

Fausse promises et stupidité organisationnelle : pourquoi l'IA n'est pas synonyme de performance

Guillaume Desjardins and Anthony Gould

Number 8, 2024

URI: <https://id.erudit.org/iderudit/1115740ar>

DOI: <https://doi.org/10.1522/radm.no8.1846>

[See table of contents](#)

Publisher(s)

Département des sciences économiques et administratives de l'Université du Québec à Chicoutimi (UQAC)

ISSN

2369-6907 (digital)

[Explore this journal](#)

Cite this article

Desjardins, G. & Gould, A. (2024). Fausse promises et stupidité organisationnelle : pourquoi l'IA n'est pas synonyme de performance. *Ad machina*, (8), 216–233. <https://doi.org/10.1522/radm.no8.1846>

Article abstract

With the arrival of ChatGPT, the term artificial intelligence (AI) seems to become a fashion where several fields of expertise must renegotiate established paradigms. In the context of work, ever-growing literature predicts major transformations in work organization resulting in jobs rendered obsolete by the use of AI. This provocative essay goes against this trend. To support the argument, the authors use the concept of organizational stupidity to demonstrate how the irrational or even stupid side of businesses will surpass the efficient rationality promised by AI. Finally, practical implications are proposed for organizational decision-makers.



Titre : Fausses promesses et stupidité organisationnelle : pourquoi l'IA n'est pas synonyme de performance

Rubrique : Perspective théorique

Auteur(s)

1 : Guillaume Desjardins, professeur

2 : Anthony Gould, professeur

Citation : Desjardins, G. et Gould, A. (2024). Fausses promesses et stupidité organisationnelle : pourquoi l'IA n'est pas synonyme de performance. *Ad Machina*, 8(1), 216-233, <https://doi.org/10.1522/radm.no8.1846>

Affiliation des auteurs

1 : Université du Québec en Outaouais

Courriel : guillaume.desjardins@uqo.ca

2 : Université Laval

Courriel : anthony.gould@rlt.ulaval.ca

Remerciements

Déclaration des conflits d'intérêts

- ☒ Aucun conflit d'intérêts à déclarer
☐ Conflit d'intérêts à déclarer (veuillez détailler)

Détails :

Résumé (250 mots)

Depuis l'apparition de ChatGPT, le terme intelligence artificielle (IA) semble devenir une mode dans laquelle plusieurs champs d'expertise doivent renégocier les paradigmes établis. Dans le contexte du travail, une littérature grandissante prédit des transformations majeures dans l'organisation du travail résultant des emplois devenus obsolètes par l'utilisation de l'IA. Cet essai provocateur ira à l'encontre de ce courant. Afin d'appuyer l'argumentation, les auteurs emploieront le concept de la stupidité organisationnelle afin de démontrer comment le côté irrationnel voire stupide des entreprises surpassera la rationalité efficiente promise par l'IA. Enfin, des implications pratiques seront proposées pour les décideurs en organisation.

Abstract

With the arrival of ChatGPT, the term artificial intelligence (AI) seems to become a fashion where several fields of expertise must renegotiate established paradigms. In the context of work, ever-growing literature predicts major transformations in work organization resulting in jobs rendered obsolete by the use of AI. This provocative essay goes against this trend. To support the argument, the authors use the concept of organizational stupidity to demonstrate how the irrational or even stupid side of businesses will surpass the efficient rationality promised by AI. Finally, practical implications are proposed for organizational decision-makers.

Mots clés

IA, stupidité organisationnelle, psychologie organisationnelle, ère numérique, stupidité fonctionnelle

Droits d'auteur

Ce document est en libre accès, ce qui signifie que le lectorat a accès gratuitement à son contenu. Toutefois, cette œuvre est mise à disposition selon les termes de la licence [Creative Commons Attribution \(CC BY NC\)](https://creativecommons.org/licenses/by-nc/4.0/).



Fausses promesses et stupidité organisationnelle : pourquoi l'IA n'est pas synonyme de performance

Guillaume Desjardins
Anthony Gould

Introduction

Les avancées technologiques sont souvent synonymes de nouvelles occasions d'affaires. Tantôt incrémentielles, ces innovations permettent l'automatisation de tâches répétitives et chronophages (Staccioli et Virgillito, 2020), la facilitation de la communication des entreprises avec ses membres ou ses clients (Wang *et al.*, 2020), et l'augmentation de la capacité de traitement des données des organisations (Bag *et al.*, 2021). Dans l'ensemble, ces éléments permettent à une organisation d'obtenir un avantage compétitif dans son industrie. Par exemple, des organisations telles qu'Amazon ou Netflix déploient leurs technologies afin de récolter et de traiter un maximum d'informations auprès de leurs utilisateurs afin de détecter les tendances futures et s'appropriier ces nouveaux marchés souvent inexploités par les concurrents – ce que Kim et Mauborgne (2005) nomment les océans bleus.

Cependant, il arrive qu'une innovation soit si forte qu'elle entraîne alors un nouveau paradigme dans une ou plusieurs industries, ce qui remet en question les façons de faire des organisations existantes – ce que Christensen (2018) nomme les innovations perturbatrices. Le téléphone intelligent en est un parfait exemple : sa prolifération sur le marché des télécommunications coïncide avec une diminution des parts de marché du service de téléphonie filaire qui, par le fait même, a fait perdre le monopole du marché à certains fournisseurs (Chang *et al.*, 2013), mais a aussi occasionné une diminution des ventes d'ordinateurs portables (Vecchiato, 2017) et d'appareils connexes tels que les lecteurs MP3 (West et Mace, 2010).

Les innovations perturbatrices sont généralement introduites par de petites ou nouvelles entreprises et ciblent initialement les segments bas de gamme ou nichés d'un marché (Christensen *et al.*, 2018; Si *et al.*, 2020). Ces innovations, bien qu'elles ne répondent pas toujours aux besoins des clients traditionnels, s'améliorent au fil du temps et finissent par supplanter les produits établis (Petzold *et al.*, 2019). Par exemple, le constructeur Tesla lança en 2008 sa première voiture entièrement électrique pour grand public avec une autonomie de 394 km. Cette percée dans l'industrie automobile a forcé les constructeurs « traditionnels » (p. ex., Ford, General Motors, Hyundai, etc.) à emboîter le pas et à offrir des solutions électriques à leur clientèle.

Partant de cette définition pour qualifier les innovations perturbatrices, il n'est pas exagéré de considérer les nouvelles formes d'intelligence artificielle (IA) telles que *ChatGPT* (OpenAI), *Bard* (Google) ou *Copilot* (Microsoft) comme étant une révolution sociétale de l'ère numérique. Cet argument n'est pas faux en soi; l'utilisation grand public des systèmes d'IA remet en question plusieurs domaines, notamment l'éducation universitaire (Grassini, 2023), le système légal (Tan *et al.*, 2023) et bien entendu, le monde du travail (Ayinde *et al.*, 2023). Ces changements se font à une vitesse si ahurissante que la recherche peine à suivre. Par exemple, en l'espace de quelques mois seulement, la plateforme *ChatGPT* s'est améliorée à un tel point qu'elle est maintenant en mesure d'être dans le top 10 % des meilleures notes de l'examen uniforme du barreau aux États-Unis (Choi *et al.*, 2022; OpenAI, 2023). Ces résultats encouragent certaines firmes d'avocats à user de logiciel d'IA pour remplacer le travail de certains parajuridiques afin de détecter les

préférences d'un juge dans le but de personnaliser et d'automatiser leur plaidoyer (Scheihauf, 2023; Stokel-Walker, 2023). Nous serions donc bien loin de l'argument traditionnel proposant qu'un changement technologique ne déplace que les emplois non qualifiés vers des emplois nécessitant des compétences spécifiques (Colbert *et al.*, 2016; Mei *et al.*, 2023).

Bien que les avancées en matière d'IA soient remarquables, l'argumentaire de cet article théorique soutiendra que les résultats de recherche en psychologie organisationnelle sont tout simplement incompatibles avec les promesses de l'IA. Pour défendre cette position, le concept de stupidité organisationnelle et ses ramifications seront mis de l'avant (Alvesson et Spicer, 2012, 2016; McComber et Jenkins, 1972; Parkinson, 1955; Peter et Hull, 1969). Ces principes, inhérents à toutes les organisations, font que la nature des interactions humaines dans une organisation est nécessairement imparfaite, voire inefficace, et ne pourra être compensée par l'utilisation d'un outil technologique. En d'autres termes, les dirigeants ne devraient pas systématiquement intégrer l'IA dans leur organisation en espérant automatiquement obtenir un gain d'efficacité puisque la stupidité l'emportera sur le rationnel.

