

القضايا الأخلاقية في استخدام الذكاء الاصطناعي (AI)

Ethical issues in the use of artificial intelligence (AI)

هشام جادالله منصور شخاترة¹ - إسراء محمد عباينة²

¹ جامعة جدارا كلية الحقوق - الأردن - عمان / dr.hishamshakhatreh@gmail.com / hshakhatreh@jadara.edu.jo

² جامعة برشلونة المستقلة - كلية الأعمال - الأردن - عمان / esraaababneh99@gmail.com

ملخص:

تقدم هذه الورقة لمحة موجزة عن القضايا الأخلاقية المتعلقة باستخدام الذكاء الاصطناعي (AI). مع التقدم السريع في تقنيات الذكاء الاصطناعي، نشأت مخاوف بشأن تأثيرها على المجتمع، بما في ذلك قضايا الخصوصية والتحيز والمساءلة واستبعاد الوظائف. هذه الاعتبارات الأخلاقية مهمة للتعامل معها في تطوير ونشر أنظمة الذكاء الاصطناعي (AI) لضمان استخدامها بطريقة مسؤولة ومفيدة، تسلط هذه الورقة الضوء على القضايا الأخلاقية في استخدام الذكاء الاصطناعي الحاسمة وتتطرق إلى الحاجة إلى استمرار الحوار والتعاون بين التقنيين وصانعي السياسات وأصحاب المصلحة الآخرين لضمان تطوير الذكاء الاصطناعي (AI) واستخدامه بطريقة تحترم المبادئ الأخلاقية وتحمي مصالح الأفراد والمجتمع ككل.

الكلمات المفتاحية: الذكاء الاصطناعي (AI)، الخوارزميات، اتخاذ القرار المتحيز، قضايا الخصوصية، التحيز، المساءلة، استبعاد الوظائف، الشفافية والمساءلة.

تصنيف JEL: C44.

Abstract:

This paper provides a brief overview of ethical issues related to the use of artificial intelligence (AI). With the rapid advancement of AI technologies, concerns have arisen about its impact on society, including issues of privacy, bias, accountability, and job exclusion. These ethical considerations are important to address in the development and deployment of artificial intelligence (AI) systems to ensure that they are used in a responsible and beneficial manner. This paper highlights the critical ethical issues in the use of artificial intelligence and addresses the need for continued dialogue and collaboration among technologists, policy makers, and other stakeholders to ensure that artificial intelligence (AI) is developed and used in a manner that respects ethical principles and protects the interests of individuals and society as a whole.

Keywords: artificial intelligence (AI), algorithms, biased decision making, privacy issues, bias, accountability, job exclusion, transparency and accountability.

Jel Classification Codes: C44.

حقق الذكاء الاصطناعي (AI) تطورات هائلة في السنوات الأخيرة، وأصبح تأثيره محسوسا في مختلف المجالات وجوانب حياتنا اليومية. في حين أن الذكاء الاصطناعي (AI) لديه القدرة على إحداث تغييرات إيجابية كبيرة، فإنه يثير أيضا مخاوف وتحديات أخلاقية (Bostrom, & Yudkowsky, 2018).

تتضمن بعض القضايا الأخلاقية المتعلقة بالذكاء الاصطناعي (AI) انتهاكات الخصوصية والتمييز والتحيز واستبعاد الوظائف والمساءلة عن أنظمة الذكاء الاصطناعي (AI) التي تسبب الضرر في مصالح الأفراد والمنظمات (Burton et al., 2017). ومن المهم للمجتمع معالجة هذه الاعتبارات الأخلاقية لتطوير ونشر تقنيات الذكاء الاصطناعي (AI) وضمان استخدامها بطريقة مسؤولة ومفيدة وصحيحة (Safdar et al., 2020).

يتزايد استخدام الذكاء الاصطناعي (AI) في التعليم لدعم وتعزيز عملية التعلم، يشير الذكاء الاصطناعي في التعليم إلى استخدام الأساليب الحسابية والقائمة على البيانات لدعم التدريس والتعلم والتقييم في البيئات التعليمية الرسمية وغير الرسمية (Munoko et al., 2020). الهدف من الذكاء الاصطناعي (AI) في التعليم هو استخدام التكنولوجيا لمساعدة الطلاب على اكتساب المعرفة وتطوير المهارات وتحقيق أهدافهم التعليمية بطريقة أكثر كفاءة وفعالية وجاذبية (Coeckelbergh, 2019).

الذكاء الاصطناعي (AI) هو مجال سريع التطور ولديه القدرة على إحداث ثورة في العديد من جوانب المجتمع، بما في ذلك الاقتصاد والرعاية الصحية والتعليم والاتصالات (Kellmeyer, 2019).

ومع ذلك، فإن استخدام الذكاء الاصطناعي يثير مخاوف أخلاقية كبيرة يجب معالجتها للتأكد من أن تطويرها ونشرها يتوافق مع القيم الإنسانية ولا يؤدي إلى عواقب سلبية (Brady & Neri, 2020).

أحد التطبيقات الرئيسية للذكاء الاصطناعي (AI) في التعليم هو أنظمة التدريس الذكية، التي تستخدم خوارزميات الذكاء الاصطناعي لتقديم ملاحظات وإرشادات مخصصة للمعلمين أثناء عملهم من خلال المواد التعليمية، تعتمد هذه الأنظمة على نماذج من معرفة الطلاب ومهاراتهم وقدراتهم، وتستخدم البيانات من تفاعلات الطلاب لتكييف المحتوى والتعليقات والتفاعلات التي يقدموها الطلاب في الوقت الفعلي (Brendel et al., 2021). لقد ثبت أن أنظمة التدريس الذكية فعالة في مجموعة متنوعة من المجالات، بما في ذلك الرياضيات والعلوم وتعلم اللغة.

تطبيق آخر للذكاء الاصطناعي (AI) في التعليم هو استخراج البيانات التعليمية، والذي يستخدم خوارزميات الذكاء الاصطناعي لتحليل مجموعات كبيرة من بيانات الطلاب، مثل ملفات السجل والتقييمات، لاكتساب رؤى حول تعلم الطلاب وأدائهم (Zhai et al., 2021). يمكن استخدام هذه المعلومات لدعم التعلم الشخصي والفردى، وتحسين التصميم التعليمي وتحديد مجالات التحسين في النظام التعليمي المستخدم.

يعد الذكاء الاصطناعي مجالا سريع النمو يتمتع بإمكانية إحداث ثورة جديدة من خلال دعمه للكثير من المجالات وتعزيزها بطرق جديدة ومبتكرة (Ashok et al., 2022)، ومع ذلك، من المهم التأكد من أن تطوير ونشر الذكاء الاصطناعي (AI) يسترشد بالمبادئ الأخلاقية والتعليمية، وأن الفوائد المحتملة لهذه التقنيات يتم توزيعها بشكل عادل ومنصف. سيتطلب ذلك بحثا وتطويرا مستمرين، بالإضافة إلى مشاركة وجهات نظر وخبرات متنوعة في تصميم وتنفيذ أنظمة الذكاء الاصطناعي (AI) في كل المجالات التعليمية والصناعية والصحية والاقتصادية (Beil et al., 2019).

تتمثل إحدى القضايا الأخلاقية الأساسية في الذكاء الاصطناعي في تأثيره على الخصوصية وحماية البيانات، يعتمد الذكاء الاصطناعي على كميات هائلة من البيانات لتدريب خوارزمياته وإجراء تنبؤات، مما قد يعرض المعلومات الشخصية

للأفراد للخطر (Krontiris et al., 2020)، أيضا، يمكن لخوارزميات الذكاء الاصطناعي أن تديم التحيز والتمييز الحاليين، مما يؤدي إلى نتائج غير عادلة، لا سيما في مجالات مثل التوظيف والائتمان والعدالة الجنائية (Mazurek & Małagocka, 2019). هناك قضية أخلاقية مهمة أخرى تتمثل في قدرة الذكاء الاصطناعي على استبدال العاملين البشريين وزيادة البطالة مما قد يؤدي إلى تفاقم عدم المساواة الاجتماعية والاقتصادية (Orr & Davis, 2020). يثير نشر الأنظمة المستقلة، مثل الطائرات بدون طيار والسيارات ذاتية القيادة والروبوتات العسكرية، مخاوف بشأن المساءلة والمسؤولية عن أفعالهم، لا سيما في المواقف التي يحدث فيها ضرر (Stahl et al., 2020).

يثير استخدام الذكاء الاصطناعي أسئلة أخلاقية معقدة يجب مراعاتها ومعالجتها بعناية لضمان تطويرها ونشرها بطريقة مسؤولة ومنصفة، وهذا يتطلب التعاون بين أصحاب المصلحة من الأوساط الأكاديمية والصناعة والحكومة، فضلا عن النقاش العام المستمر والتدقيق.

❖ المنهجية

يمكن أن تتضمن منهجية دراسة القضايا الأخلاقية في استخدام الذكاء الاصطناعي (AI) في جميع القطاعات الخطوات التالية:

- مراجعة الأدبيات: إجراء مراجعة شاملة للأدبيات الموجودة حول القضايا الأخلاقية المتعلقة بالذكاء الاصطناعي (AI)، بما في ذلك المقالات الأكاديمية والتقارير ودراسات الحالة؛
 - تطوير الإطار الأخلاقي: تطوير الأطر والمبادئ التوجيهية الأخلاقية للاستخدام المسؤول للذكاء الاصطناعي (AI) في كل قطاع، مع مراعاة التحديات الفريدة التي يطرحها الذكاء الاصطناعي (AI) في كل قطاع؛
 - تقييم الأثر: تقييم تأثير الذكاء الاصطناعي (AI) على المجتمع ومدى إتباع الأطر والمبادئ التوجيهية الأخلاقية؛
 - المراقبة المستمرة: مراقبة وتقييم تأثير الذكاء الاصطناعي (AI) على المجتمع بشكل مستمر وتعديل الأطر والمبادئ التوجيهية الأخلاقية حسب الحاجة.
- بإتباع هذه المنهجية، من الممكن اكتساب فهم شامل للقضايا الأخلاقية التي يثيرها استخدام الذكاء الاصطناعي (AI) في مختلف القطاعات وتطوير حلول فعالة تعزز رفاهية الأفراد والمجتمعات.

2. الخلفية النظرية والأدبيات السابقة

أصبح الذكاء الاصطناعي (AI) بشكل سريع حاضرا في كل مكان في المجتمع وأماكن العمل، ويمكن لتطبيقاته أن تفيد بشكل كبير الأفراد والمنظمات والمجتمعات والشركات والمؤسسات، ومع ذلك، فإن تطوير أنظمة الذكاء الاصطناعي ونشرها يثير أيضا عددا من المخاوف الأخلاقية (Taddeo, 2019).

