

Transformation numérique et intelligence artificielle : vers une gouvernance intégrée de la santé et de la sécurité au travail

Digital transformation and artificial intelligence: towards integrated occupational health and safety governance

Said AKKI, (Doctorant)

Laboratoire Management, Marketing, Logistique Internationale et Finance « **REMMALIF** »
École Supérieure de Technologie de Fès (EST).
Université Sidi Mohamed Ben Abdellah, (USMBA) Fès, Maroc.

Abderrahmane OUDDASSER, (Enseignant chercheur)

Laboratoire Management, Marketing, Logistique Internationale et Finance « **REMMALIF** »
École Supérieure de Technologie de Fès (EST).
Université Sidi Mohamed Ben Abdellah, (USMBA) Fès, Maroc

Adresse de correspondance :	Laboratoire « REMMALIF », Ecole Supérieure de Technologie (EST) ; Université Sidi Mohamed Ben Abdellah (USMBA), Maroc (Fès) Adresse : BP 2427 Route d'Imouzzer 30000 Téléphone : (+212) 6 66 21 46 30 - Email : support.est@usmba.ac.ma
Déclaration de divulgation :	Les auteurs n'ont pas connaissance de quelconque financement qui pourrait affecter l'objectivité de cette étude. Ils assument l'entière responsabilité de tout éventuel plagiat, de l'usage de l'intelligence artificielle dans la rédaction, ainsi que des résultats présentés dans cet article.
Conflit d'intérêts :	Les auteurs ne signalent aucun conflit d'intérêts.
Citer cet article	AKKI, S., & OUDDASSER, A. (2026). Transformation numérique et intelligence artificielle : vers une gouvernance intégrée de la santé et de la sécurité au travail. International Journal of Accounting, Finance, Auditing, Management and Economics, 7(2), 572–590. https://doi.org/10.5281/zenodo.18609907
Licence	Cet article est publié en open Access sous licence CC BY-NC-ND

Received: 08/01/2026

Accepted: 12/02/2026

International Journal of Accounting, Finance, Auditing, Management and Economics - IJAFAME

ISSN: 2658-8455

Volume 7, Issue 02 (2026)

Transformation numérique et intelligence artificielle : vers une gouvernance intégrée de la santé et de la sécurité au travail

Résumé

Cet article analyse les mutations profondes des paradigmes de gouvernance en santé et sécurité au travail (SST) provoquées par la transformation numérique (TN) et l'intégration massive de l'intelligence artificielle (IA) aux environnements professionnels. S'appuyant sur une revue systématique de la littérature, pluridisciplinaire et récente, l'étude identifie une problématique centrale résidant dans l'asynchrone structurelle entre la célérité de l'innovation technologique et l'inertie des cadres législatifs et normatifs. L'analyse souligne une reconfiguration systémique des typologies de risques : aux menaces professionnelles conventionnelles se greffent désormais des risques émergents tels que la surveillance panoptique, l'érosion de l'autonomie décisionnelle et l'aliénation sociotechnique découlant de l'opacité des « boîtes noires ». Si les dispositifs prédictifs et l'Internet des objets (IoT) constituent des leviers de prévention proactive inédits, leur déploiement exacerbe une tension éthique fondamentale entre impératifs de performance et préservation de la dignité humaine. Les conclusions démontrent que la transition vers une gouvernance responsable exige l'instauration d'un paradigme de régulation by design. Cette approche repose sur trois piliers : l'institutionnalisation de l'IA explicable (XAI), une convergence accrue entre les cadres juridiques (AI Act, RGPD) et la revitalisation du dialogue social, afin de garantir que la technologie demeure un auxiliaire au service de l'agence humaine.

Mots clés : Transformation numérique (TN) ; Intelligence artificielle (IA) ; Santé et sécurité au travail (SST) ; Gouvernance des risques professionnels ; Gestion algorithmique ; AI Act

Classification JEL : J24 - K32 - M14

Type du papier : Recherche théorique

Abstract

This article examines the profound shifts in occupational health and safety (OHS) governance paradigms prompted by digital transformation (DT) and the pervasive integration of artificial intelligence (AI) into professional environments. Drawing upon a recent multidisciplinary systematic literature review, the study identifies a central issue rooted in the structural asynchrony between the velocity of technological innovation and the inertia of legislative and regulatory frameworks. The analysis highlights a systemic reconfiguration of risk typologies: traditional occupational hazards are now compounded by emerging risks such as panoptic surveillance, the erosion of decision-making autonomy, and sociotechnical alienation resulting from the opacity of "black box" systems. While predictive tools and the Internet of Things (IoT) offer unprecedented opportunities for proactive prevention, their deployment exacerbates a fundamental ethical tension between performance imperatives and the preservation of human dignity. Ultimately, the findings demonstrate that the transition toward responsible governance necessitates the establishment of a "regulation by design" paradigm. This approach rests on three pillars: the institutionalization of explainable AI (XAI), enhanced convergence between legal frameworks (such as the AI Act and GDPR), and the revitalization of social dialogue to ensure that technology remains an auxiliary to human agency.

Keywords: Digital transformation (DT); Artificial intelligence (AI); Occupational health and safety (OHS); Occupational risk governance; Algorithmic management; AI Act

JEL Classification: J24 - K32 - M14

Paper type: Theoretical Research

1. Introduction

La transformation numérique (TN) et l'avènement de l'intelligence artificielle (IA) s'imposent aujourd'hui comme l'un des bouleversements paradigmatiques majeurs qui affectent en profondeur l'économie mondiale ainsi que les dynamiques du travail (Nadler, 2024). Cette disruption technologique systémique transcende la simple optimisation des outils productifs pour engendrer une reconfiguration structurelle des architectures organisationnelles. En sciences de gestion, ce phénomène implique une redéfinition structurelle des cadres de gouvernance, notamment en ce qui concerne la maîtrise des risques professionnels. Les organisations, qu'elles relèvent des sphères publiques ou privées, sont désormais contraintes d'adapter leurs stratégies pour intégrer des innovations qui redéfinissent le lien entre technologie, performance et protection humaine.

La célérité de la transition numérique catalyse l'émergence de leviers disruptifs pour l'optimisation proactive de la sécurité. Ainsi, l'IA favorise l'efficacité opérationnelle par l'automatisation des tâches à forte pénibilité physique ou cognitive, permettant une réallocation des ressources humaines vers des missions à plus haute valeur ajoutée. L'adaptabilité exigée des salariés devient une contrainte majeure, ces derniers étant soumis à une pression constante de montée en compétences numériques (Vietas, 2024). Cette nécessité de mise à jour permanente des savoirs peut devenir un obstacle à l'utilisation sereine de la technologie et générer une exposition accrue à l'incertitude professionnelle. Les individus font alors face à une perte de leurs marges d'autonomie décisionnelle et à une remise en question profonde des garanties liées à la confidentialité de leurs données personnelles (Vietas, 2024).

La numérisation des processus organisationnels et l'intégration de systèmes automatisés, de dispositifs de surveillance connectée ou d'algorithmes décisionnels sont autant de facteurs qui reconfigurent radicalement la nature de la gestion des risques au travail, quel que soit le secteur d'activité (Aloisi, 2024). Cette gestion algorithmique, bien que présentée comme un gage d'efficacité, risque de subordonner les salariés à des systèmes opaques où le contrôle est exercé par des mécanismes automatisés et impersonnels. Dans cette dynamique, les travailleurs se voient confrontés à des risques émergents, tant physiques que psychosociaux, sur fond de questionnements croissants en matière de gouvernance, d'éthique et de régulation (Awumey et al., 2024 ; Fiegler-Rudol et al., 2025). Le recours à l'IA réduit parfois les possibilités de socialisation et d'interactions humaines, rendant le travail plus aliénant malgré les gains de productivité affichés.

Face à ces mutations, la littérature scientifique récente s'efforce de conceptualiser et d'encadrer les nouveaux enjeux de santé et sécurité au travail liés à la sphère numérique. Les recherches actuelles s'interrogent sur les effets réels de l'IA sur la santé psychologique et physique des salariés, tout en soulignant le rôle crucial des politiques publiques pour réguler ces outils (Adams-Prassl et al., 2024 ; Souлами et al., 2024). Il apparaît impératif de développer une gouvernance innovante, responsable et éthique, capable de naviguer entre le potentiel de l'innovation et la nécessaire protection des droits fondamentaux des collaborateurs. La transition numérique transforme en profondeur les sociétés contemporaines, et l'un des domaines où cet impact est le plus sensible est celui de la médiation entre les impératifs de performance technologique et le bien-être humain.

