

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام – الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام – آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م – (عدد خاص)

استخدام أدوات الذكاء الاصطناعي للتحقق من التزييف العميق في المضامين الإعلامية

USING AI TOOLS TO VERIFY DEEPPKES IN MEDIA CONTENT

الدكتور محمد وهاب عبود - صحافة

Dr. Mohammed Wahhab Abboud – Journalism

جامعة بغداد – كلية العلوم – شعبة الإعلام والاتصال الحكومي

**University of Baghdad - College of Science - Department of
Media and Government Communication**

muhamedaboud19@gmail.com

الملخص

يروم هذا البحث إلى التعرف على فعالية أدوات الذكاء الاصطناعي للتحقق من التزييف العميق في المضامين الإعلامية، فضلاً عن تقييم مدى دقتها وحدود موثوقيتها، مع استكشاف وجهات نظر الصحفيين ومتخصصي الذكاء الاصطناعي حول تطبيقاتها العملية، وذلك من خلال استخدام منهج بحث نوعي يعتمد على تحليل محتوى أدوات التحقق بالذكاء الاصطناعي، وتحديدًا Sensity AI و Deepware Scanner عبر اختبار أدائهما لمجموعة بيانات من محتوى إعلامي مُتلاعب به وآخر حقيقي. بالإضافة إلى ذلك، أُجريت مقابلات شبه منظمة مع ١٥ صحفياً و ٥ متخصصين في الذكاء الاصطناعي بغية فهم تحديات كشف التزييف العميق ودور الذكاء الاصطناعي في سير العمل الصحفي. وتكشف النتائج أنه في حين تُظهر أدوات الذكاء الاصطناعي دقة عالية في الكشف عن التزييف العميق في المحتوى الاعلامي ذات المغزى السياسي إلا أنها تواجه صعوبة في التعامل مع المحتوى الصوتي المؤلّد بواسطة الذكاء الاصطناعي. وغالبًا ما تفشل هذه الأدوات في تأكيد المحتوى الأصلي على أنه حقيقي، مما يثير مخاوف وشكوك الصحفيين تجاه فعاليتها. كما تشير المقابلات مع الصحفيين إلى أن معظم المؤسسات الإخبارية تفتقر إلى برامج التدريب والدعم المؤسسي للكشف عن التزييف العميق، مما يجعل الصحفيين غير مستعدين لتحديد المحتوى المُتلاعب به. فيما يؤكد متخصصو الذكاء الاصطناعي على الحاجة إلى التحديثات المستمرة، حيث تتطور تقنية التزييف العميق بشكل أسرع من طرق الكشف. وتوصي الدراسة باعتماد نموذج تحقق هجين يجمع بين الكشف باستخدام الذكاء الاصطناعي والتحقق اليدوي، فضلاً عن تطوير أدوات ذكاء اصطناعي مفتوحة المصدر وبرامج تدريبية متخصصة للصحفيين لتعزيز الثقة بوسائل الإعلام ومواصلة الجهود الصحفية لمكافحة التزييف العميق في المحتوى الإعلامي.

Abstract

This research aims to identify the effectiveness of AI tools for detecting deepfake content in media content, as well as assess their accuracy and reliability. It also explores the perspectives of journalists and AI specialists on their practical applications. This research uses a qualitative research approach based on content analysis of AI verification tools, specifically Sensity AI and Deepware Scanner, by testing their performance on a dataset of manipulated and authentic media content. Semi-structured interviews were conducted with 15 journalists and five AI specialists to understand the challenges of deepfake detection and the role of AI in journalistic workflows. The findings reveal that while AI tools demonstrate high accuracy in detecting deepfake content in politically motivated media, they struggle with AI-generated audio content. These tools often fail to confirm the original content as authentic, raising concerns and doubts among journalists about their effectiveness. Interviews with journalists also indicate that most news organizations lack training programs and institutional support for deepfake detection, leaving journalists ill-prepared to identify manipulated content. AI experts emphasize the need for continuous updates, as deepfake technology evolves faster than detection methods. The study recommends adopting a hybrid verification model that combines AI detection and manual verification, as well as developing open-source AI tools and specialized training programs for journalists to enhance trust in the media and continue journalistic efforts to combat deepfake content.