إحدى القضايا الأخلاقية الأساسية المتعلقة بالذكاء الاصطناعي (AI) هي احتمالية اتخاذ القرار المتحيز، يمكن لأنظمة الذكاء الاصطناعي (AI) أن تستخدم تحيزات منشئها أو مصادر البيانات أو خوارزميات التدريب (Kooli & Al Muftah, 2022) ويمكن أن تستمر هذه التحيزات في تنبؤاتها وتوصياتها. يمكن أن يؤدي هذا إلى نتائج تمييزية، لا سيما في المجالات الحساسة مثل التوظيف والتعليم والأعمال الاقتصادية والعدالة الجنائية والرعاية الصحية (Kooli & Al Muftah, 2022).

تمت مناقشة إمكانية اتخاذ القرار المتحيز في أنظمة الذكاء الاصطناعي (AI) على نطاق واسع في الأدبيات العلمية باعتبارها مصدر قلق أخلاقي رئيسي، يمكن أن ينشأ التحيز في الذكاء الاصطناعي (AI) من مجموعة متنوعة من المصادر، بما في ذلك البيانات المستخدمة لتدريب الخوارزميات والأفراد الذين يصممون الأنظمة وينفذونها (Stahl et al., 2022).

أحد المصادر الشائعة للتحيز في أنظمة الذكاء الاصطناعي (AI) هو البيانات المستخدمة لتدريب الخوارزميات، تتعلم خوارزميات الذكاء الاصطناعي (AI) من البيانات التي يتم تغذيتها، وإذا كانت البيانات تحتوي على تحيزات أو عدم دقة، فيمكن أن تنعكس هذه في التنبؤات والقرارات التي يتخذها نظام الذكاء الاصطناعي (Brady & Neri, 2020). مثلاً، إذا تم تدريب نظام التعرف على الوجه على مجموعة بيانات تتكون في الغالب من أفراد ذوي بشرة فاتحة، فقد لا يعمل بشكل جيد على الأفراد ذوي البشرة الداكنة، أو بيانات تختار الرجال على النساء في الوظائف القيادية وفقاً للمدخلات مما يؤدي إلى نتائج تمييزية حسب الجنس (Siau & Wang, 2020).

مصدر آخر للتحيز في أنظمة الذكاء الاصطناعي (AI) هو الخوارزميات نفسها، قد تحتوي بعض الخوارزميات على تحيزات داخلية تعطي الأولوية لنتائج أو معايير معينة على غيرها، مما يؤدي إلى قرارات متحيزة (Ashok et al., 2020)، على سبيل المثال، قد يحتوي نظام طلب القروض المدعوم بالذكاء الاصطناعي على خوارزمية تعطي الأولوية لمقدمي الطلبات ذوي الدخل المرتفع، مما يؤدي إلى نتائج تمييزية لمقدمي الطلبات من ذوي الدخل المنخفض (Sadok et al., 2022). يمكن للأفراد الذين يصممون وينفذون أنظمة الذكاء الاصطناعي (AI) أيضاً إدخال تحيزات في الأنظمة، على سبيل المثال، إذا قام فريق من المصممين والمهندسين بإنشاء نظام ذكاء اصطناعي (AI)، فقد لا يكونون على دراية بتجارب المجموعات الأخرى ووجهات نظرهم أو يفكرون فيها. يمكن أن يؤدي هذا إلى نتائج متحيزة تعكس تحيزات وخبرات الأفراد الذين أنشأوا النظام (Brady & Neri, 2020).

تعد إمكانية اتخاذ القرار المتحيز في أنظمة الذكاء الاصطناعي (AI) مصدر قلق أخلاقي رئيسي للمستخدمين، وتتطلب اهتماماً دقيقاً ومراقبة مستمرة لضمان أن هذه الأنظمة عادلة وشفافة وخاضعة للمساءلة (Orr & Davis, 2020). يمكن تحقيق ذلك من خلال استخدام مجموعات بيانات تدريبية متنوعة، وتطوير خوارزميات غير متحيزة، وإشراك وجهات نظر وخبرات متنوعة في تصميم وتنفيذ أنظمة الذكاء الاصطناعي (Akgun & Greenhow, 2021).

مصدر قلق أخلاقي رئيسي آخر هو قدرة أنظمة الذكاء الاصطناعي (AI) على زيادة عدم المساواة الاقتصادية والاجتماعية يتمتع الذكاء الاصطناعي (AI) بالقدرة على أتمتة العديد من الوظائف والمهام، مما قد يترك بعض الأفراد بلا عمل (Habli et al., 2020). في الوقت نفسه، من المرجح أيضاً أن يؤدي الذكاء الاصطناعي (AI) إلى تفاقم أوجه عدم المساواة الحالية، حيث من المرجح أن يتمتع أولئك الذين لديهم الموارد اللازمة لاعتماد التكنولوجيا وخوارزميات الذكاء الاصطناعي والاستفادة منها بميزة كبيرة على أولئك الذين ليس لديهم مثل هذه الموارد (Kerr et al., 2020).

تمت مناقشة قدرة أنظمة الذكاء الاصطناعي (AI) على زيادة عدم المساواة الاقتصادية والاجتماعية على نطاق واسع في الأدبيات العلمية باعتبارها مصدر قلق أخلاقي رئيسي (Schönberger, 2019). يتمتع الذكاء الاصطناعي (AI) بالقدرة على أتمتة العديد من الوظائف والمهام، مما قد يؤدي إلى خسائر كبيرة في الوظائف واضطراب اقتصادي. في الوقت نفسه، من المرجح أن تتركز فوائد الذكاء الاصطناعي (AI) بين مجموعة صغيرة من الأفراد والمنظمات، مما يؤدي إلى تفاقم التفاوتات القائمة وإنشاء أخرى جديدة (Stahl, 2021).

تتمثل إحدى الطرق الرئيسية التي يمكن للذكاء الاصطناعي (AI) من خلالها زيادة عدم المساواة الاقتصادية في أتمتة الوظائف والمهام التي يؤديها البشر حالياً. وهذا يمكن أن يؤدي إلى فقدان الوظائف على نطاق واسع، لا سيما بين العاملين في الوظائف ذات المهارات المنخفضة والأجور المنخفضة وساعات العمل الطويلة (Habli et al., 2020). في الوقت نفسه، من المرجح أن يكون الأفراد والمنظمات والشركات والمؤسسات الأكثر قدرة على تبني الذكاء الاصطناعي (AI) والاستفادة منه هم أولئك الذين لديهم الموارد والخبرة اللازمة للقيام بذلك، مثل الشركات الكبيرة والأثرياء (Kumar et al., 2016).

هناك طريقة أخرى يمكن للذكاء الاصطناعي (AI) من خلالها زيادة عدم المساواة الاجتماعية وهي إدانة وتضخيم التحيزات القائمة والممارسات التمييزية (Brady & Neri, 2020).

على سبيل المثال، يمكن أن تترث أنظمة الذكاء الاصطناعي (AI) تحيزات منشئها ومصادر بياناتها، مما يؤدي إلى نتائج تمييزية في مجالات مثل التوظيف والعدالة الجنائية والرعاية الصحية (Schönberger, 2019)، بالإضافة إلى ذلك، فإن طبيعة الملكية للعديد من خوارزميات الذكاء الاصطناعي (AI) تجعل من الصعب تقييم أو التحقق من دقتها أو نزاهتها أو عدالتها، مما يسمح لهذه التحيزات بالاستمرار والتضخم والتطور (Akgun & Greenhow, 2021).

فإن قدرة أنظمة الذكاء الاصطناعي (AI) على زيادة عدم المساواة الاقتصادية والاجتماعية هي مصدر قلق أخلاقي رئيسي، وتتطلب اهتماما دقيقا ومراقبة مستمرة لضمان نشر هذه الأنظمة بطريقة تعود بالنفع على المجتمع ككل، يمكن تحقيق ذلك من خلال تطوير خوارزميات عادلة وشفافة وغير متحيزة (Safdar et al., 2020)، واستخدام مجموعات بيانات تدريبية متنوعة، وإشراك وجهات نظر وخبرات متنوعة في تصميم وتجهيز وتنفيذ أنظمة الذكاء الاصطناعي (Brendel et al., 2021).

بالإضافة إلى ذلك، سيكون من المهم معالجة الأسباب الجذرية لعدم المساواة وضمان توزيع فوائد الذكاء الاصطناعي بشكل منصف (Hagerty & Rubinov, 2019).

يعد الافتقار إلى الشفافية والمساءلة في أنظمة الذكاء الاصطناعي (AI) قضية أخلاقية رئيسية، يمكن لأنظمة الذكاء الاصطناعي (AI) اتخاذ قرارات بناء على خوارزميات معقدة يصعب على الأشخاص فهمها أو تحديها، مما يجعل من الصعب محاسبة أو اتخاذ قرارات ضد هذه الأنظمة على أفعالها وقراراتها (Santoni & Mecacci, 2021). بالإضافة إلى ذلك، فإن طبيعة الملكية للعديد من خوارزميات الذكاء الاصطناعي (AI) تجعل من الصعب تقييم أو التحقق من دقتها أو عدالتها ونزاهتها (Morgan et al., 2020).

يعد الافتقار إلى الشفافية والمساءلة في أنظمة الذكاء الاصطناعي (AI) قضية أخلاقية رئيسية تمت مناقشتها على نطاق واسع في الأدبيات العلمية (Habli et al., 2020). غالبا ما تكون أنظمة الذكاء الاصطناعي (AI) معقدة وذات ملكية خاصة، مما يجعل من الصعب على الأفراد والمؤسسات فهم كيفية اتخاذهم للقرارات ومحاسبهم على تلك القرارات (Taddeo, 2019)، يمكن أن يؤدي ذلك إلى انعدام الثقة في هذه الأنظمة ويمكن أن يكون له عواقب وخيمة، لا سيما في مجالات مثل الرعاية الصحية والعدالة الجنائية والتوظيف (Siau & Wang, 2020).

يتمثل أحد الجوانب الرئيسية لانعدام الشفافية في أنظمة الذكاء الاصطناعي (AI) في صعوبة فهم كيفية اتخاذ القرارات. تعتمد العديد من خوارزميات الذكاء الاصطناعي (AI) على نماذج معقدة ويتم تدريبها على مجموعات بيانات كبيرة ومختلفة، مما يجعل من الصعب على البشر المستخدمين فهم طريقة عملهم الداخلية (Orr & Davis, 2020)، هذا النقص في الشفافية يجعل من الصعب على الأفراد والمؤسسات تقييم دقة وعدالة أنظمة الذكاء الاصطناعي (AI)، والتأكد من توافقها مع المعايير الأخلاقية والقانونية (Stahl, 2021).

