

Paweł Nowik  <https://orcid.org/0000-0002-1824-0884>

The John Paul II Catholic University of Lublin

NEW CHALLENGES FOR TRADE UNIONS IN THE FACE OF ALGORITHMIC MANAGEMENT IN THE WORK ENVIRONMENT

Abstract

Algorithmic management is the subject of numerous scientific studies. This article attempts to answer the question of what kinds of new competencies and skills should be acquired by trade unions in the face of challenges related to algorithmic management. The author indicates two main areas of trade union activities: The first concerns the challenges associated with the process of explaining and transplanting artificial intelligence. The second concerns participation in the AI certification process. Considering that artificial intelligence algorithms' certification process is an entirely new undertaking, it should be based on a pragmatic search for peaceful solutions, encourage compliance with the law and limit the possibility of stiff administrative and criminal sanctions. For this purpose, the author considers using the theory of responsive regulation as a pragmatic approach for certification agencies and trade unions. The author considers the cooperation of artificial intelligence to be the main principle. In the working environment, there should be a principle of human importance—the focus of personalism.

Słowa kluczowe: sztuczna inteligencja, zarządzanie algorytmiczne, certyfikacja AI, Explainable AI, związki zawodowe, elastyczna regulacja

Keywords: artificial intelligence, algorithmic management, AI certification, Explainable AI, trade unions, responsive regulation

ASJC: 3308, **JEL:** H55, K31

Introductory remarks

The development of algorithmic management for the working environment is essential for research and academic debate (Aneesh 2009; Irani 2015; Lee et al. 2015b; Cherry 2016). The development of algorithmic management for the working environment is essential for research and academic debate. One view suggests that artificial intelligence (AI) and algorithmic rules are not hampered by overregulation, as this slows down their inevitable

development and, at the same time, weakens the competitiveness of a country's economy (Stewart, Stanford 2017; Deloitte 2020). The second view highlights that algorithmic management in the work environment should be regulated, emphasizing the need to provide a tight protective umbrella for employees (World Economic Forum 2018).

Reflecting on the state of the contemporary debate on this subject, Prassl correctly observes the danger of generalization being too easy when confronted with the black scenario of the algorithmic exploitation of employees under an optimistic vision of improved and efficient work (Adams-Prassl 2020). Undoubtedly, there is a shortfall of broader scientific reflection aimed at finding "ways to the middle ground" including making space for technological development and the instrument of regulation and control to protect workers.

Crucial responsibilities belong to trade unions, which play a role in regulating working conditions in areas where AI and algorithmic management are present (DeStefano 2018). Trade unions concerned with the processes of the algorithmic rule should be able to take on new skills, in addition to their traditional competencies (realization of the right of employee participation, conducting collective bargaining, and conducting collective disputes), so that there is the possibility of performing social control of the operation of algorithmic management using the technology of AI explanation. The possibility of controlling the AI evaluates the algorithm rules throughout its life cycle, starting from the initial design processes.

One interesting idea mentioned in the literature and many reports on AI activities' ethical issues certify the SI. Trade unions could play an essential role in consultation and opinion-forming in the face of such a certification procedure. Reflecting on the pragmatics of trade unions' operations, it is worth referring to responsive regulation research. Ian Ayres presented its essential elements and John Braithwaite in their famous scientific monograph entitled *Responsive Regulation: Transcending the Deregulation Debate* (Rogers, Ayres, Braithwaite 1993). An essential part of this proposal is the search for a general regulatory strategy that attempts to indicate the so-called third way between the one-sided governmental practice of control and supervision and the laissez-faire policy, which is based on market competition. According to the authors,

"responsive regulation is not a clearly defined program or a set of prescriptions concerning the best way to regulate... Responsiveness is rather an attitude that enables the blossoming of a wide variety of regulatory approaches, only some of which are canvased here. However, for the responsive regulator, there are no optimal or best regulatory solutions, just solutions that respond better than others to the plural configurations of support and opposition that exist at a particular moment in history" (Rogers, Ayres, Braithwaite 1993).

Thus, in other words, Ayres and Braithwaite proposed that regulators be flexible and tune their regulatory responses to the situations they are dealing with in practice. According to the authors, good regulation is about respecting the spirit of the law, and good regulation "promotes freedom as the absence of domination" (Braithwaite 2013).

What can Ayres and Braithwaite's ideas contribute to the regulation of algorithmic management and AI in the work environment? First, the theory of responsive regulation may become a practical tool for future AI-certified agencies' operations, the necessity of which is indicated by many thinktanks and researchers in reports (van der Heijden 2020). The potential certification process, which should also include law enforcement elements, helps to find a regulatory balance. Considering that AI algorithms' certification is a new development, it should be based on a search for flexible solutions that encourage compliance with the law and limit the possibility of applying stiff administrative and criminal sanctions. Such an approach seems to be especially advantageous if we look for a compromise between imminent technological development and providing protective guarantees for workers.

These flexible regulations are the most well-known illustrations called "the enforcement pyramids." One of them shows how a regulator can work with an individual or a business through a set of escalating regulatory interventions: clarifying the purpose of regulation to compliance (persuasion), through warning letters or civil sanctions (deterrence), criminal sanctions, or withdrawal of license (full force of law). The enforcement pyramid illustrates how compliance with "soft" regulatory interventions can be achieved, and to what extent a regulatory authority can escalate to an acute response and is willing to apply the most critical intervention (Braithwaite, Braithwaite 2001; van der Heijden 2020).

The paper's thesis states that the theory of responsive regulation can also be a proposal for trade unions and other organizations representing employees' collective interests, especially given the challenges of the broadly understood process of shaping transparency and explainability for the non-transparent algorithms that shape employee management rules. Concerning algorithmic management, trade unions should develop new skills to protect workers' collective rights and interests. Such new capabilities include, for example, conscious participation in the process of explaining AI and performing control functions in the field of law enforcement. Trade unions, particularly their transnational organizations, should also actively participate in AI certification. However, in each case, the response regulating mechanisms described by Ayres and Braithwaite might prove useful. A significant element of this concept is the so-called tripartitism that includes forced self-regulation and partial industry regulation (Rogers, Ayres, Braithwaite 1993). According to the authors, tripartitism presents the idea of strengthening civic associations' position to overcome potential pathologies (such as corruption) in traditional, bilateral interactions between the government and regulated industry. This forced self-regulation requires companies to write their own sets of corporate rules, which are then publicly ratified and enforced. Thus, this space is close to trade union competence and collective bargaining practices.

Ayres and Braithwaite introduced responsive regulation as an overall regulatory strategy that aims to build on both approaches' strengths and overcome their weaknesses. Therefore, this is worth considering in the context of new challenges related to artificial intelligence and algorithmic management regulation in the work environment.

The first thing to consider is the scope of the meaning of algorithmic management and AI concepts. In this context, referring to employment issues, the aim is to identify the main regulatory ranges concerning AI and algorithmic management. This is discussed in section 2.

Section 3 of this paper is concerned with reading the new role of trade unions and other social organizations representing employees' collective interests in the context of challenges related to the development of algorithmic management rules in the work environment. Considering traditional trade union competencies and looking at their new skills, the usefulness of Ayres' and Braithwaite's responsive theories should also be explored in this area. Such exploration specifically focuses on the trade union challenges of the algorithm clarification process.