Keywords: Artificial intelligence, fact-checking, deepfake, media content

المقدمة

شهد العقد الأخير طفرة نوعية في تطبيقات الذكاء الاصطناعي (AI)، مدفوعة بقوة الحوسبة التي وفرت إمكانات غير مسبوقة وكميات هائلة من البيانات. وقد انعكس هذا التوسع على مجال الصحافة، إذ أحدث ثورة في طرق إنتاج الأخبار وتوزيعها واستهلاكها. بيد أن هذا التقدم التكنولوجي لم يخل من تبعات سلبية، فقد حمل بين طياته تهديدا لمصداقية الإعلام، لا سيما مع تطور خوارزميات الذكاء الاصطناعي التوليدية التي أنتجت تقنية التزييف العميق

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام – الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام – آفاق الإنكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م – (عدد خاص)

(Deepfake) وتعرّف بأنها محتويات سمعية وبصرية تم التلاعب بها أو تركيبها باستخدام تقنيات الذكاء الاصطناعي، بحيث تبدو واقعية وحقيقية، وهو ما مثل واحدة من التحديات التي تواجه الصحافة المعاصرة.

كما يشكل التزييف العميق خطراً متعدد الأوجه على المجتمع، فبالإضافة إلى إمكانية استخدامه في نشر الأخبار الكاذبة والتضليل الإعلامي والتلاعب بالرأي العام، فإنه يهدد بتقويض الثقة في المؤسسات الإعلامية وبالعملية الديمقراطية نفسها. كما أن التزييف العميق يمكن أن يستخدم في التشهير بالأفراد والابتزاز والتحرش عبر الإنترنت، مما يتسبب في أضرار نفسية واجتماعية جسيمة.

وفي مواجهة هذا التحدي المتنامي، أصبح من الضروري تطوير أدوات وتقنيات جديدة للتحقق من المحتوى الإعلامي وكشف التزييف العميق. وقد بدأت بالفعل تظهر بعض الحلول القائمة على الذكاء الاصطناعي، والتي تستخدم خوارزميات متقدمة لتحليل المحتوى السمعي والبصري وكشف العلامات الدالة على التلاعب.

في هذا السياق، يأتي هذا البحث ليسلط الضوء على دور أدوات الذكاء الاصطناعي في التحقق من المحتوى الإعلامي وكشف التزييف العميق، فضلاً عن اختبار مدى كفاءة هذه الأدوات في التمييز بين المحتوى الحقيقي والمزيف. كما يتناول البحث إمكانية دمج هذه الأدوات ضمن الممارسات الصحفية لتعزيز مصداقية الإعلام، وتقديم رؤية شاملة حول فعالية تقنيات الذكاء الاصطناعي في مواجهة التزييف العميق، واقتراح حلول تساهم في الحد من آثاره السلبية على البيئة الإعلامية والمجتمعية، تحديداً في فضاء وسائل الإعلام الرقمية ومنصات التواصل الاجتماعي التي تنتشر بين ثناياها هذه التقنيات بشكل متزايد.

المبحث الأول: منهجية البحث

مشكلة البحث

في ظل التطور المتسارع لتقنيات التزييف العميق، برز تحدٍّ أمام المؤسسات الإعلامية والصحفيين يتمثل في صعوبة التمييز بين المحتوى الأصلي والمحتوى المزيف. فبينما تعتمد أساليب التحقق التقليدية على التدقيق اليدوي وتقييم المصادر، فإن التزييف العميق يتخطى قدرة هذه الأساليب على الكشف والتحقق. فمن هنا، تبرز الحاجة إلى استكشاف حلول جديدة قادرة على مواكبة هذا التطور التكنولوجي. لذا، يسعى هذا البحث إلى تقييم ما إذا كانت أدوات الذكاء الاصطناعي - بما تملكه من قدرات تحليلية متقدمة وإمكانيات للتعلم الآلي - يمكن أن تساهم في تطوير آليات أكثر فاعلية وموثوقية للتحقق من صحة المحتوى الإعلامي وكشف التزييف العميق. وبدلاً من الانطلاق من فرضية مسبقة بشأن فشل الأساليب التقليدية أو نجاح أدوات الذكاء الاصطناعي، يسعى البحث إلى فهم

الدكتور محمد وهاب عبود

أعمق لقدرات وقيود هذه الأدوات في مواجهة هذا التحدي الراهن، وتقييم إمكانية دمجها في سير العمل الصحفي لتعزيز الدقة والمصداقية.