جانب آخر من جوانب نقص المساءلة في أنظمة الذكاء الاصطناعي (AI) هو صعوبة إسناد المسؤولية عن النتائج التي تنتجها. على سبيل المثال، إذا اتخذ نظام ذكاء اصطناعي (AI) قرارا يؤدي إلى ضرر، فقد لا يكون من الواضح من الذي يجب أن يتحمل المسؤولية عن هذا الضرر نظرا لطبيعة أنظمة الذكاء الاصطناعي (AI)، يمكن أن يخلق ذلك موقفا يتكرر فيه الأفراد والمنظمات في تبني واستخدام أنظمة الذكاء الاصطناعي (AI)، لأنهم يخشون أن يتم محاسبهم على النتائج التي تصدر عن الأنظمة (Max et al., 2021).

يعد الافتقار إلى الشفافية والمساءلة في أنظمة الذكاء الاصطناعي (AI) قضية أخلاقية رئيسية تتطلب اهتماما دقيقا ومراقبة مستمرة لضمان أن هذه الأنظمة عادلة وشفافة وخاضعة للمحاسبة والمساءلة (Ashok et al., 2022)، يمكن تحقيق ذلك من خلال تطوير خوارزميات مفتوحة وشفافة، واستخدام مجموعات بيانات تدريبية متنوعة، وإشراك وجهات نظر وخبرات متنوعة في تصميم وتنفيذ أنظمة الذكاء الاصطناعي (Habli et al., 2020)، بالإضافة إلى ذلك، سيكون من المهم وضع معايير أخلاقية وقانونية واضحة للذكاء الاصطناعي (AI)، ولضمان مساءلة الأفراد والمنظمات عن النتائج الصادرة (Kooli & Al Muftah, 2022).

هناك أيضا مخاوف أخلاقية أوسع تتعلق بالاعتماد المتزايد على أنظمة الذكاء الاصطناعي (AI) والعواقب التي قد تترتب على ذلك على استقلالية الإنسان وفاعليته وكرامته. نظرا لأن أنظمة الذكاء الاصطناعي (AI) أصبحت أكثر تعقيدا وقدرة فهناك خطر يتمثل في أن البشر قد يعتمدون على هذه الأنظمة بشكل كبير، مما قد يؤدي إلى فقدان السيطرة اتخاذ القرار (Akgun & Greenhow, 2021).

فإن القضايا الأخلاقية المتعلقة باستخدام الذكاء الاصطناعي (AI) معقدة ومتعددة الأوجه، وتتطلب دراسة متأنية ومناقشة مستمرة. يعد ضمان التطوير المسؤول لأنظمة الذكاء الاصطناعي (AI) ونشرها أمرا بالغ الأهمية لتحقيق الفوائد المحتملة لهذه التكنولوجيا مع تقليل أضرارها المحتملة.

❖ دمج القيم الإنسانية في تصميم واستخدام أنظمة الذكاء الاصطناعي (AI)

مع استمرار الذكاء الاصطناعي (AI) في لعب دور أكثر أهمية في حياتنا اليومية، يصبح من المهم بشكل متزايد التأكد من أن هذه الأنظمة تتوافق مع القيم والأخلاق الإنسانية، وهذا يشمل اعتبارات مثل الخصوصية والإنصاف والشفافية والمساءلة والتأثير العام للذكاء الاصطناعي (AI) على المجتمع، من خلال دمج القيم الإنسانية في تصميم واستخدام الذكاء الاصطناعي (Safdar et al., 2020)، يمكننا المساعدة في ضمان أن هذه الأنظمة تخدم المصالح الفضلى للبشرية وتتجنب الضرر المحتمل. يتطلب نهجا متعدد التخصصات يضم خبراء من مجالات علوم الكمبيوتر والأخلاق والقانون لتطوير إرشادات وأطر أخلاقية لأنظمة الذكاء الاصطناعي (AI) التي تراعي حقوق الإنسان والقيم الإنسانية والأخلاقية (Santoni & Mecacci, 2021)، مع وضع هذه الاعتبارات في الاعتبار، يمكننا بناء أنظمة ذكاء اصطناعي (AI) تكون مسؤولة وجديرة بالثقة ومفيدة للمجتمع في نهاية المطاف.

يشير دمج القيم الإنسانية في تصميم واستخدام أنظمة الذكاء الاصطناعي (AI) إلى عملية تصميم واستخدام أنظمة الذكاء الاصطناعي بطريقة تتماشى مع القيم الإنسانية، مثل الكرامة والاحترام والإنصاف والخصوصية والاستقلالية (Brady & Neri, 2020).

❖ لدمج القيم الإنسانية في أنظمة الذكاء الاصطناعي (AI)، يمكن للمؤسسات إتباع عدة خطوات:

- تحديد وتعريف القيم الإنسانية ذات الصلة التي يجب دمجها في تصميم واستخدام أنظمة الذكاء الاصطناعي (Santoni & Mecacci, 2021)؛
- تطوير وتنفيذ المبادئ والمبادئ التوجيهية الأخلاقية التي تتماشى مع القيم الإنسانية وتقديم إرشادات لتصميم واستخدام أنظمة الذكاء الاصطناعي (Alqudah & Muradkhanli, 2021)؛
- دمج القيم الإنسانية في عملية تصميم وتطوير أنظمة الذكاء الاصطناعي (AI)، بحيث يتم بناؤها مع وضع هذه القيم في الاعتبار (Ashok et al., 2022)؛

- تقييم واختبار أنظمة الذكاء الاصطناعي (AI) للتأكد من أنها تتماشى مع القيم الإنسانية والمبادئ الأخلاقية، واتخاذ الإجراءات المناسبة لمعالجة أي قضايا يتم تحديدها (Schönberger, 2019)؛
- المراقبة والمراجعة المنتظمة لاستخدام أنظمة الذكاء الاصطناعي (AI) للتأكد من استمرارها في التوافق مع القيم الإنسانية والمبادئ الأخلاقية، وإجراء أي تغييرات ضرورية لتحسين تصميمها واستخدامها (Safdar et al., 2020)؛
- من خلال دمج القيم الإنسانية في تصميم واستخدام أنظمة الذكاء الاصطناعي (AI)، يمكن للمنظمات المساعدة في ضمان توافق أنظمة الذكاء الاصطناعي (AI) مع الأخلاقيات والقيم البشرية، وأنها تعزز الصالح العام، بالإضافة إلى ذلك، من خلال تعزيز الكرامة والاحترام والإنصاف والخصوصية والاستقلالية، يمكن للمنظمات المساعدة في بناء الثقة والتأكد من استخدام أنظمة الذكاء الاصطناعي (AI) بطريقة مسؤولة وأخلاقية.
- ❖ هناك العديد من القضايا الأخلاقية التي يجب مراعاتها عند استخدام الذكاء الاصطناعي (AI):
- التحيز والإنصاف: يمكن لخوارزميات الذكاء الاصطناعي (AI) إدانة وتضخيم التحيزات الحالية إذا تم تدريبها على البيانات المتحيزة (Kooli & Al Muftah, 2022)؛
- الخصوصية: يمكن لأنظمة الذكاء الاصطناعي (AI) جمع كميات كبيرة من البيانات الشخصية وتخزينها ومعالجتها، مما يؤدي إلى مخاوف تتعلق بالخصوصية (Ashok et al., 2022)؛
- المسؤولية والمساءلة: قد يكون من الصعب تحديد المسؤول عن أفعال نظام الذكاء الاصطناعي (AI)، لا سيما في حالات الضرر أو الخطأ (Max et al., 2021)؛
- تقليص الوظائف: يمكن لأنظمة الذكاء الاصطناعي (AI) أتمتة المهام التي كان يؤديها البشر سابقا، مما يؤدي إلى فقدان الوظائف (Safdar et al., 2020)؛
- الاستقلالية والتحكم: نظرا لأن أنظمة الذكاء الاصطناعي (AI) أصبحت أكثر تقدما، فهناك خطر من أنها يمكن أن تتخذ قرارات دون تدخل بشري ومن المحتمل أن تسبب ضررا أو خلل أو مشاكل (Brady & Neri, 2020).
- لمعالجة هذه القضايا، يجب على المنظمات النظر في تنفيذ المبادئ الأخلاقية والمبادئ التوجيهية لتطوير واستخدام أنظمة الذكاء الاصطناعي (AI). يمكن أن يشمل ذلك عمليات تدقيق منتظمة لضمان خلو أنظمة الذكاء الاصطناعي من التحيز ولديها ضوابط مناسبة للخصوصية والمساءلة.

❖ قضايا الخصوصية والتحيز والمساءلة واستبدال الوظائف في الذكاء الاصطناعي (AI)

- مع استمرار تطور الذكاء الاصطناعي (AI) وتغلغل في مختلف الصناعات، هناك عدد من القضايا الملحة التي تحتاج إلى معالجة. الخصوصية هي مصدر قلق كبير لأن أنظمة الذكاء الاصطناعي (AI) تجمع وتعالج وتخزن كميات هائلة من البيانات الشخصية، مما قد يؤدي إلى انتهاكات حقوق الخصوصية. يعد التحيز مشكلة حيث يمكن لأنظمة الذكاء الاصطناعي (AI) أن تديم التحيزات الموجودة في المجتمع بل وتضخمها إذا لم يتم تصميمها وتدريبها بشكل صحيح (Ashok et al., 2022).
- يمكن أن يؤدي هذا إلى نتائج تمييزية وإلحاق الضرر بالمجتمعات المهمشة. تعد المساءلة أيضا قضية مهمة، حيث قد يكون من الصعب تحديد المسؤول عندما تسبب أنظمة الذكاء الاصطناعي (AI) ضررا أو ترتكب أخطاء (Kooli & Al Muftah, 2022).

أخيرا، هناك مشكلة الاستغناء عن الوظائف، حيث يتم نشر أنظمة الذكاء الاصطناعي (AI) بشكل متزايد في القوى العاملة، مما قد يؤدي إلى فقدان الوظائف والحاجة إلى إعادة تدريب العمال وصقل مهاراتهم. تعتبر معالجة هذه القضايا

أمرا بالغ الأهمية للتطوير المسؤول للذكاء الاصطناعي (AI) ونشره ولضمان أن هذه الأنظمة تخدم المصالح الفضلى للمجتمع (Morgan et al., 2020).