Section 4 of this paper focuses on identifying the types of AI algorithms in a work environment that will be certified. Based on the assumption of the need for such certification concerning specific AI types, the certification authority's rules and operation methods will be analyzed. Therefore, this paper will consider to what extent Ayres and Braithwaite's theory of responsive regulation can help determine the principle of operation for a certification and inspection authority.

The final part of the paper offers summary conclusions related to the possibility of applying the responsible theory in the process of regulation and outlawing of labor law in the work environment, where algorithmic and SI management takes place. Previous to this reflexivity, this work will attempt to identify the research problems that require further study.

Authority of algorithmic management in the working environment

Algorithmic management is a varied set of technologies for employees' remote control, based on data collection and the monitoring of employees to support automated or semiautomated decision-making (Mateescu, Nguyen 2019). It may also apply to recruitment processes where SI technology is used (Mann, O'Neil 2016). Advanced surveillance technologies are used in a wide range of platforms. Algorithmic management also occurs in a traditional work environment using wearable technology and monitoring software applications. People analytics, which is defined as the systematic identification and quantification of the people who drive business results, has played a crucial role for several years (van den Heuvel, Bondarouk 2017). Access to a large variety of data is required for effective algorithmic management (Gillespie 2014).

Algorithmic management is based on an uninterrupted stream of information about employees' behavior, distinguished from conventional management models based on personal relationships (Rosenblat, Stark 2016). Despite their diversity, the technologies used in algorithmic management still have one common factor. Each of them is used as a tool to collect data that the SI can process to find patterns, trends, and correlations

to achieve a more significant workflow optimization, with particular emphasis on performance and cost aspects (Amyx n.d.).

All algorithms play an essential role in identifying useful patterns in datasets and making decisions based on these patterns. Algorithmic systems can use various methods to structure and control employee behavior. Granular data on employee behavior can also be interpreted as objective performance measures, although in reality, they can only be substitutes or incomplete measures. Finn argues that the algorithm deploys concepts from the idealized space of computation in a messy reality, with unpredictable and sometimes fascinating results (Finn 2017). He argues that we need to build an algorithmic understanding and inquisitiveness model that will serve the process leading to new experimental humanities (Finn 2017). Explainable AI (XAI) is a separate field of research knowledge. It is based on the assumption that the person who designed an algorithm is responsible for its operation's specific effects. Explainable AI is placed within the framework of a broader ethical or accountable AI. It covers ethical issues and legal regulations at the level of common law and industry regulation.

The AI's fundamental dilemma is whether its creator should bear the legal responsibility for undesirable effects. According to Cynthia Rudin, many tools used to explain the secrets of deep learning do not match the actions of the models they are supposed to explain. Rudin argues that black boxes are complex models based on deep neural networks and more straightforward SI tools protected by business secrets (Rudin 2018). According to Rudin, using interpreted SI models instead of non-transparent neural networks is possible in various areas, whereas for sensitive issues, it is necessary. Therefore, she believes that using black box-based models should be prohibited where it is possible to create an equally effective system based on classically interpreted models. A more user-friendly form of promoting transparent SI models could be the mandatory performance testing of neural network-based networks using XAI techniques (Barredo et al. 2020).

The use of modern technologies is dictated by the need to optimize the work process and reduce its costs, which involves major ethical concerns (Mittelstadt et al. 2016; Stone et al. 2016). Unquestionably, the main argument favoring algorithmic decision-making is the advantage resulting from the unprecedented speed and efficiency of systems designed in this way. However, algorithms that are increasingly difficult to understand can make decisions that raise specific ethical controversies. In this context, the integrity of the algorithms becomes an ethical, economic, and legal category (Schubert, Hütt 2019). Whether they occur within the framework of algorithmic management on various platforms or in other work environments, the algorithms themselves are not ethical or unethical. Instead, they are a formally defined sequence of logical operations that provide computers with step-by-step instructions for handling data and automated decisions (Barocas, Selbst 2016). For the personalistic dimension of work, SI systems should support human autonomy and the decision-making process according to the principle of respect for human autonomy (Ferraris, Bosco, D'Angelo 2013; Moore, Upchurch, Whittaker 2018). This requires that AI systems function as a facilitator of a democratic,

prosperous, and just society by supporting user leadership and fundamental rights and allowing for human oversight (Winfield et al. 2019; Hamon, Junklewitz, Sanchez 2020).

The general principle of user autonomy plays a central role in the function of the system. An essential aspect of this is the users' right not to be subject to a decision based solely on automated processing. Human oversight ensures that the SI system does not undermine human autonomy or cause other adverse effects. This means that people who are employed and interact in any way with AI systems should be guaranteed the freedom of self-determination (Pasquale 2015; Yeung 2018; AI HLEG 2019a), in which the human rights-based freedom of the individual plays a fundamental role (Davidow 2014; Noto La Diega 2018). Concerning AI, the individual's freedom requires the prevention of (un)indirect illegal coercion, threats to mental independence and mental health, unjustified supervision, misrepresentation, and unfair manipulation (AI HLEG 2019a). Algorithmic management refers to a commitment to enable individuals to exercise even greater control over their lives, including, among other things, the protection of the freedom to conduct business, the right to respect for private life and privacy, and the right of association and assembly. AI systems should not unduly subjugate, coerce, deceive, shape, or control people, or manipulate them (AI HLEG 2019a). Humans are all treated with the respect they deserve as moral subjects and not just as objects to be sifted, sorted, judged, collected, conditioned, or manipulated. As a result, AI systems must be developed to support and ensure respect for and protection of human beings' physical and mental integrity and personal and cultural identity, and guarantee their essential needs. The division of function between humans and AI systems should be based on human-oriented design principles and leave the humans in a real position to control the rules of the algorithm (Rosenblat, Stark 2016; Schubert, Hütt 2019). It ensures human supervision and control over the work processes of AI systems (Mittelstadt et al. 2016).

Trade unions and XAI

The intense development of SI has led to opaque decision-making systems such as deep neural networks (DNNs). The empirical success of deep learning (DL) models, such as DNNs, stems from a combination of efficient learning algorithms and their huge parametric space. The latter space comprises hundreds of layers and millions of parameters, making DNNs complex black-box models (Karanasiou, Pinotsis 2017; Barredo et al. 2020). AI models with machine learning showing the highest predictive accuracy are, paradoxically, those with the most unclear black box architecture.

Meanwhile, the unstoppable computerization of advanced industrial societies requires the use of these technologies in an increasing number of areas. The combination of both phenomena raises the problem of controlling AI (Carabantes 2020). The black box's opposite is transparency, which means a direct understanding of the model's mechanism. On the one hand, it is necessary to verify how the system works

(correctly or not), and on the other hand, it is a quest for certainty that the results of DNNs can be trusted (Ribeiro, Singh, Guestrin 2016; Preece et al. 2018). There is a shared conviction of a need to shape a so-called ethical AI that generates trust (UNI Global Union 2017; Dameski 2018; Winfield et al. 2019).