أسئلة البحث

- ١- ما مدى فعالية ودقة أدوات الذكاء الاصطناعي مثل Sensity AI و Deepware Scanner في التحقق من محتوى التزييف العميق ودرجة مصداقيتها في المضامين الإعلامية المختلفة؟
- ٢- ما هي التحديات التقنية والأخلاقية التي تواجه الصحفيين والمؤسسات الإعلامية عند استخدام أدوات الذكاء الاصطناعي للتحقق من التزييف العميق ضمن سير العمل الصحفي؟
- ٣- كيف يمكن دمج أدوات التحقق بالذكاء الاصطناعي بفعالية في الممارسات الصحفية لتعزيز القدرة على كشف التزييف العميق والمساهمة في الحفاظ على مصداقية وسائل الإعلام وثقة الجمهور؟
- ٤- ما هي وجهات نظر الصحفيين والمتخصصين في الذكاء الاصطناعي حول القدرات الحالية والمستقبلية لهذه الأدوات في مكافحة التزييف العميق؟

أهداف البحث

- تقييم دقة وموثوقية أداتي Sensity AI و Deepware Scanner التي تعتمد على الذكاء الاصطناعي في كشف أنواع مختلفة من محتوى التزييف العميق (فيديو، صوت).
- استكشاف وجهات نظر الصحفيين حول استخدام هذه الأدوات وتأثيرها المحتمل على ممارسات التحقق الصحفي.
- تحديد القيود التقنية والتحديات العملية والأخلاقية المتعلقة باستخدام أدوات الذكاء الاصطناعي للتحقق من التزييف العميق في غرف الأخبار.
- اقتراح استراتيجيات عملية لدمج أدوات التحقق بالذكاء الاصطناعي بفعالية ومسؤولية في الممارسات الصحفية لمكافحة التزييف العميق.

أهمية البحث

تنبثق أهمية هذا البحث من الحاجة الملحة لمواجهة تحدي التزييف العميق المتزايد في البيئة الإعلامية الرقمية، ويمكن تفصيلها كالتالي:

الأهمية العلمية:

- المساهمة في الأدبيات البحثية المتعلقة بتقاطع الذكاء الاصطناعي مع دراسات الإعلام والصحافة، لاسيما في مجال التحقق من المعلومات ومكافحة التضليل الرقمي.

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام – الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام – آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م – (عدد خاص)

- تقديم تقييم تجريبي ومقارن لأداء أدوات محددة لكشف التزييف العميق، Sensity AI و Deepware Scanner، مما يوفر بيانات مرجعية للباحثين في هذا المجال.
- استكشاف وتطبيق أطر نظرية مثل (نظرية ما بعد الحقيقة، حارس البوابة، الحتمية التكنولوجية) لفهم أبعاد استخدام الذكاء الاصطناعي في التحقق الصحفي.
- تطوير فهم أعمق وأشمل للتحديات المنهجية والتقنية في تصميم وتقييم أدوات كشف التزييف العميق.

الأهمية العملية:

- تزويد الصحفيين ومدققي الحقائق والمؤسسات الإعلامية بتقييم عملي لقدرات وحدود أدوات الذكاء الاصطناعي المتاحة لمساعدتهم في اتخاذ قرارات مستنيرة بشأن تبنيها.
- تقديم رؤى حول التحديات التي يواجهها الصحفيون في التعامل مع التزييف العميق، مما يساعد المؤسسات الإعلامية على تطوير برامج تدريبية وسياسات دعم مناسبة.
- المساهمة في النقاش العام ورفد صناع السياسات بالمعلومات اللازمة لوضع استراتيجيات وتشريعات تهدف إلى الحد من انتشار التزييف العميق وتأثيراته السلبية.
- اقتراح نماذج عملية (كالنموذج الهجين) يمكن للمؤسسات الإعلامية تطبيقها لتحسين عمليات التحقق من المحتوى الرقمي وتعزيز المصادقية.
- المساهمة في الجهود الرامية للحفاظ على ثقة الجمهور بوسائل الإعلام في عصر يتسم بانتشار المعلومات المضللة والتزييف المتقدم.

