

MATHEMATICAL METHODS, MODELS AND INFORMATION TECHNOLOGIES IN ECONOMY

UDC 351:004.8:331.108.45:303.732
JEL H83, J71, M53, O33

DOI: 10.26906/EiR.2025.4(99).4169

MATHEMATICAL METHODS OF THE DIGITALIZATION OF HR PROCESSES IN PUBLIC ADMINISTRATION

Oliinyk Oleksandr*, Candidate of Philosophical Sciences,
Associate Professor at the Department of Business Administration
and Management of Foreign Economic Activity

Bikulov Damir**, Doctor of Public Administration,
Professor of the Department of Business Administration
and Management of Foreign Economic Activity

Holovan Olha***, Candidate of Physical and Mathematical Sciences,
Associate Professor of the Department of Business Administration
and Management of Foreign Economic Activity

Markova Svitlana****, Doctor of Economic Sciences,
Professor of the Department of Business Administration
and Management of Foreign Economic Activity

Veritova Olha*****, Candidate of Pedagogical Sciences,
Senior Lecturer of the Department of Business Administration
and Management of Foreign Economic Activity
Zaporizhzhia National University

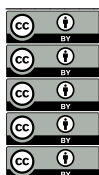
**ORCID 0000-0003-0511-7681*

***ORCID 0000-0001-9188-7310*

****ORCID 0000-0002-9410-3830*

*****ORCID 0000-0003-0675-0235*

******ORCID 0000-0001-7365-701X*



© Oliinyk O., 2025

© Bikulov D., 2025

© Holovan O., 2025

© Markova S., 2025

© Veritova O., 2025

Стаття отримана: 03.11.2025

Стаття прийнята: 22.11.2025

Стаття опублікована: 26.12.2025

Introduction. Public administration worldwide is undergoing a phase of deep digital transformation, actively implementing artificial intelligence technologies to modernize key state functions, with the human resource management sector being at the forefront of these changes. The motivation for introducing algorithmic systems into recruitment processes is based on the aspiration to increase objectivity, efficiency, and decision-making accuracy, as well as on the paradoxical expectation that machines will help neutralize the cognitive biases and prejudice inherent in humans. However, empirical reality shows that algorithmic management, especially when implemented in the form of opaque black boxes, is not a neutral instrument. Instead of eliminating

discrimination, it carries inherent risks of scaling and amplifying existing systemic imbalances and biases, which leads to new and previously unseen forms of automated discrimination. For public administration, in contrast to the private sector where failures result in reputational and legal costs, this problem has an existential character. Automated discrimination originating from the state apparatus, which by definition is supposed to guarantee equality, directly delegitimizes state power, creating a challenge for all three pillars of legitimacy, namely: 1) input legitimacy due to the non transparency of decision making; 2) procedural legitimacy due to the violation of the principle of fairness of the process; 3) output legitimacy due to the unfair distribution of outcomes or positions. Thus, the key problem lies in the fundamental conflict between the technocratic pursuit of efficiency and the inherent public values such as due process, equality, and accountability.

Analysis of recent research and publications. The conducted analysis of academic research in the field of mathematical methods for evaluating the fairness of digital HR systems in public administration demonstrates their significant methodological incompatibility, fragmentation, and the absence of an integrated normative logic. The work of Köchling A., Wehner M.C. [1] systematizes the technical sources of bias, but remains largely descriptive in nature and does not propose mechanisms for reconciling conflicting metrics. The study by Starke C., Lünich M. [2] shifts the focus toward the political legitimacy of algorithmic decisions and shows that fairness does not exist outside a value based context. Between these approaches, an epistemological gap is observed that has no instrumental solution within existing models. The work of Bansak K., Ferwerda J., Hainmueller J., Dillon A., Hangartner D., Lawrence D., Weinstein J. [3] demonstrates the effectiveness of algorithmic allocation in administrative processes, but does not provide fairness mechanisms, thus creating a gap between technical feasibility and the legitimacy of application. The Fair Machine Learning approach proposed by Mujtaba D.F., Mahapatra N.R. [4] introduces ethical constraints into mathematical models, yet remains weakly adapted to the complex regulatory environment of public administration. The study by Fregnan E., Ivaldi S., Scaratti G. [5] shows that even technically advanced models are ineffective without appropriate managerial competences, which forms a second gap between technology and administrative capacity. The analysis of HRM platforms in the work of Zeebaree S.R., Shukur H.M., Hussan B.K. [6] reveals the absence of mechanisms for assessing their compliance with non discrimination principles. The study by Acikgoz Y., Davison K.H. [7] demonstrates the divergence between mathematical and perceived fairness, creating a third gap between technical correctness and social acceptability. The work by DiRomualdo A., El-Khoury D., Girimonte F. [8] emphasizes the systemic nature of digital HR transformation, but likewise does not offer a comprehensive mathematical solution. Thus, no study provides a mechanism for aligning these dimensions within a unified mathematical framework. This cumulative imbalance explains why the research problem has not previously been solved, since existing studies lack a model capable of integrating technical metrics, ethical constraints, legal norms, and public administration requirements into a coherent quantitative system. Therefore, the “CIF-PA” model proposed by the authors will address these gaps and create an instrument that forms a unified and consistent logic of fairness governance in the public sector.