The development of DNNs has undoubtedly contributed to the reactivation of the field of knowledge known as “Explainability in artificial intelligence” (XAI), which originated in the late 1970s and early 1980s when the focus was on building expert systems to emulate human reasoning in specialist high-value domains such as medicine, engineering, and geology (Buchanan, Duda 1983; Preece et al. 2018). Recent years have been a period of landmark research on SI explainability. However, the term interpretability is now more common, indicating an emphasis on people interpreting machine learning models (Preece et al. 2018).

There is an intense discussion in the literature about the terms “interpretability” and “explainability” in the cognitive aspect of AI. Essential differences between these concepts appear to exist (Barredo et al. 2020).

Interpretability is identified as a human observer’s certain level of understanding of the model. Explainability, on the other hand, is an active feature of the model, meaning that all actions or procedures are undertaken to clarify and detail the internal function of the SI model. Standard conceptual nomenclature includes understandability (or equivalently, intelligibility), comprehensibility, interpretability, explainability, and transparency (Barredo et al. 2020). All these terms refer to different aspects of understanding AI, in which knowledge and human knowledge play a fundamental role.

Explainability is defined in various publications as an element of transparency. Sometimes, it is understood as a separate criterion from transparency or used as a substitute (ICO n.d.; Kaminski 2018; Sileno, Boer, van Engers 2019; Guo 2020; Leslie 2020). In the following reflection, I assume that transparency is a more general concept, including explainability as one element of a transparent AI system. Similarly, the term transparency is used in the ethically aligned design adopted by the European Commission and Institute of Electrical and Electronics Engineers (IEEE) (IEEE n.d.).

In the European Commission’s report, which discusses the features of a credible SI, the concept of transparency has the following requirements:

- Human agency and oversight—SI systems should support the development of a just society by strengthening the guiding role of humans and their fundamental rights, rather than reducing, limiting, or distorting human autonomy, stability, and safety.
- The algorithms used for reliable AI must be secure, reliable, and robust enough to deal with errors or inconsistencies at all stages of the SI system life cycle.
- Privacy and data protection (privacy and data governance)—citizens should have full control over their data, and the data concerning them should not be used to harm them or discriminate against them. The aim is for diversity, non-discrimination, and fairness.
- SI systems should consider the whole range of human capacities, skills, and requirements, and ensure accessibility and social and environmental well-being.

- AI systems should enhance positive social change, promote sustainable development and environmental responsibility, and support accountability.

Mechanisms should be put in place to ensure accountability for SI systems and their results (AI HLEG 2019b).

There is a broad and diverse range of stakeholders interested in the proper understanding of SI. Each of them uses different terminology, which is derived from the expectations of the XAI. These include developers (people concerned with building AI applications), users (people who use AI systems), theorists (people involved with understanding and advancing AI theory, particularly around deep neural networks), ethicists (people concerned with the fairness, accountability, and transparency of AI systems, including policy-makers and commentators and also trade unions (Preece et al. 2018).

The fundamental condition for implementing the transparency principle is to provide selected trade union activists with access to reliable information on the AI model's operation, including information on the training procedure, training data, machine learning algorithms, and testing methods for validating the AI system. The right to information is not limited. The purpose and scope of legal requirements for AI transparency may depend on the recipients' profile, understood both as the AI system operator and its end-user. Therefore, the possibility of different levels of transparency and available explanations should be considered depending on a given recipient's qualification as a member of a specific stakeholder group. Thus, the scope of causes dedicated to manager entities in AI systems is different for data scientists, developers, AI designers, executive and supervisory bodies, regulators, and certification bodies. Consequently, the explanation scope differs for end-users affected by AI decisions, such as job applicants or employees. Identifying different levels of transparency for different audiences should be preceded by assessing the risks posed by the AI system.

The impact assessment should also provide an acceptable level of autonomy in the AI system and a minimum level of explanation of its decisions or predictions regarding trade unions and workers. The results obtained after applying SI systems should not lead to discrimination against individuals or social groups based on race, religion, sex, sexual orientation, disability, ethnic origin, or other personal circumstances. The crucial criterion to consider when determining an AI system's optimum performance is its performance in terms of error optimization and how the course deals with these groups (Benjamins, Barbado, Sierra 2019).

Moreover, AI systems should not encourage the uncontrolled collection and processing of employees' personal data. Another risk element in an SI system operation is the possibility of the breach of the principle of privacy and security.

Following the above understanding of transparency, the need for privacy and security should apply to all stages of AI implementation, including design, development, personalization (customization to the needs of the user purchaser), operation, and periodic validation of AI systems' effectiveness (regular testing as well as regular external audits).

It seems appropriate to request that, at least for a certain group of AI systems in which the risks identified above can be reliably addressed, AI systems should be classified as high-risk (AI HLEG 2020). This means that specific legal requirements need to be introduced to ensure the transparency and, in particular, the clarity of decisions made by automated or semiautomated means. For example, as unions play a crucial role in building employees' confidence in AI, there is a consequential requirement that access to information about the AI system and the data it uses is guaranteed. Legal documentation obligations should be imposed on the designers, manufacturers, and operators of autonomous systems to ensure fair access to data. Such documentation should include:

1. Methods used in the design and development of AI.
2. Data management policies and procedures that consider the requirements of data availability, timeliness, integrity, and security.
3. The standards applied (technical, legal, and ethical) and certificates granted (it should be assumed here that different criteria will apply to various systems; a predictive algorithm determining the likelihood of recidivism should undergo a mandatory compliance assessment, while Netflix algorithms suggesting another film worth watching are unlikely to experience any certification).
4. A management model attributing responsibilities and liabilities within the organization to the AI.
5. Methods used for testing and validating the AI system.

One of the requirements necessary to ensure the transparency and explanation of the system is the ability to identify, at each stage of AI use, the data files (including their collection, marking, grouping, etc.) and the processes used to analyze them, including machine learning algorithms. Regarding the context of clarity, it is essential to record and archive the decisions of the SI system, together with all data and algorithms used, to ensure that, even after a certain period (within a fixed retention period), the information that leads to the decision can be analyzed. Formally, the explanation should be provided in good time, preferably in real-time, but it is also acceptable to explain within a reasonable time after the decision has been implemented. As a rule, a trustworthy AI system should explain its operations, activities, and decision-making. However, AI systems with posthoc explanations should be considered acceptable, as such systems would have to show techniques for transforming an uninterpretable model into an explainable model (Khaleghi 2019).

The AI provides the user, in our case, the trade unions, with the information needed to understand why an autonomous system behaves in a certain way under certain circumstances (or would conduct hypothetical cases) (Barredo et al. 2020). Diverse working environments, including various types of work on platforms, only form part, meaning it would be excessive to demand that all possible AI systems meet the full explanation requirement. For this reason, the level of explainability required by law is defined as:

- categories of stakeholders to whom certain information should be disclosed,
- types of information that can be disclosed to specific stakeholders,
- circumstances that specify the level of detail of the information to be announced.

One fundamental difficulty is determining the level of explicability of the AI. On the one hand, this is a problem with the procedures used in practice. This should be the subject of appropriate legal regulations guaranteeing trade unions' right to demand proper measures to protect the rights, freedoms, and legitimate interests of persons affected by AI systems. Establishing a level of clarification in trade unions may also be the subject of social dialog with the AI system's employers or operators. However, the key here is to define the categories of information that can be disclosed and the types of experts that can have access to it. Protecting the creators or producers of AI's intellectual property rights cannot be disregarded entirely.