منهجية البحث

ينتمي هذا البحث إلى صنف الدراسات الوصفية التحليلية، مُعتمداً على منهج بحثي نوعي (Qualitative Research Approach) يهدف إلى فهم ظاهرة استخدام أدوات الذكاء الاصطناعي في التحقق من التزييف العميق بعمق وسياقها. أختير المنهج النوعي كونه يسمح باستكشاف تعقيدات الظاهرة وفهم وجهات نظر المشاركين بطريقة معمقة تتجاوز محددات القياس الكمي (Berger 2018).

لقد طبق البحث المنهج من خلال الجمع بين أداتين رئيسيتين لجمع وتحليل البيانات:

١- تحليل المحتوى النوعي: (Qualitative Content Analysis)

جرى استخدام أداتين متخصصتين في الكشف عن التزييف العميق وهما **Sensity AI** و **Deepware Scanner**. تم اختيارهما نظرا لشهرتهما النسبية في هذا المجال وتوفيرهما إمكانية اختبار محدودة أو تجريبية مجانية، مما يتيح تقييمهما ضمن حدود البحث. إذ قام الباحث بتجميع مجموعة بيانات (عينة) تتضمن محتوى إعلاميا متنوعا، يشمل مقاطع فيديو حقيقية وأخرى تم التلاعب بها أو إنشاؤها بتقنية التزييف العميق. وتم إخضاع هذه العينة للتحليل باستخدام الأداتين المذكورتين. وركز التحليل على تقييم مدى قدرة كل أداة على:

- التمييز الصحيح بين المحتوى الحقيقي والمزيف.
 - تحديد أنواع معينة من التلاعب (بصري، صوتي).
 - سرعة المعالجة وتقديم النتائج.
 - وجود أو غياب تفسيرات للنتائج (قابلية التفسير).
- جرى تحليل النتائج الصادرة عن الأداتين بشكل وصفي ومقارن لتحديد نقاط القوة والضعف لكل أداة وفعاليتها الإجمالية في سياقات مختلفة.

٢- المقابلات شبه المنظمة: (Semi-Structured Interviews)

جرى تصميم دليل مقابلة شبه منظم يتضمن محاور رئيسية مع ترك مرونة لطرح أسئلة متابعة بناء على إجابات المشاركين. أجريت مقابلات مع عينة قصدية (Purposive Sample) تتكون من ٢٠ فرداً، موزعين كالتالي: ١٥ صحفياً يمارسون المهنة في مؤسسات إعلامية رقمية أو تقليدية تتعامل مع المحتوى الرقمي بشكل مكثف (مثل المواقع الإخبارية، القنوات التلفزيونية، منصات التواصل الاجتماعي)، و ٥ متخصصين في مجال الذكاء الاصطناعي أو تطوير البرمجيات لديهم خبرة بتقنيات التزييف العميق أو كشفه. أختيرت هذه العينة لتوفير وجهات نظر متنوعة تجمع بين التجربة العملية في الصحافة والخبرة التقنية في الذكاء الاصطناعي. كما أجريت المقابلات (شخصياً أو عبر الإنترنت/تطبيقات المراسلات الواتساب) وركزت على استكشاف المحاور التالية:

- تصورات الصحفيين والمتخصصين لمدى خطورة التزييف العميق.
- تجارب الصحفيين مع أدوات الكشف (إن وجدت) والتحديات التي يواجهونها.
- تقييم المتخصصين لقدرات وحدود أدوات الكشف الحالية.
- الاحتياجات التدريبية والمؤسسية لدعم الصحفيين.
- الاعتبارات الأخلاقية المتعلقة باستخدام الذكاء الاصطناعي في التحقق.
- التوصيات المستقبلية لتطوير الأدوات والممارسات.

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام – الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام – آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م – (عدد خاص)

فُرغت المقابلات وُحِلت باستخدام التحليل الموضوعي (Thematic Analysis) لتحديد الأنماط والمواضيع الرئيسية المتكررة في إجابات المشاركين وتفسيرها في ضوء أسئلة البحث وأهدافه.

يهدف هذا الأسلوب بين تحليل المحتوى والمقابلات إلى تقديم فهم شامل وأعمق ومتعدد الأبعاد لفعالية استخدام أدوات الذكاء الاصطناعي في التحقق من التزييف العميق، مع الأخذ في الاهتمام الجوانب التقنية والعملية والبشرية.