Cet article propose de commencer par poser les pourtours du concept d'IA dans sa forme générale et comment, déjà depuis son apparition, celle-ci n'a pas été en mesure d'atteindre les objectifs escomptés. Par la suite, l'argumentaire mettra en relation la stupidité organisationnelle et l'intégration, voire la démocratisation de l'IA au travail, pour élaborer sur comment les promesses d'un travail nécessairement plus efficace grâce à l'IA ne seront pas tenues. Enfin, la discussion offrira trois implications pratiques pour les décideurs, notamment en ce qui a trait aux éléments à surveiller lorsque l'IA est employée dans un contexte de travail. Des avenues pour de futures recherches seront aussi présentées.

1. Le concept d'IA : ses débuts et ses échecs

La terminologie « intelligence artificielle » (IA) est souvent un terme générique pour de nombreux autres concepts tels que l'algorithme, le modèle d'apprentissage automatique ou le réseau neuronal (IBM, s. d.). Dans sa forme la plus simple, l'intelligence artificielle est un domaine qui combine l'informatique avec des ensembles de données robustes pour permettre la résolution de problèmes. La plupart des auteurs s'entendent pour attribuer la paternité de l'IA à Allen Newell, Cliff Shaw et Herbert Simon en 1956 avec leur programme informatique du « théoriste logique » (*The Logic Theorist*) lors de la conférence de Dartmouth (Gugerty, 2006; Newell et Simon, 1956). Bien que les auteurs *du Logic Theorist* aient prôné une collaboration entre les différents spécialistes des sciences afin de faire avancer le champ de l'IA, traditionnellement celui-ci fut développé par différentes expertises (génération d'images, reconnaissance vocale, génération de sons, robotique, etc.) en silos étanches. Chaque domaine ayant sa propre méthodologie, les progrès de l'IA n'étaient bien souvent qu'itératifs (Han *et al.*, 2021). Ce n'est qu'en 2017, lorsque Google Brain a introduit le « modèle transformeur », basé sur le traitement du langage naturel, que les différents domaines d'expertise ont convergé vers un modèle unique. Cette unification de la recherche a permis une avancée exponentielle de l'IA et a démocratisé l'outil numérique de la recherche fondamentale vers des domaines plus pratiques.

Le Center for Humane Technology (2023) indique que la démocratisation de l'IA, c'est-à-dire l'utilisation grand public de celle-ci, s'est déroulée en deux phases. La première, aussi appelée premier contact, a été introduite avec les réseaux sociaux, notamment *TikTok* en 2017. Durant cette période, l'utilisateur est principalement passif en présence de l'IA : sur la base de ses entrées (p. ex., aimer ou ne pas aimer une vidéo, le temps passé à regarder un certain créateur de contenu, etc.), l'IA tente de catégoriser l'individu et d'estimer ses goûts, ses tendances et ses préférences afin de lui proposer des contenus personnalisés. Afin de convaincre les utilisateurs des bienfaits d'enregistrer et analyser leurs comportements lors de l'utilisation de ces applications, les dirigeants des réseaux sociaux promettent que cela permettra à « chacun d'exprimer son opinion et d'être entendu » ainsi que « d'entrer en contact avec des amis et des collègues de façon

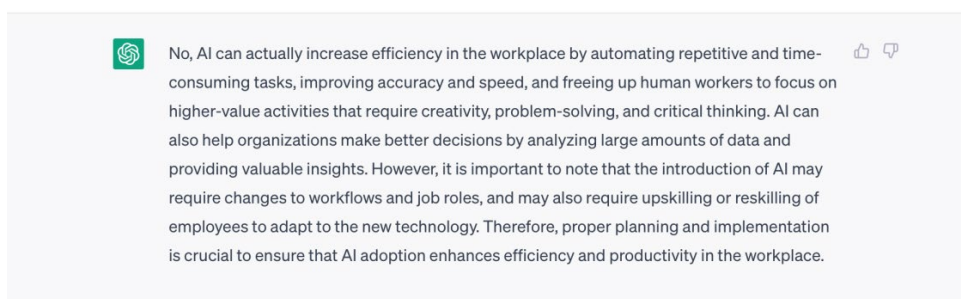
organique ». Cependant, ce premier contact a également conduit à des problèmes de dépendance en ligne (Smith et Short, 2022), à la désinformation (Basch *et al.*, 2021), au catastrophisme (Grandinetti et Bruinsma, 2023), ainsi qu'à la polarisation, qui peut conduire à des comportements extrémistes (Weinmann et Masri, 2023).

La deuxième phase de l'IA se caractérise par le fait que l'utilisateur devient un artisan – un créateur qui interagit directement avec la technologie pour moduler des résultats spécifiques. Cette période a débuté vers novembre 2022 avec l'ouverture au public de *ChatGPT* 3.0. Cette interface permet à tout utilisateur d'écrire une question/demande et de générer une réponse basée sur les données disponibles sur le Web avant 2021. Par exemple, un utilisateur peut demander à la plateforme de lui fournir un résumé de 1 000 mots sur un livre en particulier. Cette démocratisation permettrait à la communauté scientifique de résoudre de nouveaux problèmes et d'être plus efficace (Raj *et al.*, 2023). Néanmoins, plusieurs limitations sont encore présentes dans le traitement du langage naturel de l'IA, ce qui entraîne des biais dans les informations traitées (Bender et Friedman, 2018) qui peuvent parfois être discriminatoires (Heinrichs, 2022). De plus, les experts du domaine qui sont engagés pour former l'IA ne sont pas toujours en mesure d'expliquer les décisions de l'IA ou les nouvelles capacités soudaines qu'elle a développées (Zhang *et al.*, 2022). Dans une étude menée auprès de 738 chercheurs en IA, près de 50 % d'entre eux estiment qu'il existe un risque de 10 % ou plus que l'homme disparaisse en raison de notre incapacité à contrôler l'IA (Stein-Perlman *et al.*, 2022). D'ailleurs, Sam Altman, PDG de *OpenAI*, a lui-même réitéré à plusieurs reprises que l'IA allait probablement mener à la fin de l'humanité¹. Certaines études appuyant ces conclusions commencent à faire surface. Par exemple, aux États-Unis, certains propriétaires d'immeubles à logement se servent d'un programme employant l'IA afin de fixer les prix des loyers, voulant ainsi maximiser leur profit. Cependant, les données préliminaires démontreraient que l'IA ne considère pas plusieurs facteurs globaux (aspect social, économie nationale, etc.) dans sa décision finale, risquant ainsi de contrevenir à certaines lois en vigueur ou même de provoquer une crise économique (Quach, 2023). Dans un autre cas extrême, Rivera et ses collaborateurs (2024) ont remarqué que lorsque *ChatGPT* 4.0 est utilisé pour prodiguer des conseils à une nation dans le cadre de situations diplomatiques, celui-ci a une propension à escalader les tensions internationales et à proposer la guerre nucléaire comme solution.

2. L'IA, l'efficacité au travail et la stupidité

Est-ce que l'IA peut engendrer des pertes d'efficacité au travail? Selon le principal intéressé, non (voir figure 1). Néanmoins, des questions se posent quant à la véracité de cette assertion, en particulierité lorsque mise en relation avec les connaissances acquises en psychologie organisationnelle au cours des 50 dernières années.

Figure 1
Réponse de ChatGPT à la question « Can AI reduce efficiency in the workplace? »
Source : OpenAI (2023). ChatGPT. OpenAI. <https://openai.com/>



Avant toute chose, il convient de préciser que la grande majorité des études relevant une augmentation de l'efficacité au travail le font (implicitement) avec les modèles de l'IA de l'ère « prétransformeur » (Neilson *et al.*, 2008). Ainsi, les applications en milieu de travail sont relativement limitées et sont réservées à une catégorie d'employés qui détiennent des compétences spécifiques (Hudson et Cohen, 1999; Moskewicz *et al.*, 2001). Dans cet article, l'accent sera plutôt mis sur les nouvelles formes d'IA à la *ChatGPT*, qui seront ou sont déjà déployées dans les organisations sous forme de services tels que *Copilot* de la suite Office 365 (Kelly, 2023), ou *Bard*, par Google (Pichai, 2023). Ces nouvelles formes d'IA marquent une rupture, car elles sont déployées dans les organisations en vue d'une utilisation quotidienne à grande échelle et pour des utilisateurs qui n'ont pas de compétence spécifique dans ce domaine.

