to positively impact organizations through superior performance or improved efficiency (e.g., Brynjolfsson & McAfee, 2014; Daugherty & Wilson, 2018; Davenport & Kirby, 2016). In addressing these views, scholars have highlighted the difference between artificial intelligence and augmented intelligence (Zhou et al., 2021). Augmented intelligence refers to enhancing and elevating a human’s ability, intelligence, and performance with help from information technology. A key aspect of augmenting intelligence involves human-machine collaboration in which “machines perform what they do best (e.g., computing, recording, and doing routine, repetitive work) to aid humans in doing what humans do best (e.g., abstract reasoning, creating, and making in-depth discoveries about people and the world)” (Zhou et al., 2021, p. 245; Jain et al., 2021). The relationship between AI and the augmentation of human intelligence is dynamic and complex and needs to be investigated and problematized (Raisch & Krakowski, 2021; Lebovit et al., 2022). As tasks are automated with AI, the horizon of human intelligence shifts. Thus, the question of the augmentation of human intelligence depends on context. Another important issue to consider is the end goal of intelligence augmentation. In this paper, we elaborate on the context and the end goal of augmentation as both issues have surfaced as significant for the design, use, and impact of AI on professionals, organizations, and the future of work. In doing this, we draw on the discussion in a panel at the Hawaii International Conference on System Sciences (HICSS) in January 2022.

This research commentary proceeds as follows: In Section 2, we emphasize the importance of intelligence augmentation (IA) context and end goals. In Section 3, we elaborate on the human-centric approach to AI (HAI). Effective HAI systems consider human factors and ethical concerns regarding man-machine systems. We also highlight the ethical considerations that relate to these interactions. In Section 4, we explain how augmented intelligence can play a major role in shaping the future of human work. In Section 5, we discuss

research topics for intelligence augmentation and human-AI interaction. Finally, in Section 6, we conclude the research commentary.

2 Intelligence Augmentation: Context and End Goal

2.1 Augmented Intelligence in Context

We can define augmented intelligence as “enhancing and elevating human ability, intelligence, and performance with the help of information technology” (Zhou et al., 2021, p. 245). Augmented intelligence informally means computers and humans working together, by design, to enhance one another such that the intelligence of the resulting system improves. IA can pool the joint intelligence of humans and computers to transform individual work, organizations, and society (Jain et al., 2021). From the sociotechnical system perspective, one can also view augmented intelligence as a sociotechnical extension factor on human capabilities. Like any other tools that humans have developed and used in history, augmented intelligence enables individuals to become more capable, while individuals collectively make organizations and businesses more capable and society a better place. The sociotechnical systems approach recognized the relationships among people, technology, and the environment (Orlikowski & Scott, 2015). Augmented intelligence as a sociotechnical extension factor under this approach moves beyond the aspect of technology and the algorithms behind it. It also involves human and environmental factors such as supply chain management, human resources, organizational culture, social context, and more.

Taking augmented intelligence out of its application context could yield misinterpretation and technology misuse (Lebovitz et al., 2021). In many professional contexts, such as radiology, the evaluation of AI tools is challenging as there is a high degree of uncertainty involved. Many “ground truth” measures are based on knowledge claims that lack strong external validation (Lebovitz et al., 2021). It is important to investigate in various contexts how professionals respond to a new technological force that challenges the professional jurisdiction and knowledge boundaries of an existing profession, such as how professionals enact professional identity work (Lifshitz-Assaf, 2018; Tripsas, 2009) or knowledge boundary work (Anthony, 2021; Barrett et al., 2012; Levina & Vaast, 2005).

As another example, in the autonomous vehicle context, the first priority while deploying and using such vehicles is safety. Both human and machine intelligence have benefits and limitations. Under which conditions human intelligence should rely on machine intelligence and under which conditions human intelligence needs to be in absolute control depends on the scenario. Machine intelligence may not be suitable for making decisions in scenarios involving morality. The trolley problem of killing one person to save five is an example. Facing those scenarios, machine intelligence should only be used to facilitate human intelligence rather than take control. On the other hand, humans experience physical limitations. After 10 hours of driving, for instance, human intelligence may be compromised by fatigue. Let the machine intelligence be in control when humans are not fully alert would potentially avoid catastrophes.

The wide adoption of online platforms has brought up new ways to create and share new ideas, such as crowdsourcing, super minds, and democratizing innovation (Von Hippel, 2006). The notion of user-centered innovation has challenged the dominance of the expert model in the innovation and knowledge creation domain. With rapid advancements in data mining and analytics, machine performance can exceed human performance in many ways. With the help of machine intelligence, crowdsourcing could generate better results than when a task is done by experts, which has spurred heated discussions on whether organizations still need to rely on experts for innovation and problem-solving. However, researchers point out that ideas favoring one particular model carry an underlying assumption that one model, either crowdsourcing or expert, is a total substitute for another and, thus, that institutions should choose either experts or crowds. However, Kittur et al. (2019) revealed that combining all three different roles (i.e., experts, crowds, and AI) would take advantage of all three parts and overcome each part’s disadvantages in scalability, complexity, and fixation for better results. Experts can bring the complexity of the thinking, crowds can help with communication and narrow down the fixation of the domain expert, and AI can help with scalability.

2.2 The Goal of Augmented Intelligence: Automation vs. Augmentation

One may not help but wonder what the end game or goal of augmented intelligence is. Academics and practitioners have spent much effort on using machines to replicate what humans can do (also known as automation). Many researchers believe the goal of augmented intelligence, at least for some contexts, is having machines replace humans. The stage when machine intelligence augments human intelligence is a

transitive rather than a definite end goal. Considering the context of autonomous vehicles, having complete self-driving cars (i.e., achieving automation) is the goal. Such automation would free humans from the labor of driving and eliminate human errors caused by cognitive and physical limitations. Currently, autonomous vehicles are in the transitive stage of augmenting humans rather than replacing humans due to technological limitations. We cannot wholly rely on autonomous vehicles for many reasons, such as complex moral scenarios, which make it challenging to embed ethics into automated systems.

Taking augmented intelligence out of a specific context and looking at it in human history, we can gain a different perspective. Let us reflect on human history. During the early 18th century, thanks to new farming tools and methods, the Neolithic revolution (also called the Agricultural Revolution) increased the production of food, which resulted in a huge increase in population. Cities emerged, which set the stage for the industrial revolution. During the industrial age, steam engines replaced hand tools in manufacturing. Mechanization has significantly improved production but at the cost of millions of workers losing their jobs at factories. Society adapted to this change. Workers learn how to operate machines and become more specialized. In the mid-20th century, thanks to the invention of computers and later the Internet, we stepped into the Information Age. Silicon Valley became the center for high technology and innovation in the world. Nowadays, having successful online stores has become a requirement for businesses to survive. People feared that e-commerce would replace brick-and-mortar stores, especially small and local businesses. However, as technology has developed in the past few decades, we are on our way to achieving a new balance for both to co-exist.

Similarly, nowadays, augmented intelligence with automation as the goal has spurred the fear of losing jobs to machines from the public. But if we had learned from history, we would believe that any “new society-in-the-making” (Castells, 1999) would arrive in a gradual manner along with new relationships among the economy, state, and society. One change we foresee in the future is the increasing demand for innovative minds. Augmented intelligence opens the door for more people to have “digital employees”. Everyone can become a manager. Self-driving cars may cost some taxi drivers their jobs but make others the owner of a fleet of self-driving taxis.

3 Human AI Collaboration

Many IA technologies have expanded from being merely supportive tools, such as calculators, to having a more collaborative role. Human-AI collaborations have taken place in many domains. In healthcare, human-AI collaborations can assist overworked medical professionals and help provide quality healthcare (Lai et al., 2021). Human-computer vision (CV) collaboration has been employed in remote-sighted assistance to help visually impaired people (Lee et al., 2022).