❖ قضايا الخصوصية والتحيز والمساءلة واستبدال الوظائف في الذكاء الاصطناعي (AI) في هذه النقاط:

– الخصوصية: يمكن لأنظمة الذكاء الاصطناعي (AI) جمع ومعالجة وتخزين كميات كبيرة من البيانات الشخصية، مما قد يؤدي إلى مخاوف تتعلق بالخصوصية في حالة إساءة استخدام البيانات أو مشاركتها دون موافقة مناسبة، (Brady & Neri, 2020):

– التحيز: يمكن لخوارزميات الذكاء الاصطناعي (AI) إدامة وتضخيم التحيزات الموجودة إذا تم تدريبها على البيانات المتحيزة، هذا يمكن أن يؤدي إلى نتائج غير عادلة وغير منصفة وتحيزة لجنس أو عرق أو دين أو لون (Akgun & Greenhow, 2021):

– المساءلة: يمكن لأنظمة الذكاء الاصطناعي (AI) اتخاذ القرارات واتخاذ الإجراءات التي لها عواقب وخيمة، ولكن قد يكون من الصعب تحديد المسؤول عندما تسوء الأمور (Habli et al., 2020):

– إزاحة الوظائف: يمكن لأنظمة الذكاء الاصطناعي (AI) أتمتة المهام التي كان يقوم بها البشر في السابق، مما يؤدي إلى فقدان الوظائف والاضطراب الاقتصادي. يمكن أن يؤدي هذا أيضا إلى اضطرابات اجتماعية وسياسية (Alqudah & Muradkhanli, 2021):

لمعالجة هذه القضايا، من المهم وجود مبادئ وإرشادات أخلاقية واضحة لتطوير واستخدام أنظمة الذكاء الاصطناعي (AI)، بالإضافة إلى عمليات تدقيق منتظمة للتأكد من أن الأنظمة تعمل على النحو المنشود وخالية من التحيز بالإضافة إلى ذلك، يجب على المنظمات النظر في العواقب المحتملة لأنظمة الذكاء الاصطناعي (AI) الخاصة بها وتنفيذ تدابير لتخفيف الضرر وضمان المساءلة.

❖ المبادئ التوجيهية الأخلاقية التي يجب العمل بها لتطوير واستخدام أنظمة الذكاء الاصطناعي (AI):

لا توجد مجموعة واحدة من المبادئ والمبادئ التوجيهية الأخلاقية لتطوير واستخدام أنظمة الذكاء الاصطناعي (AI) ولكن بعض تلك التي يتم الاستشهاد بها بشكل شائع تشمل، هذه المبادئ والمبادئ التوجيهية الأخلاقية التي اقترحت لتطوير وتحديث استخدام أنظمة الذكاء الاصطناعي، بما في ذلك:

– المسؤولية والمساءلة: يجب تصميم أنظمة الذكاء الاصطناعي (AI) واستخدامها بطريقة تأخذ في الاعتبار تأثيرها المحتمل على المجتمع والأفراد، ويجب محاسبة المسؤولين عن تطويرها ونشرها عن أي ضرر تسببه، ويجب أن تكون هناك مسؤوليات واضحة والمساءلة عن الإجراءات والقرارات التي تتخذها أنظمة الذكاء الاصطناعي (AI)، لا سيما في حالات الضرر أو الخطأ أو عدم اليقين (Schönberger, 2019).

– الإنصاف وعدم التمييز: يجب تصميم أنظمة الذكاء الاصطناعي (AI) واستخدامها بطريقة عادلة ولا تميز ضد الأفراد أو الجماعات على أساس عوامل مثل العرق أو الجنس أو الدين أو الاتجاهات السياسية (Chen et al., 2022).

– الذكاء الاصطناعي (AI) الشفاف والقابل للتفسير: يجب أن تكون أنظمة الذكاء الاصطناعي (AI) شفافة في كيفية عملها واتخاذ القرارات، ويجب أن يكون الأفراد قادرين على فهم الأسباب الكامنة وراء القرارات التي تتخذها أنظمة الذكاء الاصطناعي (AI) التي تؤثر عليهم، ويجب أن تكون أنظمة الذكاء الاصطناعي (AI) شفافة في عمليات صنع القرار الخاصة بهم، السماح للمستخدمين بفهم كيفية اتخاذ القرارات (Schönberger, 2019).

- الخصوصية: يجب حماية البيانات الشخصية التي تجمعها أنظمة الذكاء الاصطناعي (AI) واستخدامها وفقا لقوانين ولوائح الخصوصية (Brady & Neri, 2020).
 - السلامة والأمن: يجب تصميم أنظمة الذكاء الاصطناعي (AI) وتنفيذها بطريقة تضمن السلامة والأمن والحماية من الأذى والاستخدام الضار، ويجب أن تحترم أنظمة الذكاء الاصطناعي خصوصية الأفراد وألا تجمع المعلومات الشخصية أو تخزينها أو استخدامها بطرق غير أخلاقية أو ضارة (Orr & Davis, 2020).
 - احترام القيم الإنسانية: يجب تصميم أنظمة الذكاء الاصطناعي (AI) واستخدامها بطريقة تتماشى مع القيم الإنسانية مثل الكرامة والاحترام والاستقلالية (Larsson & Heintz, 2020).
 - تقليل التحيز: يجب تصميم أنظمة الذكاء الاصطناعي (AI) واستخدامها بطريقة تقلل من التحيز وتعزز اتخاذ القرار المحايد.
 - التحكم البشري: يجب تصميم أنظمة الذكاء الاصطناعي (AI) واستخدامها بطريقة تسمح بالتحكم البشري والتدخل لضمان توافقها مع القيم الإنسانية والمفاهيم والمبادئ الأخلاقية للمستخدمين (Kooli & Al Muftah, 2022).
- لم يتم الاتفاق على هذه المبادئ التوجيهية الأخلاقية عالميا ولا تزال تتطور مع استمرار تقدم تقنية الذكاء الاصطناعي ومع ذلك، فهي بمثابة نقطة انطلاق للمناقشات الأخلاقية وصنع القرار فيما يتعلق بتطوير واستخدام أنظمة الذكاء الاصطناعي (Max et al., 2021).

هذه المبادئ ليست شاملة وقد يكون للمنظمات المختلفة أولويات وتفسيرات مختلفة لما يشكل الذكاء الاصطناعي الأخلاقي (Kooli & Al Muftah, 2022)، ومع ذلك، يمكن أن توفر هذه المبادئ نقطة انطلاق للمنظمات والشركات والمؤسسات التي تتطلع إلى تطوير مبادئ توجيهية أخلاقية لأنظمة الذكاء الاصطناعي (AI) الخاصة بها.

❖ الشفافية:

الشفافية: يجب أن تكون الشفافية شفافة في عمليات صنع القرار، مما يسمح للمستخدمين بفهم كيفية اتخاذ القرارات (Walmsley, 2021). الشفافية جانب مهم في أي منظمة، لا سيما عندما يتعلق الأمر بعمليات صنع القرار، يشير إلى فعل الكشف بصراحة وصدق عن المعلومات حول إجراءات وسياسات وممارسات المنظمة أو المؤسسة (Larsson & Heintz, 2020). في هذا السياق، تعني الشفافية أن عمليات صنع القرار في منظمة ما يجب أن تكون سهلة الفهم ويمكن الوصول إليها من قبل المستخدمين، مما يسمح لهم بمعرفة كيفية اتخاذ القرارات ومساءلة المنظمة، هذا النوع من الانفتاح يعزز الثقة والمساءلة والإنصاف، ويمكن أن يؤدي إلى نتائج أفضل لجميع أصحاب المصلحة المعنيين (Kumar et al., 2016).

الشفافية في عمليات صنع القرار في مجال الذكاء الاصطناعي (AI) هي قضية حاسمة حيث يتم استخدام أنظمة الذكاء الاصطناعي (AI) في عدد متزايد من التطبيقات التي تؤثر على حياة الناس، مثل العدالة والرعاية الصحية والخدمات المالية (Max et al., 2021)، غالبا ما تتخذ أنظمة الذكاء الاصطناعي (AI) قرارات لها عواقب وخيمة على الأفراد، ومن المهم التأكد من أن هذه القرارات التي تتخذ يتم اتخاذها بشكل عادل ودقيق وشفاف (Safdar et al., 2020).

ومع ذلك، تعمل العديد من أنظمة الذكاء الاصطناعي (AI) بطريقة "الصندوق الأسود"، حيث يصعب فهم كيفية اتخاذ القرارات، هذا النقص في الشفافية يمكن أن يثير مخاوف بشأن التحيز والمساءلة، لأن المستخدمين غير قادرين على رؤية كيفية اتخاذ القرارات ومن المسؤول عنها. يمكن أن يؤدي أيضا إلى عدم الثقة في التكنولوجيا، مما قد يحد من اعتمادها على نطاق واسع (Schönberger, 2019).

لمعالجة هذه المخاوف، من المهم أن تكون المنظمات شفافة في عمليات صنع القرار الخاصة بالذكاء الاصطناعي (Chen et al., 2022). يتضمن ذلك توفير معلومات حول كيفية اتخاذ القرارات، والبيانات التي يتم استخدامها، وكيفية معالجة التحيزات المحتملة، يمكن للمستخدمين استخدام هذه المعلومات لفهم كيفية وصول نظام الذكاء الاصطناعي (AI) إلى قراراته ولإخضاع المؤسسة أو المنظمة للمساءلة عن النتائج (Siau & Wang, 2020).

بالإضافة إلى ذلك، يمكن للشفافية في صنع القرار بالذكاء الاصطناعي (AI) تحسين دقة وعدالة القرارات التي يتم اتخاذها. على سبيل المثال، من خلال جعل عملية صنع القرار شفافة، يمكن لأصحاب المصلحة تقديم التغذية الراجعة والمساعدة في تحديد التحيزات المحتملة أو النقاط العمياء المخزنة في النظام (Robinson, 2020). يمكن أن يؤدي هذا النوع من المشاركة إلى اتخاذ قرارات أفضل وزيادة الدعم للقرارات التي يتم اتخاذها (Walmsley, 2021).

تعد الشفافية في اتخاذ القرار بشأن الذكاء الاصطناعي (AI) أمراً ضرورياً لبناء الثقة والمساءلة والإنصاف في استخدام الذكاء الاصطناعي (McKinney et al., 2022)، من خلال جعل عملية صنع القرار شفافة، يمكن للمنظمات ضمان استخدام الذكاء الاصطناعي (AI) بطريقة مسؤولة وأخلاقية، ويمكن أن تساعد في بناء ثقة الجمهور في التكنولوجيا (Felzmann et al., 2020).