Pour appuyer l'argumentaire de cet article, le concept de stupidité organisationnelle sera la pierre angulaire afin d'illustrer les défis auxquels les organisations risquent de se heurter lors de l'intégration de l'IA dans leur milieu de travail, et comment les promesses d'efficacité de l'IA seront ultimement défaits par la réalité de la psychologie de l'humain dans les organisations. Bref, un peu comme la venue des ordinateurs dans les milieux de travail, l'ajout d'un outil rationnel n'augmente pas nécessairement les performances d'une organisation, qui, elle, est irrationnelle.

L'une des plus grandes erreurs que peut faire un chercheur en psychologie organisationnelle est de supposer qu'un travailleur prendra toujours la décision la plus rationnelle et bénéfique pour son organisation. En effet, les milieux organisationnels, par leurs rationalités limitées ainsi que leurs intérêts multiples et divergents, sont propices à être victimes de ce que plusieurs auteurs nomment la stupidité organisationnelle. Ce concept fait référence à une stupidité collective systémique qui peut potentiellement imprégner des organisations entières (Kerfoot, 2003; Paulsen, 2017; Peter et Hull, 1969). La stupidité organisationnelle est définie comme l'absence de réflexivité, de raisonnement substantiel et de justification dans les organisations (Karimi-Ghartemani *et al.*, 2020). Cela ne signifie pas que les travailleurs de l'organisation sont dépourvus d'intelligence. Au contraire, comme le mentionne Albrecht (1980, p. 236) : « *Intelligent people who can think and act very effectively as individuals can, as a group, exhibit the most profoundly stupid and counterproductive responses to the demands of their environment* ».

La recherche sur la stupidité organisationnelle a été fondée par quelques publications provenant de McComber et Jenkins (1972), Parkinson (1955) et Peter et Hull (1969) avant de retrouver un second souffle au cours de la dernière décennie (Alvesson et Spicer, 2012; Paulsen, 2017; Starkey *et al.*, 2019; Azevedo, 2020).

2.1 L'IA et la stupidité organisationnelle

Fondamentalement, la stupidité organisationnelle empêche les organisations d'accumuler des connaissances et d'améliorer leur prise de décision. Cette stupidité procure un sentiment de confiance qui permet aux organisations de fonctionner sans friction et conduit à préserver l'organisation et ses membres des incertitudes causées par le doute (Alvesson et Spicer, 2017). Dans une telle situation, l'organisation se crée en quelque sorte une réalité, à la fois parallèle et incohérente par rapport à son environnement externe. Ce déphasage provoque une prise de décision qui est suboptimale pour l'organisation. Trois dimensions peuvent alimenter la stupidité organisationnelle. Tout d'abord, la stupidité fonctionnelle est tributaire de l'attitude des gestionnaires – en particulier lorsque ceux-ci imposent une discipline qui contraint la relation entre les employés (Alvesson et Spicer, 2012). La seconde dimension, soit l'incompétence organisationnelle, se manifeste plutôt au niveau de la structure de l'organisation (Peter et Hull, 1969). Il s'agit ici de toutes les politiques/pratiques internes qui empêchent l'organisation d'apprendre. La dernière dimension s'intéresse à ce qui est dit dans l'organisation, soit la *bullshit* organisationnelle (Christensen *et al.*, 2019). Les prochaines sections s'affaireront à décortiquer chacune de ces dimensions et à les mettre en lien avec la démocratisation de l'IA dans les milieux de travail.



2.2 Une stupidité fonctionnelle

La stupidité fonctionnelle tend à être plus présente quand les gestionnaires forcent une discipline qui contraint la relation entre les employés, la créativité et la réflexion organisationnelle (Alvesson et Spicer, 2012). Dans ces organisations où la stupidité fonctionnelle est élevée, les gestionnaires refusent le raisonnement rationnel, les nouvelles idées, et résistent au changement, ce qui a pour effet d'augmenter la stupidité organisationnelle (Paulsen, 2018).

Ce genre de culture engendre une stupidité de groupe qui encourage les employés à ne pas travailler en équipe, ainsi qu'à ne pas partager leurs connaissances avec leurs collègues (Karimi-Ghartemani et al., 2020). Ces comportements entraîneront une diminution de l'apprentissage organisationnel et de l'intelligence totale dans l'organisation (Doaei, 2012). En d'autres termes, ce type de culture pousse les employés à conserver leurs ressources et connaissances pour leurs bénéfices personnels au détriment de l'organisation.

L'intégration de l'IA dans les milieux de travail ne viendra qu'exacerber la stupidité de groupe. Traditionnellement, le développement de capacités par l'employé au sein de son organisation était chronophage, et nécessitait un certain degré d'expertise. Par exemple, un employé ne pouvait pas simultanément s'approprier les résultats d'une planification financière, détecter les nouvelles tendances de la clientèle à partir d'une feuille Excel et préparer le résumé d'un rapport de recherche. C'est pourquoi l'organisation fait appel à différents « experts » qui ont une tâche précise à accomplir, pour ensuite transmettre l'information aux équipes de travail. L'IA, par son instantanéité, permet la réalisation de ces tâches d'une manière rapide. Cependant, cela produit l'effet pervers de rendre ces actions disponibles à tous les employés de l'organisation, et ce, sans contextualisation de l'extrait. Ainsi, dans leur quête de conservation de leurs ressources, les employés auront tendance à décupler le nombre de demandes faites à partir de l'IA, pour un travail qui n'était pas initialement le leur. Bref, nous assisterons à la multiplication d'une même action (p. ex., résumé d'un rapport fiscal) par plusieurs travailleurs ce qui, ultimement, réduira l'efficacité organisationnelle. En effet, comme le dénote Boudreau (2010), l'accès numérique aux données par l'entremise de l'IA n'améliore pas les performances organisationnelles si les bonnes questions ne sont pas préalablement posées par une ressource experte. Le fait d'analyser des données sans se poser préalablement les bonnes questions risque de provoquer l'accumulation de données inutiles « *warehouses of data* » au mieux, ou la mésinterprétation des données recueillies, au pire (Rasmussen et al., 2024).

2.3 L'IA et l'incompétence organisationnelle

Alors que la première apparition de la stupidité organisationnelle se situe au niveau du gestionnaire, la seconde manifestation se traduit au niveau de la structure de l'organisation. L'incompétence organisationnelle peut être définie comme le « schéma répété d'une organisation incapable ou désireuse d'apprendre de son environnement, de ses échecs ou de ses succès » (Ott et Shafritz, 1994, p. 370). Ce concept, bien que proche de la stupidité fonctionnelle, se distingue néanmoins sous deux angles. Premièrement, l'incompétence a une connotation un peu plus procédurale et moins axée sur le comportement des gestionnaires (Peter et Hull, 1969). Deuxièmement, l'incompétence tend à être perçue comme étant plus susceptible d'être corrigée par une restructuration de l'organisation (Peter et Hull, 1969) plutôt que par des changements individuels (Ott et Shafritz, 1994). L'un des premiers auteurs ayant tenté d'expliquer la stupidité organisationnelle (à partir de l'incompétence) à l'aide de principes est Cyril N. Parkinson. La loi de Parkinson est l'observation que la bureaucratie organisationnelle se développe, quelle que soit la quantité de travail à accomplir (Parkinson, 1955). En d'autres termes, le travail (et la charge cognitive nécessaire à celle-ci) sera modulé pour remplir la totalité du délai prescrit. Ce principe détient des preuves empiriques dans la littérature scientifique (Locke et Bryan, 1967; Gutierrez et Kouvelis, 1991), et s'applique tout autant à l'organisation privée (Latham et Locke, 1975). Parkinson (1995) attribue ce phénomène à deux facteurs, soit le fait que les gestionnaires veulent multiplier leurs subordonnés, et que

les gestionnaires se créent mutuellement du travail entre eux. Ces deux facteurs provoquent un milieu organisationnel où des tâches à peu de valeur ajoutée sont créées périodiquement et attribuées à un bassin de main-d'œuvre grandissant. Ce genre d'emplois, ne servant qu'à flatter l'égo de certains gestionnaires, est nommé *bullshit job*, selon Graeber (2018). Ce dernier avance que les progrès technologiques de l'ère numérique ont réduit le besoin de travail manuel et répétitif. Cependant, au lieu de combler ce vide par une augmentation des activités de loisirs, Graeber affirme que pour soutenir la croissance économique alimentée par les consommateurs, des emplois – aussi inutiles soient-ils – doivent être créés en permanence. Graeber (2018) estime que jusqu'à 50 % de la main-d'œuvre de l'ère numérique occupe des emplois ayant ces caractéristiques.