Cette revue de littérature scientifique propose une analyse structurée et approfondie des principaux axes, des problématiques émergentes et des avancées majeures qui façonnent le nouveau paradigme de la gestion des risques professionnels à l'ère de l'intelligence artificielle.

La problématique centrale de cet article réside dans l'asynchrone structurelle entre l'expansion fulgurante de l'intelligence artificielle (IA) et l'inertie des cadres de régulation de la santé et sécurité au travail (SST). Elle interroge la capacité des organisations à concilier l'optimisation proactive des risques professionnels avec la préservation des droits fondamentaux face aux

menaces d'aliénation numérique et d'opacité algorithmique. S'appuyant sur les apports des publications académiques les plus citées ainsi que sur les cadres normatifs internationaux de référence, l'étude vise à fournir un état de l'art exhaustif permettant d'éclairer la compréhension des enjeux actuels et prospectifs de cette transformation technologique. L'objectif fondamental est de mettre en évidence les réponses organisationnelles et réglementaires nécessaires à une intégration éthique, efficace et responsable de l'IA dans le champ de la prévention des risques professionnels (Park & Kang, 2024 ; Capasso et al., 2025). Il s'agit d'explorer comment l'IA peut renforcer les dispositifs de contrôle interne sans toutefois remplacer totalement l'expertise humaine indispensable pour interpréter des environnements incertains.

L'analyse mobilise une méthodologie rigoureuse d'examen systématique de la littérature scientifique, avec une attention particulière portée aux travaux récents et aux sources à forte visibilité académique. Cette approche systématique est indispensable, car le domaine de l'IA en gestion des risques est actuellement marqué par une fragmentation des cadres théoriques. En cartographiant les contributions pertinentes, cette recherche permet d'identifier les dynamiques émergentes, de révéler les lacunes conceptuelles et méthodologiques actuelles, et d'ouvrir de nouvelles pistes de recherche dans un domaine en plein renouvellement (Yadav & Ganesan, 2025 ; El Bouchikhi et al., 2024). L'étude culmine par la formulation de directives stratégiques visant à concilier efficacité systémique et résilience éthique.

L'architecture de cet article s'articule autour de cinq axes principaux. La première section appréhende l'émergence de nouveaux risques professionnels consécutifs à la transformation numérique des environnements de travail. Le deuxième volet approfondit l'analyse des enjeux critiques, en mettant l'accent sur les risques psychosociaux (RPS) et les problématiques inhérentes à la surveillance algorithmique. La troisième section examine les dynamiques de gouvernance ainsi que les cadres de régulation impératifs à une intégration éthique et responsable de ces outils. La quatrième partie procède à une évaluation critique des dispositifs actuels, interroge les tensions éthiques identifiées et formule des recommandations stratégiques pour une gestion responsable. Enfin, l'étude se clôt par une synthèse des contributions théoriques et managériales de cette recherche.

2. Mutation des risques professionnels à l'ère numérique

2.1. Évolution des risques et nouveaux paradigmes de sécurité

La transformation numérique (TN) ne saurait être réduite à une simple adoption technologique, elle constitue une mutation structurelle des milieux de travail qui transfigure la nature des risques professionnels. Ce changement de paradigme exige de passer d'une posture de sécurité réactive à une gestion prospective, centrée sur la convergence des systèmes physiques et numériques. Tandis que l'IA offre des leviers de prévention sophistiqués, elle introduit des risques systémiques complexes que les cadres de régulation classiques peinent à appréhender. Ce processus n'est pas une évolution incrémentale, mais une rupture exigeant de nouveaux cadres conceptuels pour définir la sécurité. En intégrant capteurs et algorithmes au cœur de l'activité, la technologie redéfinit ainsi les limites de la protection corporelle et de l'intégrité numérique, imposant une vision holistique de la sécurité au travail.

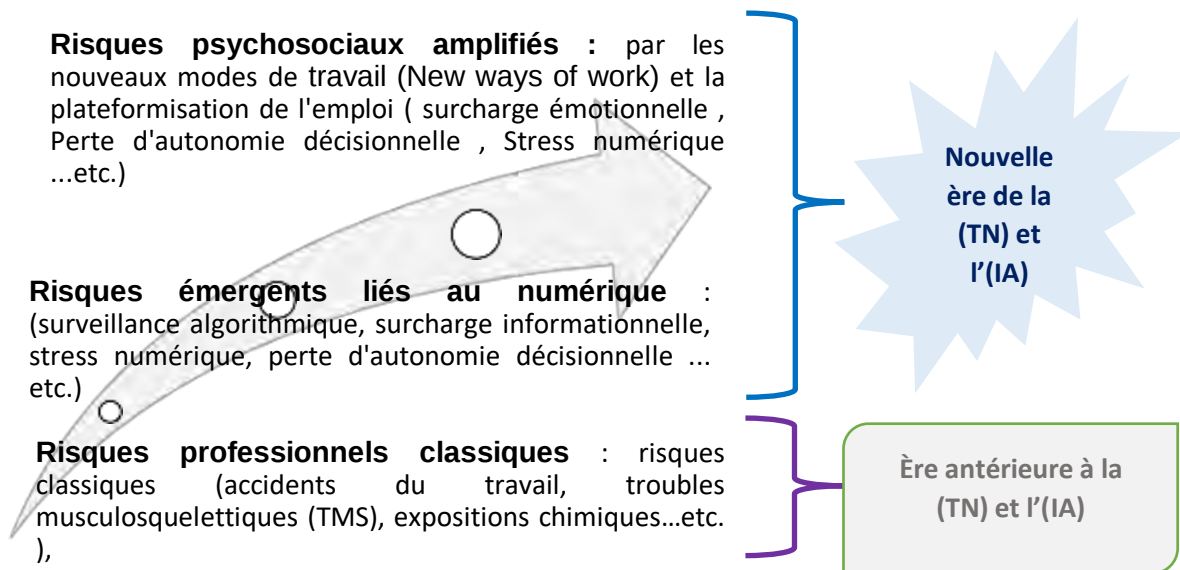
Nadler (2024) démontre que l'intégration de l'apprentissage automatique en Santé et Sécurité au Travail (SST) se heurte à une absence de cadres normatifs standardisés, entravant ainsi la capacité des organisations à assurer un traitement éthique des données de sécurité. Parallèlement, les travaux de (Vietas, 2024) mettent en exergue l'équilibre délétaire entre les perspectives d'optimisation de la sécurité et les risques d'insécurité professionnelle, d'érosion de l'autonomie et d'atteinte à la confidentialité des données.

L'intelligence artificielle regroupe un ensemble de technologies d'automatisation décisionnelle, d'apprentissage automatique et d'analyse prédictive qui transforment radicalement

l'environnement de travail (Park & Kang, 2024). Ces technologies incluent les dispositifs intelligents, les objets connectés (Internet of Things, IoT), les capteurs biométriques et les plateformes collaboratives qui génèrent des volumes massifs de données personnelles et professionnelles (Jiang & Zhang, 2024).

L'illustration ci-après offre une synthèse structurée des évolutions paradigmatiques liées aux risques professionnels à l'ère digitale. Elle souligne notamment le rôle pivot de l'intelligence artificielle (IA) dans le renouvellement des dispositifs de pilotage de la santé-sécurité au travail (SST) :

Figure 1 : La mutation des risques professionnels à l'épreuve de la transition numérique et l'(IA)



Source : Fiegler-Rudol et al., (2025) "Exploring Human-AI Dynamics in Enhancing Workplace Safety and Health: A Narrative Review"

Les recherches de Parmar et Tiwari (2025) illustrent cette complexification par leur framework dual pour la détection des actes et conditions dangereuses utilisant la reconnaissance faciale dans l'industrie. Leur approche révèle les tensions inhérentes entre efficacité sécuritaire et respect de la vie privée, mettant en lumière les défis éthiques posés par l'intégration de l'IA dans la gestion des risques. Ces préoccupations sont renforcées par l'analyse de Naidoo et al. (2025) qui identifie l'IA comme un facteur de transformation majeur de la gestion de la sécurité au travail, nécessitant une approche multidisciplinaire combinant expertise technique, juridique et sociologique.