تشير الشفافية في صنع القرار للذكاء الاصطناعي (AI) إلى فكرة أن أنظمة الذكاء الاصطناعي (AI) يجب أن تكون منفتحة وشفافة بشأن كيفية توصيلها إلى قراراتها، وهذا يشمل جعل الخوارزميات والبيانات المستخدمة من قبل الأنظمة في متناول المستخدمين ومفهومة لهم (Morgan et al., 2020).

من خلال توفير الشفافية، يمكن للمستخدمين فهم كيفية اتخاذ القرارات بشكل أفضل ولماذا يتم اتخاذها بطريقة معينة. يمكن أن يساعد ذلك في بناء الثقة في النظام، وزيادة المساءلة، ومنع العواقب الضارة أو غير المقصودة (Robinson, 2020).

لتحقيق الشفافية يمكن للمنظمات تنفيذ العديد من الممارسات، مثل:

- تقديم تفسيرات واضحة لكيفية وصول نظام الذكاء الاصطناعي (AI) إلى قراراته، مثل استخدام تقنيات الذكاء الاصطناعي القابلة للتفسير (Felzmann et al., 2020)؛
- جعل الخوارزميات والبيانات المستخدمة من قبل نظام الذكاء الاصطناعي متاحة للجمهور وشفافة، بحيث يمكن للمستخدمين فهم كيفية اتخاذ النظام للقرارات (Robinson, 2020)؛
- المراجعة المنتظمة لنظام الذكاء الاصطناعي (AI) للتأكد من أنه يعمل على النحو المنشود وخال من التحيز (Larsson & Heintz, 2020)؛
- تسهيل وصول المستخدمين إلى معلومات حول كيفية عمل نظام الذكاء الاصطناعي (AI) وكيفية اتخاذ القرارات، من خلال التوثيق الواضح أو الواجهات سهلة الاستخدام أو غيرها من الوسائل؛
- المشاركة مع أصحاب المصلحة، بما في ذلك المستخدمين والخبراء والجمهور والمختصين، لجمع الملاحظات ودمج وجهات نظرهم المختلفة في تطوير وتحديث واستخدام أنظمة الذكاء الاصطناعي (Larsson & Heintz, 2020).
- من خلال إعطاء الأولوية للشفافية في عملية صنع القرار في مجال الذكاء الاصطناعي (AI)، يمكن للمؤسسات المساعدة في بناء الثقة وزيادة المساءلة ومنع النتائج الضارة.

❖ المسؤولية والمساءلة:

يجب أن تكون المسؤولية والمساءلة واضحين بالنسبة للإجراءات والقرارات التي تتخذها أنظمة الذكاء الاصطناعي (AI) لا سيما في حالات الضرر أو الخطأ. نظراً لأن الذكاء الاصطناعي (AI) أصبح أكثر اندماجاً في حياتنا اليومية، فمن الأهمية بمكان إنشاء خطوط واضحة للمسؤولية والمساءلة عن الإجراءات والقرارات التي تتخذها هذه الأنظمة (Stahl, 2021). هذا مهم بشكل خاص في الحالات التي تسبب فيها أنظمة الذكاء الاصطناعي (AI) ضرراً أو ترتكب أخطاء، حيث يجب أن تكون هناك آلية واضحة لتحديد المسؤول والمساءلة (Orr & Davis, 2020). وهذا يتطلب دراسة متأنية للأثار القانونية والأخلاقية للذكاء الاصطناعي (AI)، فضلاً عن تطوير الأنظمة والعمليات التي تضمن المساءلة والشفافية. من خلال تحديد المسؤولية والمساءلة الواضحة، يمكننا المساعدة في بناء الثقة في أنظمة الذكاء الاصطناعي (AI) والتأكد من استخدامها بطرق مسؤولة وأخلاقية تخدم المصالح الفضلى للمجتمع (Siau & Wang, 2020).

تشير المسؤولية والمساءلة إلى فكرة أن أولئك الذين يصممون ويطورون ويستخدمون أنظمة الذكاء الاصطناعي (AI) يجب أن يكونوا مسؤولين عن الإجراءات والقرارات التي تتخذها هذه الأنظمة. هذا مهم بشكل خاص في الحالات التي يحدث فيها ضرر أو خطأ، حيث قد يكون من غير الواضح من المسؤول عن هذه النتائج (Loi & Spielkamp, 2021). لضمان المسؤولية والمساءلة الواضحة لأنظمة الذكاء الاصطناعي (AI)، يمكن للمنظمات تنفيذ العديد من الممارسات، مثل:

- تحديد الأدوار والمسؤوليات الواضحة للمشاركين في تطوير واستخدام أنظمة الذكاء الاصطناعي (AI)، بما في ذلك المصممون والمطورون والمستخدمون (Orr & Davis, 2020).
 - إنشاء سلسلة واضحة للمساءلة، بحيث يسهل تحديد المسؤول في حالة حدوث ضرر أو خطأ؛
 - المراجعة المنتظمة لنظام الذكاء الاصطناعي (AI) للتأكد من أنه يعمل على النحو المنشود وخالٍ من التحيز (Orr & Davis, 2020)؛
 - تنفيذ تدابير لضمان تشغيل نظام الذكاء الاصطناعي (AI) بطريقة آمنة ومأمونة، مع وجود ضوابط مناسبة لمنع الضرر (Dignum, 2020)؛
 - وضع إجراءات واضحة للإبلاغ عن حوادث الأذى أو الخطأ والتحقيق فيها، واتخاذ الإجراءات المناسبة لمنع الحوادث المستقبلية (Habli et al., 2020).
- من خلال ضمان المسؤولية والمساءلة الواضحة لأنظمة الذكاء الاصطناعي (AI)، يمكن للمؤسسات المساعدة في بناء الثقة وزيادة الشفافية ومنع النتائج الضارة (Stahl, 2021). بالإضافة إلى ذلك، يمكن أن تساعد المساءلة الواضحة في ضمان تحفيز المسؤولين عن نظام الذكاء الاصطناعي (AI) لإجراء تحسينات ومعالجة أي مشكلات تنشأ.

تصميم عدم التمييز والإنصاف لاستخدامه بطريقة تتجنب التمييز وتعزز العدالة:

عند تطوير أنظمة الذكاء الاصطناعي (AI) ونشرها، من الأهمية بمكان إعطاء الأولوية لعدم التمييز والإنصاف لضمان استخدام هذه الأنظمة بطريقة تتجنب التمييز وتعزز العدالة لجميع الأفراد والجماعات. يتطلب ذلك النظر في مجموعة متنوعة من العوامل، مثل البيانات المستخدمة لتدريب أنظمة الذكاء الاصطناعي (Kooli & Al Muftah, 2022) والخوارزميات المستخدمة في اتخاذ القرارات، والتأثير العام للذكاء الاصطناعي (AI) على المجتمعات المختلفة، يتطلب ضمان عدم التمييز والإنصاف في الذكاء الاصطناعي (AI) نهجاً متعدد التخصصات يجمع خبراء من علوم الكمبيوتر والأخلاق والقانون لتطوير

المبادئ التوجيهية والأطر الأخلاقية لأنظمة الذكاء الاصطناعي. من خلال إعطاء الأولوية لعدم التمييز والإنصاف، يمكننا المساعدة في ضمان استخدام أنظمة الذكاء الاصطناعي (AI) لتعزيز العدالة الاجتماعية وخدمة المصالح الفضلى لجميع أفراد المجتمع (Orr & Davis, 2020).

يشير عدم التمييز والإنصاف في الذكاء الاصطناعي إلى فكرة وجوب تصميم أنظمة الذكاء الاصطناعي (AI) واستخدامها بطريقة تتجنب التمييز وتعزز الإنصاف لجميع المستخدمين، بغض النظر عن العرق أو الجنس أو الدين أو أي خصائص محمية أخرى (Alqudah & Muradkhanli, 2021)، هذا مهم للتأكد من أن أنظمة الذكاء الاصطناعي (AI) لا تديم أو تزيد من التحيزات والتمييز في المجتمع.

لتعزيز عدم التمييز والإنصاف في الذكاء الاصطناعي (AI)، يمكن للمنظمات تنفيذ العديد من الممارسات، مثل:

- التأكد من أن البيانات المستخدمة لتدريب أنظمة الذكاء الاصطناعي (AI) متنوعة وتمثيلية، بحيث لا يديم نظام الذكاء الاصطناعي (AI) التحيزات الموجودة (Alqudah & Muradkhanli, 2021)؛
 - المراجعة المنتظمة لنظام الذكاء الاصطناعي (AI) لتحديد ومعالجة أي حالات تمييز أو تحيز؛
 - تنفيذ تدابير لضمان عدم تمييز نظام الذكاء الاصطناعي (AI) على أساس الخصائص المحمية، مثل العرق أو الجنس أو الدين؛
 - المشاركة مع أصحاب المصلحة، بما في ذلك المستخدمون والخبراء والجمهور، لجمع التعليقات ودمج وجهات نظرهم في تطوير واستخدام أنظمة الذكاء الاصطناعي (Schönberger, 2019)؛
 - وضع سياسات وإجراءات واضحة لمعالجة حالات التمييز أو التحيز في نظام الذكاء الاصطناعي (AI)، واتخاذ الإجراءات المناسبة لمنع الحوادث المستقبلية.
- من خلال تعزيز عدم التمييز والإنصاف في الذكاء الاصطناعي (AI)، يمكن للمنظمات المساعدة في بناء الثقة وزيادة الشفافية ومنع النتائج الضارة. بالإضافة إلى ذلك، يمكن لأنظمة الذكاء الاصطناعي (AI) العادلة وغير التمييزية أن تساعد في إنشاء مجتمع أكثر إنصافاً وعدالة، من خلال ضمان حصول جميع المستخدمين على مزايا التكنولوجيا على قدم المساواة.
- ❖ الخصوصية:

يجب حماية البيانات الشخصية التي تجمعها أنظمة الذكاء الاصطناعي (AI) واستخدامها وفقاً لقوانين ولوائح الخصوصية.

نظراً لأن أنظمة الذكاء الاصطناعي (AI) تجمع كميات هائلة من البيانات الشخصية وتعالجها وتخزنها، فمن الضروري ضمان حماية هذه البيانات واستخدامها وفقاً لقوانين ولوائح الخصوصية. يتضمن ذلك اعتبارات مثل التخزين الآمن للبيانات والشفافية في استخدام البيانات الشخصية والحصول على موافقة من الأفراد قبل استخدام بياناتهم (Habli et al., 2020). يؤدي عدم حماية البيانات الشخصية إلى عواقب وخيمة، بما في ذلك انتهاكات حقوق الخصوصية وإلحاق الضرر بالأفراد والمجتمعات (Kooli & Al Muftah, 2022).