Humans and AI have complementary strengths. Tasks that involve complex and equivocal processes may present challenges to humans with limited cognitive ability. Solving tasks with uncertain outcomes calls for examining a large volume of historical data, which increases the level of information processing. This may pose a problem for humans having limited cognitive abilities. On the other hand, tasks that involve intuitive reasoning and reliance on memory, especially episodic memories, may not be automated effectively using AI systems. Thus, it is necessary to rely on the complementarity of humans and AI in performing tasks that are characterized by uncertainty, complexity, and equivocality. AI, with a high level of information processing and analytical capacity, can help humans to meet the cognitive loads of performing tasks that are complex yet structured. Humans can offer a more holistic, intuitive approach to dealing with uncertainty and equivocality in unstructured and semi-structured tasks. Thus, the partnership between humans and AI can form the foundation of augmented intelligence (Jarrahi, 2018). Table 1 shows how the combination of humans and AI can address the complexity, uncertainty, and equivocality of organizational tasks.

Closely related to the concept of complementarities of human and machine intelligence is the goal of human-centered AI (HAI). HAI research strategies focus on AI to enhance humans rather than replace them. Xu (2019) proposes an extended framework of HAI that has three components:

- **Technology that reflects human intelligence**, which provides good complementarity of humans and AI.
- **Ethically aligned design**, which focuses on providing AI solutions that are fair and that avoid discrimination. We discuss ethical issues in Section 3.2.

- **Human factors design**, which ensures that AI solutions are explainable, comprehensible, useful, and usable. An explainable AI enables users to understand the parameters and algorithms used to arrive at AI solutions. For example, based on the medical history of a patient, an AI system may predict the possibility of having serious health problems (such as heart disease and cancer). An explainable AI will provide reasoning for the prediction. However, while a domain expert might understand the reasoning that an AI provides, other users may not. The goal of HAI is to improve comprehensibility so the common users can understand the solution and the reasoning for arriving at the solution. Another aspect of human factors design is the usability of AI which implies that an AI system should be easy to learn and use. This calls for designing an effective user interface.

Table 1. Complementary Human and Machine Intelligence (Adapted from Jarrahi, 2018)

	Complexity	Uncertainty	Equivocality
Human intelligence	Select data sources and models to explore different solutions. Choose the solution that meets the objective function.	Explore the riskiness of different situations. Apply intuition to select a solution that is less risky.	Explore diverse interpretations. Arrive at a consensus.
Machine intelligence	Curate and process diverse data. Analyze data.	Identify anomalies.	Analyze sentiment. Present diverse interpretations.

There has been considerable interest in explainable and interpretable AI. Contemporary AI tools, such as deep-learning algorithms, are often designed as “black boxes” to users, which makes it very difficult or even impossible to examine how the algorithm arrived at a particular output (Christin, 2020; Diakopoulos, 2020; Pasquale & Cashwell, 2015). While experiencing opacity and using “black box” technologies (e.g., cars or computers) is ubiquitous (Anthony, 2021), problems arise when one needs to integrate diverse knowledge claims into a single decision that a human expert can stand behind. This is the case for many scenarios of AI use for critical decisions, such as in medicine (Lebovitz et al., 2022), human resource management, and criminal justice, where opacity associated with AI use is particularly problematic (Christin, 2020; Van Den Broek et al., 2021; Waardenburg et al., 2018). The topic of black boxed AI and its implications is creating a heated debate in the field these days. While many researchers focus on developing “explainable AI” or “interpretable AI” (e.g., Barredo Arrieta & Del Ser, 2020; Hooker & Kim, 2019; Rudin & Radin, 2019; Samek & Müller, 2019; Teodorescu et al., 2021), some leading scholars (Cukier, 2021; Simonite, 2018) and AI designers believe there is no need for explanations. Some argue that explainable AI represents the lowest level of ‘bi-directional’ human-AI collaboration (van den Bosch & Bronkhorst, 2018). It does not elicit how AI initiates and sustains effective interactions, which involves detecting the conditions of interactions and triggering appropriate actions. This line of research suggests that at the next level, HAI should be able to reason human actions and intentions so that it can act in an adaptive, intelligent manner (van den Bosch & Bronkhorst, 2018). Trust in AI will evolve as the system becomes adaptive and intelligent. A recent line of research brings issues of opacity into use and shifts the analytical focus from what appears as an innate and fixed property of technology to the broad socio-material practice that produces opacity as a specific technology is used in a particular context (Lebovitz et al., 2022). This line of research focuses on the process of how AI opacity emerges in practice and how, in some cases, professionals can deal with it.

3.1 Trust in Human-AI Interactions

While HAI focuses on emulating and augmenting human intelligence, its success depends largely on effective human-AI interaction (HAI). Humans demonstrate different personalities while facing AI (Mou & Xu, 2017). Research has found users to be less open, agreeable, extroverted, and conscientious when interacting with AI than when interacting with other humans (Mou & Xu, 2017). Explanations help to increase trust in AI (Meske & Bunde, 2020; Zhang et al., 2021). The use of AI-based virtual assistants by human resources in organizations can engage employees at a very personalized level and enhance the climate of trust and fairness (Dutta & Mishra, 2021). Trust in human-AI systems is an important aspect of sustained HAI. Lee and Moray (1992) suggest four dimensions of trust in man-machine systems: 1) foundation, 2) performance, 3) process, and 4) purpose of the man-machine system. Foundation is the assumption of the natural and social order that makes other dimensions possible. Performance is the “expectation of consistent, stable, and desirable performance” (p. 124) of the system. The process is the underlying characteristic that governs performance. The purpose is the motive or intent of using the man-machine system. Later, Lee and See (2004) focused on automation and elaborated on the dimensions of

performance, process, and purpose. We draw on this literature to present the dimensions that are relevant for trust in HAI in Table 2.

Table 2. Dimensions of HAI Trust

HAI trust dimensions	Trust elements
Performance	• Competence of HAI ◦ Functional ◦ Human-AI interaction • Timeliness of solution in real-time HAI • Reliability ◦ Context-specific reliability
Process	• Openness • Consistency • Understandability • Predictability • Data integrity • Accessibility
Purpose	• Authorized responsibility for administering and using HAI • Intention of machines and users in HAI • Faith in HAI

Lee and See (2004) also discuss the importance of calibration of trust, which is the “correspondence between a person’s trust in the automation and the automation’s capabilities” (Lee & See, 2004, p. 55). Ideally, trust in HAI should match the capabilities of AI, which will help to achieve the optimum use of HAI. Neither over-trust (i.e., trust exceeding the capabilities of AI) nor distrust (i.e., trust falling short of the capabilities of AI) are desirable. Over-trust leads to abuse of AI and distrust will result in the underuse of AI. Trust calibration becomes necessary as AI is pushed towards higher levels of human-AI collaboration. When a system shows adaptive behavior, trust should evolve with it. HAI facilitates the dynamic process of trust calibration as it enables humans to “continuously experience, interrogate, and judge the functioning of the AI” (Van den Bosch & Bronkhorst, 2018, p. 9).

Another important consideration in HAI is the concept of situation awareness (SA), especially in time critical decision-making (Wei et al., 2020). SA refers to “the perception of the elements in the environment within a volume of time and space, the comprehension of their meaning, and the projection of their status in the near future” (Endsley, 1987, 1988). Jiang et al. (2002) suggest that users’ enactment of SA will mitigate some negative impacts of AI systems on user experience, improve human agency during AI system use, and promote more efficient and effective in-situ decision-making. Endsley elaborates the definition of SA by suggesting three hierarchical phases, which are: perception of the elements of the environment (Phase 1), comprehension of the current situation (Phase 2), and projection of future status (Phase 3). Intelligent agents are aimed at meeting the requirements of SA. Jiang et al. (2022) suggest how considerations of the SA perspective in designing HAI can help in alleviating tensions between 1) automation and human agency, 2) AI system uncertainty and user confidence, and 3) the system’s objective complexity and users’ perceived complexity. To overcome these tensions, it is necessary to increase users’ sense of control, boost users’ confidence by presenting relevant plans and projections of HAI, and reduce users’ perceived complexity in the system (Jiang et al., 2022).