Dans un contexte organisationnel à forte incompetence, c'est-à-dire où les tâches à réaliser sont modulées par rapport au temps disponible pour les accomplir, l'intégration de l'IA ne permettra pas une augmentation de l'efficacité, bien au contraire. La loi de Parkinson nous renseigne sur le fait que si les délais ne sont pas judicieusement ajustés par les gestionnaires, les tâches à réaliser par les employés – même si celles-ci peuvent se faire pratiquement instantanément à l'aide de l'IA – s'étireront néanmoins aussi loin que les échéanciers. Par exemple, un gestionnaire pourrait demander à un employé de systématiquement se servir de l'IA afin de créer un rapport exhaustif à partir des notes prises au cours de chaque réunion – démontrant ainsi la quantité de travail réalisée – afin d'envoyer ces documents à une autre équipe de travail qui risque fort bien de réutiliser l'IA afin de créer un résumé de ce rapport afin d'en faciliter la lecture. Qui plus est, les égos des gestionnaires, mesurés par le nombre de subordonnés sous leur aile, ne seront pas tentés de réduire leurs effectifs, mais de plutôt trouver des tâches, voire créer des postes entiers à peu de valeur ajoutée, pour maintenir en emploi leurs effectifs. De plus, puisque les données générées par l'IA sont souvent testées sur des modèles prédictifs qui proposent d'interpréter les données historiques plutôt que d'offrir de nouvelles possibilités pour le futur (Rasmussen et al., 2024), un risque subsiste pour l'organisation de se baser uniquement sur leurs expériences passées pour prédire les prochaines décisions. En d'autres termes, l'IA ne semble pas être en mesure de penser *outside the box*, ce qui augmente le risque pour l'organisation de répéter les mêmes comportements dans un environnement qui, lui, change.

Un second élément de l'incompétence organisationnelle provient de la façon dont une organisation promeut ses gestionnaires. Selon Mohammed (2020), de plus en plus d'organisations de l'ère du numérique se retrouvent dans une *kakistocracy*, où l'individu qui semble avoir le moins de compétences pour occuper un poste de gestion se retrouve néanmoins dans cette position. Ce processus prend forme dans la manière dont les structures hiérarchiques octroient les promotions, souvent en considérant la performance actuelle de l'employé dans son poste actuel, plutôt que de considérer ses habiletés à répondre aux exigences du nouveau rôle. Le principe de Peter (Peter et Hull, 1969) propose que lorsque ce phénomène est reproduit à grande échelle, l'organisation se retrouve éventuellement avec une structure hiérarchique où les gestionnaires, réalisant qu'ils ne sont pas en mesure de remplir leur rôle de façon satisfaisante, ressentent du stress, de l'anxiété et une perte de confiance. Ainsi, les organisations qui offrent déjà des promotions aux employés sous l'égide d'une bonne performance actuelle risquent d'exacerber le principe de Peter. Par exemple, un employé, en mesure de maîtriser l'IA plus rapidement que ses collègues, aura un avantage sur ceux-ci, en étant en mesure d'écrire un code de programmation en un temps record à l'aide de l'IA pour régler plusieurs problèmes mineurs au travail. Cette aptitude lui sera favorable lors de son évaluation de rendement et son éventuelle promotion. Néanmoins, son expertise en IA ne lui permettra pas de relever les défis de résolution de conflit et de leadership que les postes de gestion amènent. Le nouveau gestionnaire risque alors de souffrir d'*injelitance* (une combinaison d'incompétence et de jalousie) (Brimelow, 1989) par rapport à ses nouveaux défis. Dans son sentiment d'incompétence, le gestionnaire aura besoin de plus de temps pour prendre une décision (Harrison, 1996; Spanier, 2011), et aura tendance à trouver des solutions à des problèmes inexistantes (voir *The Garbage-Can Model*, Cohen et al., 1972). Cela pourrait se traduire, par exemple, par un gestionnaire qui demande à ses subordonnés la création de multiples rapports au moyen



de l'IA afin de « s'assurer » de bien avoir toute l'information disponible avant de prendre une décision; après tout, l'IA permet l'élaboration de ces documents en si peu de temps! Néanmoins, selon Agrawal et ses collaborateurs (2020), une utilisation en silo de l'IA risque de donner des résultats (et une prise de décision y découlant) qui seront déconnectés des besoins organisationnels.

2.4 L'IA et la *bullshit*

Alors que l'incompétence organisationnelle s'attarde à ce qui est fait, la *bullshit* s'intéresse à ce qui est dit. Bien que les deux concepts se soldent par une perte de temps et de ressources pour l'organisation, la *bullshit* est plus souvent associée à des pratiques en management (Christensen *et al.*, 2019). Selon Frankfurt (2005, p. 60), la *bullshit* est plus dangereuse que le mensonge, car elle ignore totalement les faits :

Someone who lies and someone who tells the truth are playing on opposite sides, so to speak, in the same game. Each responds to the facts as he understands them, although the response of the one is guided by the authority of the truth, while the response of the other defies that authority and refuses to meet its demands. The bullshitter ignores these demands altogether. He does not reject the authority of the truth, as the liar does, and oppose himself to it. He pays no attention to it at all. By virtue of this, bullshit is a greater enemy of the truth than lies are.

Deux fonctions sont proposées par la *bullshit* organisationnelle. La fonction première est de gérer l'impression que les autres ont de soi-même (Allen *et al.*, 2012). En ce sens, elle détient un fondement d'utilité dans l'organisation en modulant les relations sociales entre les travailleurs. C'est pourquoi elle est fréquemment maintenue par design. Elle survient fréquemment, par exemple, lorsqu'une personne est appelée à prodiguer un conseil ou une opinion sur un sujet qui dépasse son expertise, ou lorsque les faits sont tout simplement inconnus (Spicer, 2018). Dans ces situations, le *bullshitter* tente néanmoins de maintenir sa crédibilité et son autorité grâce à un dialogue ou des actions vides de sens. La *bullshit* organisationnelle visant à gérer les impressions peut prendre de nombreuses formes, telles que des réunions interminables, des mesures et de la paperasserie inutiles, et l'accent mis sur l'apparence d'être occupé plutôt que sur l'accomplissement d'un travail significatif (Spicer, 2018). Ces pertes de temps proviennent d'un style de gestion qui désire donner l'impression de réaliser plusieurs tâches plutôt que d'obtenir de véritables résultats. Dans son langage, le *bullshitter* utilisera des termes comme « gagnant-gagnant » ou « effet de levier » qui semblent souvent convaincants pour l'autre partie, car ils sont présentés de manière qu'ils paraissent inévitables et hautement souhaitables (Argyris, 1990). Dans de tels cas, le message vide de sens peut étouffer la critique constructive qui est nécessaire à l'apprentissage organisationnel (Senge, 1990). Qui plus est, dans certaines situations, les représentations répétées du *bullshitter* peuvent avoir un effet transformateur sur ce dernier, modifiant ainsi la perception qu'il a de lui-même et de son importance dans l'organisation (Allen *et al.*, 2012). Mears (2002) indique que la réalité fabriquée par le *bullshitter* est à risque de se diffuser à toute l'organisation lors d'une situation de crise, alors que les remises en question sont effectuées sur les pratiques actuelles. Les employés cherchent alors une réponse à leurs interrogations, souvent en faisant fi de la réalité et en acceptant, sans trop de résistance, la proposition émise par le *bullshitter*.

La seconde utilité de la *bullshit* organisationnelle est de permettre aux gestionnaires de commander leur main-d'œuvre, sans en donner l'impression. Comme le souligne Jackall (1988), dans les organisations modernes, les employés – du fait de leurs compétences et expertises – ont moins tendance à accepter des ordres explicites venant de leur supérieur. La *bullshit* offre alors une solution de rechange aux gestionnaires afin de maintenir ce pouvoir (Christensen *et al.*, 2019). En effet, bien que les propos soient vagues et vides de sens, l'euphémisme créé par la *bullshit* permet au gestionnaire de commander sans donner d'ordres directs. En tant que tel, il maintient les positions de pouvoir existantes tout en permettant à toutes les personnes impliquées de maintenir une image positive d'elles-mêmes (Eisenberg, 2007). Ainsi donc, la

bullshit organisationnelle permet la décontextualisation d'une situation afin de permettre au *bullshitter* de rendre une décision sans en prendre la responsabilité. Dans l'une des rares études sur le sujet, Bartosiak et Modlinski (2022) ont démontré empiriquement que lors d'une décision disciplinaire à rendre envers un employé, les individus prendront moins de temps à prendre une décision (et donc à moins appréhender le contexte) s'ils sont assistés par l'IA, et ce, peu importe leur biais à l'égard de cette technologie.