الدراسات السابقة:

دراسة: (Cazzamatta & Sarısakaloglu (2025)

هدفت الدراسة إلى التعرف ورصد مجال الذكاء الاصطناعي المتطور في مجال التحقق من صحة المعلومات الصحفية، مستخدماً نهجاً متعدد الأساليب لاستكشاف مدى تكامله مع البحث العلمي والتطبيقات العملية. وبناءً على مراجعة منهجية لـ ٣٤٨ منشوراً حول الصحافة القائمة على الخوارزميات، تحدد الدراسة الاتجاهات الرئيسية من خلال تحليل كمي لمحتوى ٣١٥٤ مقالة تحقق في ٢٣ منظمة في ثماني دول (الأرجنتين والبرازيل وتشيلي وفنزويلا والمملكة المتحدة وألمانيا والبرتغال وإسبانيا). يُستكمل هذا بتحليل نوعي لمواقع المنظمات الإلكترونية والمواد الترويجية. كشفت النتائج عن خلل في الإنتاج العلمي وندرة في الأبحاث التي تربط بين الصحافة القائمة على الخوارزميات والتحقق من صحة المعلومات، إلا أن الأدوات الخاصة بكل منطقة للكشف عن المعلومات المضللة تكشف عن تباينات التي غالباً ما تمولها جهات مثل جوجل والاتحاد الأوروبي. كشفت الدراسة عن وجود فجوة بين البحث العلمي والتطبيقات العملية للذكاء الاصطناعي في التحقق الصحفي، وندرة في الأبحاث التي تربط بين الصحافة الخوارزمية والتحقق. كما أظهرت تبايناً في الأدوات المستخدمة إقليمياً، والتي غالباً ما تكون ممولة من جهات خارجية مثل جوجل والاتحاد الأوروبي، مما يشير إلى تفاوت في القدرات والتوجهات.

دراسة: ابراهيم، سامح محمد محمد السيد (٢٠٢٥).

سعت الدراسة إلى رصد مخاطر التزييف العميق (Deepfake) الذي يعتمد على الذكاء الاصطناعي لإنتاج محتوى رقمي مزيف بالغ الواقعية، ويشمل فيديوهات وصوراً وأصواتاً مُصنَّعة بدقة، وشخصت الدراسة التهديدات الناجمة عنه: من انتهاك الخصوصية والتلاعب بالرأي العام إلى التشهير والاحتيال المالي المعقد، بالإضافة إلى تآكل الثقة بالمعلومات الرقمية، وكذلك يُستخدم التزييف العميق للابتزاز ونشر الشائعات وتقمص الهويات الموثوقة في عمليات الاحتيال. كما أوصت الدراسة بضرورة تطوير أدوات للكشف عن المحتوى المزيف وسن تشريعات رادعة وتعزيز الوعي المجتمعي وتقوية الإجراءات الأمنية في المؤسسات إذ إن التعاون الدولي وتبادل الخبرات يمثلان عنصراً حاسماً لضمان حماية المجتمعات، كما ان الحد من مخاطر التزييف العميق يتطلب جهوداً مشتركة بين الحكومات والمؤسسات التقنية والمجتمع المدني لضمان استخدام

الدكتور محمد وهاب عبود

التكنولوجيا بشكل أخلاقي وآمن. شخّصت الدراسة مخاطر متعددة للتزييف العميق تشمل انتهاك الخصوصية والتلاعب بالرأي العام والتشهير والاحتيال المالي وتآكل الثقة الرقمية، موصية بضرورة تطوير أدوات كشف وسن تشريعات وتعزيز الوعي وتقوية الأمن المؤسسي والتعاون الدولي كآليات للمواجهة.