3.2 The Ethics of Using AI Technology

Robots are a common embodiment of AI technology. The ethical and unethical use of robots is a key challenge. To be clear, the use of automation to replace humans is not new. What is unique about robots is that they are not necessarily eliminating roles done by humans but instead filling those roles. When this occurs, previous human-to-human interactions in the workplace are transformed into human-to-robot interactions. In essence, replacing the role with a robot rather than eliminating it necessitates the need to design a robot that can engage in social interactions with humans (You & Robert, 2022). To accomplish this, humanoid robots are being designed to encourage humans to anthropomorphize them in order to promote social interactions. For example, in healthcare settings, robots are being developed and deployed as care providers in the home and in hospitals to fill both medical and social interaction roles (Esterwood & Robert, 2021).

Questions about the ethical use of robots to fill in for humans have been largely avoided. An important question to ask is whether it is ethical or unethical to replace humans with robots. There is an ongoing debate about the benefits and cost of robots to our society but typically this debate concerns the fear of job loss (Eglash et al., 2020). There are clear and legit concerns about the ethics associated with displacing workers. However, another less studied problem is the implications associated with the dehumanization of our society. Dehumanization of our society occurs when we replace meaningful relationships between humans and humans with humans and robots. Ethically, the cost to replace human-to-human interactions with human-to-robot interactions in our society remains unclear. Is it ethical to encourage humans to build relationships with technology in the way they have become accustomed to building relationships with humans? Several scholars are beginning to express concerns and, in some cases, outrage at some actions (Liang et al., 2021; Shneiderman, 2020).

Ethical concerns around the dehumanization of our society are of particular importance for organizations. Relationships between humans in organizations can be characterized as one of two types of relationships: 1) work related and 2) non-work related or social. Yet, research in organizations has repeatedly demonstrated that social relationships are often the best predictor of employee outcomes (e.g., see Lykourantzou et al., 2021). Therefore, replacing robots with humans requires us to understand the cost of replacing human-to-human social interactions with human-to-robot interactions in our society.

3.3 Embedding Ethics in Technology to Complement Humans

The second ethical challenge is the embedding of ethics within technology. The idea is to develop robots that can make or help humans make ethical decisions. On the one hand, issues with totally relying on such systems to make ethical decisions independently of humans have drawn increasing attention. For example, questions about ethical driving decisions by autonomous vehicles have received increased attention (Robert, 2019). On the other hand, questions about whose ethics are to be embedded into such a system have not received the same attention (Robert et al., 2020a). Finally, it is very difficult at times to technically embed ethics into such systems.

One solution is to design hybrid systems that rely on both the robot and the human. These hybrid systems can be designed to allow both agents to complement one another (You & Robert, 2018). Although such systems help to address the technical problems of embedded ethics within the system, they do little to address whose ethics should be the normative standard. Relying on the human in the human-robot hybrid system to make the ethical decision ignores the role the robot has in helping the human make such ethical decisions and carrying out their corresponding actions. More specifically, a robot designed to help humans make and carry out fair actions defined by equality is likely to differ from a robot designed to be fair and defined by equity.

Nonetheless, such hybrid systems are rarely the end game. For example, the end game in the development of autonomous vehicles is not to have the human and the vehicle drive together but instead to develop autonomous vehicles that can eventually drive without human input or supervision. Therefore, the hybrid in the human-robot hybrid system will rarely be the goal and is always going to be a transitive situation. In many cases, the hybrid system is designed as a stop-gap approach to allow the human and the robot to work together long enough to figure out how to design a robot to work without the human. Ethical questions about when the robot is likely to be much more capable than the human persist and are not easily addressed (Liang et al., 2021; Robert et al., 2020a). For example, when do human's ad hoc problem-solving skills outweigh the robot's reliability and precision? Determining this handoff point is likely to be plagued by ethical questions.

3.4 Mitigating Bias in Human-AI Interaction

Organizations are increasingly employing AI systems to manage their workforce which has the potential to manage their employees efficiently and effectively. Unfortunately, these systems have been shown to be biased and may not be fair to their employees (Harini, 2018). AI bias often results from human biases encoded directly into the AI system or can result from learned behavior over time. Therefore, finding ways to mitigate those biases has become one of the most pressing issues in our society (Robert et al., 2020b). Below, we present a framework for mitigating bias in human-AI interaction based on Robert et al.'s (2020b) framework.

The first step in mitigating bias in human-AI interaction is to identify the bias. It begins with clear definitions of AI bias. Generally, AI bias can be viewed as decisions and/or actions that are an "unfair assessment in

favor of or against some person or group” (Robert et al., 2020a, p. 100). Robert et al. (2020b) identified three types of biases: distributive, procedural, and interactional bias. Distributive AI bias is the unfair allocation of outcomes such as pay and other resources. Procedural AI bias is the unfairness associated with a process employed to reach or decide the outcome. Interactional AI bias is unfairness in the degree of respect, dignity, and information shared.

The second step is drawing attention to AI bias. It requires the aid of systems designed to specifically identify and draw attention to biases, which could be accomplished through interfaces that can highlight potential problems (e.g., an interface that could alert the user when an AI hiring system’s selection criteria produces a list of applicants with no women or minorities). More sophisticated systems can work to make bias transparent through explainability. On the one hand, transparency helps to make the underlying AI mechanics visible and known to the employee, while explainability helps to describe the AI’s decisions/actions to the user in terms that they can understand. The third step would be to allow users a voice in the AI process. According to Robert et al. (2020b), voice provides users with an opportunity to communicate and provide feedback to the AI and is measured by input influence. Voice empowers users, allowing them to push back against the system. Finally, if harm is done, there should be a way to redress biased actions. Redressing AI bias focuses on taking actions needed to remedy or set right unfairness. Taken together, these steps help provide a framework for mitigating bias in human-AI interaction.

4 Impact of Artificial Intelligence and Intelligence Augmentation on Future of Work

There have been rapid advances in AI technology aided by significant improvements in processing and storage technologies, the availability of huge amounts of data, and advances in algorithms such as deep learning algorithms and reinforcement learning. These developments are starting to impact businesses in a significant way (e.g., Agrawal et al., 2018) and have significantly increased anxiety related to the number of job losses. The estimates vary from 10 to more than 50 percent of jobs lost due to the implementation of AI and related technologies. Specifically, these losses are expected to impact white-collar work and middle management. Frey and Osborne (2017) categorized tasks by their susceptibility to automation; linked these tasks to occupation, employment, and wage data; and found that 47 percent of US employment was at high risk of automation. One assumption embedded in the Frey and Osborne model is that all workers in the same occupational category face the same threat of automation. On the other hand, an OECD Report (Arntz et al., 2016) argued that there might be task variation between individuals within the same occupation. For example, managers of different firms may treat shop floor labor differently depending on the company culture. Thus, the OECD Report used individual-level data to predict how susceptible occupations may be to automation and found that only nine percent of jobs in the United States and across OECD countries will be highly susceptible to automation.

Thus, there has been significant disagreement on the true impact of AI on jobs. Acemoglu et al. (2022) studied the impact of AI on labor markets using online vacancies in the United States from 2010 onwards. They concluded that “while visible at the establishment level, the aggregate impacts of AI-labor substitution on employment and wage growth in more exposed occupations and industries is currently too small to be detectable” (p. 293). Other researchers have reached a similar conclusion; namely, that although AI will have a potentially huge impact on job losses, the impact so far has been minimal.