Si la transformation numérique redessine les contours globaux de l'activité productive, elle cristallise des tensions spécifiques et immatérielles qu'il convient désormais d'approfondir, notamment à travers le prisme des risques psychosociaux et de l'influence croissante des algorithmes sur le travailleur.

2.2. Défis critiques : surveillance, biais et autonomie

L'incorporation des dispositifs algorithmiques au sein de la gestion de la Santé et Sécurité au Travail (SST) opère une reconfiguration paradigmatique : Le modèle préventif conventionnel subit une mutation vers une surveillance algorithmique panoptique, induisant une réification du travailleur par la quantification systématique des données comportementales. Cette objectivation technologique induit une rupture axiologique, exacerbant les tensions éthiques au

sein de l'écosystème organisationnel. In fine, la gouvernance de la SST se heurte à l'opacité intrinsèque des systèmes décisionnels automatisés, dont le déploiement massif fait émerger des risques transversaux excédant les taxonomies conventionnelles de la sécurité professionnelle. L'étude systématique d'Awumey et al. (2024), portant sur 129 publications, révèle que les technologies de monitoring biométrique alimentées par l'IA exposent les travailleurs à divers préjudices sociotechniques, notamment l'obligation d'effectuer un travail émotionnel supplémentaire pour naviguer dans des systèmes de surveillance affective peu fiables.

Les biais algorithmiques constituent une problématique centrale identifiée par Patil (2025) dans son analyse des défis éthiques des applications d'IA industrielle. Ces biais, résultant d'erreurs de conception algorithmique, peuvent perpétuer voire amplifier les discriminations existantes dans les domaines critiques du recrutement, de l'évaluation des performances et de la gestion des risques. Cette problématique est particulièrement préoccupante dans les contextes de haute responsabilité où les décisions automatisées peuvent avoir des conséquences durables sur la carrière et la sécurité des travailleurs (Huang, 2025).

La question de l'autonomie au travail se trouve également au cœur des transformations induites par l'IA. Les recherches de Soulamy et al. (2024), basées sur une analyse bibliométrique de 92 articles et une revue systématique de 25 études, identifient quatre clusters principaux d'impact de l'adoption de l'IA sur le bien-être des employés : l'éthique, l'autonomie au travail, le stress des employés et la santé mentale. Ces travaux soulignent la nécessité pour les organisations d'implémenter des stratégies de soutien pour atténuer les effets négatifs potentiels de l'IA tout en préservant l'autonomie et la créativité des travailleurs.

L'intensification de la surveillance constitue un autre défi majeur soulevé par Carter (2024) dans son analyse de la surveillance IA et du contrôle informationnel. L'auteur démontre comment la prolifération de la surveillance IA sur le lieu de travail transforme les travailleurs en entités statistiques, objectivant leur valeur comme simples points de données pour l'analyse algorithmique. Cette co-modification exacerbe les déséquilibres informationnels existants entre travailleurs et employeurs, nécessitant des mécanismes renforcés de protection de la vie privée et de contrôle informationnel.

2.3. Perspectives sectorielles et approches empiriques

L'analyse des applications sectorielles révèle une grande diversité dans l'adoption et l'impact de l'IA selon les domaines d'activité. Les travaux de Yadav et Ganesan (2025) démontrent l'efficacité des plateformes de conscience contextuelle alimentées par l'IA dans les industries à haut risque, où ces systèmes ont significativement réduit les incidents enregistrables et les quasi-accidents. Leur recherche met en évidence la transition des méthodes de sécurité réactives traditionnelles vers des systèmes de gestion des risques proactifs.

Dans le secteur manufacturier, les recherches de Khurram et al. (2025) explorent le potentiel de l'IA pour améliorer la sécurité des travailleurs par la maintenance prédictive, l'automatisation et l'atténuation des risques en temps réel. Leurs travaux soulignent la nécessité d'établir des standards et réglementations pour gouverner l'intégration des technologies IA, tout en prévenant leur usage adversarial. Cette approche est complétée par l'analyse de Bispo et Amaral (2024) qui propose un modèle théorique d'interaction pour illustrer les relations causales entre les exigences et facteurs de risque de l'Industrie 4.0.

Le secteur de la construction bénéficie également des innovations IA, comme le démontre l'étude d'Agapiou (2024) sur l'optimisation des processus de planification et de gestion des risques. Ces applications permettent une meilleure anticipation des situations dangereuses et une allocation plus efficace des ressources de sécurité.

Les recherches de Maheshwari et al. (2024) explorent l'impact de l'IA sur la vie privée des employés, évaluant les compromis entre l'efficacité organisationnelle et la protection de la vie privée. Leur revue systématique met en évidence la nature duale de l'impact de l'IA sur la

confidentialité en milieu de travail, soulignant la nécessité d'équilibrer les bénéfices organisationnels avec la protection des droits individuels.

Dans le contexte des plateformes numériques, les travaux de Potocka-Sionek (2024) analysent le premier instrument réglementaire de l'UE pour réguler la supervision et la prise de décision automatisée dans le contexte du travail, spécifiquement la Directive sur l'amélioration des conditions de travail dans le travail de plateforme. Cette analyse révèle que les dispositions légales sur la gestion algorithmique constituent les aspects les plus progressistes et les mieux conçus de la Directive.

L'étude de Jangid (2024) explore le potentiel de l'IA pour surveiller et améliorer la santé mentale des employés sur le lieu de travail, utilisant des technologies comme l'analyse de sentiment et les chats bots. Ces recherches soulignent l'importance du consentement des employés et de la protection de la vie privée pour maintenir la confiance et assurer une implémentation éthique des systèmes de surveillance de la santé mentale.

2.4. Synthèse critique : apports théoriques et lacunes identifiées par la recherche

L'analyse croisée des contributions révèle une littérature en pleine mutation, oscillant entre optimisme technologique et vigilance éthique.

Sur le plan des contributions, les travaux récents soulignent deux prérequis à l'intégration de l'IA : le contrôle individuel des données (Das Swain et al., 2024) et la transparence algorithmique comme gage d'acceptabilité (Mapungwana, 2025). La recherche identifie désormais l'émergence de modèles hybrides de prévention. Selon Marsden et Steyer (2024), la SST de demain repose sur une dualité : d'un côté, un modèle proactif-prédictif (IA, IoT, Big Data) et de l'autre, un modèle coconstruit-participatif fondé sur le dialogue social.

Toutefois, des lacunes conceptuelles majeures subsistent. Nadler (2024) pointe l'absence de cadres standardisés et de lignes directrices éthiques claires pour l'apprentissage automatique appliqué à la sécurité. Surtout, la littérature converge vers un risque de « désintermédiation » de l'expert humain. Fiegler-Rudol et al. (2025) rappellent alors l'urgence d'une approche multidisciplinaire — croisant droit, gestion et ingénierie — pour combler ce fossé.

En synthèse, l'enseignement fondamental de cette revue réside dans le constat que l'efficacité technologique ne saurait garantir la sécurité si elle s'exerce au préjudice de l'agence humaine. L'identification de cette potentielle « atrophie de l'agence humaine » au sein des protocoles automatisés impose une refonte des structures de régulation. Dès lors, comment bâtir un cadre de gouvernance capable de résorber cette asynchrone structurelle tout en protégeant la dignité des travailleurs ? C'est l'objet de la section suivante.

3. Gouvernance et régulation des systèmes d'IA en faveur de la SST

Une approche holistique est requise pour harmoniser les cadres juridiques, notamment par une interopérabilité normative entre le Règlement (UE) 2024/1689 (AI Act), le RGPD et les législations nationales du travail. Cette approche vise à combler les "zones de non-droit" créées par les asymétries de définitions et à garantir que la protection de la santé mentale et physique soit une condition sine qua non de la conformité technologique. La gouvernance doit s'opérationnaliser par la mise en place d'instances de régulation dotées de compétences hybrides (juridiques et algorithmiques). Il s'agit de créer des mécanismes de recours effectifs face à la gestion algorithmique, permettant aux travailleurs de contester non seulement l'issue d'une décision automatisée, mais également la légitimité de l'*heuristique décisionnelle* (la logique interne de l'algorithme) ayant conduit à ladite décision.