Objectives of the article. Develop and provide scientific justification for a composite mathematical toolkit for the quantitative evaluation and auditing of fairness and non-discrimination in personnel selection algorithms within the public administration system. In contrast to existing approaches that focus on the analysis of individual and often mathematically incompatible fairness metrics, this study proposes a synergistic “CIF-PA” model that integrates technical metrics, ethical requirements for transparency and oversight, as well as regulatory and legal requirements reflected in the provisions of the EU Regulation “AI Act”, into a single multi-criteria “MCDA” evaluation system adapted for managerial decision making in conditions of high risk and uncertainty inherent in the public sector.

The main material of the study. Modern HR technologies are fundamentally transforming the entire personnel management cycle by combining automated resume screening, neural networks, and natural language processing tools for analyzing job descriptions and candidate responses, as well as algorithmic systems for evaluating video interviews. Gamified methods for identifying cognitive and personality characteristics and predictive analytics that forecast candidates’ potential performance based on large volumes of historical data are becoming increasingly widespread. The scale of such implementation is growing extremely rapidly, since most large organizations already use various forms of recruitment automation. However, such rapid technological development creates a noticeable gap between the speed of their emergence and the availability of adequate control mechanisms. New tools are introduced much faster than standards, regulatory requirements, and clear rules for their use are developed, especially in the public sector. As a result, this space of uncertainty is

often filled by commercial systems whose operation remains opaque. For this reason, mathematical methods for evaluating fairness cease to be a purely theoretical task and become a necessary instrument for ensuring accountability in the procurement and operation of high-risk algorithms in public administration.

Therefore, in this context, there emerges a need to systematize approaches to the quantitative measurement of fairness through their classification and the characterization of key mathematical methods for evaluating AI fairness, which are presented in Table 1.

We classify and characterize the key mathematical methods for evaluating AI fairness in Table 1.

Table 1

Classification and characterization of key mathematical methods for assessing the fairness of AI

Method/metric name	Method description	Interpretation of the method in the context of personnel selection
Statistical Parity	Proportion of selected (Group A) = Proportion of selected (Group B)	The proportion of candidates selected should be the same for all protected groups, regardless of their actual qualifications.
Equal Opportunity	Chance of selection (for Qualified Group A) = Chance of selection (for Qualified Group B)	Qualified candidates from all groups should have the same chance of being selected (equal True Positive Rate, TPR).
Equalized Odds	Equality of chances (TPR and FPR) for Group A = Equality of chances (TPR and FPR) for Group B	Equal chances for both qualified (equal TPR) and unqualified (equal False Positive Rate, FPR).
Predictive Parity	Prediction accuracy (for Selected Group A) = Prediction accuracy (for Selected Group B)	Among those recommended by the algorithm, the proportion of truly qualified candidates is the same for all groups (equal Positive Predictive Value).
Adverse Impact Ratio	(Proportion of selected Group A) / (Proportion of selected Group B) > 0.8	Legal concept of the "4/5 Rule" ("EEOC"). The selection rate for a protected group should not be lower than 80% of the rate for the dominant group.
Counterfactual Fairness	Outcome(candidate, Group A) = Outcome (same candidate if he were in Group B)	The decision about a particular candidate would not change if the only change were his protected characteristic (e.g., gender), all other factors being held constant. Requires a causal model.
Individual Fairness	"Similar" candidates (x, z) get "Similar" decisions (f(x), f(z))	"Similar individuals should receive similar treatment". The main difficulty is the mathematical definition of the "similarity" (d) of individuals x and z.
Accuracy Parity	Overall accuracy (Group A) = Overall accuracy (Group B)	The overall accuracy (percentage of correct solutions) of the algorithm is the same for all protected groups.