The actual impact of AI on jobs may depend on how and where the technology is implemented, its cost, and overall economics. The situation presents a golden opportunity for researchers and other industry leaders to guide AI and Augmented AI development for the benefit of society. We can draw some lessons regarding job displacement from the vast literature on automation and its impact on jobs. Acemoglu and Restrepo (2019) presented a task-based framework related to the impact of automation. Their framework considered that production requires various tasks that can be either allocated to capital or labor. The development of new automation technology potentially changes the allocation of tasks to factors of production, from labor to capital, thus impacting labor demand and productivity. The allocation decision is generally based on economics and potentially productivity and quality gains. The result of this trend has been the sharp slowdown of U.S. wage growth, especially in blue-collar work, in the last three decades (Acemoglu & Restrepo, 2019).

The impact of AI on jobs can be partly understood by its automation potential. Agrawal et al. (2019) argue that understanding the impact of AI requires comprehending the capabilities of the technology. However, the capabilities of AI technologies are a moving target. These capabilities change with every new

development. For example, most of the recent achievements in AI are the results of advances in machine learning. However, machine learning does not represent an increase in artificial general-purpose intelligence, which could substitute machines for all aspects of human cognition, but generally addresses one aspect of intelligence (i.e., prediction) (Agrawal et al., 2018). Based on a task-based framework, advances in prediction technology may affect labor in four ways: 1) substituting capital for labor in the prediction task; 2) automating the decision task, which may increase relative return on capital; 3) enhancing the labor where automating the routine prediction may enhance productivity in related decision tasks; and 4) creating new decision tasks since automation of prediction significantly reduces uncertainty, which may enable new decisions that were not feasible previously (Agrawal et al., 2019). Thus, in the longer term, automating prediction can make it more attractive to automate complementary decision tasks resulting in complete automation possibly because machines may be faster than humans. For example, machines may predict potential car accidents faster than human reaction time, which may make it more attractive to employ automated braking in cars.

There are a few tangible examples where machine learning-based prediction enhances complementary tasks not feasible without machine learning-based prediction. For example, in brain cancer surgery, to ensure that all cancerous tissues are removed, surgeons frequently end up removing more brain matter than necessary. An ODS medical device, which resembles a connected pen-like camera, uses artificial intelligence to predict whether an area of brain tissue has cancer cells or not, which helps surgeons in real time to decide which area should be removed. By predicting with more than 90 percent accuracy whether a cell is cancerous, the device enables the surgeon to reduce both type I errors (removing non-cancerous tissue) and type II errors (leaving cancerous tissue). The effect is to augment human capabilities to improve the overall outcome (Agrawal et al., 2019). Thus, it is difficult to assess the net effect of AI on labor even in the short run since multiple factors impact jobs that both increase and decrease demand. The net effect will vary across countries, industries, and applications. There is general agreement that it takes significant efforts to adopt technologies. Entrepreneurs and innovators take time to adopt new technologies, reconfigure existing work, discover new business processes, and co-invent complementary technologies (Bresnahan et al., 1996).

Brynjolfsson et al. (2018) argued that unleashing the full potential of AI will require unbundling of tasks in jobs and a significant redesign of the task content of jobs. They found that the focus needs to shift from full automation in most cases to redesign of jobs and reengineering of business processes. For example, to understand how drafting emails might affect different types of jobs differently, one can use the O*NET database, which offers detailed descriptions of the tasks involved in almost 1,000 occupations (<https://www.onetcenter.org>). This data includes a task described as “prepare responses to correspondence containing routine inquiries”. The executive assistant role (along with eight other occupations such as clerks, tellers, receptionists, and so on) includes this task. Executive assistants would typically draft possible responses for someone else to decide whether to send them, so a system such as Gmail’s smart reply can potentially fully automate the executive assistant’s decision. On the other hand, the executive assistant might use this technology but still retain the decision task of what to ultimately send to the senior manager for review. So, in the former case, AI replaces labor, while, in the latter case, it enhances labor. While the current level of technology suggests a human should remain in the loop for many jobs, it is plausible that, over time, artificial intelligence will improve, lead to full automation, and, thus, reduce the demand for labor (Agrawal et al., 2019).

Babina et al. (2021) studied the impact of a firm’s investment in AI technologies on product innovation and growth. They found a significant increase in investment in AI technologies, especially by larger firms. Their study reveals that AI investing firms see increased product innovations, sales, employment, and valuation of firms. Babina et al. (2022) found a change in the workforce composition of firms using AI technologies. The firms investing in AI technologies tend to move to the highly educated workforce with undergraduate and graduate degrees in STEM disciplines. Additionally, the organization tends to become flat with a larger number of junior workers and fewer workers in middle management and senior roles. Felten et al. (2021) linked advancement in different categories of AI to different types of abilities and used it to correlate advances in AI to actual changes to occupational descriptions.

Another major factor that will moderate the impact of AI on jobs is an increased movement toward augmented intelligence. In their editorial to a special section of *Information Systems Research* on augmented intelligence and the future of work, Jain et al. (2021) argue that humans always have the desire to overcome their physical and intellectual limitations by developing technologies and infrastructure, such as transportation systems to overcome limitations in travel speed and communication infrastructure to

overcome long-distance communications limitations. They argue that developments in AI should be guided by a desire to overcome the intellectual limitations of humans and not to replace human intelligence with machines. They suggest that intelligence augmentation infrastructure has the potential to pool the intelligence of human beings and computers to transform individual work, businesses, institutions, and even society in an unprecedented manner and at scale. Since the capabilities of computers and human brains are different, the best of these abilities can be combined to optimize the system. To understand the impact of augmented intelligence, we need to focus on the ability of humans and the ability of AI. Peterson et al. (2001) describe the core concept of human abilities. They propose a job classification system called O*Net that basically divides human abilities into 52 distinct groups and then groups them broadly into four categories: cognitive, psychomotor, physical, and sensory. The O*Net framework can be used to develop a conceptual model to study the complementary nature of human and AI capabilities and guide the augmentation of humans' intellectual capabilities. Figure 1 shows a framework for conceptualizing the future of work in an augmented intelligence world.