La régulation ne saurait être monolithique. La promotion de standards de transparence radicale et de participation praxéologique (basée sur l'expérience réelle de travail) doit s'adapter aux spécificités de chaque secteur. Ces standards doivent institutionnaliser le "droit à l'explication"

et transformer l'auditabilité en un processus continu et démocratique, plutôt qu'une simple certification ponctuelle.

3.1. Cadres normatifs et enjeux éthiques

3.1.1. Régulations européennes et internationales

L'encadrement juridique de l'IA dans le contexte professionnel connaît une évolution rapide, particulièrement au niveau européen. L'AI Act de l'Union européenne, analysé par Popa et Pascariu (2024), introduit des obligations spécifiques pour les "systèmes à haut risque", incluant les dispositifs d'IA déployés en milieu professionnel. Cette réglementation établit un cadre de classification des risques et impose des obligations particulières aux fournisseurs et utilisateurs de systèmes d'IA à haut risque.

Les travaux de Carinci et Ingraio (2025) sur l'impact de l'AI Act sur le droit du travail révèlent que la réglementation interdit certains systèmes d'IA à haut risque, tels que la notation sociale, la reconnaissance d'émotions et la catégorisation biométrique des travailleurs. Leurs recherches critiquent cependant le manque de soutien normatif pour la participation syndicale et la négociation collective en relation avec les systèmes d'IA.

L'analyse de Pardo López (2024) du cadre normatif de l'IA et de la protection des droits fondamentaux souligne l'importance de l'harmonisation des normes pour la conception, le développement et l'utilisation des systèmes d'IA. Cette harmonisation vise particulièrement à protéger les droits fondamentaux, notamment la vie privée, dans le contexte de l'implémentation de l'IA sur le lieu de travail.

Les recherches de Cavas Martínez (2024) sur l'intelligence artificielle et les relations de travail examinent les limites de la gestion algorithmique à la lumière de la nouvelle législation européenne sur l'IA et le travail sur plateformes digitales. Cette analyse révèle la nécessité de conceptualiser juridiquement l'impact de l'IA sur le travail, en reconnaissant les droits et obligations tant pour les travailleurs que pour les employeurs.

Pérez Amorós (2025) analyse les aspects du travail du Règlement (UE) 2024/1689 sur l'intelligence artificielle, mettant l'accent sur les obligations des employeurs qui utilisent des systèmes et modèles d'IA, ainsi que sur les droits individuels et collectifs des travailleurs affectés par l'implémentation de l'IA sur le lieu de travail.

3.1.2. Standards d'auditabilité et transparence algorithmique

La question de la transparence algorithmique constitue un défi majeur identifié par Adams-Prassl et al. (2024) dans leur blueprint pour un standard international de régulation de la gestion algorithmique. Leurs recherches soulignent que les risques significatifs associés à la gestion algorithmique incluent l'augmentation des atteintes à la vie privée, l'élargissement des asymétries informationnelles et la perte d'agence humaine dans la prise de décision sur le lieu de travail.

Les travaux d'Özkızıltan et Landini (2025) sur la gouvernance des technologies d'IA sur le lieu de travail sous l'AI Act de l'UE révèlent que, malgré l'objectif de l'AI Act d'assurer une IA centrée sur l'humain, plusieurs lacunes réglementaires peuvent favoriser l'autonomie des entreprises technologiques d'IA et les déséquilibres de pouvoir. Ces lacunes risquent de rendre les lieux de travail européens vulnérables aux technologies d'IA biaisées et intrusives.

L'analyse de Chivato Pérez (2025) des défis éthiques et légaux de l'IA dans le monde du travail met l'accent sur l'importance de normes harmonisées pour la conception, le développement et l'utilisation des systèmes d'IA. Cette recherche souligne la nécessité de pratiques appropriées en matière de gestion et protection des données parmi les opérateurs juridiques.

Les recherches de Reddy et al. (2024) sur les implications éthiques et légales de l'IA sur les affaires et l'emploi proposent un cadre pour aborder la vie privée, les biais et la responsabilité par le biais de chiffrement avancé, d'algorithmes de détection de biais et de processus de prise

de décision transparentes. Leurs résultats démontrent des avancées significatives dans la résolution des violations de la vie privée, des biais algorithmiques et des mécanismes de responsabilité au sein des systèmes d'IA.

La problématique de l'auditabilité est également abordée par Peruzzi (2023) qui identifie les réponses réglementaires que l'ordre juridique de l'UE peut offrir pour protéger les travailleurs face aux processus algorithmiques. Son analyse souligne l'importance du RGPD et de la proposition en cours de réglementation sur l'IA, construite sur des principes de transparence et de supervision humaine.

En somme, l'architecture réglementaire actuelle, centrée sur l'IA Act et le RGPD, tente de circonscrire les risques professionnels par une approche de conformité statique et descendante (*top-down*). Toutefois, l'analyse factuelle de ces textes révèle un décalage persistant entre la rigidité normative et la fluidité évolutive des systèmes algorithmiques en milieu de travail. Ce constat d'une régulation principalement descriptive impose de passer d'une simple observation du cadre légal à une réflexion critique sur les modalités concrètes de sa mise en œuvre. La section suivante explore ainsi les limites opérationnelles de cette régulation *ex-ante* et propose un basculement vers un paradigme de gouvernance dynamique, fondé sur l'auditabilité permanente et la responsabilité sociotechnique.

3.2. Gouvernance algorithmique et technologies intelligentes au service de la SST

3.2.1. Systèmes prédictifs et analyse de données massives

Les avancées technologiques en matière d'IA et d'objets connectés permettent l'émergence de nouvelles stratégies de prévention et de gestion des risques professionnels. Les recherches de Yadav et Ganesan (2025) sur les plateformes de conscience contextuelle alimentées par l'IA dans les industries à haut risque démontrent une transition significative des méthodes de sécurité réactives traditionnelles vers des systèmes de gestion des risques proactifs. Leurs travaux révèlent des améliorations substantielles de la sécurité, incluant une réduction des incidents enregistrables et des quasi-accidents, ainsi qu'une amélioration des capacités de détection des dangers.

L'intégration de capteurs biométriques, d'analyses en temps réel et de plateformes centralisées permet une surveillance continue des conditions de travail et une détection précoce des situations dangereuses (Park & Kang, 2024). Ces systèmes utilisent des méthodes avancées d'apprentissage automatique, incluant les réseaux de neurones profonds et l'apprentissage par renforcement, pour améliorer les capacités des systèmes de sécurité.

Les algorithmes prédictifs avancés, analysés par Marsden et Steyer (2024), évaluent la probabilité d'incidents, adaptent les consignes de sécurité en temps réel et personnalisent les notifications de prévention selon les profils individuels et les contextes spécifiques. Cette approche permet une gestion plus fine et personnalisée des risques, tenant compte des variations individuelles et situationnelles.

L'étude de Jiang et Zhang (2024) sur les technologies numériques en santé au travail dans les industries manufacturières identifie les technologies numériques prédominantes telles que les dispositifs portables, les capteurs, la collaboration homme-robot et l'analytique d'apprentissage profond. Leurs recherches révèlent que les facilitateurs de l'implémentation des technologies numériques incluent la fabrication intelligente, les conditions concurrentielles, les outils de prise de décision basés sur les données et les considérations de bien-être et de santé.

3.2.2. Internet des objets (IoT) et monitoring intelligent

L'Internet des objets (IoT) transforme radicalement les approches traditionnelles de surveillance et de gestion des risques professionnels. Les travaux d'El Bouchikhi et al. (2024) identifient six catégories principales de préoccupations éthiques liées à l'utilisation des dispositifs IoT pour la

santé et sécurité au travail : les opportunités éthiques, la surveillance et la réutilisation problématique des données, la difficulté d'informer, consulter et obtenir le consentement des employés, les effets adverses non intentionnels et imprévisibles, la gestion sous-optimale des données, et les facteurs externes qui sont propices aux problèmes éthiques.