Source: formed on the basis of the following sources [9–13]

The system of artificial intelligence fairness evaluation metrics presented in Table 1 begins with an examination of demographic parity, which is based on the principle of equality of outcomes and requires the same proportion of approved candidates for all groups, while ignoring the actual level of their professional qualifications [9, 10]. This approach carries the risk of reducing staff quality in favor of statistical balance, therefore a more advanced alternative is the concept of equality of opportunity, which focuses exclusively on qualified applicants and guarantees them identical chances of successful selection regardless of social characteristics [9, 11]. A logical extension of this meritocratic principle is the method of equalized odds, which balances not only correct decisions but also the algorithm's errors, thereby ensuring the same probability of both correct selection and erroneous rejection for all categories of participants [10, 11]. In parallel with this, predictive parity is applied, which tests the reliability of the model's predictions by comparing the proportion of truly competent specialists among all individuals recommended by the system across different demographic groups [10, 12, 13]. The legal aspect of fairness is regulated by the adverse impact ratio, which transforms mathematical calculations into the four-fifths rule and determines the permissible limits of selection disproportionality for protected population categories. A deeper level of verification is provided by causal fairness, which uses causal models to analyze whether the decision regarding a candidate would change if only their protected attribute were altered while other parameters

remained constant [9, 11]. In contrast to group-based approaches, individual fairness requires equal treatment of candidates with similar professional profiles and is based on a mathematical definition of similarity between individuals [9, 10]. The final evaluation element is accuracy parity, which controls the overall correctness of the algorithm's performance and prevents situations in which the system operates more effectively for one demographic group compared to another [10, 11, 13].

The next step will be the classification of sources and manifestations of algorithmic bias at the stages of personnel selection in the public sector, presented in Table 2.

Table 2

**Classification of sources and manifestations of algorithmic bias
at the stages of personnel selection in the public sector**

HR Stage	Bias Type	Source of Bias	Implications for public administration
Sourcing	Exposure Bias	Algorithms of advertising platforms optimize display by cost/clicks and “decide” that a civil servant vacancy is “irrelevant” for certain demographic groups.	Violation of the principle of equal access to the civil service even before the application is submitted.
Screening	Sexism	A model trained on historical data, where most managers are men, and penalizes words associated with women.	Violation of the principle of equal access to the civil service even before the application is submitted. Systemic discrimination against women, undermining the policy of gender equality in government bodies.
Screening	Racism	The model penalizes names associated with minorities or educational institutions	Reduction of ethnic diversity in the civil service, strengthening of systemic discrimination.
Screening	Ageism	Data-based training, where a “recent graduate” is an ideal candidate, because the algorithm is unable to evaluate nonlinear experience.	Violation of labor rights, loss of valuable administrative experience in public administration.
Video Interview	Disability Bias	Pseudoscientific analysis of “emotions”, “confidence” or “body language”. Algorithm penalizes for atypical facial expressions, accent, or speech disorders	Direct discrimination against people with disabilities (speech disorders, hyperkinetic manifestations of the body, autism spectrum, cerebral palsy).
Screening	Cultural Capital Bias	Evaluation of “cultural capital” through prestigious universities, volunteering in the “right” organizations, use of specific vocabulary	Creation of a “glass ceiling” for talented candidates with low socio-economic status, reduction of social mobility through the civil service.
Assessment	Measurement Bias	Psychometric model validated on one cultural group, but not valid for other population groups	Systematic screening of candidates whose thinking style or cultural background does not correspond to the “gold standard” built into the system.
Ranking	Bias Laundering	The algorithm gives the recruiter a “score” because a person is inclined to trust this score and not re-test a candidate with a low score, giving the bias an “objective” appearance.	The algorithm does not replace human bias, but legitimizes it, removing responsibility from the person.
The whole process	Representation Bias	A specific group of candidates, underrepresented in the training data, will systematically receive incorrect scores	The model works well for the “typical” candidate, but systematically fails for the “atypical” one.
The whole process	Feedback Loop Bias	The model learns from the decisions it helped make. (E.g., the model recommends type A candidates, who are then hired. They then become a “successful example” of an effective candidate and begin to recommend even more type A candidates)	“Self-reinforcing inequality”. The system becomes increasingly homogeneous and biased over time.

Source: formed on the basis of the following sources [9–16]