لمعالجة هذه المخاوف، من المهم دمج اعتبارات الخصوصية في تصميم ونشر أنظمة الذكاء الاصطناعي (AI) منذ البداية والتأكد من امتثال هذه الأنظمة لقوانين ولوائح الخصوصية ذات الصلة. من خلال إعطاء الأولوية للخصوصية يمكننا المساعدة في بناء الثقة في أنظمة الذكاء الاصطناعي (AI) والتأكد من استخدامها بطرق مسؤولة وأخلاقية وبلا تمييز تخدم المصالح الفضلى للمجتمعات والمؤسسات (Santoni & Mecacci, 2021).

- تشير الخصوصية في الذكاء الاصطناعي (AI) إلى حماية البيانات الشخصية التي تجمعها أنظمة الذكاء الاصطناعي (AI). البيانات الشخصية هي معلومات حساسة يمكن استخدامها لتحديد هوية الأفراد، مثل الاسم أو العنوان أو المعلومات الصحية (Larsson & Heintz, 2020). لضمان حماية البيانات الشخصية واستخدامها وفقا لقوانين ولوائح الخصوصية يمكن للمؤسسات تنفيذ العديد من الممارسات، مثل:
- الالتزام بقوانين ولوائح الخصوصية، مثل اللائحة العامة لحماية البيانات (GDPR) في أوروبا وقانون خصوصية المستهلك في كاليفورنيا (CCPA) في الولايات المتحدة؛
 - الحصول على موافقة مستنيرة من الأفراد قبل جمع واستخدام بياناتهم الشخصية؛
 - تنفيذ التدابير الفنية والتنظيمية المناسبة لحماية البيانات الشخصية من الوصول أو الاستخدام أو الكشف غير المصرح به؛
 - تدقيق نظام الذكاء الاصطناعي بانتظام للتأكد من أن البيانات الشخصية يتم جمعها ومعالجتها وتخزينها وفقا لقوانين وأنظمة الخصوصية؛
 - تزويد الأفراد بمعلومات واضحة حول كيفية استخدام بياناتهم الشخصية ومعالجتها بواسطة نظام الذكاء الاصطناعي بما في ذلك حقوقهم في الوصول إلى بياناتهم وتصحيحها وحذفها.
- من خلال حماية البيانات الشخصية وفقا لقوانين ولوائح الخصوصية، يمكن للمؤسسات المساعدة في بناء الثقة وزيادة الشفافية ومنع النتائج الضارة، بالإضافة إلى ذلك، يمكن أن يساعد احترام حقوق الخصوصية في ضمان قدرة الأفراد على التحكم في بياناتهم الشخصية والحفاظ على حقوق الخصوصية الخاصة بهم في العصر الرقمي.
- ❖ تصميم أنظمة الذكاء الاصطناعي (AI) بطريقة تسمح بالتدخل والتحكم البشري.
- نظرا لأن أنظمة الذكاء الاصطناعي (AI) أصبحت أكثر تقدما وتكاملا في حياتنا اليومية، فمن الأهمية بمكان التأكد من أن هذه الأنظمة مصممة بطريقة تسمح بالتدخل والتحكم البشري. هذا مهم لعدد من الأسباب، بما في ذلك ضمان المساءلة والشفافية (Orr & Davis, 2020)، ومعالجة الشواغل الأخلاقية والمعنوية، وتجنب إلحاق الضرر بالأفراد والمجتمعات، يمكن أن يتخذ التدخل البشري والتحكم عدة أشكال، مثل قدرة الأفراد على فهم البيانات التي تجمعها أنظمة الذكاء الاصطناعي (AI) وتستخدمها والتحكم فيها، فضلا عن القدرة على التدخل وتجاوز قرارات الذكاء الاصطناعي (AI) في ظروف محددة، من خلال دمج التدخل البشري والتحكم في تصميم أنظمة الذكاء الاصطناعي (AI)، يمكننا المساعدة في ضمان استخدام هذه الأنظمة بطرق مسؤولة وأخلاقية تخدم المصالح الفضلى للمجتمع.
- يشير التدخل والتحكم البشري في الذكاء الاصطناعي (AI) إلى فكرة أن أنظمة الذكاء الاصطناعي (AI) يجب أن تصمم بطريقة تسمح بالتدخل البشري والإشراف، هذا مهم لضمان توافق أنظمة الذكاء الاصطناعي (AI) مع القيم والأخلاق الإنسانية، ومنع الضرر أو العواقب غير المقصودة (Schönberger, 2019).
- للسماح بالتدخل البشري والتحكم في أنظمة الذكاء الاصطناعي (AI)، يمكن للمنظمات تنفيذ العديد من الممارسات مثل:
- تصميم أنظمة الذكاء الاصطناعي (AI) مع مراعاة الشفافية والمساءلة، بحيث يكون من الواضح كيفية اتخاذ القرارات ومن المسؤول عنها (Larsson & Heintz, 2020)؛

- تنفيذ تدابير لضمان إمكانية مراجعة أنظمة الذكاء الاصطناعي (AI) ومراقبتها، وأن هناك ضوابط مناسبة قائمة لمنع الضرر (Alqudah & Muradkhanli, 2021) :
 - السماح بالتدخل البشري في عمليات صنع القرار الرئيسية، بحيث يمكن للأفراد مراجعة القرارات التي يتخذها نظام الذكاء الاصطناعي (AI) وتجاوزها، إذا لزم الأمر؛
 - تصميم أنظمة الذكاء الاصطناعي (AI) مع القدرة على تقديم تفسيرات واضحة لقراراتهم، بحيث يمكن للأفراد فهم كيف ولماذا يتم اتخاذ القرارات (Orr & Davis, 2020)؛
 - التأكد من أن أنظمة الذكاء الاصطناعي (AI) مصممة بطريقة تسمح بتعديلها وتحسينها بسهولة، بحيث يمكن أن تتطور وتتكيف بمرور الوقت (Taddeo, 2019).
- من خلال تصميم أنظمة الذكاء الاصطناعي (AI) مع وضع التدخل البشري والتحكم في الاعتبار، يمكن للمؤسسات المساعدة في بناء الثقة وزيادة الشفافية ومنع النتائج الضارة (Ashok et al., 2022)، بالإضافة إلى ذلك، فإن ضمان خضوع أنظمة الذكاء الاصطناعي (AI) للرقابة البشرية يمكن أن يساعد في ضمان توافقها مع القيم والأخلاق الإنسانية، وأنها تعزز الصالح العام (Schönberger, 2019).
- يجب تصميم أنظمة الذكاء الاصطناعي (AI) وتنفيذها بطريقة تضمن السلامة والأمن والحماية من الأذى والاستخدام الضار.
- نظراً لأن أنظمة الذكاء الاصطناعي أصبحت أكثر انتشاراً في كل مكان، فمن الأهمية بمكان ضمان تصميم هذه الأنظمة وتنفيذها بطريقة تضمن السلامة والأمن والحماية من الأذى والاستخدام الضار، يتطلب ذلك دراسة متأنية لمجموعة متنوعة من العوامل الفنية والأخلاقية، مثل التخزين الآمن للبيانات الحساسة، وتطوير خوارزميات قوية ومقاومة للتلاعب، والحماية من الجهات الخبيثة التي قد تسعى إلى إساءة استخدام أنظمة الذكاء الاصطناعي لأغراض ضارة (Hagerty & Rubinov, 2019).
- يعد ضمان السلامة والأمن في أنظمة الذكاء الاصطناعي (AI) أمراً بالغ الأهمية لبناء الثقة في هذه التقنيات ولضمان استخدامها بطرق مسؤولة وأخلاقية تخدم المصالح الفضلى للمجتمع. من خلال إعطاء الأولوية للسلامة والأمن في الذكاء الاصطناعي (AI)، يمكننا المساعدة في ضمان استخدام هذه الأنظمة لتعزيز رفاهية الأفراد والمجتمعات وعدم التسبب في ضرر (Alqudah & Muradkhanli, 2021).
- يشير الأمان والأمان في الذكاء الاصطناعي (AI) إلى فكرة أن أنظمة الذكاء الاصطناعي (AI) يجب تصميمها وتنفيذها بطريقة تحمي من الأذى والاستخدام الضار، هذا مهم لمنع العواقب غير المقصودة، مثل الأذى الجسدي أو فقدان المعلومات الحساسة (Taddeo, 2019).
- لضمان السلامة والأمن في أنظمة الذكاء الاصطناعي (AI)، يمكن للمنظمات تنفيذ العديد من الممارسات، مثل:
- إجراء تقييمات شاملة للمخاطر لتحديد مخاطر السلامة والأمن المحتملة المرتبطة بنظام الذكاء الاصطناعي (AI)، واتخاذ التدابير المناسبة للتخفيف منها؛
 - تنفيذ التدابير الفنية والتنظيمية المناسبة للحماية من الوصول غير المصرح به أو استخدام أو الكشف عن المعلومات الحساسة؛