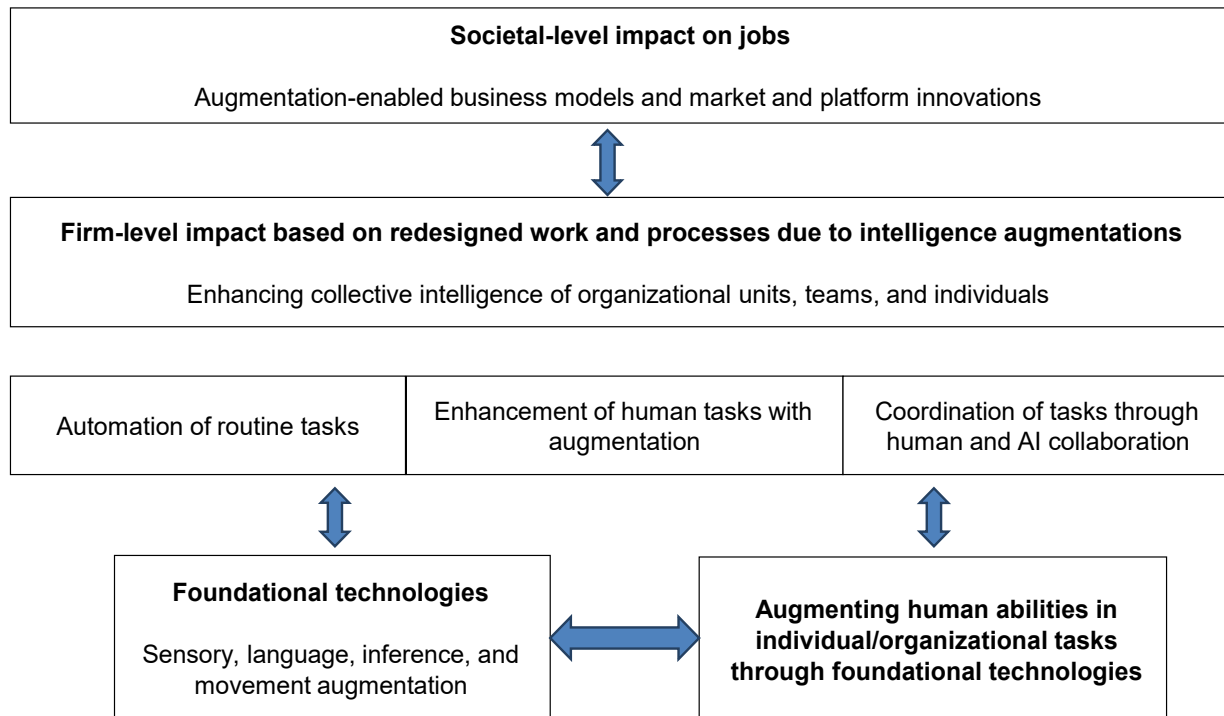


Figure 1. Conceptualizing the Future of Work in an Augmented Intelligence World

The foundation of augmented technologies, which are available in the computer area now, is that there are sensor technologies that essentially can sense what humans cannot sense. We have language technology, inference technology, and movement augmentation using robotic technology. These technologies, which we call foundational technologies, can be combined with human abilities (cognitive, psychomotor, physical, and sensory abilities) to redesign tasks at the firm level. Some of the routine tasks can be completely automated while some other tasks may be enhanced, and new tasks may be created which are done by augmented human capabilities. The coordination of automated and enhanced human tasks can be performed through human and AI collaboration. This will result in a firm-level impact on jobs based on redesigned work and processes. The aggregation of all firm-level impacts on jobs may help us arrive at the societal-level impacts on jobs.

5 Future Research Directions

In this paper, we discuss the augmentation of HAI, human factors and ethical considerations in HAI, and the future of work in the realm of augmented intelligence. Several research topics may be pursued in each of these areas. We list some topics in Table 3. This list is not comprehensive. It highlights some research topics that emanate from the discussions in the paper.

Table 3. Sample IA Research Areas and Topics

HAI research: sociotechnical areas	Sample research topics
Augmentation: context and goal	Performances measurements in various contexts, human perception toward AI
AI scalability and empowerment	Human-in-the-loop AI, computing, algorithm design
Human factors in HAI	Roles of human intelligence and machine intelligence in group settings and organizational settings Relationship between situation awareness and trust in HAI Relationship between situation awareness and privacy concerns
Trust in HAI	Factors lead to trust/distrust Trust calibration and levels of HAI
Bias in HAI	Discrimination (e.g., racial, gender) toward AI agents Bias in distinct levels of HAI and mitigation of bias
Future of work	The structure of work Allocation of tasks between AI and humans How allocation of tasks gets impacted by local economic and social conditions Collaboration between AI and humans Issue of ultimate control Ethics of using AI Training humans to work with AI systems Developing AI systems that can seamlessly work with humans AI with metaverse
AI and crowdwork	AI trends on crowdwork platforms
Skill and augmented intelligence	Reforms in education to manage AI

Additionally, if there is to be a science of intelligence augmentation (IA), then what would it primarily measure? One candidate measurement is the so-called "socio-technical extension factor" (Kline, 1995). For example, the distance at which two people can communicate simultaneously using spoken language was at most a few hundred feet throughout most of human history. However, with the technology of satellites and telephones, now two people can communicate easily using spoken language even if separated by thousands of miles. The socio-technical extension factor for this communication task has increased by more than six orders of magnitude along the spatial dimension of distance. Kline (1995) shows the exponential improvement of human socio-technical systems for tasks such as communications, transportation, the radius of destruction, and others. For example, the number of computation steps that a computer can perform per second has been increasing exponentially, which relates to Moore's Law, and has recently been increasing a million-fold every 20 years. It is a socio-technical extension factor that has been directly related to the potential for intelligence augmentation with respect to the speed of symbolic processing of information (Engelbart, 1962).

A fundamentally important cognitive task is learning, so exploring the socio-technical extension factor for learning would be a foundational contribution to an emerging science of intelligence augmentation. The challenge of individual learning to perform a task like a human expert or the challenge of a team of people learning to perform a task at the performance of an expert team is of great scientific interest and economic significance. Reducing the amount of time it takes individuals and teams to learn to perform more like experts and less like novices is important for economic progress. Technology is not the only factor that contributes to socio-technical extension factors; organizations are another way to extend the capabilities of an individual (Norman, 1994). Ultimately, reducing the amount of time and energy that are required in a variety of contexts to accomplish a type of beneficial outcome as measured with respect to socio-technical extension factors is relevant to the science of intelligence augmentation.

6 Conclusion

The evaluations of AI technologies depend on the context they are used in. Incorporating discussions from the autonomous vehicles, healthcare, and innovation contexts, we believe putting AI back in context would benefit researchers to better answer the question of whether, when, and how to apply AI technology. Along with technology development, the question of what the end goal of AI emerges. Is automation, fully replacing

human intelligence with machine intelligence, the end goal of AI? In this paper, we emphasize a human-centric approach to AI (HAI) that considers human factors and ethical concerns regarding man-machine systems. Through our elaboration on the intricacies of human-machine interaction that form the foundation of augmented intelligence, we discuss how AI would impact the future of work, the challenges of embedding ethics in developing AI, and how to build trust in HAI. We believe the AI area has abundant research opportunities and issues in refining and shaping the future of human work that invite cross-disciplinary collaboration.

Acknowledgments

We gratefully thank Dr. Sarah Lebovitz (University of Virginia) for serving as one of the panelists in the workshop on “Augmented Intelligence and Future of Human Work” at the 55th Hawaii International Conference on System Sciences (HICSS-55). We also thank Dr. Lina Zhou (University of North Carolina Charlotte) for assistance in organizing and facilitating the panel discussion of “Augmented Intelligence and Future of Human Work” at HICSS-55.