The classification of sources and manifestations of algorithmic bias at the stages of personnel selection in the public sector presented in Table 2 exposes deep structural distortions in the recruitment process, which are initiated already at the sourcing stage, where commercial advertising platform algorithms, guided by the logic of minimizing cost per click, autonomously determine the irrelevance of public service vacancies for certain

demographic groups, thus effectively usurping an administrative function and violating the fundamental right to equal access to public service even before a potential application is submitted [10, 12–14]. This discriminatory dynamic is further reinforced at the stage of automated resume screening, when historical biases embedded in data sets, reflecting the past dominance of men in managerial positions, are transformed into penalty mechanisms for lexical markers associated with women, thereby directly undermining state strategies of gender equality. At the same time, the use of discriminatory variables such as personal names or specific educational institutions automates latent ethnic segregation and blocks cultural diversity [9, 10, 13, 15]. A separate vector of risk is created by age bias in algorithms which, being trained on the profiles of recent graduates as ideal candidates, prove incapable of adequately assessing the non-linear experience of mature professionals, leading to the loss of valuable administrative capital and the violation of labor legislation. At the same time, bias related to cultural capital, through the prioritization of elite universities and specific professional vocabulary, produces an insurmountable “glass ceiling” for talented representatives of lower socio-economic strata, thereby slowing social mobility [10, 11, 15, 16]. The implementation of video interview and gamification technologies is often based on pseudo-scientific interpretations of emotions and body language, which turns into a tool of direct discrimination against persons with disabilities or neurodivergent individuals whose facial expressions or speech deviate from statistical norms, while incorrect validation of psychometric models on a single cultural group leads to the systemic exclusion of candidates with alternative cognitive patterns [12, 13, 16]. This destructive cycle is completed by the phenomenon of bias laundering at the ranking stage, when recruiters under the influence of automation perceive subjective machine scores as objective truth, thereby relieving themselves of responsibility for decisions. At the same time, closed feedback loops, in which the model is trained on its own previously distorted outputs, transform the selection system into a mechanism of self-reinforcing inequality, which over time makes the state apparatus increasingly homogeneous and isolated from society [9, 10, 11, 13, 14].

Next, we consider the regulatory and ethical framework for the governance of “high-risk” AI in HR under the EU “AI Act”, presented in Table 3.

Table 3

Regulatory and ethical framework for regulating “high-risk” AI in HR (EU “AI Act”)

Regulatory Requirement	Source	Essence of Requirement	Implications for Mathematical Assessment
1	2	3	4
High-Risk Classification	EU “AI Act”, Annex AI	HR selection systems are high-risk “by definition” (unless proven otherwise).	Fairness and Risk Audit (AIA) is not an optional but a mandatory legal condition for admission to operation.
Data Governance	EU “AI Act” (Art. 10)	Training and test data must be “relevant, representative, error-free and complete”.	The need for quantitative audit of datasets before training the model (e.g., measuring data bias).
Transparency	EU “AI Act” (Art. 13) European Convention on Human Rights	Providing instructions to users. Ensuring “transparency” and “right to explanation”.	Mathematical methods of XAI become legal requirements to ensure procedural fairness.
Human Oversight	EU “AI Act” (Art. 14)	The system must be designed to “allow” the introduction of human oversight. The principle of “a human makes the final decision”.	The model should be fully autonomous (“human-in-the-loop”). Math. evaluation should include an assessment of interpretability for a public management manager.
Accuracy & Robustness	EU “AI Act” (Art. 15)	Ensuring “appropriate levels of accuracy, reliability and cybersecurity”.	Fairness assessment (should be balanced with classical accuracy metrics)
Non-discrimination	European Convention on Human Rights	An explicit ban on discrimination in AI decisions on protected grounds	A direct requirement to apply group fairness metrics and conduct an audit on “disparate impact”.
Protection of personal data	General “GDPR” regulations	The processing of personal data must be lawful, minimized and have a clear purpose.	Fairness audit requires data on protected features, but their collection and use are strictly regulated.
Record-keeping / Logging	EU “AI Act” (Art. 12)	The system must “automatically record events” (logs) relevant for risk identification and post-audit.	Mathematical evaluation should be reproducible. Audit results (metric values) should be stored as an “audit trail”

Continuation of table 3

1	2	3	4
Impact Assessment	Application recommendations (AIA) Fundamental Rights Impact Assessment	Conducting Algorithmic Impact Assessments (AIA) prior to implementation is necessary to form a proactive assessment of social and human rights risks.	Is the quantitative part and methodology for conducting such AIA.
Accountability	EU “AI Act” (fines) OECD Principles	Determining who is responsible for damage caused by AI. High fines for violations.	The availability of a quantitative documented audit is a key factor in proving due diligence.