Les dispositifs portables et les capteurs intégrés permettent une collecte de données en temps réel sur les paramètres physiologiques, les conditions environnementales et les comportements de travail. Cette approche est illustrée par les recherches de Das Swain et al. (2024) qui révèlent que les travailleurs sont plus enclins à accepter les systèmes d'IA de détection passive (PSAI) lorsqu'ils se concentrent sur la détection de l'utilisation du temps numérique ou de l'activité physique, plutôt que sur des modes visuels.

L'analyse de Huang (2024) sur la protection des informations personnelles des travailleurs sous surveillance algorithmique souligne les risques croissants de violations de données, de surveillance et de discrimination potentielle. Cette recherche met en évidence l'équilibre que les entreprises doivent maintenir entre une gestion efficace et la protection des informations personnelles des travailleurs.

Les technologies immersives, notamment la réalité virtuelle et les simulateurs, révolutionnent la préparation aux situations à risque et la formation continue des travailleurs. Ces approches permettent une formation dans des environnements sans risque tout en préparant efficacement aux situations dangereuses réelles (Khurram et al., 2025).

3.3. Déterminants organisationnels, culturels, sociaux et législatifs de l'adoption des innovations technologiques au milieu professionnel

Les modalités d'adoption et d'acceptation des technologies d'IA varient considérablement selon les contextes organisationnels et sectoriels. Les recherches de Cefaliello et al. (2023) sur la sécurisation, de la gestion algorithmique pour les travailleurs soulignent que les cadres réglementaires européens existants et proposés sont inadéquats pour aborder les risques de santé et sécurité au travail, particulièrement les risques psychosociaux, associés à la gestion algorithmique. Leurs travaux préconisent la nécessité d'une législation autonome au niveau de l'UE ciblant spécifiquement la gestion algorithmique.

Le rôle du dialogue social dans l'acceptation des innovations technologiques est souligné par De Stefano et Taes (2022) dans leur analyse de la gestion algorithmique et de la négociation collective. Leurs recherches identifient les défis posés par la gestion algorithmique et l'intelligence artificielle sur le lieu de travail, particulièrement concernant les droits fondamentaux et les principes. Ils argumentent que la négociation collective constitue l'instrument réglementaire le plus approprié pour répondre à ces défis.

Les résistances culturelles face à la numérisation intensive sont documentées par Souлами et al. (2024), particulièrement dans les secteurs traditionnellement peu technologisés. Leur analyse bibliométrique révèle une augmentation significative des publications liées à l'adoption de l'IA sur le lieu de travail à partir de 2020, avec une concentration de recherche principalement aux États-Unis et en Chine.

L'étude de Corvite et al. (2023) sur les perspectives des sujets de données concernant l'utilisation de l'IA émotionnelle sur le lieu de travail révèle que les participants, tout en reconnaissant les avantages potentiels, expriment des préoccupations significatives concernant les risques associés à l'utilisation de l'IA émotionnelle, incluant les impacts négatifs potentiels sur leurs bien-être, environnement de travail et statut d'emploi.

Les enjeux de formation continue et d'accompagnement personnalisé à la transformation émergent comme facteurs critiques de succès. Les travaux de Santoso et al. (2024) sur les réglementations légales pour anticiper les travailleurs basés sur l'IA soulignent la nécessité de formations spécifiques liées à l'IA et de droits du travail pertinents aux avancées technologiques.

L'acceptabilité sociale des innovations varie en fonction des contextes culturels, organisationnels et individuels. Les recherches de Kuzminac et Reljanović (2024) sur les algorithmes dans l'emploi et les relations de travail révèlent que les algorithmes ont significativement affecté les processus de travail mondialement depuis plus d'une décennie, mais qu'il n'existe que des tentatives pionnières dans la plupart des pays pour analyser l'impact des algorithmes sur les droits des travailleurs.

L'instauration d'un cadre de gouvernance rigoureux constitue désormais une condition sine qua non pour atténuer les risques inhérents au déploiement de l'intelligence artificielle (IA), tout en maximisant ses externalités préventives. Cette régulation doit s'affranchir des modèles conventionnels pour privilégier une approche dynamique, holistique et multipartite. Le consensus doctrinal souligne, à cet égard, l'impératif que la technologie demeure un auxiliaire d'aide à la décision, préservant ainsi la responsabilité discrétionnaire et l'arbitrage de la ligne managériale.

Dans cette perspective, l'action régulatrice ne saurait se limiter à une simple atténuation réactive des préjudices ; elle doit réaffirmer la primauté de la finalité politique et humaine sur la seule performance technique. Ce changement de paradigme marque le passage d'une régulation ex post vers une gouvernance by design, au sein de laquelle la protection du travailleur est intégrée dès la conception structurelle des systèmes algorithmiques. Il s'agit, en somme, d'évoluer d'une logique instrumentale de gestion du risque vers une véritable éthique de la responsabilité organisationnelle.

Néanmoins, l'efficacité des cadres normatifs et juridiques se heurte à l'asynchrone structurelle entre la célérité de l'innovation technologique et l'inertie des processus régulateurs. Ce décalage impose une analyse critique des limites des approches contemporaines afin de dégager des leviers stratégiques opérationnels pour les organisations. Ces défis, ainsi que les recommandations pour une gouvernance responsable, feront l'objet de la section suivante.

4. Défis, Recommandations de gouvernance responsable et perspectives de recherche

L'intégration massive de l'intelligence artificielle au sein des environnements de travail transcende la simple optimisation technique pour exiger une refonte profonde des cadres de régulation. Si ces technologies offrent des leviers de prévention inédits, leur déploiement impose de concilier impératifs de performance et protection des droits fondamentaux. Cette section synthétise les enseignements majeurs de la revue de littérature selon un triptyque analytique : l'examen des limites intrinsèques et des défaillances systémiques des modèles actuels, la formulation de recommandations pratiques et managériales à l'usage des décideurs, et l'établissement d'un agenda de recherche prospective destiné à orienter les travaux futurs sur cette mutation technologique.

4.1. Défaillances systémiques et limites intrinsèques des cadres actuels

Cette section examine les carences structurelles qui entravent une régulation efficiente de l'intelligence artificielle (IA) dans le domaine de la santé et de la sécurité au travail (SST). L'analyse se focalise sur la dualité entre la vacuité normative et l'opacité technique des systèmes.

4.1.1. Fragmentation réglementaire et lacunes normatives

La littérature scientifique contemporaine met en exergue une aporie normative majeure : le décalage croissant entre la vélocité de la gestion algorithmique et l'inertie des cadres juridiques traditionnels. Cette défaillance structurelle, qualifiée par certains auteurs de véritable «

démission régulatrice », traduit l'incapacité des normes actuelles à encadrer les mutations du travail à l'ère numérique.

- ***L'insuffisance des cadres réglementaires existants***

Bien que des avancées législatives soient notables, Adams-Prassl et al. (2024) démontrent que les réglementations en vigueur échouent à traiter de manière holistique les risques multidimensionnels associés à l'automatisation. Ce constat s'applique même aux juridictions dotées d'une protection des données robuste, révélant que le droit de la vie privée (RGPD) ne saurait pallier, à lui seul, les lacunes du droit du travail face à l'hégémonie algorithmique.

- ***Fragmentation et autonomie discrétionnaire des firmes***

Cette fragmentation n'est pas une simple lacune technique ; elle favorise activement la discrétionnarité managériale des firmes technologiques. À cet égard, Özkızıltan et Landini (2025) soulignent que les zones d'ombre de l'IA Act de l'Union Européenne pourraient permettre aux entreprises d'imposer leurs propres standards de conception, exacerbant ainsi les déséquilibres de pouvoir préexistants.

- ***Disparités sectorielles et insécurité juridique***

L'absence de coordination internationale cristallise des inégalités profondes. Les recherches de Coloma-Armijos et al. (2024) illustrent comment l'usage de l'IA, en l'absence de cadre spécifique, peut conduire à des violations caractérisées du droit au travail.

En somme, l'opacité algorithmique sature les dispositifs juridiques classiques. Comme l'analyse Pinedo (2025), cette situation est particulièrement préjudiciable dans le domaine du droit du licenciement, où le cadre légal actuel semble insuffisamment outillé pour relever les défis de la preuve et de la responsabilité algorithmique. Ce vide normatif prive la SST d'un socle protecteur contraignant, exposant les travailleurs à une insécurité juridique structurelle.