References

- Acemoglu, D., & Restrepo, P. (2019). Automation and new tasks: How technology displaces and reinstates labor. *Journal of Economic Perspectives*, 33(2), 3-30.
- Acemoglu, D., Autor, D., Hazell, J., & Restrepo, P. (2022). Artificial intelligence and jobs: evidence from online vacancies. *Journal of Labor Economics*, 40(S1), S293-S340.
- Agrawal, A., Gans, J. S., & Goldfarb, A. (2018). *Prediction machines: The simple economics of artificial intelligence*. Harvard Business Review Press.
- Agrawal, A., Gans, J. S., & Goldfarb, A. (2019). Artificial intelligence: The ambiguous labor market impact of automating prediction. *Journal of Economic Perspectives*, 33(2), 31-50.
- Anthony, C. (2021). When knowledge work and analytical technologies collide: The practices and consequences of black boxing algorithmic technologies. *Administrative Science Quarterly*, 66(4), 1173-1212.
- Arntz, M., Gregory, T., & Zierahn, U. (2016). *The risk of automation for jobs in OECD countries*. Retrieved from <https://doi.org/10.1787/5jlz9h56dvq7-en>
- Babina, T., Fedyk, A., He, A. X., & Hodson, J. (2021). *Artificial intelligence, firm growth, and product innovation*. Retrieved from <https://ssrn.com/abstract=3651052>
- Babina T., Fedyk A., He, A., & Hodson, J. (2022). *Firm investments in artificial intelligence technologies and changes in workforce composition*. Retrieved from <https://ssrn.com/abstract=4060233>
- Barredo-Arrieta, A., & Del Ser, J. (2020). Plausible counterfactuals: Auditing deep learning classifiers with realistic adversarial examples. In *Proceedings of the International Joint Conference on Neural Networks*.
- Barrett, M., Oborn, E., Orlikowski, W. J., & Yates, J. (2012). Reconfiguring boundary relations: Robotic innovations in pharmacy work. *Organization Science*, 23(5), 1448-1466.
- Bresnahan, T. F., Greenstein, S., Brownstone, D., & Flamm, K. (1996). Technical Progress and Co-Invention in Computing and in the Uses of Computers. *Brookings Papers on Economic Activity: Microeconomics*, 1-83.
- Brynjolfsson, E., Mitchell, T., & Rock, D. (2018). What can machines learn, and what does it mean for occupations and the Economy? *AEA Papers and Proceedings*, 108, 43-47.
- Castells, M. (1999). *Information technology, globalization and social development* (vol. 114). UNRISD.
- Christin, A. (2020). The ethnographer and the algorithm: Beyond the black box. *Theory and Society*, 49(5), 897-918.
- Cukier, K. (2021). Commentary: How AI shapes consumer experiences and expectations. *Journal of Marketing*, 85(1), 152-155.
- Daugherty, P. R., & Wilson, H. J. (2018). *Human+machine: Reimagining work in the age of AI*. Harvard Business Press.
- Davenport, T. H., & Kirby, J. (2016). *Only humans need apply: Winners and losers in the age of smart machines*. Harper Business.
- Diakopoulos, N. (2020). "Transparency". In M. S. Dubber, F. Pasquale, & S. Das (Eds.), *The Oxford handbook of ethics of AI*. Oxford University Press.
- Dutta, D., & Mishra, S. K. (2021). Chatting with the CEO's virtual assistant: Impact on climate for trust, fairness, employee satisfaction, and engagement. *AIS Transactions on Human-Computer Interaction*, 13(4), 431-452.
- Eglash, R., Robert, L. P., Bennett, A., Robinson, K. P., Lachney, M., & Babbitt, W. (2020). Automation for the artisanal economy: enhancing the economic and environmental sustainability of crafting professions with human-machine collaboration. *AI & Society*, 35, 595-609.

- Endsley, M. R. (1987). *SAGAT: A methodology for the measurement of situation awareness* (NORDOC87-83). Northrop Corp.
- Endsley, M. R. (1988). Situation awareness global assessment technique (SAGAT). In *Proceedings of the National Aerospace and Electronics Conference*.
- Engelbart, D. (1962). *Augmenting human intellect: A conceptual framework*. Retrieved from <https://www.dougenelbart.org/pubs/augment-3906.html>
- Esterwood, C., & Robert, L. P. (2021). Robots and COVID-19: Re-imagining human–robot collaborative work in terms of reducing risks to essential workers, *Robonomics: The Journal of the Automated Economy*, 1, 1-5.
- Faraj, S., Pachidi, S., & Sayegh, K. (2018). Working and organizing in the age of the learning algorithm. *Information and Organization*, 28(1), 62-70.
- Felten, E., Raj, M., & Seamans, R. (2021). Occupational, industry, and geographic exposure to artificial intelligence: A novel dataset and its potential uses. *Strategic Management Journal*, 42(12), 2195-2217.
- Frey, C. B., & Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254-280.
- Furman, J., & Seamans, R. (2019). AI and the economy. *Innovation Policy and the Economy*, 19(1), 161-191.
- Harini, V. (2018). A.I. “bias” could create disastrous results, experts are working out how to fight it. *Yahoo!* Retrieved from <https://finance.yahoo.com/news/apos-bias-apos-could-create-054900771.html>
- Hooker, J., & Kim, T. W. (2019). Truly autonomous machines are ethical. *AI Magazine*, 40(4), 66-73.
- Jain, H., Padmanabhan, B., & Pavlou, P. A., & Raghu, T. S. (2021). Editorial for the special section on humans, algorithms, and augmented intelligence: The future of work, organizations, and society. *Information Systems Research*, 32(3), 675-687.
- Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making. *Business Horizons*, 61(4), 577-586.
- Jiang, J., Karran, A. J., Coursaris, C. K., Léger, P. M., & Beringer, J. (2022). A situation awareness perspective on human-AI interaction: Tensions and opportunities. *International Journal of Human-Computer Interaction*. Retrieved from <https://www.tandfonline.com/doi/full/10.1080/10447318.2022.2093863>
- Jordan, M. I., & Mitchell, T. M. (2015). Machine learning: Trends, perspectives, and prospects. *Science*, 349(6245), 255-260.
- Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Academy of Management Annals*, 14(1), 366-410.
- Kittur, A., Yu, L., Hope, T., Chan, J., Lifshitz-Assaf, H., Gilon, K., & Shahaf, D. (2019). Scaling up analogical innovation with crowds and AI. *Proceedings of the National Academy of Sciences*, 116(6), 1870-1877.
- Kline S. J. (1995). *Conceptual foundation of multidisciplinary thinking*. Stanford University Press.
- Lai, Y., Kankanhalli, A., & Ong, D. (2021). Human-AI collaboration in healthcare: A review and research agenda. In *Proceedings of the 54th Hawaii International Conference on System Sciences*.
- Lebovitz S., Lifshitz-Assaf H., & Levina N. (2022). To engage or not to engage AI for critical judgments: How professionals deal with opacity when using AI for medical diagnosis. *Organization Science*, 33(1), 126-14.
- Lebovitz, S., & Levina, N., Lifshitz-Assaf, H. (2021). Is AI ground truth really “true”? The dangers of training and evaluating AI tools based on experts’ know-what. *MIS Quarterly*, 45(3), 1501-1526.
- Lee, J., & Moray, N. (1992). Trust, control strategies and allocation of function in human-machine systems. *Ergonomics*, 35(10), 1243-1270.

- Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50-80.
- Lee, S., Yu, R., Xie, J., Billah, S. M., & Carroll, J. M. (2022). Opportunities for human-AI collaboration in remote sighted assistance. In *Proceedings of the 27th International Conference on Intelligent User Interfaces*.
- Levina, N., & Vaast, E. (2005). The emergence of boundary spanning competence in practice: Implications for implementation and use of information systems. *MIS Quarterly*, 29(2), 335-363.
- Liang, T.-P., Robert, L. P., Sarkar, S., Christy, C., Matt, C., Trenz, M., & Turel, O. (2021). Artificial intelligence and robots in individuals' lives: How to align technological possibilities and ethical issues. *Internet Research*, 31(1), 1-10.
- Lifshitz-Assaf, H. (2018). Dismantling knowledge boundaries at NASA: The critical role of professional identity in open innovation. *Administrative Science Quarterly*, 63(4), 746-782.
- Lykourantzou, I., Robert, L. P., Barlatier, J.-P. (2021). Unleashing the crowd work's potential: The need for a Post-Taylorism crowdsourcing model. *M@n@gement*, 24(4), 64-69.
- Meske, C., & Bunde, E. (2020). Transparency and trust in human-AI-interaction: The role of model-agnostic explanations in computer vision-based decision support. In H. Degen & L. Reinerman-Jones (Eds.), *Artificial intelligence in HCI* (LNCS vol. 12217). Springer.
- Mou, Y., & Xu, K. (2017). The media inequality: Comparing the initial human-human and human-AI social interactions. *Computers in Human Behavior*, 72, 432-440.
- Nah, F. F.-H., & Benbasat, I. (2004). Knowledge-based support in a group decision making context: An expert-novice comparison. *Journal of the Association for Information Systems*, 5(3), 125-150.
- Nah, F. H., Mao, J., & Benbasat, I. (1999). The effectiveness of expert support technology for decision making: individuals versus small groups. *Journal of Information Technology*, 14(2), 137-147.
- Norman, D. A. (1994). *Things that make us smart: Defending human attributes in the age of the machine*. Basic Books.
- Orlikowski, W. J., & Scott, S. V. (2015). The algorithm and the crowd. *MIS Quarterly*, 39(1), 201-216.
- Pasquale, F., & Cashwell, G. (2015). Four futures of legal automation. *UCLA Law Review Discourse*, 63, 26-48.
- Peterson, N. G., Mumford M. D., Borman W. C., Jeanneret P. R., Fleishman E. A., Levin K. Y., Campion M. A. (2001). Understanding work using the Occupational Information Network (O*NET): Implications for practice and research. *Personnel Psychology*, 54(2), 451-492.
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation–augmentation paradox. *Academy of Management Review*, 46(1), 192-210.
- Robert, L. P. (2019). Are automated vehicles safer than manually driven cars? *AI & Society*, 34, 687-688.
- Robert, L. P., Gaurav, B., & Lü tge, C. (2020a). ICIS 2019 SIGHCI workshop panel report: Human-computer interaction challenges and opportunities for fair, trustworthy and ethical artificial intelligence. *AIS Transactions on Human-Computer Interaction*, 12(2), 96-108.
- Robert, L. P., Pierce, C., Marquis, E., Kim, S., & Alahmad, R. (2020b). Designing fair AI for managing employees in organizations: A review, critique, and design agenda. *Human-Computer Interaction*, 35(5-6), 545-575.
- Rudin, C., & Radin, J. (2019). Why are we using black box models in AI when we don't need to? A lesson from an explainable AI competition. *Harvard Data Science Review*. Retrieved from <https://hdr.mitpress.mit.edu/pub/f9kuryi8/release/8>
- Samek, W., & Müller, K. R. (2019). Towards explainable artificial intelligence. In W. Samek, G. Montavon, A. Vedaldi, L. K. Hansen, & K.-R. Müller (Eds.), *Explainable AI: Interpreting, explaining and visualizing deep learning* (LNCS vol. 11700, pp. 5-22). Springer.
- Shneiderman, B. (2020). Human-centered artificial intelligence: Three fresh ideas. *AIS Transactions on Human-Computer Interaction*, 12(3), 109-124.

- Simonite, T. (2018). AI is the future—but where are the women. *Wired*. Retrieved from <https://www.wired.com/story/artificial-intelligence-researchers-gender-imbalance>
- Teodorescu, M. H., Morse, L., Awwad, Y., & Kane, G. C. (2021). Failures of fairness in automation require a deeper understanding of human-ML augmentation. *MIS Quarterly*, 45(3), 1-18.
- Tripsas, M. (2009). Technology, identity, and inertia through the lens of “The Digital Photography Company”. *Organization Science*, 20(2), 441-460.
- Van Den Bosch, K., & Bronkhorst, A. (2018). *Human-AI cooperation to benefit military decision making*. In *Proceedings of the NATO IST-160 Specialist Meeting on Big Data and Artificial Intelligence for Military Decision Making*.
- Van den Broek, E., Sergeeva, A., & Huysman, M. (2021). When the machine meets the expert: An ethnography of developing AI for hiring. *MIS Quarterly*, 45(3), 1557-1580.
- Von Hippel, E. (2006). *Democratizing innovation*. MIT Press.
- Waardenburg, L., Sergeeva, A., & Huysman, M. (2018). Hotspots and blind spots. In *Proceedings of the IFIP WG 8.2 Working Conference on the Interaction of Information Systems and the Organization*.
- Wang, P., Angarita, R., & Renna, I. (2018). *Is this the era of misinformation yet: Combining social bots and fake news to deceive the masses*. Retrieved from <https://hal.archives-ouvertes.fr/hal-01722413>
- Wei, W., Wu, J., & Zhu, C. (2020). Special issue on situation awareness in intelligent human-computer interaction for time critical decision making. *IEEE Intelligent Systems*, 35(1), 3-5.
- Xu, W. (2019). Toward human-centered AI: A perspective from human-computer interaction. *Interactions*, 26(4), 42-46.
- You, S., & Robert, L. P. (2018). Emotional attachment, performance, and viability in teams collaborating with embodied physical action (EPA) robots. *Journal of the Association for Information Systems*, 19(5), 377-407.
- Zhang, Q., Yang, X. J., & Robert, L. P. (2021). What and when to explain? A survey of the impact of explanations on attitudes towards adopting automated vehicles. *IEEE Access*, 9, 159533-159540.
- Zhou, L., Paul, S., Demirkan, H., Yuan, L., Spohrer, J., Zhou, M., & Basu, J. (2021). Intelligence augmentation: Towards building human-machine symbiotic relationship. *AIS Transactions on Human-Computer Interaction*, 13(2), 243-264.

About the Authors

Souren Paul is a Professor of Information Systems at the School of Computing and Analytics of Northern Kentucky University. His research interests are in areas of virtual teams, collaboration systems, behavioral information security, and augmented intelligence. He has published research articles in *Journal of Management Information Systems*, *Decision Support Systems*, and *Information & Management*. He has served as Conference Co-Chair for the 2015 Americas Conference on Information Systems (AMCIS) and the 2020 International Conference on Information Systems (ICIS).

Lingyao (Ivy) Yuan is an Assistant Professor at the Department of Information Systems and Business Analytics at Debbie and Jerry Ivy College of Business at Iowa State University. Her research interests include the impact of non-rational cognition and automatic processes on individuals' communication and behavior with virtual agents and digital humans, and on decision-making within computer-mediated environments. She has conducted research in the fields of electronic commerce and social media. She has been published in journals such as the *Journal of Management Information Systems (JMIS)*, *the Journal of the Association for Information Systems*, and *Decision Sciences*. She serves as the mini track chair of Actors, Agents, and Avatars: Visualizing Digital Humans in E-Commerce and Social Media at *Hawaii International Conference on System Sciences (HICSS)*.

Hemant Jain is W. Max Finally Chair and Professor of Data Analytics, in Rollins College of Business at University of Tennessee Chattanooga. His work has appeared in *ISR*, *MISQ*, *IEEE Transactions on Software Engineering*, *JMIS*, *IEEE Transactions on Systems Man and Cybernetics*, *Naval Research Quarterly*, *Decision Sciences*, *Decision Support Systems*, *Communications of ACM*, and *Information & Management*. He served as Associate Editor-in-Chief of *IEEE Transactions on Services Computing* and as Associate Editor of *JAIS & ISR*. He received his Ph. D from Lehigh University, a M. Tech. from IIT Kharagpur, and B. E. University of Indore, India.

Lionel P. Robert Jr. is a Professor and Associate Dean of Faculty Development & Affairs in the School of Information at the University of Michigan. He is an AIS Distinguished Member Cum Laude and an IEEE Senior Member. He completed his PhD in Information Systems from Indiana University where he was a BAT Fellow and a KPMG Scholar. He is the director of the Michigan Autonomous Vehicle Research Intergroup Collaboration (MAVRIC) and also a core faculty member of the Robotics Institute, an affiliate faculty of the National Center for Institutional Diversity all at the University of Michigan, and an affiliate faculty of the Center for Computer-Mediated Communication at Indiana University. He is currently serving on the editorial boards of the *Journal of the Association for Information Systems*, *AIS Transactions on Human-Computer Interaction*, *ACM Transactions on Social Computing*, *Collective Intelligence*, and *Management Information Systems Quarterly*. He has appeared in print, radio, and/or television for ABC, CNN, CNBC, Michigan Radio, *The New York Times*, and the Associated Press.