Source: generated by the authors

The conducted analysis of the regulatory landscape based on the EU “AI Act” (Table 3) defines AI-based personnel selection systems as high-risk technologies. This status automatically transforms fairness auditing from a voluntary internal practice into a mandatory legal prerequisite for the deployment of such a product. The legislative provisions impose strict requirements on data quality, making preliminary analysis of training datasets for representativeness necessary even before model training begins. Transparency requirements and the right to explanation transform mathematical interpretability methods from purely technical tools into a legal mechanism for ensuring fairness. The principle of human oversight makes the use of fully autonomous “black boxes” impossible and obliges developers to ensure that the results are understandable for public administration managers. At the same time, the direct prohibition of discrimination requires the application of group fairness metrics, which creates a complex dilemma between the need to process sensitive data for auditing purposes and privacy requirements. Systematic record keeping and event logging form a reproducible evidentiary trace that confirms the developer’s due diligence. The presence of such documented auditing becomes a key factor in minimizing legal risks and avoiding significant penalties for violations of regulatory requirements. Concluding the analytical and theoretical review of the key sources and mechanisms of algorithmic bias, it is determined that structural distortions arise already at the sourcing stage through the automatic restriction of access to public vacancies. These distortions are further intensified during resume screening due to historical data inequalities and the use of discriminatory variables. They are deepened by age and cultural biases in the evaluation of professional profiles, complicated by incorrect psychometric validation in video interviews and gamified tests, and finally completed by the phenomenon of bias laundering at the ranking stage, when the algorithm reproduces its own distorted decisions. This indicates the formation of an integrated chain of systemic vulnerabilities that interact with each other and gradually make the recruitment mechanism increasingly homogeneous and alienated from society. Such a multi-level and interconnected structure of violations makes it impossible to assess fairness on the basis of isolated technical indicators alone, which determines the necessity of transition to an integral quantitative measurement methodology capable of covering the full spectrum of risks and transforming them into a unified manageable system.

Therefore, the next logical step is the development of a composite fairness and non-discrimination index, which is presented in Formula 1.

For calculating the composite fairness and non-discrimination index in order to identify critical vulnerability points of algorithmic systems in personnel selection, see Formula 1.

$$CIF - PA = (w_1 + IDQ) + (w_2 + IAE) + (w_3 + IGF) + (w_4 + ICF) + (w_5 + ITE) + (w_6 + IDP) + (w_7 + IRS) + (w_1 + IHO) \quad (f.1)$$

where $CIF - PA$ – composite index of fairness and non-discrimination;

w_i – weight coefficient of the i -th subindex, which is determined by the political regulator;

IDQ – “Data quality and representativeness” normalized subindex indicator;

IAE – “Model efficiency and accuracy” is a normalized subindex indicator;

IGF – “Group justice” is a normalized subindex indicator;

ICF – “Individual and causal justice” is a normalized subindex indicator;

ITE – “Transparency and Clarity” is a normalized subindex indicator;

IDP – “Data protection compliance” is a normalized subindex indicator;

IRS – “Reliability and stability” is a normalized subindex;

IHO – “Support for human supervision” is a normalized subindex score.

Having formed a composite index of fairness and non-discrimination to identify critical points of vulnerability of algorithmic systems in personnel selection (f.1), it would be advisable to develop a methodology for its calculation (Table 4) and an assessment scale.

Table 4

Methodology for calculating a composite index of fairness and non-discrimination to identify critical vulnerability points of algorithmic systems in personnel selection

Subindex	Description	Calculation Methodology (Conceptually)
1	2	3
IDQ	EU “AI Act” (Art. 10) Dataset assessment for representativeness, completeness, model training	1. Calculate the percentage of omissions (in %) 2. Calculate the measure of discrepancy between the data and the target population 3. $IDQ = (1 - \text{Omission right}) * (1 - \text{The measure of divergence})$ 4. If $IDQ < 0.8$ then the data is unsuitable for training. High risk of representation bias
IAE	EU “AI Act” (Art. 15) Evaluation of the basic usefulness of the model. The model must be accurate	1. Calculate the standard integral metric “F1-Score” based on the Precision and Recall indicators, where Precision (selection accuracy) reflects the proportion of actually suitable candidates among all candidates recommended by the algorithm, and Recall (completeness of coverage) characterizes the proportion of suitable candidates identified by the algorithm from the total number of all suitable candidates in the labor market? 1.1. $Precision = \frac{TP(\text{Good candidates})}{TP(\text{Good candidates}) + FP(\text{Redundant candidates})}$ 1.2. $Recall = \frac{TP(\text{Good candidates})}{TP(\text{Good candidates}) + FN(\text{Missed candidates})}$ 1.3. $F_1 = 2 \times \frac{Precision \times Recall}{Precision + Recall}$ 2. Calculate precision parity (AP) $AP = 1 - (Precision(Gp.A) - (Precision(Gp.B)))$ 3. $IAE = F1 \text{ Score} - AP$
IGF	Disparate impact assessment using a set of metrics	1. The government body politically chooses a priority metric (e.g., equality of opportunity (EO)) 2. Violation $V_{EO} = (\text{Chance of selection } (Qual. Gr. A) - \text{Chance of selection } (Qual. Gr. B))$ 3. Calculate Adverse Impact Ratio (AIR). $AIR = \frac{(\text{Min. selection rate})}{(\text{Max. selection rate})}$ 4. $= (1 - V_{EO}) * \left(\min(1, \frac{AIR}{0,8}) \right)$
ICF	Deep validation: Assessing whether the model is making decisions based on protected features	1. Use counterfactual audit Create N pairs of candidate lookalikes 2. Calculate the Violation Rate – % of cases where changing only the protected feature changed the decision 3. $ICF = 1 - (\text{Violation rate в } \%)$
ITE	EU “AI Act” (Art. 13) Conformity assessment (transparency) and human-in-the-loop capabilities	1. Availability of technical documentation (TDS) $V_1 = (\text{Yes} = 1, \text{No} = 0)$ 2. Availability of XAI explanations (SHAP/LIME) $V_2 = (\text{Yes} = 1, \text{No} = 0)$ 3. Quantitative assessment (“stability” of XAI) V_3 (from 0 to 1) 4. $ITE = (w_1 + V_1) + (w_2 + V_2) + (w_3 + V_3)$
IDP	General Data Protection Regulation (GDPR).	1. Does the model use “sensitive data” for training (and not just for auditing): $V_1 = (\text{Yes} = 1, \text{No} = 0)$ 2. The existence of the “right to be forgotten” mechanism $V_2 = (\text{Yes} = 1, \text{No} = 0)$ 3. $IDP = (0,5 + V_1) + (0,5 + V_2)$