- التأكد من أن أنظمة الذكاء الاصطناعي (AI) مصممة مع مراعاة السلامة، مع مراعاة العواقب المحتملة لفشل أو أخطاء النظام (Schönberger, 2019)؛
 - إجراء عمليات تدقيق أمنية منتظمة واختبارات الاختراق لتحديد ومعالجة نقاط الضعف في نظام الذكاء الاصطناعي (AI)؛
 - وضع سياسات وإجراءات واضحة للاستجابة للحوادث الأمنية، واتخاذ الإجراءات المناسبة لمنع الحوادث المستقبلية.
- من خلال ضمان سلامة وأمن أنظمة الذكاء الاصطناعي (AI)، يمكن للمؤسسات المساعدة في بناء الثقة وزيادة الشفافية ومنع النتائج الضارة. بالإضافة إلى ذلك، من خلال الحماية من الأذى والاستخدام الضار، يمكن للمؤسسات المساعدة في ضمان استخدام أنظمة الذكاء الاصطناعي (AI) بطريقة مسؤولة وأخلاقية، لصالح الجميع (Taddeo, 2019).
- يجب تصميم أنظمة الذكاء الاصطناعي (AI) واستخدامها بطريقة تتماشى مع القيم الإنسانية، مثل الكرامة والاحترام والاستقلالية.
- نظراً لأن أنظمة الذكاء الاصطناعي (AI) أصبحت أكثر اندماجاً في حياتنا اليومية، فمن المهم التأكد من تصميم هذه الأنظمة واستخدامها بطريقة تتفق مع القيم الإنسانية، مثل الكرامة والاحترام والاستقلالية وعدم التحيز (Orr & Davis, 2020). وهذا يتطلب دراسة متأنية للاعتبارات الأخلاقية والمعنوية، مثل ضمان أن أنظمة الذكاء الاصطناعي (AI) لا تديم التحيزات القائمة أو تسبب ضرراً للأفراد والمجتمعات والمؤسسات (Schönberger, 2019). كما يتطلب أيضاً أن تحترم أنظمة الذكاء الاصطناعي (AI) استقلالية ووكالة الأفراد، مما يسمح لهم باتخاذ قراراتهم الخاصة والتحكم في حياتهم. من خلال تصميم واستخدام أنظمة الذكاء الاصطناعي (AI) بطريقة تتماشى مع القيم الإنسانية، يمكننا المساعدة في ضمان استخدام هذه الأنظمة بطرق مسؤولة وأخلاقية تخدم المصالح الفضلى للمجتمع. من خلال إعطاء الأولوية لقيم مثل الكرامة والاحترام والاستقلالية في الذكاء الاصطناعي (AI)، يمكننا المساعدة في ضمان استخدام هذه الأنظمة لتعزيز رفاهية الأفراد والمجتمعات وأصحاب المصلحة وعدم التسبب في أي ضرر (Habli et al., 2020).
- يشير الاتساق مع القيم الإنسانية في الذكاء الاصطناعي (AI) إلى فكرة أن أنظمة الذكاء الاصطناعي (AI) يجب تصميمها واستخدامها بطريقة تتفق مع القيم الإنسانية، مثل الكرامة والاحترام والاستقلالية، هذا مهم لضمان أن أنظمة الذكاء الاصطناعي تعزز الصالح العام ولا تسبب ضرراً أو تنتهك المبادئ الأخلاقية.
- لضمان الاتساق مع القيم الإنسانية في أنظمة الذكاء الاصطناعي (AI)، يمكن للمنظمات والشركات والمؤسسات تنفيذ العديد من الممارسات، مثل:
- دمج القيم الإنسانية، مثل الكرامة والاحترام والاستقلالية، في تصميم أنظمة الذكاء الاصطناعي (AI) وتطويرها.
 - التأكد من أن أنظمة الذكاء الاصطناعي (AI) مصممة لاحترام حقوق خصوصية الأفراد، ولحماية المعلومات الحساسة من الوصول أو الاستخدام أو الكشف عن الاستخدام غير المصرح به (Alqudah & Muradkhanli, 2021)؛
 - تجنب التمييز وتعزيز الإنصاف في تصميم واستخدام أنظمة الذكاء الاصطناعي (AI)، بحيث لا تستند القرارات إلى عوامل مثل العرق أو الجنس أو العمر (Orr & Davis, 2020)؛
 - تقديم تفسيرات واضحة للقرارات التي تتخذها أنظمة الذكاء الاصطناعي (AI)، بحيث يمكن للأفراد فهم كيف ولماذا يتم اتخاذ القرارات (Schönberger, 2019)؛

— السماح بالتدخل البشري والرقابة، بحيث يمكن للأفراد مراجعة القرارات التي يتخذها نظام الذكاء الاصطناعي (AI) وتجاوزها، إذا لزم الأمر (Ashok et al., 2022).

من خلال دمج القيم الإنسانية في تصميم واستخدام أنظمة الذكاء الاصطناعي (AI)، يمكن للمنظمات المساعدة في ضمان توافق أنظمة الذكاء الاصطناعي (AI) مع الأخلاقيات والقيم البشرية، وأنها تعزز الصالح العام. بالإضافة إلى ذلك، من خلال تعزيز الكرامة والاحترام والاستقلالية، يمكن للمنظمات المساعدة في بناء الثقة والتأكد من استخدام أنظمة الذكاء الاصطناعي (AI) بطريقة مسؤولة وأخلاقية.

3. النتيجة:

استخدام الذكاء الاصطناعي (AI) لديه القدرة على تحقيق فوائد كبيرة في مجموعة متنوعة من القطاعات المختلفة بما في ذلك الرعاية الصحية، التمويل، التعليم، وأكثر من ذلك. ومع ذلك، فإن تطبيق الذكاء الاصطناعي (AI) في هذه القطاعات يؤثر أيضا قضايا أخلاقية مهمة يجب دراستها ومعالجتها. يمكن أن تختلف هذه القضايا اعتمادا على القطاع والسياق المحددين، ولكنها يمكن أن تشمل مخاوف تتعلق بالخصوصية والتحيز والمساءلة وتهجير الوظائف وعدم التمييز والإنصاف والسلامة والأمن والاتساق مع القيم الإنسانية مثل الكرامة والاحترام والاستقلالية. من الضروري إجراء دراسات شاملة ومتعددة التخصصات لهذه القضايا الأخلاقية من أجل تطوير مبادئ توجيهية وأطر للاستخدام المسؤول للذكاء الاصطناعي (AI) في كل قطاع. يجب أن تجمع هذه الدراسات خبراء من المجالات ذات الصلة، بما في ذلك علوم الكمبيوتر والأخلاق والقانون وغيرها، لضمان فهم التحديات الأخلاقية الفريدة التي يطرحها الذكاء الاصطناعي (AI) في كل قطاع ومعالجتها بشكل كامل، من خلال القيام بذلك، يمكننا المساعدة في ضمان تحقيق فوائد الذكاء الاصطناعي (AI) بطرق تتوافق مع قيمنا وتعزز رفاهية الأفراد والمجتمعات.

يثير استخدام الذكاء الاصطناعي (AI) عددا من القضايا الأخلاقية المهمة التي يجب دراستها ومعالجتها بعناية. وتشمل هذه القضايا المتعلقة بالخصوصية، والتحيز، والمساءلة، والاستغناء عن الوظائف، وعدم التمييز والإنصاف، والسلامة والأمن، والاتساق مع القيم الإنسانية مثل الكرامة والاحترام والاستقلالية، يعد ضمان تصميم أنظمة الذكاء الاصطناعي (AI) واستخدامها بطريقة مسؤولة وأخلاقية أمرا بالغ الأهمية لبناء الثقة في هذه التقنيات ولضمان استخدامها بطرق تخدم المصالح الفضلى للمجتمع. يتطلب ذلك نهجا متعدد التخصصات يجمع خبراء من علوم الكمبيوتر والأخلاق والقانون ومجالات أخرى لمعالجة هذه القضايا الأخلاقية وتطوير إرشادات وأطر للاستخدام المسؤول للذكاء الاصطناعي.

4. مناقشة النتائج:

في الختام، تعد دراسة القضايا الأخلاقية في استخدام الذكاء الاصطناعي (AI) عبر مختلف القطاعات قضية مهمة وملحة، إن تطبيق الذكاء الاصطناعي (AI) لديه القدرة على تحقيق فوائد كبيرة في مجموعة متنوعة من القطاعات ولكنه يثير أيضا أسئلة وتحديات أخلاقية مهمة يجب معالجتها. يمكن أن تشمل هذه القضايا مخاوف تتعلق بالخصوصية، والتحيز والمساءلة، واستبدال الوظائف، وعدم التمييز والإنصاف، والسلامة والأمن، والاتساق مع القيم الإنسانية مثل الكرامة والاحترام والاستقلالية. من الضروري إجراء دراسات متعددة التخصصات لهذه القضايا الأخلاقية من أجل فهم التحديات الفريدة التي يطرحها الذكاء الاصطناعي (AI) في كل قطاع ووضع مبادئ توجيهية وأطر للاستخدام المسؤول للذكاء الاصطناعي (AI). من خلال القيام بذلك، يمكننا المساعدة في ضمان استخدام الذكاء الاصطناعي (AI) بطرق تعزز رفاهية الأفراد والمجتمعات وتتفق مع قيمنا ومبادئنا الأخلاقية.

الذكاء الاصطناعي (AI) هو تقنية سريعة التطور ولديها القدرة على تحويل العديد من جوانب حياتنا، من الرعاية الصحية والتمويل إلى التعليم وما بعده، ومع ذلك، فإن تطوير واستخدام الذكاء الاصطناعي (AI) يثير أيضا أسئلة وتحديات أخلاقية مهمة يجب معالجتها.

إحدى القضايا الرئيسية هي الخصوصية، غالبا ما تجمع أنظمة الذكاء الاصطناعي (AI) كميات كبيرة من البيانات الشخصية وتعالجها، مما قد يثير مخاوف بشأن الخصوصية ويعرض المعلومات الحساسة للخطر، من المهم التأكد من حماية البيانات الشخصية التي تجمعها أنظمة الذكاء الاصطناعي (AI) واستخدامها وفقا لقوانين ولوائح الخصوصية. قضية أخرى مهمة هي التحيز، يمكن لأنظمة الذكاء الاصطناعي (AI) أن تديم التحيزات الحالية بل وتضخمها، والتي يمكن أن يكون لها آثار سلبية على الأفراد والمجتمعات. من الضروري تصميم أنظمة الذكاء الاصطناعي بطريقة تتجنب التمييز وتعزز العدالة.

تعد المساءلة أيضا قضية مهمة في استخدام الذكاء الاصطناعي (AI)، من المهم التأكد من أن المسؤولية والمساءلة واضحة بالنسبة للإجراءات والقرارات التي تتخذها أنظمة الذكاء الاصطناعي (AI)، خاصة في حالات الضرر أو الخطأ. يعد الاستغناء عن الوظائف مصدر قلق آخر يثيره استخدام الذكاء الاصطناعي (AI)، يتمتع الذكاء الاصطناعي (AI) بالقدرة على أتمتة العديد من الوظائف، مما قد يؤدي إلى فقدان الوظائف والمصاعب الاقتصادية للأفراد والمجتمعات، من المهم النظر في تأثير الذكاء الاصطناعي (AI) على التوظيف ووضع استراتيجيات للتخفيف من إزاحة الوظائف.

5. خاتمة:

1.5. التوصيات

يثير استخدام الذكاء الاصطناعي (AI) أيضا أسئلة حول عدم التمييز والإنصاف، يجب تصميم أنظمة الذكاء الاصطناعي بطريقة تضمن استخدامها بطريقة تتجنب التمييز وتعزز العدالة. أخيرا، تعد السلامة والأمن مسألة حاسمة في استخدام الذكاء الاصطناعي (AI)، يجب تصميم أنظمة الذكاء الاصطناعي (AI) وتنفيذها بطريقة تضمن السلامة وتحمي من الأذى والاستخدام الضار. من خلال معالجة هذه القضايا والتحديات الأخلاقية، يمكننا المساعدة في ضمان استخدام الذكاء الاصطناعي (AI) بطرق مسؤولة وأخلاقية تخدم المصالح الفضلى للمجتمع وتتفق مع القيم الإنسانية مثل الكرامة والاحترام والاستقلالية.