دراسة: Rosca, Stancu, Iovanovici (2025)

سعت هذه الدراسة إلى تقديم نهجًا جديدًا للكشف عن نصوص التزييف العميق، في مواجهة التحدي الشامل المتمثل في صحة النصوص، آخذة في الاعتبار مخاطر التلاعب الكامنة في المحتوى المؤلّد بالذكاء الاصطناعي. وإدراكًا لقيود أدوات الكشف عن الذكاء الاصطناعي الحالية، وخاصةً في الحالات التي يُحاكي فيها المؤلفون أنماط كتابة الذكاء الاصطناعي، تقترح الدراسة منهجية قائمةً على أنماط الكتابة الفردية. وباستخدام مقاييس مثل متوسط عدد الكلمات وطول الكلمة ونسب الكلمات الفريدة وتحليل المشاعر، إذ قاموا بتطوير نموذج نص تزييف عميق مخصص يحقق دقة ٩٩% ودقة ٩٧.٨٣%. علاوةً على ذلك، تم إنشاء نموذج نص تزييف عميق لتحديد الانحرافات عن الخصائص النصية المحددة للمؤلف، مما يُظهر دقة ٨٨.٩% ودقة ١٠٠% وتذكر ٨٩.٩%. تُظهر هذه النتائج إمكانية تحديد التزييف العميق على المستوى النصي، مما يجعل هذه الدراسة خطوةً أساسيةً في مكافحة المعلومات المضللة القائمة على النصوص. ونجحت الدراسة في تطوير نموذج للكشف عن نصوص التزييف العميق بناءً على أنماط الكتابة الفردية، محققة دقة عالية (تصل إلى ٩٩% في النموذج المخصص و ٨٨.٩% في نموذج تحديد الانحرافات). و هذا يثبت إمكانية الكشف عن التزييف على المستوى النصي بفعالية.

دراسة: مصطفى، محمود أحمد. (٢٠٢٥).

سعت هذه الدراسة إلى الكشف عن دور التربية الإعلامية الرقمية في تحصين وعي طلاب الإعلام التربوي بمخاطر تطبيقات التزييف العميق، معتمدة على نظرية دافع الحماية وثلاثة نماذج نظرية أخرى، والمنهج الوصفي الميداني من خلال استبانة طبقت على عينة عشوائية قوامها ٤٠٠ طالب إعلام تربوي. أظهرت النتائج تجانسًا في استجابات العينة تجاه اقتراحات تحصين الوعي، وجود ارتباط طردي بين مهارات التربية الإعلامية الرقمية ومستوى تحصين الوعي، وجود علاقة ارتباطية طردية بين كثافة استخدام التزييف العميق ومستوى الثقة الزائفة، عدم وجود فروق ذات دلالة إحصائية في الوعي بمخاطر التزييف العميق وفقًا للمتغيرات الديموغرافية، ووجود فرق دال إحصائيًا في مهارات التربية الإعلامية الرقمية وفقًا لمتغيرات النوع والجامعة والمستوى الاقتصادي. أظهرت الدراسة وجود ارتباط إيجابي بين مهارات التربية الإعلامية الرقمية ومستوى تحصين وعي الطلاب ضد مخاطر التزييف العميق. كما وجدت علاقة بين كثافة استخدام التزييف العميق ومستوى الثقة الزائفة لدى الطلاب، ووجود فروق في مهارات التربية الإعلامية الرقمية تعزى للنوع والجامعة والمستوى الاقتصادي.

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام – الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام – آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م – (عدد خاص)

دراسة: قادم، جميلة لصوان (٢٠٢٤).

تناولت الدراسة مسألة "التأثير السلبي لتقنية التزييف العميق على سمعة الشخصيات البارزة على منصات التواصل الاجتماعي"، إذ تمثل قضية بالغة الأهمية نظرًا للتطور السريع لهذه التقنية وتأثيرها المتزايد على الحياة الرقمية. وهدفت الدراسة إلى فهم كيفية تأثير التزييف العميق على صورة الشخصيات البارزة واستكشاف سبل التصدي لهذا التأثير، وذلك من خلال تحليل محتوى الفيديوهات المزيفة والكشف عن علامات التزييف العميق. توصلت النتائج إلى أن هذه الفيديوهات تؤثر بشكل كبير على سمعة الأفراد، مما يستدعي تعزيز الوعي الرقمي وتطوير آليات للمواجهة، بالإضافة إلى تبني سياسات أمان رقمي قوية وتشديد الرقابة على المحتوى المشبوه. ختمت الدراسة بالقول أن تسارع التحديات الرقمية يتطلب تحركًا عاجلاً للحفاظ على سمعة الأفراد المعرضين لهذه التقنية. توصلت الدراسة إلى أن فيديوهات التزييف العميق تؤثر سلبيًا وبشكل كبير على سمعة الشخصيات البارزة المستهدفة. وأكدت الحاجة لتعزيز الوعي الرقمي وتطوير آليات مواجهة وتبني سياسات أمان وتشديد الرقابة لمواجهة هذا التحدي المتسارع.