1	2	3
IRS	EU “AI Act” (Art. 15) reliability and cybersecurity	1. Conduct adversarial attack testing (Adversarial Attacks) 2. Calculate the proportion of successful attacks (in %) 3. $IRS = 1 - (\text{Share of successful attacks in \%})$
IHO	EU “AI Act” (Art. 14) Conformity assessment (supervision) and the principle of “human decision-making”	1. Presence of an explicit function in the recruiter interface to cancel/review the decision AI $V_1 = (\text{Yes} = 1, \text{No} = 0)$ 2. Does the system provide information for making a cancellation decision? (XAI): $V_2 = (\text{Yes} = 1, \text{No} = 0)$ 3. $IHO = (0,5 + V_1) + (0,5 + V_2)$ $IHO = 0 \text{ unacceptable for civil service}$

Source: generated by the authors

According to the proposed methodology for calculating the composite fairness and non discrimination index (Table 4), the following scale is provided for categorizing the obtained “CIF-PA” scores.

1. A range from 0.90 to 1.00. “Green” means that the “CIF-PA” has a minimal level of risk for application, which leads to the managerial decision “Permitted for full operation”. The system complies with all ethical and efficiency standards; 2. A range from 0.70 to 0.89. “Yellow” means that the “CIF-PA” has a medium level of risk for application, which leads to the managerial decision “Permitted with enhanced supervision”. The system is effective but requires mandatory human involvement in the loop; 3. A range from 0.50 to 0.69. “Orange” means that the “CIF-PA” has a high level of risk for application, which leads to the managerial decision “Prohibited for operation until deficiencies are eliminated”. The system has critical vulnerabilities for example bias. Return to the developer; 4. A range below 0.50. “Red” means that the “CIF-PA” has an unacceptable level of risk for application, which leads to the managerial decision “Complete prohibition”. Refusal of procurement or immediate termination of use. A direct threat to human rights.

Thus, the conceptual methodology for calculating the Composite Fairness and Non-Discrimination Index “CIF-PA” made it possible to identify critical vulnerability points of algorithmic systems in personnel selection and demonstrated that the mathematical neutrality of algorithms is a fiction because it requires strict external auditing. The developed evaluation model proved that fairness in digital HR is not an abstract ethical category but a measurable technical parameter that can be quantified through a system of weighted sub-indices. The application of this methodology makes it possible to transform the digitalization of public service from a high-risk zone into a controlled process in which artificial intelligence technologies are directed toward strengthening competence-based selection systems rather than scaling hidden discrimination.

Conclusions. As the analysis has shown, the implementation of artificial intelligence in human resource management is not merely a technical issue but primarily a socio-technical and politico-legal problem that directly affects the foundations of the legitimacy of state power. Attempts to solve this problem by searching for a single ideal mathematical metric of fairness are doomed to failure due to their fundamental incompatibility and conflicting definitions. The scientific novelty of this study lies in rejecting such a metrically limited approach through the application of the developed comprehensive “CIF-PA” model, which serves as a methodological response to this multidimensional problem and transforms unbalanced mathematical metrics and abstract legal requirements into a concrete quantitative and step-by-step audit tool that enables public officials to make well-founded managerial decisions. The practical significance of the proposed model lies in the fact that it allows public authorities to move from passive responses to discrimination risks toward proactive management of fairness, which is a necessary precondition both for fulfilling obligations within the framework of international legislative integration under the EU “AI Act” and the “European Convention on Human Rights”, and for increasing real rather than merely declarative trust in the digital state. Further directions of research will focus on the practical validation of the “CIF-PA” model, the development of standardized benchmarks, the study of psychological aspects of the perception of algorithmic fairness, and the analysis of the long-term effects of applying these technologies to the structure and quality of public administration.