2.5. النتائج:

بناء على دراسة القضايا الأخلاقية في استخدام الذكاء الاصطناعي (AI) عبر مختلف القطاعات، يمكن تقديم التوصيات التالية:

- نهج متعدد التخصصات: إجراء دراسات متعددة التخصصات للقضايا الأخلاقية المتعلقة بالذكاء الاصطناعي (AI) لفهم التحديات الفريدة التي يطرحها الذكاء الاصطناعي في كل قطاع ولتطوير حلول شاملة وفعالة؛
- حماية الخصوصية: تأكد من حماية البيانات الشخصية التي تجمعها أنظمة الذكاء الاصطناعي (AI) واستخدامها وفقا لقوانين ولوائح الخصوصية؛
- تقليل التحيز: تصميم أنظمة الذكاء الاصطناعي (AI) لتقليل التحيز والتمييز، وتعزيز العدالة وعدم التمييز؛

- المساءلة الواضحة: تحديد المسؤولية والمساءلة بشكل واضح عن الإجراءات والقرارات التي تتخذها أنظمة الذكاء الاصطناعي (AI)، لاسيما في حالات الضرر أو الخطأ؛
 - التخفيف من إزاحة الوظائف: النظر في تأثير الذكاء الاصطناعي (AI) على التوظيف ووضع استراتيجيات للتخفيف من إزاحة الوظائف؛
 - موائمة القيم الإنسانية: تصميم واستخدام أنظمة الذكاء الاصطناعي (AI) بطريقة تتماشى مع القيم الإنسانية مثل الكرامة والاحترام والاستقلالية؛
 - السلامة والأمن: تأكد من تصميم أنظمة الذكاء الاصطناعي (AI) وتنفيذها بطريقة تعطي الأولوية للسلامة والأمن وتحمي من الأذى والاستخدام الضار؛
 - تطوير الإطار الأخلاقي: تطوير الأطر والمبادئ التوجيهية الأخلاقية للاستخدام المسؤول للذكاء الاصطناعي (AI)، مع مراعاة التحديات الفريدة التي يطرحها الذكاء الاصطناعي (AI) في كل قطاع؛
 - المراقبة المستمرة: مراقبة وتقييم تأثير الذكاء الاصطناعي (AI) على المجتمع بشكل مستمر وتعديل الأطر والمبادئ التوجيهية الأخلاقية حسب الحاجة.
- بإتباع هذه التوصيات، يمكننا المساعدة في ضمان استخدام الذكاء الاصطناعي (AI) بطرق مسؤولة وأخلاقية تعزز رفاهية الأفراد والمجتمعات وتتوافق مع قيمنا ومبادئنا الأخلاقية.

6. قائمة المراجع:

1. Akgun, S., & Greenhow, C. (2021). Artificial intelligence in education: Addressing ethical challenges in K-12 settings. *AI and Ethics*, 1-10.
2. Alqudah, M. A., & Muradkhanli, L. (2021). Artificial Intelligence in Electric Government; Ethical Challenges and Governance in Jordan. *Electronic Research Journal of Social Sciences and Humanities*, 3, 65-74.
3. Ashok, M., Madan, R., Joha, A., & Sivarajah, U. (2022). Ethical framework for Artificial Intelligence and Digital technologies. *International Journal of Information Management*, 62, 102433.
4. Beil, M., Proft, I., van Heerden, D., Sviri, S., & van Heerden, P. V. (2019). Ethical considerations about artificial intelligence for prognostication in intensive care. *Intensive Care Medicine Experimental*, 7(1), 1-13.
5. Bostrom, N., & Yudkowsky, E. (2018). The ethics of artificial intelligence. In *Artificial intelligence safety and security* (pp. 57-69). Chapman and Hall/CRC.
6. Brady, A. P., & Neri, E. (2020). Artificial intelligence in radiology—ethical considerations. *Diagnostics*, 10(4), 231.
7. Brendel, A. B., Mirbabaie, M., Lembcke, T. B., & Hofeditz, L. (2021). Ethical management of artificial intelligence. *Sustainability*, 13(4), 1974.
8. Burton, E., Goldsmith, J., Koenig, S., Kuipers, B., Mattei, N., & Walsh, T. (2017). Ethical considerations in artificial intelligence courses. *AI magazine*, 38(2), 22-34.
9. Chen, D. K., Modi, Y., & Al-Aswad, L. A. (2022). Promoting transparency and standardization in ophthalmologic artificial intelligence: a call for artificial intelligence model card. *The Asia-Pacific Journal of Ophthalmology*, 10-1097.
10. Coeckelbergh, M. (2019). Artificial intelligence: Some ethical issues and regulatory challenges. *Technology and regulation*, 2019, 31-34.
11. Dignum, V. (2020). Responsibility and artificial intelligence. *The oxford handbook of ethics of AI*, 4698, 215.
12. Felzmann, H., Fosch-Villaronga, E., Lutz, C., & Tamò-Larriex, A. (2020). Towards transparency by design for artificial intelligence. *Science and Engineering Ethics*, 26(6), 3333-3361.

13. Habli, I., Lawton, T., & Porter, Z. (2020). Artificial intelligence in health care: accountability and safety. *Bulletin of the World Health Organization*, 98(4), 251.
14. Hagerty, A., & Rubinov, I. (2019). Global AI ethics: a review of the social impacts and ethical implications of artificial intelligence. *arXiv preprint arXiv:1907.07892*.
15. Kellmeyer, P. (2019). Artificial intelligence in basic and clinical neuroscience: opportunities and ethical challenges. *Neuroforum*, 25(4), 241-250.
16. Kerr, A., Barry, M., & Kelleher, J. D. (2020). Expectations of artificial intelligence and the performativity of ethics: Implications for communication governance. *Big Data & Society*, 7(1), 2053951720915939.
17. Kooli, C., & Al Muftah, H. (2022). Artificial intelligence in healthcare: a comprehensive review of its ethical concerns. *Technological Sustainability*.
18. Krontiris, I., Grammenou, K., Terzidou, K., Zacharopoulou, M., Tsikintikou, M., Baladima, F., ... & Kaouras, K. (2020, December). Autonomous vehicles: data protection and ethical considerations. In *Proceedings of the 4th ACM Computer Science in Cars Symposium* (pp. 1-10).
19. Kumar, N., Kharkwal, N., Kohli, R., & Choudhary, S. (2016, February). Ethical aspects and future of artificial intelligence. In *2016 International Conference on Innovation and Challenges in Cyber Security (ICICCS-INBUSH)* (pp. 111-114). IEEE.
20. Larsson, S., & Heintz, F. (2020). Transparency in artificial intelligence. *Internet Policy Review*, 9(2).
21. Loi, M., & Spielkamp, M. (2021, July). Towards accountability in the use of artificial intelligence for public administrations. In *Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society* (pp. 757-766).
22. Max, R., Kriebitz, A., & Von Websky, C. (2021). Ethical considerations about the implications of artificial intelligence in finance. *Handbook on Ethics in Finance*, 577-592.
23. Mazurek, G., & Małagocka, K. (2019). Perception of privacy and data protection in the context of the development of artificial intelligence. *Journal of Management Analytics*, 6(4), 344-364.
24. McKinney, S. M., Karthikesalingam, A., Tse, D., Kelly, C. J., Liu, Y., Corrado, G. S., & Shetty, S. (2020). Reply to: Transparency and reproducibility in artificial intelligence. *Nature*, 586(7829), E17-E18.
25. Morgan, F. E., Boudreaux, B., Lohn, A. J., Ashby, M., Curriden, C., Klima, K., & Grossman, D. (2020). Military applications of artificial intelligence: ethical concerns in an uncertain world. *RAND PROJECT AIR FORCE SANTA MONICA CA SANTA MONICA United States*.
26. Munoko, I., Brown-Liburd, H. L., & Vasarhelyi, M. (2020). The ethical implications of using artificial intelligence in auditing. *Journal of Business Ethics*, 167, 209-234.
27. Orr, W., & Davis, J. L. (2020). Attributions of ethical responsibility by Artificial Intelligence practitioners. *Information, Communication & Society*, 23(5), 719-735.
28. Robinson, S. C. (2020). Trust, transparency, and openness: How inclusion of cultural values shapes Nordic national public policy strategies for artificial intelligence (AI). *Technology in Society*, 63, 101421.
29. Sadok, H., Sakka, F., & El Maknouzi, M. E. H. (2022). Artificial intelligence and bank credit analysis: A review. *Cogent Economics & Finance*, 10(1), 2023262.
30. Safdar, N. M., Banja, J. D., & Meltzer, C. C. (2020). Ethical considerations in artificial intelligence. *European journal of radiology*, 122, 108768.
31. Santoni de Sio, F., & Mecacci, G. (2021). Four responsibility gaps with artificial intelligence: Why they matter and how to address them. *Philosophy & Technology*, 34, 1057-1084.
32. Schönberger, D. (2019). Artificial intelligence in healthcare: a critical analysis of the legal and ethical implications. *International Journal of Law and Information Technology*, 27(2), 171-203.
33. Shreve, J. T., Khanani, S. A., & Haddad, T. C. (2022). Artificial Intelligence in Oncology: Current Capabilities, Future Opportunities, and Ethical Considerations. *American Society of Clinical Oncology Educational Book*, 42, 842-851.

34. Siau, K., & Wang, W. (2020). Artificial intelligence (AI) ethics: ethics of AI and ethical AI. *Journal of Database Management (JDM)*, 31(2), 74-87.
35. Stahl, B. C., (2021). Ethical issues of AI. *Artificial Intelligence for a better future: An ecosystem perspective on the ethics of AI and emerging digital technologies*, 35-53.
36. Stahl, B. C., Antoniou, J., Ryan, M., Macnish, K., & Jiya, T. (2022). Organisational responses to the ethical issues of artificial intelligence. *AI & SOCIETY*, 37(1), 23-37.
37. Taddeo, M. (2019). Three ethical challenges of applications of artificial intelligence in cybersecurity. *Minds and machines*, 29, 187-191.
38. Walmsley, J. (2021). Artificial intelligence and the value of transparency. *AI & SOCIETY*, 36(2), 585-595.
39. Zhai, X., Chu, X., Chai, C. S., Jong, M. S. Y., Istenic, A., Spector, M., ... & Li, Y. (2021). A Review of Artificial Intelligence (AI) in Education from 2010 to 2020. *Complexity*, 2021, 1-18.