Ces comportements de *bullshit* sont déjà visibles dans les grandes organisations. En 2018, Amazon a renvoyé 300 employés de son entrepôt de Baltimore, aux États-Unis, pour manque de rendement, lequel a été mesuré exclusivement par une IA. Selon l'enquête journalistique menée par Lecher (2019), l'IA, en étant inflexible sur les indicateurs de performance, ne considérait pas les tâches supplémentaires demandées par le poste, ni les *time off task* comme aller à la salle de bain. Bien qu'Amazon assure que le gestionnaire peut infirmer la décision, l'organisation refuse de donner le nombre de fois qu'un gestionnaire a manuellement vérifié un dossier de mesure disciplinaire. Pour répondre à ces allégations, Amazon a publié un document de huit pages qui présente plusieurs caractéristiques relevées par les auteurs mentionnés plus haut se rapportant à la *bullshit* organisationnelle – laquelle se reflétant par des mots vides de sens, accompagnés d'un manque de responsabilisation des décisions prises par l'organisation. À cet égard, Amazon ne semble pas avoir appris de ses expériences passées. L'utilisation de l'IA pour mesurer le rendement de ses travailleurs du programme de livraison de colis « Flex » entre 2020 et 2021 a mené au renvoi de plusieurs livreurs pour des raisons contextuelles hors de leur contrôle (Soper, 2021). En effet, l'IA ne semblait pas en mesure de considérer des facteurs circonstanciels comme des complexes d'appartements dont l'accès est restreint. L'IA envoyait automatiquement une lettre de terminaison d'emploi au travailleur si leur cote d'efficacité tombait à un certain niveau arbitraire. Bien que l'organisation offre un délai pour faire appel de la décision, il semble qu'Amazon n'est pas enclin à infirmer la décision de l'IA : « I am asking for specific details on how this decision for deactivation of my account was reached [...] But then they replied: "We understand that every delivery partner has difficult days and that you may sometimes experience delays, and we have already taken this into account." I didn't get my job back. » (Soper, 2021, p. 6).

3. Discussion

L'objectif de cet article était de démontrer l'inadéquation entre les promesses d'efficacité de l'IA et la réalité (souvent stupide) des organisations. En effet, les 50 dernières années d'études en psychologie organisationnelle démontrent que les décideurs ne devraient pas espérer que l'intégration d'un outil aussi rationnel que l'IA au travail générerait automatiquement les bénéfices escomptés en temps et ressources. L'utilisation du concept de stupidité organisationnelle a été en mesure de relever plusieurs répercussions négatives qui semblent être ignorées par la littérature (voir tableau 1).



Tableau 1
Conséquences de l'IA mises en relation avec les concepts de la stupidité organisationnelle

Stupidité organisationnelle	Concepts reliés	Conséquences
La stupidité fonctionnelle	<i>La stupidité de groupe.</i>	• Décuplement des demandes faites à partir de l'IA, redondance entre les « experts ».
L'incompétence organisationnelle	<i>Loi de Parkinson. Principe de Peter.</i>	• Création de postes pour flatter l'égo de certains gestionnaires. • Élasticité du travail, pour le compresser par la suite. • Augmentation du temps pour prendre une décision.
La <i>bullshit</i> organisationnelle	<i>Commander sans donner d'ordre.</i>	• Focus du gestionnaire sur l'output de l'IA aux dépens du processus. • Déresponsabilisation du processus décisionnel et justification faisant fi de la réalité.

Les conséquences de l'intégration de l'IA dans les milieux de travail relevées dans cet article tournent autour de deux thématiques, soit la multiplication des tâches et des postes à peu de valeur ajoutée, et la perte de réflexivité dans le processus de prise de décision. Traditionnellement, plusieurs tâches chronophages pouvaient être jugées comme non essentielles dans l'organisation; le ratio temps investi versus les bénéfices n'en valant pas la peine. Cependant, l'IA, par sa rapidité de traitement de données, rend ce calcul plus attrayant pour les gestionnaires. Ainsi, ceux-ci auront tendance à demander à leurs subordonnés une panoplie de tâches sous prétexte que ces informations pourraient éventuellement servir. Dans le même ordre d'idées, ces employés justifieront leur temps de travail (via – en se référant au principe de Peter) par l'adéquation de plusieurs tâches anodines donnant l'impression à l'équipe de gestion qu'ils atteignent plusieurs résultats.

La perte de réflexivité, quant à elle, est attribuée au fait que les organisations, par leurs processus d'attribution de promotion, ont une tendance à mettre un travailleur dans une position de gestion qui nécessite des capacités fort différentes comparativement à son précédent emploi; celui-ci est donc mal outillé pour y faire face. Dans ce contexte, il y a fort à parier que le nouveau gestionnaire s'appuiera de façon pathologique sur l'IA dans sa prise de décision afin de ne « rien oublier ». Qui plus est, par son peu d'expérience en leadership, le gestionnaire aura tendance à s'appuyer sur la conclusion de l'IA, hormis les éléments circonstanciels. Dès lors, il devra expliquer ses décisions sans maîtriser le processus, ce qui, inévitablement, se fera par un discours déconnecté de la réalité.

3.1 Implications pratiques

L'utilisation du concept de la stupidité organisationnelle comme prisme théorique, par lequel l'intégration (et la démocratisation) de l'IA dans le monde du travail est observée, met en relief plusieurs défis pour les organisations. Ainsi donc, il serait pertinent pour les décideurs de considérer certaines dimensions de leur organisation avant d'espérer obtenir des gains d'efficacité en ajoutant l'IA au flux de travail actuel. L'argumentaire de cet article permet de relever trois éléments.

Plus ne veut pas dire mieux

L'un des premiers pièges à éviter avec l'IA est sa capacité à gérer un nombre impressionnant de données de manière quasi instantanée. Certes, dans certaines circonstances cela peut s'avérer utile. Cependant, une

trop forte tendance d'une organisation à générer systématiquement des données « au cas où » ou pour être sûr de ne rien manquer, risque de devenir rapidement pathologique et donner lieu à des pertes d'efficacité. Selon Van Oudenhoven et ses collaborateurs (1998), cette tendance risque de s'accroître dans les organisations où la tolérance envers l'incertitude est faible – les dirigeants cherchent alors toutes les informations disponibles dans leur environnement avant de prendre la bonne décision. Ce faisant, ceux-ci augmentent le temps de prise de décision, ce qui peut entraîner une perte de fenêtre d'opportunité – en particulier dans une industrie volatile (Ceschi *et al.*, 2017).

Afin d'aider les décideurs à choisir à quel moment la sollicitation de l'IA serait pertinente, Rasmussen et ses collaborateurs (2024) suggèrent que les données doivent permettre aux gestionnaires de hiérarchiser les ressources qui seraient autrement dépensées ailleurs. Pour ce faire, les auteurs proposent de présenter les résultats de l'analyse suivis du marqueur « et alors? » afin de démontrer les interventions suggérées, le coût et l'impact attendu par le changement. Cette présentation devrait se tenir dans une seule diapositive, ce qui forcerait son créateur à ne présenter que les points clés, en évitant ainsi de tomber dans le piège d'une présentation longue et formelle. Cette suggestion est aussi entérinée par Wright et ses collaborateurs (2005) qui indiquent que l'analyse doit permettre aux gestionnaires de comprendre et interpréter facilement les résultats, mais surtout les implications du changement proposé pour l'organisation.

L'erreur fondamentale des analyses prédictives

Alors que les gestionnaires s'intéressent particulièrement à l'output, c'est-à-dire le résultat de l'IA, il importe néanmoins de comprendre le fonctionnement interne de l'algorithme afin de contextualiser le processus. L'IA devient « intelligente » en analysant des données existantes (p. ex., des sites web) afin de raffiner son modèle d'analyse. Ainsi, c'est par les analyses prédictives que l'IA se développe. Cette méthode peut être appropriée sur des sujets immuables comme les lois de la physique; cependant, elle est discutable dans un milieu organisationnel (Delen et Ram, 2018). En effet, ce qui permet à une organisation de se démarquer, à un moment donné, dans une industrie particulière, n'est peut-être pas dû aux données probantes (Rasmussen *et al.*, 2024). Fondamentalement, les analyses prédictives se concentrent seulement sur ce qui a conduit aux performances passées et non sur ce qui permettra d'augmenter les résultats futurs d'une organisation (Edwards et Edwards, 2019). En d'autres termes, une dépendance accrue aux données probantes offertes par l'IA ne garantit en rien le succès d'une organisation.