REFERENCES:

1. Köchling A., Wehner M.C. (2020). Discriminated by an algorithm: a systematic review of discrimination and fairness by algorithmic decision making in the context of HR recruitment and HR development. *Business Research*, vol. 13, pp. 795–848.

2. Starke C., Lünich M. (2020). Artificial intelligence for political decision making in the European Union: effects on citizen's perceptions of input, throughput and output legitimacy. *Data and Policy*, no. 2, pp. 1–17.
3. Bansak K., Ferwerda J., Hainmueller J., Dillon A., Hangartner D., Lawrence D., Weinstein J. (2018). Improving refugee integration through data driven algorithmic assignment. *Science*, vol. 359, pp. 325–329.
4. Mujtaba D.F., Mahapatra N.R. (2019). Ethical considerations in AI based recruitment. *IEEE International Symposium on Technology and Society*, pp. 1–7.
5. Fregnan E., Ivaldi S., Scaratti G. (2020). HRM 4.0 and new managerial competences profile: the COMAU case. *Frontiers in Psychology*, vol. 11, pp. 1–16.
6. Zeebaree S.R., Shukur H.M., Hussan B.K. (2019). Human resource management systems for enterprise organizations: a review. *Periodicals of Engineering and Natural Sciences*, vol. 7, no. 2, pp. 660–669.
7. Acikgoz Y., Davison K.H., Compagnone M., Laske M. (2020). Justice perceptions of artificial intelligence in selection. *International Journal of Selection and Assessment*, vol. 28, pp. 399–416.
8. DiRomualdo A., El Khoury D., Girimonte F. (2018). HR in the digital age: how digital technology will change HR organization structure, processes and roles. *Strategic HR Review*, vol. 17, pp. 234–242.
9. Verma S., Rubin J. (2018). Fairness definitions explained. *Proceedings of the IEEE ACM International Workshop on Software Fairness*, vol. 18, pp. 1–7.
10. Mehrabi N., Morstatter F., Saxena N., Lerman K., Galstyan A. (2021). A survey on bias and fairness in machine learning. *ACM Computing Surveys*, vol. 54, no. 6, pp. 1–35.
11. Alves G., Bernier F., Couceiro M., Makhoul K., Palamidessi C., Zhioua S. (2023). Survey on fairness notions and related tensions. *EURO Journal on Decision Processes*, vol. 11, pp. 1–14.
12. Mujtaba D.F., Mahapatra N.R. (2024). Fairness in AI driven recruitment: challenges, metrics, methods, and future directions. *arXiv preprint*, vol. 1, pp. 1–17.
13. Fabris A., Baranowska N., Dennis M.J. et al. (2024). Fairness and bias in algorithmic hiring: a multidisciplinary survey. *ACM Transactions on Intelligent Systems and Technology*, vol. 16, no. 1, pp. 1–54.
14. Chen Z. (2023). Ethics and discrimination in artificial intelligence enabled recruitment practices. *Humanities and Social Sciences Communications*, vol. 10, pp. 1–12.
15. Capasso M., Arora P., Sharma D., Tacconi C. (2024). On the right to work in the age of artificial intelligence: ethical safeguards in algorithmic human resource management. *Business and Human Rights Journal*, vol. 9, issue 3, pp. 346–360.
16. Alon Barkat S., Busuioc M. (2021). Human AI interactions in public sector decision making: automation bias and selective adherence to algorithmic advice. *Journal of Public Administration Research and Theory*, vol. 33, issue 1, pp. 153–169

UDC 351:004.8:331.108.45:303.732

JEL H83, J71, M53, O33

Oleksandr Oliinyk, Candidate of Philosophical Sciences, Associate Professor of the Department of Business Administration and Management of Foreign Economic Activity, Zaporizhzhia National University. **Damir Bikulov**, Doctor of Public Administration, Professor of the Department of Business Administration and Management of Foreign Economic Activity, Zaporizhzhia National University. **Olha Holovan**, Candidate of Physical and Mathematical Sciences, Associate Professor of the Department of Business Administration and Management of Foreign Economic Activity, Zaporizhzhia National University. **Svitlana Markova**, Doctor of Economic Sciences, Professor of the Department of Business Administration and Management of Foreign Economic Activity, Zaporizhzhia National University. **Olha Veritova**, Candidate of Pedagogical Sciences, Senior Lecturer of the Department of Business Administration and Management of Foreign Economic Activity, Zaporizhzhia National University. **Mathematical methods of the digitalization of HR processes in public administration.**