دراسة: بلواضح، بن ابراهيم (٢٠٢١).

استهدفت هذه الدراسة استكشاف استخدام تقنيات الذكاء الاصطناعي، وتحديدًا التزييف العميق، في عمليات الفبركة الإعلامية المتعلقة بالانتخابات الرئاسية الأمريكية لعام ٢٠٢٠، سعيًا منها للتعرف على كيفية توظيف هذه التقنيات في تزييف الصوت والصورة والحركات بدقة، وفهم مساهمتها في التضليل الإعلامي. اعتمدت الدراسة على منهج تحليل المضمون، باستخدام عينة قصدية من منصة تويتر وأداة استمارة تحليل المضمون لجمع البيانات. خلصت النتائج إلى أن تقنيات الذكاء الاصطناعي في التزييف العميق تستخدم بدقة في فبركة عناصر متعلقة بالشخصيات البارزة، مما يصعب كشفه بالعين المجردة، ويشكل خطرًا على مستقبل الإعلام في غياب قوانين تنظمها. الكلمات المفتاحية: الذكاء الاصطناعي، التزييف العميق، الفبركة الإعلامية، الانتخابات الرئاسية الأمريكية ٢٠٢٠. خلصت الدراسة إلى أن تقنيات التزييف العميق تُستخدم بدقة في فبركة محتوى يتعلق بالشخصيات البارزة (خاصة في سياق الانتخابات)، مما يجعل كشفها صعبًا بالعين المجردة ويشكل خطرًا على مستقبل الإعلام في غياب التنظيم القانوني.

المبحث الثاني: الإطار النظري والمفاهيمي

أثار انتشار تقنية التزييف العميق تحديات جسيمة للصحافة، مما أوجع المخاوف بشأن مصداقية وسائل الإعلام والتضليل الإعلامي وتآكل ثقة الجمهور. لقد برز الذكاء الاصطناعي كأداة محورية في مكافحة انتشار محتوى التزييف العميق إذ يوفر آليات للكشف والتحقق تتجاوز أساليب التحقق التقليدية من الحقائق. ومع ذلك لا تزال فعالية الذكاء الاصطناعي في تحديد الوسائط المزيفة موضع نقاش حيث تتطور تقنية التزييف العميق بوتيرة تتجاوز في كثير من الأحيان قدرات الكشف.

الدكتور محمد وهاب عبود

ويُرسى هذا القسم الأسس النظرية والمفاهيمية للدراسة كما يتناول بشكل نقدي النظريات ذات الصلة في دراسات الإعلام والذكاء الاصطناعي وأبحاث التضليل الإعلامي لتحديد دور الذكاء الاصطناعي في التحقق من التزييف العميق.

نظرية "ما بعد الحقيقة" والتضليل الإعلامي

يرتبط صعود تقنية التزييف العميق ارتباطاً وثيقاً بظاهرة سياسات ما بعد الحقيقة الأوسع حيث تغطي الروايات المشحونة عاطفياً والتحيزات الأيديولوجية على الحقائق الموضوعية بشكل متزايد إذ يتميز عصر ما بعد الحقيقة بالتلاعب المتعمد بالمعلومات لتشكيل التصور العام، غالباً على حساب نزاهة الصحافة (McIntyre 2018). حيث تُجسّد تقنية التزييف العميق هذه الحالة من خلال تمكينها من تصنيع محتوى سمعي بصري بسلسلة، مما يزيد من صعوبة التمييز بين الوسائط الأصلية والمصنوعة على الجمهور.

توضح النظرية بمزيد من التفصيل الآليات التي تُساهم بها تقنية التزييف العميق في نشر الأكاذيب إذ إن التضليل والتضليل الإعلامي والمعلومات الخاطئة كظواهر منفصلة ومتراصة، يلعب كل منها دوراً في تآكل الثقة في وسائل الإعلام التقليدية. كما تعمل تقنية التزييف العميق، لا سيما عند استخدامها لأغراض الدعاية السياسية أو تشويه السمعة، كمعلومات مضللة تشير إلى محتوى مُضلل عمداً مُصمم لخداع الجماهير. لذلك، يجب ألا تقتصر أدوات التحقق المُدارة بالذكاء الاصطناعي على كشف الوسائط المُصنوعة فحسب، بل يجب أن تُعالج أيضاً الأزمة المعرفية الأوسع التي تُشكلها المعلومات المُضللة المُنتَجة من التزييف العميق (Abbood 2024).