4.1.2. Opacité algorithmique et déficit de transparence

L'opacité algorithmique, inhérente à la nature de « boîte noire » des modèles d'apprentissage profond (*deep learning*), ne constitue pas seulement un défi technique ; elle représente une rupture fondamentale du contrat informationnel en milieu de travail. Selon Patil (2025), cette caractéristique entrave la traçabilité des processus décisionnels et la redevabilité des acteurs, neutralisant les mécanismes classiques de responsabilité et altérant la confiance organisationnelle.

- ***De l'asymétrie de pouvoir à l'entrave réglementaire***

Cette invisibilité des critères décisionnels institutionnalise une asymétrie de pouvoir durable. Carter (2024) démontre que l'opacité exacerbe le déséquilibre informationnel entre l'employeur et le salarié, dépossédant ce dernier — ainsi que ses représentants — de la capacité d'auditer ou de contester des décisions automatisées impactant sa trajectoire professionnelle.

- ***Obstacles à l'action publique***

Au-delà de la sphère de l'entreprise, le manque de transparence limite l'efficacité des politiques publiques. Huang (2025) souligne que les obstacles techniques freinent l'identification des responsabilités liées aux biais, empêchant le développement de régulations fondées sur des preuves empiriques rigoureuses pour lutter contre la discrimination algorithmique.

- ***Vers une impérative gouvernance éthique***

Lorsque les systèmes d'IA rendent les décisions de gestion « intraduisibles », ils menacent directement la sécurité et l'évolution de carrière des individus. Face à ce risque, Capasso et al.

(2025) soutiennent que l'impact des biais sur les droits fondamentaux exige l'instauration de sauvegardes éthiques impératives. La préservation d'une gouvernance intégrée au sein des ressources humaines repose désormais sur la mise en place de garde-fous capables de restaurer la maîtrise des flux informationnels.

4.2. Piliers d'une gouvernance responsable et recommandations stratégiques

Pour pallier les défaillances systémiques précédemment identifiées, la littérature converge vers l'élaboration de cadres de gouvernance plus inclusifs. Ces approches visent à réconcilier l'efficacité technologique avec les impératifs de protection des travailleurs à travers deux leviers cardinaux : l'intelligibilité technique et la participation organisationnelle.

4.2.1. Vers une intelligence artificielle explicable (XAI) et une conception participative

L'Intelligence Artificielle Explicable (*eXplainable AI* - XAI) s'affirme comme un levier critique pour restaurer **l'agence humaine** et garantir la transparence des processus décisionnels en SST. Selon Nadler (2024), le développement de systèmes explicables constitue une trajectoire de recherche prioritaire pour optimiser la prise de décision tout en préservant l'intégrité de la santé au travail.

- L'explicabilité comme vecteur de redevabilité

L'intelligibilité algorithmique requiert un accès documenté au fonctionnement des modèles, aux critères de pondération et à la traçabilité des données. Cela suppose l'implémentation de mécanismes de médiation — tels que des « interfaces de transparence » — permettant aux utilisateurs finaux de décrypter les logiques décisionnelles (Reddy et al., 2024). Une telle transparence n'est pas qu'une exigence technique ; elle est la condition de la confiance organisationnelle.

- L'auditabilité par des tiers indépendants

La certification et l'audit externe constituent des mécanismes de régulation indispensables pour garantir l'absence de biais et la conformité aux normes de sécurité. Les travaux de Chivato Pérez (2025) soulignent que la robustesse technique doit impérativement s'accompagner de pratiques rigoureuses en matière de gestion et de protection des flux de données personnelles, conformément aux standards internationaux émergents.

- La co-conception sociotechnique et participative

Le déploiement de l'IA ne saurait être un processus purement descendant (*top-down*). La co-conception émerge comme une approche privilégiée pour impliquer les parties prenantes — managers, salariés et représentants — dès la phase embryonnaire du projet. Monod (2025) démontre que l'acceptabilité des systèmes de gestion automatisés dépend intrinsèquement des dynamiques de pouvoir et de la culture de l'entreprise. Une implémentation éthique doit donc être « psychologiquement durable » et inclusive pour neutraliser les sentiments d'aliénation. En l'absence de dispositifs d'IA explicable (XAI) et d'une démarche participative, la technologie risque de demeurer un instrument de contrôle vertical. Le déficit d'intelligibilité prive les responsables SST des ressources critiques nécessaires pour assurer une surveillance éthique et garantir la sécurité psychologique des travailleurs face à l'automatisation.

4.2.2. Renforcement du dialogue social et de la formation interdisciplinaire : vers une gouvernance multiniveau

Au-delà des solutions techniques, la littérature identifie le renforcement du dialogue social comme le pivot stratégique de la gouvernance algorithmique. De Stefano et Taes (2022) soulignent l'impératif d'une **implication institutionnalisée** des partenaires sociaux dans les

cycles de sélection, d'évaluation et de supervision des systèmes d'IA. Cette participation ne saurait se limiter à une dimension consultative ; elle exige une reconnaissance formelle de la négociation collective comme vecteur de régulation au sein des cadres normatifs, garantissant ainsi une surveillance effective des dispositifs connectés et une limitation de la discrétionnarité managériale.

Littératie numérique et inclusion éthique

Le développement des compétences constitue le corollaire indispensable à une transition numérique socialement responsable. Cette montée en compétences s'articule autour de deux axes majeurs :

- Équité et protection des droits fondamentaux

Oyekunle et al. (2024) démontrent la nécessité de processus de conception inclusifs pour préserver l'équité. Ils préconisent une mutation des politiques publiques vers l'établissement de lignes directrices morales robustes, visant à prémunir les travailleurs contre les abus systémiques et l'aliénation numérique.

- Prévention des biais par l'interdisciplinarité

La formation des acteurs (responsables RH, informaticiens, managers) à l'éthique numérique et à la gestion des risques psychosociaux (RPS) demeure un défi managérial de premier plan. Caldas et Costa (2025) identifient des leviers opérationnels pour prévenir la discrimination algorithmique, insistant sur la représentativité des jeux de données et la diversité des profils au sein des équipes de développement.

Cadres normatifs et principe de précaution

L'harmonisation législative requiert une approche granulaire, capable de moduler les standards sectoriels selon la criticité des risques identifiés. Valentini (2024) analyse les risques d'exploitation liés aux données produites par les travailleurs et appelle à une consolidation des cadres juridiques pour protéger l'intégrité de la vie privée en milieu professionnel.

Par ailleurs, l'incertitude quant aux impacts de l'IA à long terme impose l'application d'un principe de précaution renforcé. Todolí Signes (2019) soutient que l'usage d'algorithmes de direction devrait être proscrit sans une évaluation préalable rigoureuse des risques professionnels. Une perspective strictement technocentrée, omettant la dimension humaine, risque de réduire le travailleur à une simple unité de donnée, compromettant ainsi sa souveraineté individuelle et sa sécurité psychologique.

Recommandations pour une régulation intégrée et systémique

Pour les régulateurs et les décideurs, l'enjeu prioritaire réside dans la construction d'une articulation systémique entre l'AI Act, le RGPD et le droit du travail. L'objectif est de résorber les asymétries juridiques où la protection des données s'efface souvent devant le lien de subordination. Cette convergence doit favoriser :

- **L'interopérabilité juridique** : Harmoniser les obligations de transparence pour éviter les vides réglementaires entre les différentes strates législatives.
- **Des mécanismes de remédiation spécifiques** : Instituer des droits de recours dédiés, permettant aux salariés de contester les décisions automatisées impactant leur santé ou leur trajectoire de carrière.
- **La promotion de standards sectoriels** : Substituer à la conformité générique une participation obligatoire des parties prenantes, transformant la régulation en un levier de protection adapté aux réalités de chaque métier.

4.3. Implications managériales et perspectives de recherche

L'investigation transversale de la littérature met en exergue que la gouvernance de la SST ne traverse pas une simple phase d'ajustement, mais une véritable mutation paradigmatique. Si le déploiement de l'IA offre des leviers de prévention proactive incontestables — illustrés par l'efficacité des systèmes de vigilance contextuelle dans les industries à haute fiabilité (Yadav et Ganesan, 2025) — il n'en demeure pas moins porteur de risques d'aliénation sociotechnique. La réussite de cette transition ne repose plus exclusivement sur l'innovation technique, mais sur la résolution du clivage entre monitoring biométrique (Awumey et al., 2024) et intégrité de la dignité humaine. Cette ambition impose de consolider trois piliers stratégiques : la transparence radicale, l'inclusivité participative et une régulation publique proactive.