Pour pallier l'erreur fondamentale des analyses prédictives, Ulrich et ses collaborateurs (2021) proposent que les équipes de travail qui sondent l'IA doivent être chapeautées par une ressource dotée de compétences en gestion et d'un sens des affaires. En effet, c'est par une vision globale de la situation que la bonne question *prompt* peut être soumise à l'IA et être contextualisée par la suite (Kim *et al.*, 2023).

La place des RH par rapport à l'IA

Traditionnellement, l'utilisation de l'IA (modèle prétransformeur) en organisation était réservée au département d'informatique ou à quelques employés détenant des compétences particulières (Moskewicz *et al.*, 2001). Cela pouvait s'expliquer par le fait que les retombées des analyses de l'IA y étaient pour le moins limitées (Saner et Wallach, 2015). Cependant, avec la démocratisation de l'IA, plusieurs parties prenantes en entreprise commencent maintenant à utiliser cette technologie pour accomplir plusieurs de leurs tâches quotidiennes. Cette transition se doit d'être vue dans une approche holistique de l'IA et non au simple laisser-aller selon les compétences individuelles des employés. En somme, des politiques organisationnelles encadrant l'usage de l'IA se doivent d'être proactivement développées par la direction des ressources humaines (conjointement avec le département d'informatique, lorsque présent) afin d'être diffusées de façon uniforme à l'ensemble des travailleurs.



Un dialogue doit être enclenché afin de baliser les différents paramètres d'utilisation de l'IA (p. ex., donnons-nous à l'IA l'accès à nos données confidentielles? Est-ce que certaines activités RH [dotation, gestion du rendement, etc.] sont exclues dans notre utilisation de l'IA? Quelle est la marge de manœuvre d'un gestionnaire par rapport à une conclusion de l'IA?). Les RH doivent faire partie intégrante de ces questionnements et assumer leur rôle de partenaire stratégique envers l'organisation (Jo *et al.*, 2023). Dans le cas contraire, si ces responsabilités ne sont pas prises, le discours risque d'être dirigé par les spécialistes en informatique. N'ayant pas nécessairement une vision complète des acteurs en jeu, certaines décisions quant à l'IA pourraient alors être prises sans égard aux répercussions sur la main-d'œuvre. Ainsi, rejoignant les conclusions de plusieurs auteurs (Fenwick *et al.*, 2023; Yawalkar, 2019; Pandey *et al.*, 2024), les RH sont, à l'heure de l'implantation de l'IA dans les milieux organisationnels, plus importantes que jamais.

3.2 Piste de recherches futures

Malgré ce portrait quelque peu négatif dessiné du futur de l'intégration de l'IA dans le monde du travail, des avenues intéressantes pour les futures recherches se dressent. Tout d'abord, il serait pertinent que des études se penchent sur les formations nécessaires au bon fonctionnement de l'utilisation de l'IA dans un contexte organisationnel. À l'heure actuelle, l'utilisation de cet outil est réservée principalement aux travailleurs relevant du service en informatique (Moskewicz *et al.*, 2001). Sa démocratisation et son utilisation dans des contextes variés, notamment dans le domaine de la gestion des ressources humaines, nécessiteront l'appropriation de nouvelles compétences de la part des gestionnaires. Dans ce contexte, la réflexivité du gestionnaire devra être mise de l'avant; autrement, le dogme façonné par l'IA prendra le dessus.

Dans un second temps, le domaine de l'éthique en IA se doit de prendre une place de choix dans la littérature scientifique. Déjà, certaines études traitent du sujet (Jobin *et al.*, 2019; Hagendorff, 2020; Mittelstadt, 2019). Cependant, depuis l'intégration du modèle transformeur dans les modèles neuronaux, les avancées de l'IA sont exponentielles à un tel point que même ceux qui y travaillent quotidiennement ne comprennent pas tous les développements (Center for Humane Technology, 2023). Néanmoins, il convient de noter que la plupart des études en éthique de l'IA traitent des répercussions sur la société en général, telle que la désinformation (Kertysova, 2018), notamment lors d'élections (Marsden *et al.*, 2020). Il est primordial que la recherche développe un solide pan de l'éthique dans le monde des affaires.

Dans un troisième temps, les pistes de recherches futures devront se pencher sur les caractéristiques organisationnelles qui peuvent influencer la performance de celle-ci lors de l'intégration de l'IA dans les postes de travail. Les études antérieures en psychologie organisationnelle tendent à démontrer que les risques seront les plus grands dans des firmes ayant une forte bureaucratie; ce sont ces mêmes organisations qui risquent d'intégrer l'IA en premier. Nous faisons l'hypothèse que les organisations mettant l'accent sur une forte bureaucratie, qui ont une culture axée sur les résultats plutôt que sur les processus, et où la structure hiérarchique est constituée d'individus promus selon leur niveau de performance plutôt que par leur potentiel, seront celles qui auront le plus à perdre en intégrant l'IA.

Enfin, les chercheurs devraient répondre à l'appel lancé par Azevedo (2023) et approfondir le champ d'études de la stupidité organisationnelle. Plus particulièrement, au lieu de poursuivre futilement l'éradication de la stupidité organisationnelle, les chercheurs devraient apprendre à connaître la stupidité et à l'étudier, par exemple, comment sa compréhension permet aux individus de progresser et aux entreprises de prospérer. Il apparaît clair, selon nous, que ce concept sera en mesure d'expliquer plusieurs comportements qui se manifesteront prochainement dans les organisations.

Conclusion

Bien qu'une littérature grandissante louange les bénéfices de l'IA, l'objectif de cet article était de démontrer que les promesses d'efficacité de l'IA sont simplement incompatibles avec la recherche à ce jour sur la psychologie organisationnelle. Certes, tout comme l'arrivée des ordinateurs dans le milieu de travail, l'intégration de l'IA aura certains effets positifs; cette technologie ne sera toutefois pas la révolution espérée par plusieurs. En utilisant le concept de la stupidité organisationnelle ainsi que ses embranchements, ce texte a fait la démonstration, à l'aide d'exemples théoriques et pratiques, des risques afférents à la démocratisation de l'IA dans les milieux de travail. Une organisation qui fera fi de ses défis à l'interne en espérant être sauvée par l'objectivité d'une assistance artificielle risque d'avoir un effet catalyseur sur la stupidité fonctionnelle, l'incompétence organisationnelle et la *bullshit* organisationnelle.

Le fantasme techno-utopiste de certains gestionnaires se doit d'être recadré dans la réalité – souvent complexe et imprédictible – de la psychologie humaine. L'IA reste un outil pour l'organisation, au même titre qu'un marteau pour le menuisier; c'est la capacité de ce dernier à manier l'outil et savoir dans quelles circonstances il est pertinent (ou non) de l'employer qui fait sa force. En somme, dans sa forme actuelle, l'intégration de l'IA n'apportera pas les résultats espérés par tant de décideurs en organisation.

NOTE

- 1 Voir <https://twitter.com/ygrowthco/status/1760794728910712965>

RÉFÉRENCES

- Allen, DE, Allen, RS. et McGoun, EG. (2012). Bull markets and bull sessions. *Culture and Organization* 18(1), 15-31.
- Agrawal, V., Khanna, P. et Singhal, K. (2020). The role of research in business schools and the synergy between its four subdomains. *Academy of Management Review*, 45(4), 896-912. <https://doi.org/10.5465/amr.2019.0223>
- Albrecht K. (1980). *Brain Power: Learn to Develop Your Thinking Skills*. Englewood Cliffs, NJ: Prentice-Hall.
- Alvesson, M. et Spicer, A. (2012). A Stupidity-Based Theory of Organizations. *Journal of Management Studies*, 49(7), 1194-1220. <https://doi.org/10.1111/j.1467-6486.2012.01072.x>
- Alvesson, M. et Spicer, A. (2017). *The Stupidity Paradox: The Power and Pitfalls of Functional Stupidity at Work*. Profile Books.
- Argyris, C. (1990). *Overcoming organizational defenses: Facilitating organizational learning*. Allyn & Bacon.
- Ayinde, L., Prabu Wibowo, M., Bin Emdad, F. et Ravuri, B. (2023). ChatGPT as an important tool in organizational management: A review of the literature. *Business Information Review*, 43(3), 137-149. <https://doi.org/10.1177/02663821231187>
- Azevedo, G. (2023). Into the realm of organizational folly: A poem, a review, and a typology of organizational stupidity. *Management Learning*, 54(2), 267-281.
- Bag, S., Pretorius, J. H. C., Gupta, S. et Dwivedi, Y. K. (2021). Role of institutional pressures and resources in the adoption of big data analytics powered artificial intelligence, sustainable manufacturing practices and circular economy capabilities. *Technological Forecasting and Social Change*, 163. <https://doi.org/10.1016/j.techfore.2020.120420>
- Basch, C. H., Meleo-Erwin, Z., Fera, J., Jaime, C. et Basch, C. E. (2021). A global pandemic in the time of viral memes: COVID-19 vaccine misinformation and disinformation on TikTok. *Human vaccines & immunotherapeutics*, 17(8), 2373-2377.
- Bartosiak, M.L. et Modlinski, A. (2022), "Fired by an algorithm? Exploration of conformism with biased intelligent decision support systems in the context of workplace discipline", *Career Development International*, 27(6/7), 601-615. <https://doi.org/10.1108/CDI-06-2022-0170>
- Bender, E. M. et Friedman, B. (2018). Data Statements for Natural Language Processing: Toward Mitigating System Bias and Enabling Better Science. *Transactions of the Association for Computational Linguistics*, 6, 587-604. https://doi.org/10.1162/tacl_a_00041