The study analyzes the dual effect of AI in public sector recruitment: it increases productivity but risks structural inequalities and algorithmic discrimination, undermining public authority. Current approaches to fairness remain fragmented, creating gaps between efficiency and societal perception. Bias emerges at all recruitment stages. Under the EU “AI Act”, such systems are classified as high-risk technologies requiring mandatory audits. To address this, a composite fairness and non-discrimination index is proposed, integrating technical, ethical, and legal metrics into a unified assessment system. This model enables the identification of vulnerabilities and supports evidence-based decisions. It is substantiated that digitalization requires a systemic methodology for ensuring fairness as a key condition for trust in government algorithms.

Key words: public administration, artificial intelligence in personnel selection, algorithmic bias, algorithmic discrimination, high-risk HR systems, EU “AI Act”, data audit, mathematical methods.

УДК 351:004.8:331.108.45:303.732

JEL H83, J71, M53, O33

Олійник Олександр Миколайович, кандидат філософських наук, доцент кафедри бізнес-адміністрування і менеджменту зовнішньоекономічної діяльності, Запорізький національний університет. **Бікулов Дамір Тагірович**, доктор наук з державного управління, професор кафедри бізнес-адміністрування і менеджменту зовнішньоекономічної діяльності, Запорізький національний університет. **Головань Ольга Олексіївна**, кандидат фізико-математичних наук, доцент, кафедри бізнес-адміністрування і менеджменту зовнішньоекономічної діяльності, Запорізький національний університет. **Маркова Світлана Вікторівна**, доктор економічних наук, професор кафедри бізнес-адміністрування і менеджменту зовнішньоекономічної діяльності, Запорізький національний університет. **Верітова Ольга Сергіївна**, кандидат педагогічних наук, старший викладач кафедри бізнес-адміністрування і менеджменту зовнішньоекономічної діяльності, Запорізький національний університет. **Математичні методи цифровізації HR-процесів у публічному управлінні.**

У дослідженні було виявлено, що впровадження штучного інтелекту в процеси добору кадрів державного сектору має подвійний ефект бо з одного боку алгоритмічні системи розширюють можливості оптимізації, стандартизації й підвищення продуктивності при одночасному скороченні ресурсних витрат, а з іншого боку вони генерують нові ризики відтворення структурних нерівностей і алгоритмічної дискримінації, що безпосередньо впливає на легітимність публічної влади. Проаналізовані сучасні наукові підходи до виявлення джерел упередженості показали несбалансованість технічних, етичних і правових рамок, що зумовлює розрив між ефективністю моделей, регуляторними вимогами та сприйняттям справедливості кандидатами. Доведено, що алгоритмічна упередженість проявляється на всіх етапах рекрутингового циклу, включно з таргетингом вакансій, автоматизованим скринінгом резюме, аналізом відеоінтерв'ю, психометричною оцінкою та ранжуванням, унаслідок чого алгоритми здатні масштабувати наявні соціальні нерівності й формувати нові приховані асиметрії. Систематизовано нормативно-правові вимоги щодо використання ШІ у сфері зайнятості відповідно до EU "AI Act", «Європейської конвенції з прав людини» та "GDPR", і обґрунтовано чому вони віднесені до високоризикових систем із необхідністю обов'язкового аудиту справедливості, прозорості, якості даних, людського нагляду та недискримінаційності. Для подолання виявлених суперечностей запропоновано композитний індекс справедливості та недискримінаційності, який інтегрує технічні метрики, етичні вимоги й правові критерії в єдину кількісну систему оцінювання, дозволяючи ідентифікувати критичні точки вразливості, визначати рівень ризику для публічного сектору та формувати доказову базу управлінських рішень. Вказано, що така модель переводить справедливість із абстрактного етичного принципу в вимірюваний параметр стандартизованого аудиту й регуляторного контролю та створює практичну основу для підвищення довіри до державних алгоритмів за умови системного забезпечення недискримінаційності та суспільної орієнтації цифрових кадрових рішень.

Ключові слова: публічне управління, штучний інтелект у відборі кадрів, алгоритмічна упередженість, алгоритмічна дискримінація, високоризикові HR-системи, EU "AI Act", аудит даних, математичні методи.