Axes de recherches futures

Les futures investigations scientifiques devront approfondir plusieurs chantiers critiques pour stabiliser ce nouveau paradigme :

- **Évaluations longitudinales** : L'étude des effets à long terme des innovations numériques sur la santé holistique (physique et mentale) des travailleurs constitue une priorité. Cela nécessite le développement de méthodologies robustes d'évaluation des systèmes d'IA *in situ*.
- **Analyse comparative et culturelle** : Une approche comparative des modèles de gouvernance selon les contextes culturels et sectoriels permettrait d'identifier les variables de succès et les meilleures pratiques organisationnelles.
- **Modèles d'équité numérique** : L'exploration de modèles économiques et sociaux favorisant une transition numérique équitable représente un défi majeur pour la pérennité du contrat social en entreprise.

Synthèse épistémologique : de la technique à l'éthique de la responsabilité

En dernière analyse, la gouvernance des risques professionnels à l'ère de l'intelligence artificielle constitue un défi global qui transcende les considérations purement instrumentales. L'intégration de l'IA au sein de la SST cristallise une **antinomie épistémologique majeure** : si le postulat d'une sécurité proactive promet une réduction systémique de l'aléa, la matérialité de sa mise en œuvre induit corrélativement un risque de surveillance panoptique.

Cette gouvernance se trouve ainsi fragilisée par une **asynchrone structurelle** entre la célérité de l'innovation technologique et l'inertie des cadres normatifs et régulateurs. *In fine*, l'objet de la problématique ne réside pas tant dans l'ontologie de l'algorithme que dans la reconfiguration des rapports de force et l'exacerbation des asymétries informationnelles qu'il génère. Ces dynamiques rendent obsolètes les mécanismes traditionnels de contrôle et imposent une collaboration renforcée entre chercheurs, praticiens et décideurs pour bâtir un avenir du travail où l'innovation demeure, en tout point, un auxiliaire de l'humanité.

5. Conclusion :

Cette recherche a examiné la reconfiguration paradigmatique de la gouvernance de la santé et de la sécurité au travail (SST) à l'aune de la transformation numérique et de l'intégration de l'intelligence artificielle. À travers une revue systématique de la littérature, nous avons mis en exergue que l'introduction de l'IA au sein des organisations ne constitue pas une simple évolution technique, mais une rupture profonde des cadres de régulation traditionnels.

L'analyse révèle que si l'IA offre des perspectives inédites en matière de prévention proactive et de réduction de la pénibilité, elle engendre simultanément des risques émergents — de nature psychosociale et organisationnelle — cristallisés par l'opacité algorithmique et l'érosion de l'autonomie des travailleurs. Le constat majeur de cette étude réside dans l'identification d'une

asynchrone structurelle entre la célérité de l'innovation technologique et l'inertie des dispositifs normatifs et législatifs.

Face à cette aporie, l'article soutient qu'une gouvernance responsable ne peut se limiter à une approche réactive ou strictement juridique. Elle exige l'instauration d'un paradigme de **régulation by design**, fondé sur la transparence algorithmique (XAI), le renforcement du dialogue social et une interdisciplinarité accrue entre les sciences de gestion, le droit et les technologies de l'information.

En définitive, l'avenir de la SST dans un écosystème numérisé dépendra de notre capacité à préserver la **primauté de l'agence humaine** sur la seule performance technique. Concilier innovation et humanité n'est plus seulement un impératif éthique, mais une condition *sine qua non* de la durabilité des organisations contemporaines. Cette étude ouvre ainsi la voie à de futurs travaux qui devront tester l'efficacité de ces nouveaux modèles de gouvernance participative face à la complexité croissante des systèmes autonomes.

Références :

- (1). Adams-Prassl, J., Rakshita, S., & Abraha, H. H. (2024). Towards an international standard for Regulating Algorithmic Management: A Blueprint. Social Science Research Network. <https://doi.org/10.2139/ssrn.4684947>
- (2). Agapiou, A. (2024). Construction project management optimization using AI. Construction Innovation, 24(3), 456-472.
- (3). Aloisi, A. (2024). Regulating Algorithmic Management at Work in the European Union: Data Protection, Non-discrimination and Collective Rights. International Journal of Comparative Labour Law and Industrial Relations, 40(1), 123-145.
- (4). Awumey, E., Das, S., & Forlizzi, J. (2024). A Systematic Review of Biometric Monitoring in the Workplace: Analyzing Socio-technical Harms in Development, Deployment and Use. Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems. <https://doi.org/10.1145/3630106.3658945>
- (5). Bispo, L. G. M., & Amaral, F. G. (2024). The impact of Industry 4.0 on occupational health and safety: A systematic literature review. Journal of Safety Research, 89, 234-248. <https://doi.org/10.1016/j.jsr.2024.04.009>
- (6). Caldas, C. O. L., & Costa, G. S. D. (2025). A discriminação algorítmica nas relações trabalhistas e as ferramentas para prevenção e solução extrajudicial de conflitos. Revista Brasileira de Direito, 1(1), 45-67. <https://doi.org/10.62140/ccgc1902025>
- (7). Capasso, M., Arora, P., & Sharma, D. (2025). On the Right to Work in the Age of Artificial Intelligence: Ethical Safeguards in Algorithmic Human Resource Management. Business and Human Rights Journal, 10(1), 78-95. <https://doi.org/10.1017/bhj.2024.26>
- (8). Carinci, M. T., & Ingraio, A. (2025). L'impatto dell'AI Act sul diritto del lavoro. Giornale di Diritto del Lavoro e di Relazioni Industriali, 184(2), 189-210. <https://doi.org/10.3280/gdl2024-184002>
- (9). Carter, C. (2024). AI surveillance: Reclaiming privacy through informational control. European Labour Law Journal, 15(4), 456-478. <https://doi.org/10.1177/20319525241306327>
- (10). Cavas Martínez, F. (2024). Inteligencia artificial y relaciones laborales: límites a la gestión algorítmica del trabajo. Revista Española de Derecho del Trabajo, 218, 123-145.
- (11). Cefaliello, A., Moore, P. V., & Donoghue, R. C. (2023). Making algorithmic management safe and healthy for workers: Addressing psychosocial risks in new legal provisions. European Labour Law Journal, 14(2), 234-256. <https://doi.org/10.1177/20319525231167476>