Jim Spohrer is a member of ISSIP.org, a UIDP Senior Fellow, and a retired IBM Executive. He led IBM Global University Programs, co-founded Almaden Service Research, and was CTO of Venture Capital Group. After his MIT BS in Physics, he developed speech recognition systems at Verbex (Exxon) before receiving his Yale Ph.D. in Computer Science/AI. In the 1990's, he attained Apple Computers' Distinguished Engineer Scientist and Technologist role for next generation learning platforms. With over ninety publications and nine patents, he received the Gummesson Service Research award, Vargo and Lusch Service-Dominant Logic award, Daniel Berg Service Systems award, and a PICMET Fellow for advancing service science.

Hila Lifshitz-Assaf is a Professor of Management at Warwick Business School and a visiting faculty at the Lab for Innovation Science at Harvard University. Her research focuses on developing an in-depth empirical and theoretical understanding of the micro-foundations of scientific and technological innovation and knowledge creation processes in the digital age. She explores how the ability to innovate is being transformed, as well as the challenges and opportunities the transformation means for R&D organizations, professionals and their work. She conducted an in-depth 3-year longitudinal field study of NASA's experimentation with open innovation online platforms and communities, resulting in a scientific breakthrough. This study received the best dissertation Grigor McClelland Award at the European Group for Organizational Studies (EGOS) 2015, Best Administrative Science Quarterly (ASQ) paper based on dissertation (2018), and best published paper elected by the Organizational Communication and Information Systems Division of Academy of Management (2018).

She investigates new forms of organizing for the production of scientific and technological innovation such as crowdsourcing, open source, open online innovation communities, Wikipedia, hackathons, makeathons, etc. Her work received the prestigious INSPIRE grant from the National Science Foundation and has been presented and taught at a variety of institutions including MIT, Harvard, Stanford, INSEAD, Wharton, London Business School, Bocconi, IESE, UCL, UT Austin, Columbia, and Carnegie Mellon. Her work was recognized to have a strong impact on the industry. She received the Industry Studies Association Frank Giarrantani Rising Star award and the Industry Research Institute grant for research on R&D.

Copyright © 2022 by the Association for Information Systems. Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and full citation on the first page. Copyright for components of this work owned by others than the Association for Information Systems must be honored. Abstracting with credit is permitted. To copy otherwise, to republish, to post on servers, or to redistribute to lists requires prior specific permission and/or fee. Request permission to publish from: AIS Administrative Office, P.O. Box 2712 Atlanta, GA, 30301-2712 Attn: Reprints or via e-mail from publications@aisnet.org.



Editor-in-Chief

<https://aisel.aisnet.org/thci/>

Fiona Nah, City University of Hong Kong, Hong Kong SAR

Advisory Board

Izak Benbasat, University of British Columbia, Canada
John M. Carroll, Penn State University, USA
Phillip Ein-Dor, Tel-Aviv University, Israel
Dennis F. Galletta, University of Pittsburgh, USA
Shirley Gregor, National Australian University, Australia
Elena Karahanna, University of Georgia, USA
Paul Benjamin Lowry, Virginia Tech, USA
Jenny Preece, University of Maryland, USA

Gavriel Salvendy, University of Central Florida, USA
Suprateek Sarker, University of Virginia, USA
Ben Shneiderman, University of Maryland, USA
Joe Valacich, University of Arizona, USA
Jane Webster, Queen's University, Canada
K.K. Wei, Singapore Institute of Management, Singapore
Ping Zhang, Syracuse University, USA

Senior Editor Board

Torkil Clemmensen, Copenhagen Business School, Denmark
Fred Davis, Texas Tech University, USA
Gert-Jan de Vreede, University of South Florida, USA
Soussan Djasasbi, Worcester Polytechnic Institute, USA
Traci Hess, University of Massachusetts Amherst, USA
Shuk Ying (Susanna) Ho, Australian National University, Australia
Matthew Jensen, University of Oklahoma, USA
Richard Johnson, Washington State University, USA
Atreyi Kankanhalli, National University of Singapore, Singapore
Jinwoo Kim, Yonsei University, Korea
Eleanor Loiacono, College of William & Mary, USA
Anne Massey, University of Massachusetts Amherst, USA
Gregory D. Moody, University of Nevada Las Vegas, USA

Stacie Petter, Baylor University, USA
Lionel Robert, University of Michigan, USA
Choon Ling Sia, City University of Hong Kong, Hong Kong SAR
Heshan Sun, University of Oklahoma, USA
Kar Yan Tam, Hong Kong U. of Science & Technology, Hong Kong SAR
Chee-Wee Tan, Copenhagen Business School, Denmark
Dov Te'eni, Tel-Aviv University, Israel
Jason Thatcher, Temple University, USA
Noam Tractinsky, Ben-Gurion University of the Negev, Israel
Viswanath Venkatesh, University of Arkansas, USA
Mun Yi, Korea Advanced Institute of Science & Technology, Korea
Dongsong Zhang, University of North Carolina Charlotte, USA

Editorial Board

Miguel Aguirre-Urreta, Florida International University, USA
Michel Avital, Copenhagen Business School, Denmark
Gaurav Bansal, University of Wisconsin-Green Bay, USA
Ricardo Buettner, Aalen University, Germany
Langtao Chen, Missouri University of Science and Technology, USA
Christy M.K. Cheung, Hong Kong Baptist University, Hong Kong SAR
Tsai-Hsin Chu, National Chiayi University, Taiwan
Cecil Chua, Missouri University of Science and Technology, USA
Constantinos Coursaris, HEC Montreal, Canada
Michael Davern, University of Melbourne, Australia
Carina de Villiers, University of Pretoria, South Africa
Gurpreet Dhillon, University of North Texas, USA
Alexandra Durcikova, University of Oklahoma, USA
Andreas Eckhardt, University of Innsbruck, Austria
Brenda Eschenbrenner, University of Nebraska at Kearney, USA
Xiaowen Fang, DePaul University, USA
James Gaskin, Brigham Young University, USA
Matt Germonprez, University of Nebraska at Omaha, USA
Jennifer Gerow, Virginia Military Institute, USA
Suparna Goswami, Technische U.München, Germany
Camille Grange, HEC Montreal, Canada
Juho Harami, Tampere University, Finland
Khaled Hassanein, McMaster University, Canada
Milena Head, McMaster University, Canada
Netta Iivari, Oulu University, Finland
Zhenhui Jack Jiang, University of Hong Kong, Hong Kong SAR
Weiling Ke, Southern University of Science and Technology, China

Sherrie Komiak, Memorial U. of Newfoundland, Canada
Yi-Cheng Ku, Fu Chen Catholic University, Taiwan
Na Li, Baker College, USA
Yuan Li, University of Tennessee, USA
Ji-Ye Mao, Renmin University, China
Scott McCoy, College of William and Mary, USA
Tom Meservy, Brigham Young University, USA
Stefan Morana, Saarland University, Germany
Robert F. Otondo, Mississippi State University, USA
Lingyun Qiu, Peking University, China
Sheizaf Rafaeli, University of Haifa, Israel
Rene Riedl, Johannes Kepler University Linz, Austria
Khawaja Saeed, Wichita State University, USA
Shu Schiller, Wright State University, USA
Christoph Schneider, IESE Business School, Spain
Theresa Shaft, University of Oklahoma, USA
Stefan Smolnik, University of Hagen, Germany
Jeff Stanton, Syracuse University, USA
Chee-Wee Tan, Copenhagen Business School, Denmark
Horst Treiblmaier, Modul University Vienna, Austria
Ozgur Turetken, Ryerson University, Canada
Wietske van Osch, HEC Montreal, Canada
Weiquan Wang, City University of Hong Kong, Hong Kong SAR
Dezhi Wu, University of South Carolina, USA
Fahri Yetim, FOM U. of Appl. Sci., Germany
Cheng Zhang, Fudan University, China
Meiyun Zuo, Renmin University, China

Managing Editor

Gregory D. Moody, University of Nevada Las Vegas, USA