- Boudreau, J. W. (2010). *Retooling HR: Using proven business tools to make better decisions about talent*. Harvard Business Press.
- Brimelow P (1989) How do you cure ineptitude? *Forbes* 144(3), 42-44.
- Center for Humane Technology. (2023). *The A.I. Dilemma—March 9, 2023*. <https://www.youtube.com/watch?v=xoVJKj8lcNQ>
- Ceschi, A., Demerouti, E., Sartori, R. et Weller, J. (2017). Decision-making processes in the workplace: How exhaustion, lack of resources and job demands impair them and affect performance. *Frontiers in psychology*, 8, 1-14.
- Chang, R., Kauffman, R. et Kim, K. (2013). How Strong Are the Effects of Technological Disruption? Smartphones' Impacts on Internet and Cable TV Services Consumption. *46th Hawaii International Conference on System Sciences*. <https://doi.org/10.1109/HICSS.2013.252>
- Choi, J., Hickman, K., Monahan, A. et Schwarcz, D. (2022). ChatGPT Goes to Law School. *Journal of Legal Education*, 71(387), 1-16. <https://dx.doi.org/10.2139/ssrn.4335905>
- Christensen, C. M., McDonald, R., Altman, E. J. et Palmer, J. E. (2018). Disruptive Innovation: An Intellectual History and Directions for Future Research. *Journal of Management Studies*, 55(7), 1043-1078. <https://doi.org/10.1111/joms.12349>
- Colbert, A., Yee, N., et George, G. (2016). The Digital Workforce and the Workplace of the Future. *Academy of Management Journal*, 59(3), 731-739. <https://doi.org/10.5465/amj.2016.4003>
- Cohen MD, March JG et Olsen JP. (1972). A garbage can model of organizational choice. *Administrative Science Quarterly* 17(1), 1-25.
- Delen, D. et Ram, S. (2018). Research challenges and opportunities in business analytics. *Journal of Business Analytics*, 1(1), 2-12.
- Doaci, M. (2012). *Identify and assess the causes of group stupidity in Pars Aviation Services Company* (Doctoral dissertation, Ms dissertation. Islamic Azad University of Khorasgan).
- Edwards, M. et Edwards, K. (2019). Predictive HR analytics: Mastering the HR metric. Kogan Page. 536 pages.
- Eisenberg, J. (2007). *Group Cohesiveness*, in Baumeister R.F. and Vohs K.D. (eds.), *Encyclopaedia of Social Psychology*, Sage, Thousand Oaks, 386-388.
- Fenwick, A., Molnar, G., et Frangos, P. (2023). Revisiting the role of HR in the age of AI: bringing humans and machines closer together in the workplace. *Frontiers in Artificial Intelligence*, 6. <https://doi.org/10.3389/frai.2023.1272823>
- Frankfurt HG (2005) *On Bullshit*. Princeton, NJ: Princeton University Press.
- Grandinetti, J. et Bruinsma, J. (2023). The affective algorithms of conspiracy TikTok. *Journal of Broadcasting & Electronic Media*, 67(3), 274-293.
- Grassini, S. (2023). Shaping the Future of Education: Exploring the Potential and Consequences of AI and ChatGPT in Educational Settings. *Education Sciences*, 13(7), 1-13. <https://doi.org/10.3390/educsci13070692>
- Gugerty, L. (2006). Newell and Simon's Logic Theorist: Historical Background and Impact on Cognitive Modeling. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, 50(9), 880-884. <https://doi.org/10.1177/154193120605000904>
- Gutierrez, G. J. et Kouvelis, P. (1991). Parkinson's law and its implications for project management. *Management Science*, 37(8), 990-1001.
- Graeber, D. (2018). *Bullshit Jobs: A Theory*. Simon & Schuster.
- Hagendorff, T. (2020). The ethics of AI ethics: An evaluation of guidelines. *Minds and machines*, 30(1), 99-120.
- Harrison, E. F. (1996). A process perspective on strategic decision making. *Management decision*, 34(1), 46-53.
- Han, X., Zhang, Z., Ding, N., Gu, Y., Liu, X., Huo, Y., Qiu, J., Yao, Y., Zhang, A., Zhang, L., Han, W., Huang, M., Jin, Q., Lan, Y., Liu, Y., Liu, Z., Lu, Z., Qiu, X., Song, R., ... Zhu, J. (2021). Pre-trained models: Past, present and future. *AI Open*, 2, 225-250. <https://doi.org/10.1016/j.aiopen.2021.08.002>
- Heinrichs, B. (2022). Discrimination in the age of artificial intelligence. *AI & SOCIETY*, 37(1), 143-154. <https://doi.org/10.1007/s00146-021-01192-2>

- Hudson, D. L. et Cohen, M. E. (1999). *Neural Networks and Artificial Intelligence for Biomedical Engineering*. John Wiley & Sons.
- IBM. (n.d.). *What is Artificial Intelligence (AI)?* Retrieved March 1, 2024, from <https://www.ibm.com/topics/artificial-intelligence>
- Jackall, R. (1988). Moral mazes: The world of corporate managers. *International Journal of Politics, Culture, and Society*, 1, 598-614.
- Jo, J., Chadwick, C. et Han, J. (2023). How the human resource (HR) function adds strategic value: A relational perspective of the HR function. *Human Resource Management*, 63(1), 5-23. <https://doi.org/10.1002/hrm.22184>
- Jobin, A., Ienca, M. et Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature machine intelligence*, 1(9), 389-399.
- Karimi-Ghartemani, S., Khani, N. et Nasr Isfahani, A. (2022). A qualitative analysis and a conceptual model for organizational stupidity. *Journal of Organizational Change Management*, 35(3), 441-462.
- Kelly, S. M. (2023, March 16). *Microsoft is bringing ChatGPT technology to Word, Excel and Outlook* | CNN Business. CNN. <https://www.cnn.com/2023/03/16/tech/openai-gpt-microsoft-365/index.html>
- Kerfoot K (2003) Organizational intelligence/organizational stupidity: The leader's challenge. *Nursing Economic*, 21(2), 91-93.
- Kertysova, K. (2018). Artificial intelligence and disinformation: How AI changes the way disinformation is produced, disseminated, and can be countered. *Security and Human Rights*, 29(1-4), 55-81.
- Kim, C. et Mauborgne, R. (2005). Blue Ocean Strategy: From Theory to Practice. *California Management Review*, 47(3), 105-121. <https://doi.org/10.1177/000812560504700301>
- Kim, K., Clark, K. et Messersmith, J. (2023). High performance work systems and perceived organizational support: The contribution of human resource department's organizational embodiment. *Human Resource Management Journal*, 62(2), 181-196. <https://doi.org/10.1002/hrm.22142>
- Latham, G. P. et Locke, E. A. (1975). Increasing productivity and decreasing time limits: A field replication of Parkinson's law. *Journal of Applied Psychology*, 60(4), 524-526. <https://doi.org/10.1037/h0076916>
- Lecher, C. (2019, April 25). How Amazon automatically tracks and fires warehouse workers for 'productivity.' The Verge. <https://www.theverge.com/2019/4/25/18516004/amazon-warehouse-fulfillment-centers-productivity-firing-terminations>
- Locke, E. A. et Bryan, J. F. (1967). Performance goals as determinants of level of performance and boredom. *Journal of Applied Psychology*, 51(2), 120-130.
- Marsden, C., Meyer, T. et Brown, I. (2020). Platform values and democratic elections: How can the law regulate digital disinformation?. *Computer law & security review*, 36, 2-18.
- McComber, C., et Jenkins, N. (1972). The organizational stupidity factor. *Management Review*, 61(5), 44-46.
- Mears, D. (2002). The ubiquity, functions, and contexts of bullshitting. *Journal of Mundane Behavior*, 3(2), 233-256.
- Mei, L., Feng, X., et Cavallaro, F. (2023). Evaluate and identify the competencies of the future workforce for digital technologies implementation in higher education. *Journal of Innovation & Knowledge*, 8(4), 2-19. <https://doi.org/10.1016/j.jik.2023.100445>
- Mittelstadt, B. (2019). Principles alone cannot guarantee ethical AI. *Nature machine intelligence*, 1(11), 501-507.
- Mohammed S (2020) Book review: Business Bullshit. *Management Learning* 51(2), 244-246.
- Moskewicz, M. W., Madigan, C. F., Zhao, Y., Zhang, L. et Malik, S. (2001). Chaff: Engineering an efficient SAT solver. *Proceedings of the 38th Annual Design Automation Conference*, 530-535. <https://doi.org/10.1145/378239.379017>
- Neilson, GL, Martin, KL et Powers, E. (2008). The secrets to successful strategy execution. *Harvard Business Review*, 86(6), 60-70.
- Newell, A. et Simon, H. (1956). The logic theory machine—A complex information processing system. *IRE Transactions on Information Theory*, 2(3), 61-79. <https://doi.org/10.1109/TTT.1956.1056797>