In addition to setting up procedures for the initial phase of the clarification process and the challenges of communicating specialized information to AI systems, the scope of this information remains a critical issue. Business guidelines (ICO n.d.), as well as literature (Spreeuwenberg 2020), indicate that the explanation one should expect from the SI should include:

1. An explanation of the decision-making process, that is, an indication of the reasons that lead to a decision of certain content, which is provided in an accessible and non-technical way.
2. Clarification of responsibilities, that is, who was involved in the development, implementation, management, and operation of the SI system.
3. Explanation of the data, that is, which data were used in the decision, which data were used for training and testing the AI system and how, whether the data were reliable (not causing algorithmic bias), and whether the quantity of data used was sufficient.
4. A safety explanation, that is, proof of the SI system's accuracy, reliability, security, and robustness.
5. An explanation of the impact that using the SI system and its decisions has or may have on an individual or, more broadly, on a specific social group.
6. The justification of the outcome; it is crucial to explain why a specific decision was taken and justify that the development of the AI is objective and fair (ICO n.d.).

The essential element of transparency highlighted by High-Level Expert Group on Artificial Intelligence (HLEG) in The Assessment List on Trustworthy Artificial Intelligence (ALTAI) is communication (AI HLEG 2020). The system should inform the user from the beginning that they are interacting with AI. Trade unions should ensure that each employee has the right to know whether AI is the only source of information for the system operator and whether it provides a supporting function, or whether it is a fully autonomous system that functions without human intervention.

The assumption is that not all users, including trade unions, have the appropriate knowledge to understand how the data and codes collected translate into specific benefits or disadvantages. However, the adequately limited scope of the information disclosed must be sufficient in each case to determine whether the AI meets the standards adopted by law or compliance markings held, particularly in terms of reliability and security. The right to explain AI decisions does not necessarily mean the “black box.” The black box will be opened, that is, the end-user does not need to know how the algorithm works (the system may be so complicated that a layman will not be able to explain its operation). Nevertheless, they should be given feedback that allows them to understand the decision, change their behavior (to obtain a different conclusion), or appeal the outcome if separate provisions provide for the right of appeal. The explanation should be delivered understandably in written or visual form, adapted to the stakeholder’s level of knowledge. The simplest form of presentation of the reason is visualization, which highlights the relationship between input and output. A more advanced form involves testing hypotheses, where a well-formulated argument is tested based on information and production. However, it seems that for a non-expert user, the best way to present an explanation is to use natural language—both verbal and written communication—and to indicate which data features and algorithmic functions led to the decision. However, this solution is probably the most technically complex (Arrieta et al. 2020; Bhatt et al. 2020).

The explanation the user receives should present the specific context in which the SI operates. Temporary life situations (for example, childhood or illness), market factors (for example, information asymmetry or market power), economic factors (for example, poverty), identity factors (for example, gender, religion, or culture), and other factors may be relevant in the working environment. There is also a context for the field of application and the life-cycle phase of the AI. This also includes information on the data used to train the model. The explanation should indicate when a case is an “outlier” and very different from the data used to train the model (this will identify any situation where an AI gives an incorrect result and requires human intervention). The AI should present an alternative effect showing how the result obtained (as predictions or decisions) differs from other potential consequences. If a lack of sufficient knowledge or data occurs, the AI system should report this, potentially blocking the system operator’s ability to decide.

Transparency as a requirement for AI is indicated in almost every document on ethical and trustworthy AI. The term covers data use, human-SI interaction, and automated decision-making. In the report “Responsible AI by Design in Practice,” in the context of transparent and explainable SI principles, the authors point out that it is crucial to ensure a certain level of understanding of the SI system’s decisions. Everyone should be aware of when they are communicating with a person and when they are communicating with the AI system (Benjamins et al. 2019). In addition to identifiability, the concept of transparency also includes explaining and communicating clearly about the limitations of the AI system. These issues are essential for the data subjects, employees, and trade unions. Transparency reduces the risk of liability for damages and enhances trust in AI.

Considering the potential role of unions in participating in the development of a trustworthy AI for employees, it should be noted that organizations are increasingly publishing AI principles, declaring that they want to avoid unintended negative consequences. Nevertheless, there is minimal experience with the actual implementation of these principles in organizations. Many organizations already have standards and procedures for protecting the privacy, safety, and security of their data. However, an issue that involves operationalizing parts specific to the AI while using existing processes for more general principles remains. Therefore, unions can define co-worker agreements or available documents on the proposed values and limits of AI. Various training and information campaigns can increase knowledge of potential issues, both technical and non-technical, related to AI. Such a tool may be a questionnaire proposed by trade unions, which forces people to reflect on the AI system (impact clarification) (B. Jo Lea 2018). The questionnaire should provide specific guidance on what to do if certain adverse effects are detected. Union activists should be increasingly aware of the tools that help answer some of the questions that arise and mitigate any problems identified. Implementation of the principles of transparent and explainable AI requires unions to acquire new responsibilities. A significant challenge is the practical knowledge of XAI techniques (Benjamins, Barbado, Sierra 2019).

AI certification and responsive regulation theory

Ethical issues constitute a fairly common consensus in the debate on the development of AI in the context of the problems of developing future work. In 2020, Deloitte reported that 85% of the organizations surveyed believe that the future of work is linked to ethical challenges. Only 27% believe that they have clear rules and leaders who can manage these challenges (Deloitte 2020). When asked what makes the management of ethical issues related to the future of work more important, the organizations surveyed most often pointed to several main factors: legal and regulatory requirements (38%); the rate of implementation of AI technology in the workplace (34%); changes in the composition of the workforce (for example, growth of the alternative workforce) (32%), pressure from external stakeholders (for example, investors, clients, inter groups, etc.) (29%), pressure from employees (18%), and directions from management and leaders (12%) (Deloitte 2020).

For respondents, the leading factors were legal and regulatory requirements. The lack of crucial legal solutions that consider ethical issues, in terms of using AI in the working environment, is one of the main challenges. AI and other technologies make future work ethics more relevant, as the spread of technology leads to a redefinition of the impact of technology on the human role at work. With technology becoming ever more rooted in work, its design and application must be assessed in terms of fairness and equity. The key questions remain as to whether the use of technology reduces or increases discriminatory attitudes, what procedures are in place to protect workers'

privacy, whether technological decisions are transparent and explainable, and which rules are made to hold people responsible for the results of these decisions (Deloitte 2020).

Efforts to develop standards, both technical and ethical, for AI are essential. The work on standards for reliable AI is carried out, among others, by the Joint Technical Committee of the International Organization for Standardization (ISO), the International Electrotechnical Commission (IEC), and the Institute of Electrical and Electronics Engineers (IEEE). In 2019, the IEEE published a document entitled *Ethically Aligned Design*, containing an ethical framework for AI and recommendations to serve as a reference for legislators, engineers, developers, and companies that are implementing, selling, and using AI systems (How 2017). The above reports and scientific literature raise the demand for control over algorithms and AI through specialized national and international institutions.