- (12). Chivato Pérez, T. (2025). Desafíos éticos y legales de la inteligencia artificial en el mundo laboral. *Revista de Innovación Legal*, 3(1), 89-112. <https://doi.org/10.69592/979-13-7011-111-3-cap-10>
- (13). Coloma-Armijos, J. A., Maldonado-Valle, K. J., & Chicaiza-Flores, M. J. (2024). Vulneración del derecho al trabajo por el uso de la IA en el sistema judicial. *Verdad y Derecho*, 12(2), 156-178. <https://doi.org/10.62574/bs4gv663>
- (14). Corvite, S., Roemmich, K., & Andalibi, N. (2023). Data Subjects' Perspectives on Emotion Artificial Intelligence Use in the Workplace: A Relational Ethics Lens. *Proceedings of the ACM on Human-Computer Interaction*, 7(CSCW1), 1-28. <https://doi.org/10.1145/3579600>
- (15). Das Swain, V., Gao, L., Mondal, A., Jain, R., & McDuff, D. (2024). Sensible and Sensitive AI for Worker Wellbeing: Factors that Inform Adoption and Resistance for Information Workers. *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3613904.3642716>
- (16). De Stefano, V., & Taes, S. (2022). Algorithmic management and collective bargaining. *Transfer*, 28(4), 567-584. <https://doi.org/10.1177/10242589221141055>
- (17). El Bouchikhi, M., Weerts, S., & Clavien, C. (2024). Behind the good of digital tools for occupational safety and health: a scoping review of ethical issues surrounding the use of the internet of things. *Frontiers in Public Health*, 12, Article 1468646. <https://doi.org/10.3389/fpubh.2024.1468646>
- (18). Fiegler-Rudol, J., Lau, K., & Mroczek, A. (2025). Exploring Human–AI Dynamics in Enhancing Workplace Safety and Health: A Narrative Review. Preprints, 2025010483. <https://doi.org/10.20944/preprints202501.0483.v1>
- (19). Huang, M. (2025). Protection of workers' rights and interests in the problem of algorithmic discrimination: legal challenges and response strategies. *Education and Social Work*, 2(1), 148–155. ISSN Print: 3079-515X; ISSN Online: 3079-5168. <https://doi.org/10.63313/esw.9069>
- (20). Huang, R. (2024). Protection of Personal Information of Workers under Algorithmic Monitoring. *Communications in Humanities Research*, 33, 94-102. <https://doi.org/10.54254/2753-7064/33/20240094>
- (21). Jangid, A. (2024). AI and employee wellbeing: how artificial intelligence can monitor and improve mental health in the workplace. *International Journal of Advanced Research*, 12(10), 456-467. <https://doi.org/10.21474/ijar01/19693>
- (22). Jiang, L., & Zhang, J. (2024). Digital technology in occupational health of manufacturing industries: a systematic literature review. *SN Applied Sciences*, 6, Article 349. <https://doi.org/10.1007/s42452-024-06349-4>
- (23). Khurram, M., Wu, C., & Muhammad, S. M. (2025). Artificial Intelligence in Manufacturing Industry Worker Safety: A New Paradigm for Hazard Prevention and Mitigation. *Processes*, 13(5), Article 1312. <https://doi.org/10.3390/pr13051312>
- (24). Kuzminac, M., & Reljanović, M. (2024). New actors or new tools – algorithms in employment and labour relations. *Regional Law Review*, 5, 340-365. https://doi.org/10.56461/iup_rlrc.2024.5.ch20
- (25). Maheshwari, S., Kaur, A., & Bose, I. (2024). Watch Out, You are Live! Toward Understanding the Impact of AI on Privacy of Employees. *Communications of the AIS*, 55, Article 23. <https://doi.org/10.17705/1cais.05523>
- (26). Mapungwana, P. (2025). Addressing the Ethical and Data Privacy Concerns Related to AI in Occupational Health and Safety. In *Advances in Computational Intelligence and Robotics* (pp. 1-18). IGI Global. <https://doi.org/10.4018/979-8-3693-9301-7.ch001>

- (27). Marsden, É., & Steyer, V. (2024). Artificial intelligence and safety management: an overview of key challenges. *Industrial Safety Review*, 15(3), 234-251. <https://doi.org/10.57071/iae290>
- (28). Monod, E. (2025). Robots as Managers? Organizational Justice, Power Asymmetries, and the Future of AI Supervision. *Management Research Quarterly*, 8(2), 123-145. <https://doi.org/10.63029/rqg5z883>
- (29). Nadler, D. (2024). Machine Learning in Occupational Health and Safety: A Review of Knowledge Gaps. *Preprints*, 2024110464. <https://doi.org/10.20944/preprints202411.0464.v1>
- (30). Naidoo, C. M., Lawrence, C., & Mkolo, N. M. (2025). Future Trends and Innovations. In *Advances in Computational Intelligence and Robotics Book Series* (pp. 156-178). IGI Global. <https://doi.org/10.4018/979-8-3693-9301-7.ch006>
- (31). Oyekunle, D., Boohene, D., & Preston, D. (2024). Ethical Considerations in AI-Powered Work Environments: A Literature Review and Theoretical Framework for Ensuring Human Dignity and Fairness. *International Journal of Scientific Research and Management*, 12(3), 234-251. <https://doi.org/10.18535/ijssrm/v12i03.em18>
- (32). Özkızıltan, D., & Landini, F. (2025). Trustworthy and human-centric? The new governance of workplace AI technologies under the EU's Artificial Intelligence Act. *Transfer*, 31(2), 189-207. <https://doi.org/10.1177/10242589251336193>
- (33). Pardo López, M. M. (2024). Marco normativo de la IA: protección de derechos fundamentales en el contexto de la IA. *Revista de Derecho Digital*, 8(2), 123-145. <https://doi.org/10.69592/978-84-1194-692-6-cap1>
- (34). Park, J., & Kang, D. (2024). Artificial Intelligence and Smart Technologies in Safety Management: A Comprehensive Analysis Across Multiple Industries. *Applied Sciences*, 14(24), Article 11934. <https://doi.org/10.3390/app142411934>
- (35). Parmar, J., & Tiwari, V. (2025). Enhancing the detection of unsafe Acts and Conditions using Face Recognition Technique in the Industry. *Journal of Advances in Science and Technology*, 10(4), 456-467. <https://doi.org/10.29070/bp7hj708>
- (36). Patil, D. (2025). Ethical Challenges In Industrial Artificial Intelligence Applications: Bias, Privacy, And Accountability. *Social Science Research Network*. <https://doi.org/10.2139/ssrn.5057418>
- (37). Pérez Amorós, F. (2025). Aspectos laborales del Reglamento (UE) 2024/1689 sobre inteligencia artificial. *e-Revista Internacional de la Protección Social*, 1(1), 1-23. <https://doi.org/10.12795/e-rips.2025.i01.01>
- (38). Peruzzi, M. (2023). La protection des travailleurs dans l'ordre juridique de l'UE face à l'intelligence artificielle. *Revue de Droit Comparé du Travail et de la Sécurité Sociale*, 2, 78-95. <https://doi.org/10.4000/rdctss.4939>
- (39). Pinedo, D. (2025). AI-effect: is het arbeids- en ontslagrecht klaar voor een 'robocalyps'? *Tijdschrift voor Ontslagrecht*, 9(1), 15-32. <https://doi.org/10.5553/tvo/254253152025009001003>
- (40). Popa, A. A., & Pascariu, L. (2024). Impact of the EU's Artificial Intelligence Regulation on Workers. *European Journal of Law and Public Administration*, 11(2), 234-251. <https://doi.org/10.18662/eljpa11.2234>
- (41). Potocka-Sionek, N. (2024). Gestión de algoritmos. El caso del trabajo en plataformas. *Labos*, 5(2), 123-145. <https://doi.org/10.20318/labos.2024.9031>
- (42). Reddy, K. J., Kethan, M., & Basha, S. M. (2024). Ethical and Legal Implications of AI on Business and Employment: Privacy, Bias, and Accountability. In *Proceedings of the international Conference on Knowledge Engineering and Communication Systems* (pp. 234-241). IEEE. <https://doi.org/10.1109/ickecs61492.2024.10616875>

- (43). Santoso, I. B., Triyunarti, W., & Nurhalimah, S. (2024). Legal Regulations for Anticipating Artificial Intelligence-Based Workers through Institutional Transformation of Job Training and the Human Resources Revolution. *Devotion*, 5(6), 743-759. <https://doi.org/10.59188/devotion.v5i6.743>
- (44). Soulami, M., Bencheckroun, S., & Galiulina, A. V. (2024). Exploring how AI adoption in the workplace affects employees: a bibliometric and systematic review. *Frontiers in Artificial Intelligence*, 7, Article 1473872. <https://doi.org/10.3389/frai.2024.1473872>
- (45). Todolí Signes, A. (2019). En cumplimiento de la primera Ley de la robótica: Análisis de los riesgos laborales asociados a un algoritmo/inteligencia artificial dirigiendo el trabajo. *Labour & Law Issues*, 5(2), 1-18. <https://doi.org/10.6092/ISSN.2421-2695/10237>
- (46). Valentini, R. S. (2024). The use of Workers-Produced Data in Artificial Intelligence-Based Systems. *Brazilian Journal of Law, Technology and Innovation*, 2(2), 34-53. <https://doi.org/10.59224/bjlti.v2i2.34-53>
- (47). Vietas, J. (2024). Artificial intelligence and the Future of Work; opportunities and risks associated with worker health and safety. *Annals of Work Exposures and Health*, Volume 68, Issue Supplement_1, June 2024, Page 1, <https://doi.org/10.1093/annweh/wxae035.021>
- (48). Yadav, S. and Ganesan, H. (2025) AI-powered Contextual Awareness for Next-Gen Safety Platforms in High-Risk Industries. *World Journal of Advanced Research and Reviews*, 26(03), 1701-1709. Article DOI: <https://doi.org/10.30574/wjarr.2025.26.3.2303>.