- OpenAI. (2023). *GPT-4 Technical Report* (arXiv:2303.08774; Computation and Language, Technical Report. 100p. <https://doi.org/10.48550/arXiv.2303.08774>
- Ott S. et Shafritz, J. (1994). Toward a Definition of Organizational Incompetence: A Neglected Variable in Organization Theory. *Public Administration Review*, 54(4), 370-377. <https://doi.org/10.2307/977385>
- Paulsen, R. (2017). Slipping into functional stupidity: The bifocality of organizational compliance. *Human Relations*, 70(2), 185-210.
- Pandey, A., Grima, S., Pandey, S. et Balusamy, B. (Eds.). (2024). *The Role of HR in the Transforming Workplace: Challenges, Technology, and Future Directions*. CRC Press. 186p.
- Parkinson, C. (1955). Parkinson's Law. *The Economist*.
- Peter, R. et Hull, L. J. (1969). *The Peter Principle* (First Edition). William Morrow.
- Petzold, N., Landinez, L. et Baaken, T. (2019). Disruptive innovation from a process view: A systematic literature review. *Creativity and Innovation Management*, 28(1), 157-174. <https://doi.org/10.1111/caim.12313>
- Pichai, S. (2023, February 6). *An important next step on our AI journey*. Google. <https://blog.google/technology/ai/bard-google-ai-search-updates/>
- Quach, K. (2023). *US Justice Dept checking AI rent-pricing biz* RealPage. https://www.theregister.com/2022/11/28/doi_realpage_ai/
- Raj, R., Singh, A., Kumar, V. et Verma, P. (2023). Analyzing the potential benefits and use cases of ChatGPT as a tool for improving the efficiency and effectiveness of business operations. *BenchCouncil Transactions on Benchmarks, Standards and Evaluations*, 3(3), 1-10. <https://doi.org/10.1016/j.tbench.2023.100140>
- Rasmussen, T. H., Ulrich, M. et Ulrich, D. (2024). Moving People Analytics From Insight to Impact. *Human Resource Development Review*, 23(1), 11-29.
- Rivera, J-P., Mukobi, G., Reuel, M., Lamparth, M., Smith, C. et Schneider, J. (2024). Escalation Risks from Language Models in Military and Diplomatic Decision-Making. <https://doi.org/10.48550/arXiv.2401.03408>
- Saner, M. et Wallach, W. (2015). Technological unemployment, AI, and workplace standardization: The convergence argument. *Journal of Ethics and Emerging Technologies*, 25(1), 74-80.
- Scheihauf, R. (2023). While courts still use fax machines, law firms are using AI to tailor arguments for judges. *CBC*. <https://www.cbc.ca/news/opinion/opinion-artificial-intelligence-courts-legal-analytics-1.6762257>
- Senge, P.M. (1990). *The Fifth Discipline: The Art and Practice of the Learning Organization*, Doubleday/Currency, New York, NY.
- Si, S., Zahra, S., Wu, X. et Jeng, D. J.-F. (2020). Disruptive innovation and entrepreneurship in emerging economics. *Journal of Engineering and Technology Management*, 58, 1-12. <https://doi.org/10.1016/j.jengtecman.2020.101601>
- Smith, T. et Short, A. (2022). Needs affordance as a key factor in likelihood of problematic social media use: Validation, latent Profile analysis and comparison of TikTok and Facebook problematic use measures. *Addictive Behaviors*, 129. <https://doi.org/10.1016/j.addbeh.2022.107259>
- Soper, S. (2021, June 28). Fired by Bot at Amazon: 'It's You Against the Machine.' *Bloomberg.Com*. <https://www.bloomberg.com/news/features/2021-06-28/fired-by-bot-amazon-turns-to-machine-managers-and-workers-are-losing-out>
- Spanier, N. (2011). *Competence and incompetence training, impact on executive decision-making capability: Advancing theory and testing* (Doctoral dissertation, Auckland University of Technology).
- Spicer A (2018) *Business Bullshit*. London: Routledge.
- Staccioli, J. et Virgillito, M. E. (2020). The Present, Past, and Future of Labor-Saving Technologies. In K. F. Zimmermann (Ed.), *Handbook of Labor, Human Resources and Population Economics* (pp. 1-16). Springer International Publishing. https://doi.org/10.1007/978-3-319-57365-6_229-1

- Starkey K, Tempest S, et Cinque S (2019) Management education and the theatre of the absurd. *Management Learning* 50(5), 591-606.
- Stein-Perlman, Z., Weinstein-Raun, B., et Grace K. (2022). 2022 Expert Survey on Progress in AI. *AI Impacts*. <https://aiimpacts.org/2022-expert-survey-on-progress-in-ai/>
- Stokel-Walker, C. (2023). Generative AI Is Coming for the Lawyers. *WIRED*. <https://www.wired.co.uk/article/generative-ai-is-coming-for-the-lawyers>
- Soper, S. (2021, June 28). *Fired by Bot at Amazon: It's You Against the Machine*. Bloomberg.Com. <https://www.bloomberg.com/news/features/2021-06-28/fired-by-bot-amazon-turns-to-machine-managers-and-workers-are-losing-out>
- Tan, J., Westermann, H. et Benyekhlef, K. (2023). ChatGPT as an Artificial Lawyer? *Workshop on Artificial Intelligence for Access to Justice*, 1-9.
- Ulrich, M., Ignam Mathon, A., Pettman, J., Malo, S. et Dias, N. (2021). What do career paths in HR analytics look like? [Research report]. HR Analytics ThinkTank.
- Van Oudenhoven, J. P., Mechelse, L. et De Dreu, C. K. (1998). Managerial conflict management in five European countries: The importance of power distance, uncertainty avoidance, and masculinity. *Applied psychology*, 47(3), 439-455.
- Vecchiato, R. (2017). Disruptive innovation, managerial cognition, and technology competition outcomes. *Technological Forecasting and Social Change*, 116, 116-128. <https://doi.org/10.1016/j.techfore.2016.10.068>
- Wang, B., Liu, Y. et Parker, S. K. (2020). How Does the Use of Information Communication Technology Affect Individuals? A Work Design Perspective. *Academy of Management Annals*, 14(2), 695-725. <https://doi.org/10.5465/annals.2018.0127>
- Weimann, G. et Masri, N. (2023). Research note: Spreading hate on TikTok. *Studies in conflict & terrorism*, 46(5), 752-765.
- West, J. et Mace, M. (2010). Browsing as the killer app: Explaining the rapid success of Apple's iPhone. *Telecommunications Policy*, 34(5-6), 270-286. <https://doi.org/10.1016/j.telpol.2009.12.002>
- Wright, P., Gardner, T., Moynihan, L. et Allen, M. (2005). The relationship between HR practices and firm performance: Examining causal order. *Personnel Psychology*, 58(2), 409-446. <https://doi.org/10.1111/j.1744-6570.2005.00487.x>
- Yawalkar, M. V. V. (2019). A study of artificial intelligence and its role in human resource management. *International Journal of Research and Analytical Reviews (IJRAR)*, 6(1), 20-24
- Zhang, B., Dreksler, N., Anderljung, M., Kahn, L., Giattino, C., Dafoe, A. et Horowitz, M. C. (2022). Forecasting AI progress: Evidence from a survey of machine learning researchers. <https://doi.org/10.48550/arXiv.2206.04132>