Established international institutions must follow the technical progress and protect the interests and secrets of entrepreneurs using SI systems. Such institutions' task is to ensure that the SI operation rules are transparent and explainable. In addition, they provide the necessary support to protect the rights of employed persons. Traditional instruments arising from labor law, personal data protection, and consumer protection are insufficient for this aim. One of the problems with this lies in the algorithms themselves, especially the reading and understanding of them.

Such international institutions are mainly intended to ensure the cybersecurity and safety of users and systems during an AI algorithm's full life cycle. Algorithmic management failures and prejudices can potentially harm employees whose rights and obligations arise directly from the system's proper functioning. The AI certification system should be based on two pillars: first, a thorough assessment of the risk of the impact of AI on its environment and the risks arising from it, and second, a comprehensive artificial intelligence testing system aimed at achieving transparency and explainability (Hamon, Junklewitz, Sanchez 2020). The fundamental question here is what kind of algorithms related to the working environment's management process should be certified by international institutions. Indeed, this question applies to the type of algorithmic management whose system is outside the framework of individual countries. Particularly, this includes crowdworking and all the platforms that standardize the algorithmic management process at the supranational level. Otherwise, the regulatory authorities' algorithmic governance process should be controlled at the national level with the unions' active involvement. A thorough expert analysis should precede the certification process at the international level to examine whether a type of artificial intelligence algorithm should be certified. Undoubtedly, this requires the development of a methodology for assessing the impact of an AI system on society, in our case, on employees managed by a dedicated algorithm. Following on from the European Parliament resolution of 16 February 2017 with recommendations to the Commission on Civil Law Rules on Robotics, (2015/2103(INL)), OJ C 2018/252, pp. 239–257, the European Parliament underlined the importance of standards and ensuring interoperability, and called on the Commission to continue its work on the harmonization of technical standards at the international level,

in particular, in cooperation with the European Organization for Standardization and the ISO, to promote innovation, avoid fragmentation of the internal market, guarantee a high level of product and consumer safety, and ensure safety standards in the working environment (European Parliament 2017).

The question of which international organizations will be responsible is not easy to answer. There is the ILO for labor law, the World Intellectual Property Organization for intellectual property law, and the World Trade Organization for international trade law. The Organisation for Economic Co-operation and Development (OECD) could be a relevant actor, since it has adopted corporate governance guidelines and multinational companies, and has a digital agenda (Rosenblat, Stark 2016). There is certainly a lot to consider in this regard. This is an interesting space for further scientific research.

The development of new technologies and machine learning systems generally progresses faster than AI testing methods developed by specialists. There are many serious challenges, such as the problem of legal infringement, which are not entirely capable of adopting further and ongoing technological changes. For this reason, it is essential to establish a system that minimizes the risks associated with the phenomenon of undermining the law of technology (Rosenblat, Stark 2016). This means that the postulated institutional support for control mechanisms is possible and feasible in practice. Such a tool exists in Art. 22 of the General Data Protection Regulation (GDPR). The GDPR, however, does not apply to partially automated decisions (made with human participation) or AI decisions concerning groups of people (residents of a specific area, people of a given sex, various types of minorities). The regulation of the requirement to ensure SI transparency (including explainability) should, therefore, be broader and require entities using autonomous systems to document the parameters and instructions provided as input data for the SI and factors influencing the forecasts made by AI appropriate explanations.

The book *Responsive Regulation: Transcending the Deregulation Debate*, published in 1992, has become a central work in the field of regulatory scholarship. It was written in collaboration with PJan Ayres (Yale University) and John Braithwaite (Australian National University) and was built on Braithwaite's previous research on regulation, enforcement, and compliance. Since its inception, responsive regulation has been an approach based on evidence from economics and social science on how best to regulate individuals and organizations. Dissuasive measures are important economic compliance principles that discourage individual choices and provide incentives to make other options more attractive. This model aims to implement legislation with soft instruments to encourage its adoption. A gradual transition to incentives or penalties is recommended, based on their cost and level of management difficulties. Using mechanisms to stimulate entities, on the one hand, shows them the benefits of implementing regulation and, on the other, applies substantial penalties for those who knowingly violate the regulatory principles introduced. This regulatory concept is supported by the model developed by Ayres and Braithwaite, termed the "implementation pyramid" (Rogers, Ayres, Braithwaite 1993).

The term “pyramid” indicates that these instruments cannot be used in any order. The strength of the regulated entities’ effectiveness and sensitivity to the stimulation tools adopted must be preserved. This model indicates the regulator’s central position, which cannot limit itself to drawing up the regulation and implementing it as quickly as possible but must also focus on control (quality of enforcement) and building relationships with regulated entities. Moreover, it points out that legislation cannot be implemented effectively and efficiently if only sanctions and intimidation options for regulated entities are adopted. These principles can be read from the perspective of SI certification and regulation, especially in the context of algorithmic management in the workplace. To facilitate the work of prioritizing stimulus tools and to try to simplify the complexity of this concept, Braithwaite proposed the application of some heuristic principles (Braithwaite 2011).

First, the regulator should consider the broader context of regulation without considering any theoretical concept. To a large extent, the regulation process should refer to historical experience to better understand the range of forces (technological, economic, and socio-political) at work. It shows that all main features of algorithmic management were visible in earlier periods of capitalism, but became less visible with the emergence of the “standard employment relationship” in the 20th century. The increase and decline of the standard employment relationship existed in relation to the changing context of private employers’ efforts to attract labor. Improving understanding of the full range of the drivers of change in a work organization and rejecting the assumption that change is technologically determined and, therefore, inevitable can influence regulatory and political responses to the development of algorithmic management work (Stanford 2017). Braithwaite also emphasizes the need to think in line with the spirit of the times and recognize the importance of history in drafting regulations. In his opinion, historians act like journalists and ask many questions, creating clearer premises for regulation and better prospects for its implementation.

Second, the regulator should be active and conduct dialog that allows the stakeholders to express themselves, define the agreed effects of the regulation and how they should be monitored, build commitment by helping regulated entities to find incentives to improve the achievement of the regulatory objectives and maintain a communication system until the problem is resolved. The certification process needs to involve trade unions, which could play an essential consultative role. Unions are also becoming increasingly active in research and cooperation concerning algorithmic management. For platforms, the aim is to update information on the dynamically changing platform work phenomenon and develop ways of responding.

In Italy, the *Unione Italiana Lavoratori Turismo Commercio Servizi* (UILTuCS) has set up an observatory for contracts based on the mediation of digital platforms, which looks at platform workers’ working conditions. In addition, committees have been set up at the local level to monitor platform work. In France, the *Institute Research Économiques Et Sociales* (IRES) and *Astrees* set up the *Sharers&Workers* network in 2015 to bring together the stakeholders (trade unions, researchers, experts, and public bodies) of the

digital economy to reflect together on the future of work and social relations. At the EU level, in 2019, the European Trade Union Confederation (ETUC), IRES, and ASTREES established the European Observatory of Digital Work Platforms. It aims to map and assess the existing practices of worker representation and social dialog within the platforms and develop new representation and dialog methods with stakeholders and, in particular, with platform staff at the European level.

Trade unions have developed a range of services aimed directly at platform workers or at a broader group of independent workers (for example, self-employed, civil law, contract workers) whose *modus operandi* is similar to everyday work. Online services such as *faircrowd.work*—an online platform launched by IG Metall, the Austrian Chamber of Labour, the Austrian Confederation of Trade Unions, and Unionen to evaluate (rank) digital platforms from the perspective of employees. The evaluation was based on questionnaires filled in by employees and the exchange of experience—have been developed. The portal also contains other useful information, such as basic knowledge about platform work, contacts with platform workers' trade unions, and a link to the "Frankfurt Declaration" (*selbstaendigen.info*). Another example is the online platform operated by *Ver.di*, aimed at independent contractors (self-employed, platform workers, etc.), providing legal, tax, and contract-related advice to both trade union members and other independent contractors. The latter group is charged a fee for advice, while trade unionists receive support through membership dues *turespuestasindical.es*, an online platform run by the *Unión General de Trabajadores* to provide information and advice platform workers in Spain. The platform's activities include information, enforcement of rights, organization and termination issues, and the portal serves as a tool to promote trade union membership. For example, *precaritywar.es*, an online platform run by the Workers' Commissions (CCOO), is aimed at atypical workers to provide legal support to platform workers and offer evidence in Madrid; while a trade union freelance platform, created by the French Democratic Confederation of Labour (CFDT) offers support in accounting, social and professional insurance and supplementary health insurance, and legal advice and mutual assistance to members. The project was not very active between 2016 and 2019, but in 2019, it gained momentum with the creation of the association, supported by the CFDT but remaining independent of this federation.

According to Braithwaite, the regulatory process should be characterized by cooperation in which the regulator assumes that regulated entities have everything necessary to make changes. The regulator considers the values, motivations, skills, and resources of regulated entities to help them achieve the desired changes (J. Braithwaite and Director 2011). According to Grabosky, Ayres and Braithwaite show how external stakeholders' activities can stimulate regulators and reduce the likelihood of regulation being "captured" by the private sector. Third, parties can also directly influence the addressees of regulation by mobilizing negative publicity following a serious breach or, conversely, by acting constructively beyond the certification body's regulatory gaze (Grabosky n.d.). Similarly, Christine Parker writes about the participation of external stakeholders by introducing a porous self-regulation structure. The management of a company is

open to a deliberative dialog with stakeholders to assess a proposed regulation's value (Parker 2002). It seeks to arouse and use regulated entities' concerns to make changes to help them resolve problems and achieve their goals, and the management should focus on regulating entities' opinions and conducting a dialog that strengthens their motivation to implement regulations.

The certification system's scope of tasks as an instrument of legal control of algorithms is associated with ensuring a stable and reliable testing method. It is not without learning that the legal environment itself can undertake certification activities and take proper account of the legality of operating individual testing tools (Schubert, Hütt 2019). Undoubtedly, this is an important area for further scientific research, especially in scientific cognition, such as computer science and the data analysis of social and legal science.

Another question concerns whether, when establishing such an agency and defining its competences and responsibilities, ex-post algorithm control procedures should be applied, or whether an ex-ante certification process should be applied. This second solution would lead to a significant change in the commonly used control paradigm, such as in competition law, anti-discrimination law, or data protection law based on ex-post control. It seems that the legitimacy of applying the ex-ante rule can only be achieved when there is a very high risk of the algorithm violating human rights. This is a highly repressive procedure, including a preventive or categorical ban on the use of specific algorithms for people at risk. As part of ex-ante control, there is also the option of a prophylactic ban with the conditional option of authorization (consent) if the algorithm's specific direction is reasonable and generally desirable. This could be the case in areas of technological development, especially in cases of potential product failures that could threaten to cause irreversible damage to health or life. Accordingly, the preventive ban may only apply to algorithms that have a potential negative impact on the health, safety, or life of the persons concerned.

For algorithmic management algorithms for machine learning, as long as it is still widely accepted that, despite numerous controversies, they promote the good of society and the economy, the application of severe legal rigor associated with ex-ante control should not apply. On the other hand, the fundamental problems of the lack of transparency of algorithmic rules seriously complicate ex-post control enforcement. For this reason, it is essential to develop XAI techniques in which social organizations, including trade unions, can play an important role. Unions should, therefore, not only make increasingly conscious and specialized use of XAI techniques but, just as importantly, should support the certification process for ex-post controls substantively (Schubert, Hütt 2019).

Conclusion

Developing AI and algorithmic management forces trade unions to become more active and acquire new qualifications. In the working environment, there should be a principle of human importance, with a focus on personalism. Among the ethical controversies regarding algorithmic management's functioning, some whole groups of problems can be identified. These include the issues resulting from the false assumption of AI's faultlessness and legal responsibility for their consequences. Its manifestations may consist of phenomena such as algorithmic discrimination, algorithmic exclusion, and dehumanization of work. Other challenges include uncontrolled data collection and processing by AI and the problems of AI's non-transparent decision-making processes. For the practical defense of individual and collective labor rights, trade unions should, as a matter of priority, work toward transparency at all stages of the implementation of the SI, starting from the design stage, and moving through development, personalization (customization to the needs of the user purchaser), and operation (the periodic validation of the effectiveness of systems using the AI). This should include both regular tests and cyclical external audits.

The purpose and scope of the legal requirements for AI transparency may vary depending on the recipients' profiles. In the case of trade unions, AI transparency should be considered from the employees' perspectives. The essence of algorithmic management is to support decision-making. If so, the algorithmic rule may help with the management of humans. It may also constitute a fundamental element of the working environment, where middle-level managers are not involved, and the entire management system is based on AI.

In such a situation, trade unions can stabilize and control roles in the work process, mainly in traditional working environments based on the bilateral employment model, where traditionally, in an earlier period, unions played an important role, or where they could overcome specific difficulties during their activities. In work environments with global internet platforms, trade unions can play a reactive role by gathering information and processing initiatives. However, they cannot take proactive measures to influence the scope of decisions supported by algorithmic management. In the latter case, they can maintain the certification procedures indicated by various institutions investigating the ethical aspects of AI. Such an institution could be the OECD.

The certification system, as an instrument for the legal control of algorithms, requires a stable and reliable testing method. Learning is needed for the legal environment itself to undertake certification activities and take proper account of the legality of individual testing tools. The theory of responsive regulation may become a useful pragmatism for the operation of future AI certification bodies and trade unions. Ayres and Braithwaite argue that strictly government-led regulatory policies, mandates, and control are often not the best ways to solve social problems. However, neither are laissez-faire policies based on market competition. Ayres and Braithwaite introduced responsive regulation as an overall regulatory strategy that seeks to build on both approaches' strengths and overcome their weaknesses. This is "an attitude that enables a wide range of regulatory

approaches to flourish and not a clearly defined programme or set of recommendations on how best to regulate.” This is undoubtedly an essential space for further scientific research, especially for fields like computer science and social and legal science data analysis. However, it must be understood that effective certification requires the transformation of traditional standards into XAI techniques.

References

- Adams-Prassl J. (2020) *When Your Boss Comes Home: Three Fault Lines for the Future of Work in the Age of Automation, AI, and COVID-19*, “Ethics of AI in Context.”
- AI HLEG [High-Level Independent Group on Artificial Intelligence] (2019a) *Ethics Guidelines for Trustworthy AI*, <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai> (access: 18 December 2020).
- AI HLEG (2019b) *A Definition of AI: Main Capabilities and Disciplines*, <https://digital-strategy.ec.europa.eu/en/library/definition-artificial-intelligence-main-capabilities-and-scientific-disciplines> (access: 18 December 2020).
- AI HLEG (2020) *Assessment List for Trustworthy AI (ALTAI)*, <https://digital-strategy.ec.europa.eu/en/library/assessment-list-trustworthy-artificial-intelligence-altai-self-assessment> (access: 18 December 2020).
- Amyx S. (n.d.) *Wearing Your Intelligence: How to Apply Artificial Intelligence in Wearables and IoT*, Wired. Com, <https://www.wired.com/insights/2014/12/wearing-your-intelligence/> (accessed: 2 December 2020).
- Aneesh A. (2009) *Global Labor: Algocratic Modes of Organization*, “Sociological Theory,” Vol. 27, Issue 4.
- Barocas S., Selbst A.D. (2016) *Big Data’s Disparate Impact*, “California Law Review,” Vol. 104, No. 671.
- Barredo A.A., Díaz-Rodríguez N., Del Ser J., Bennetot A., Tabik S., Barbado A., Garcia S. et al. (2020) *Explainable Artificial Intelligence (XAI): Concepts, Taxonomies, Opportunities and Challenges toward Responsible AI*, “Information Fusion,” Vol. 58.
- Beer D. (2017) *The Social Power of Algorithms*, “Information, Communication & Society,” Vol. 20, Issue 1.
- Benjamins R., Barbado A., Sierra D. (2019) *Responsible AI by Design in Practice*, ArXiv, no. Ec.
- Bhatt U., Xiang A., Sharma S., Weller A., Taly A., Jia Y., Ghosh J., Puri R., Moura J.M.F., Eckersley P. (2020) *Explainable Machine Learning in Deployment* [in:] *FAT* 2020: Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency*, New York.
- Braithwaite J. (2011) *The Essence of Responsive Regulation*, *Fasken Lecture*, “University of British Columbia Law Review,” Vol. 44, Issue 3.
- Braithwaite J. (2013) *Relational Republican Regulation*, “Regulation and Governance,” Vol. 7, Issue 1.
- Braithwaite V., Braithwaite J. (2001) *An Evolving Compliance Model for Tax Enforcement* [in:] N. Shover, J.P. Wright (eds.), *Crimes of Privilege: Readings in White-Collar Crime*, New York.

- Brandy J.L. (2018) *Artificial Intelligence ('AI') in the Legal Profession*, "Law Audience Journal", Vol. 2, Issue 3.
- Buchanan B.G., Duda R.O. (1983) *Principles of Rule-Based Expert Systems*, "Advances in Computers" 22(C).
- Carabantes M. (2020) *Black-Box Artificial Intelligence: An Epistemological and Critical Analysis*, "AI and Society", Vol. 35 (2).
- Cherry M. (2016) *Beyond Misclassification: The Digital Transformation of Work*, "Comparative Labor Law and Policy Journal", Vol. 37, No. 3.
- Dameski A. (2018) *A Comprehensive Ethical Framework for AI Entities: Foundations*, "Lecture Notes in Computer Science", Vol. 10999 LNAI.
- Davidow B. (2014) *Welcome to Algorithmic Prison: The Use of Big Data to Profile Citizens is Subtly, Silently Constraining Freedom*, Washington, DC.
- De Stefano V. (2020a) *Algorithmic Bosses and What to Do About Them: Automation, Artificial Intelligence and Labour Protection*, "Studies in Systems, Decision and Control", https://doi.org/10.1007/978-3-030-45340-4_7.
- De Stefano V. (2020b) *'Master and Servers': Collective Labour Rights and Private Government in the Contemporary World of Work*, "International Journal of Comparative Law Law and Industrial Relations," Vol. 4, Issue 36.
- Deloitte (2020) *Trustworthy AI™: Bridging the ethics gap surrounding AI*, <https://www2.deloitte.com/us/en/pages/deloitte-analytics/solutions/ethics-of-ai-framework.html> (access: 18 December 2020).
- European Parliament (2017) European Parliament resolution of 16 February 2017 with recommendations to the Commission on Civil Law Rules on Robotics (2015/2103(INL)).
- Ferraris V., Bosco F., D'Angelo E. (2013) *The Impact of Profiling on Fundamental Rights*, Working Paper No. 3, https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2366753.
- Finn E. (2017) *What Algorithms Want: Imagination in the Age of Computing*. What Algorithms Want: Imagination in the Age of Computing. The MIT Press.
- Gillespie T (2014) *The Relevance of Algorithms* [in:] T. Gillespie, P.J. Boczkowski, K.A. Foot KA (eds.), *Media Technologies: Essays on Communication, Materiality and Society*, 167–194. MIT Press: Cambridge, MA.
- Goldstein B., Laurenand D., Nemani (eds.) (2013) *Beyond Transparency: Open Data and the Future of Civic Innovation*. AI Matters, "AI Matters," Vol. 5, Issue 2.
- Grabosky P. (n.d.) *Meta-Regulation* [in:] *Regulatory Theory*, <https://about.jstor.org/terms> (access: 18 December 2020).
- Guo W. (2020) *Explainable Artificial Intelligence for 6G: Improving Trust between Human and Machine*, "IEEE Communications Magazine," Vol. 58, Issue 6.
- Hamon R., Junklewitz H., Sanchez I. (2020) *Robustness and Explainability of Artificial Intelligence*, <https://doi.org/10.2760/5749>.
- High-Level Independent Group on Artificial Intelligence (AI HLEG) (2019) *A Definition of AI: Main Capabilities and Disciplines*, "European Commission," 7, <https://ec.europa.eu/digital-single->.

- High-Level Independent Group on Artificial Intelligence (AI HLEG) (2020) *Assessment List for Trustworthy AI (ALTAI)*, <https://ec.europa.eu/digital-single-market/en/news/assessment-list-trustworthy-artificial-intelligence-altai-self-assessment> (access: 18 December 2020).
- How J.P. (2017) *Ethically Aligned Design: A Vision for Prioritising Human Well-Being with Autonomous and Intelligent Systems–Version 2*, “IEEE Control Systems Magazine,” Vol. 38, Issue 3.
- ICO (n.d.) *Guidance on AI and Data Protection*, <https://ico.org.uk/for-organisations/guide-to-data-protection/key-data-protection-themes/guidance-on-ai-and-data-protection> (access: 18 December 2020).
- IEEE (n.d.) *IEEE Ethics In Action in Autonomous and Intelligent Systems*, <https://ethicsinaction.ieee.org/> (access: 18 December 2020).
- ILO (2018) *Digital Labour Platforms and the Future of Work: Towards Decent Work in the Online World*. International Labour Office, https://www.ilo.org/wcmsp5/groups/public/---dgreports/-ddcomm/---publ/documents/publication/wcms_645337.pdf, (access: 18 December 2020).
- Irani L. (2015) *The Cultural Work of Microwork*, “New Media & Society,” Vol. 17, Issue 5.
- Kaminski M.E. (2018) *The Right to Explanation, Explained*, “SSRN Electronic Journal,” <https://doi.org/10.2139/ssrn.3196985>.
- Karanasiou A., Pinotsis D. (2017) *Towards a Legal Definition of Machine Intelligence: The Argument for Artificial Personhood in the Age of Deep Learning* [in:] *Proceedings of the International Conference on Artificial Intelligence and Law*, New York.
- Khaleghi B. (2019) *The How of Explainable AI: Explainable Modelling*, Medium, <https://medium.com/towards-data-science/the-how-of-explainable-ai-post-modelling-explainability-8b4cbc7adf5f> (access: 18 December 2020).
- Kim N. (2014) *The Wrap Contract Morass*, “Southwestern University Law Review.”
- Lee M.K., Kusbit D., Metsky E., Dabbish L. (2015a) *Working with Machines: The Impact of Algorithmic and Data-Driven Management on Human Workers* [in:] *CHI '15: Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems*, New York.
- Lee M.K., Kusbit D., Metsky E., Dabbish L. (2015b) *Working with Machines: The Impact of Algorithmic and Data-Driven Management on Human Workers* [in:] *CHI '15: Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems*, New York.
- Leslie D. (2020) *Understanding Artificial Intelligence Ethics and Safety: A Guide for the Responsible Design and Implementation of AI Systems in the Public Sector*, “SSRN Electronic Journal,” <https://doi.org/10.2139/ssrn.3403301>.
- Maltseva K. (2020) *Wearables in the Workplace: The Brave New World of Employee Engagement*, “Business Horizons,” Vol. 63, Issue 4.
- Mann G., O’Neil C. (2016) *Hiring Algorithms Are Not Neutral*, “Harvard Business Review.”
- Mateescu A., Nguyen A. (2019) *Algorithmic Management in the Workplace*, “Data & Society Research Institute.”
- Mazmanian M., Orlikowski W.J., Yates J.A. (2013) *The Autonomy Paradox: The Implications of Mobile Email Devices for Knowledge Professionals*, “Organization Science,” Vol. 24, Issue 5.
- Mittelstadt B.D., Allo P., Taddeo M., Wachter S., Floridi L. (2016) *The Ethics of Algorithms: Mapping the Debate*, “Big Data and Society,” Vol. 2, Issue 2.

- Moore Ph.V. (2020) *Data Subjects, Digital Surveillance, AI and the Future of Work*.
- Moore Ph.V., Upchurch M., Whittaker X. (2018) *Humans and Machines at Work: Monitoring, Surveillance and Automation in Contemporary Capitalism*, Cham.
- Murray W.C., Rostis A. (2007). 'Who's Running the Machine?' A Theoretical Exploration of Work Stress and Burnout of Technologically Tethered Workers, "Journal of Individual Employment Rights," Vol. 12, Issue 3.
- Noto La Diega G. (2018) *Against the Dehumanisation of Decision-Making: Algorithmic Decisions at the Crossroads of Intellectual Property, Data Protection, and Freedom of Information*, "Journal of Intellectual Property, Information Technology and Electronic Commerce Law," Vol. 9, Issue 1.
- Page T. (2015) *A Forecast of the Adoption of Wearable Technology*, "International Journal of Technology Diffusion," Vol. 6, Issue 2.
- Parker Ch. (2002) *The Open Corporation: The Open Corporation*.
- Pasquale F. (2015) *THE BLACK BOX SOCIETY The Secret Algorithms That Control Money and Information*, <http://raley.english.ucsb.edu/wp-content/Engl800/Pasquale-blackbox.pdf> (access: 18 December 2020).
- Preece A., Harborne D., Braines D., Tomsett R., Chakraborty S. (2018) *Stakeholders in Explainable AI*, *ArXiv*.
- Renan Barzilay A., Ben-David A. (2018) *Platform Inequality: Gender in the Gig-Economy*, "Seton Hall Law Review," Vol. 47, No. 393.
- Ribeiro M.T., Singh S., Guestrin C. (2016) 'Why Should i Trust You?' Explaining the Predictions of Any Classifier [in:] *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, New York.
- Rogers J., Ayres Y., Braithwaite J. (1993) *Responsive Regulation: Transcending the Deregulation Debate*, "Contemporary Sociology," Vol. 22, No. 3.
- Rosenblat A., Stark L. (2016) *Algorithmic Labor and Information Asymmetries: A Case Study of Uber's Drivers*, "International Journal of Communication," Vol. 10.
- Rudin C. (2018) *Please Stop Explaining Black-Box Models for High Stakes Decisions*, *NIPS*.
- Schubert C., Hütt M.T. (2019) *Economy-on-Demand and the Fairness of Algorithms*, "European Labour Law Journal," Vol. 10, Issue 1.
- Sileno G., Boer A., van Engers T. (2019) *The Role of Normware in Trustworthy and Explainable AI* [in:] *CEUR Workshop Proceedings*.
- Spreeuwenberg S. (2020) *Choose for AI and for Explainability* [in:] Ch. Debruyne, H. Panetto, W. Guédria, P. Bollen, I. Ciuciu, G. Karabatis, R. Meersman (eds.), *On the Move to Meaningful Internet Systems: OTM 2019 Workshops*, Cham (series: "Lecture Notes in Computer Science," Vol. 11878).
- Stanford J. (2017) *The Resurgence of Gig Work: Historical and Theoretical Perspectives*, "The Economic and Labour Relations Review," Vol. 28, Issue 3.
- Stewart A., Stanford J. (2017) *Regulating Work in the Gig Economy: What Are the Options?*, "The Economic and Labour Relations Review," Vol. 28, Issue 3.
- Stone P., Brooks R., Brynjolfsson E., Calo R., Etzioni O., Hager G., Hirschberg J. et al. (2016) *Artificial Intelligence and Life in 2030: One Hundred Year Study on Artificial Intelligence*, "Stanford University," Vol. 52.

- Till A.L., Ratcheva V.S., Zahidi S. (2018) *Future of Jobs Report 2018*, "World Economic Forum;," <http://reports.weforum.org/future-of-jobs-2018/> (access: 18 December 2020).
- UNI Global Union (2017) *Top 10 Principles for Ethical Artificial Intelligence*, 10, www.uniglobalunion.org.
- van der Heijden J. (2020) *Responsive Regulation in Practice*, "Regulatory Governance Research."
- van den Heuvel S., Bondarouk T. (2017) *The Rise (and Fall?) Of HR Analytics*, "Journal of Organizational Effectiveness: People and Performance," Vol. 4, No. 2.
- Winfield A.F., Michael K., Pitt J., Evers V. (2019) *Machine Ethics: The Design and Governance of Ethical Ai and Autonomous Systems*, "Proceedings of the IEEE," vol. 107, Issue 3.
- World Economic Forum (2018) <http://reports.weforum.org/future-of-jobs-2018/shareable-infographics/2018>.
- Yeung K. (2018) *Algorithmic Regulation: A Critical Interrogation*, "Regulation and Governance," Vol. 12, Issue 4, <https://doi.org/10.1111/rego.12158>.
- Zuiderveen Borgesius F.J., Poort J. (2017) *Online Price Discrimination and EU Data Privacy Law*, "Journal of Consumer Policy," Vol. 40.