تفترض هذه النظرية أن الخطاب العام المعاصر يتسم بتراجع أهمية الحقائق الموضوعية أمام التأثيرات العاطفية والمعتقدات الشخصية، مما يسهل التلاعب بالرأي العام

وتساعد هذه النظرية في تأطير مشكلة التزييف العميق كأحد أبرز تجليات "ما بعد الحقيقة"، حيث يوفر أداة قوية لإنتاج "حقائق" زائفة مقنعة عاطفياً. بالتالي فإن الحاجة إلى أدوات الذكاء الاصطناعي للتحقق لا تنبع فقط من التحدي التقني، بل من الأزمة المعرفية الأوسع التي تهدد الثقة بالحقيقة نفسها. وتستفيد الدراسة من هذا الإطار لفهم دوافع استخدام التزييف العميق وتحليل أثره على تصور الجمهور للمعلومات، وتقييم ما إذا كانت أدوات الذكاء الاصطناعي قادرة على المساهمة في استعادة بعض الموضوعية في بيئة "ما بعد الحقيقة".

مصادقية وسائل الإعلام ونظرية حارس البوابة

تعتمد مصادقية الصحافة على قدرتها على توفير معلومات دقيقة وقابلة للتحقق في عصرٍ ينتشر فيه التلاعب الرقمي حيث تفترض نظرية حراسة البوابة التي صاغها كيرت ليفين (1943) ووسّعها لاحقاً شوميكور وفوس (2009)، أن الصحفيين والمؤسسات الإعلامية يعملون كمرشحات للمعلومات، مُحددين أي الروايات تصل إلى الجمهور (أبو الحمام & عزام 2017). ويُعطل ظهور

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام – الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام – آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م – (عدد خاص)

تقنية التزييف العميق وظيفة حراسة البوابة هذه من خلال إدخال محتوى مُضلل يُمكنه تجاوز الضمانات التحريرية التقليدية.

في هذا السياق، تُعد أدوات الكشف القائمة على الذكاء الاصطناعي امتدادًا للرقابة الصحفية، حيث تُوّمت عملية التحقق لمواجهة انتشار الوسائط المُصنّعة، فإن الاعتماد على التحقق الخوارزمي يُثير معضلات أخلاقية ومعرفية جديدة، لا سيما فيما يتعلق بغموض عملية صنع القرار باستخدام الذكاء الاصطناعي واحتمالية ظهور نتائج خاطئة (Tandoc et al., 2021). ويتطلب دمج الذكاء الاصطناعي في سير العمل الصحفي إعادة تقييم أطر مصداقية وسائل الإعلام بما يضمن تعزيز تقنيات التحقق لثقة الجمهور بدلًا من تقويضها.

وتفترض النظرية أن وسائل الإعلام والعاملين فيها يقومون بدور "حارس البوابة" حيث يختارون وينظمون تدفق المعلومات إلى الجمهور، مما يؤثر على ما يعرفه الناس وما يعتبرونه مهمًا.

ويمثل التزييف العميق تحديًا مباشرًا لوظيفة حراسة البوابة التقليدية، حيث يمكن للمحتوى المزيف أن يتجاوز آليات التحقق التحريرية وينتشر بسرعة، وتُعد أدوات الذكاء الاصطناعي للكشف بمثابة "حارس بوابة تكنولوجي" جديد أو أداة مساعدة للحارس البشري حيث تستفيد الدراسة من هذه النظرية لتحليل كيف يمكن لهذه الأدوات أن تدعم أو تعيد تشكيل دور حراسة البوابة في العصر الرقمي، وكيف يؤثر استخدامها أو عدم استخدامها على مصداقية المؤسسة الإعلامية. كما تساعد في فهم التوترات بين الأتمتة (AI) والحكم البشري الصحفي في عملية التحقق.

الحمية التكنولوجية وأخلاقيات الذكاء الاصطناعي

يُثير استخدام الذكاء الاصطناعي في الصحافة تساؤلات جوهرية حول الحمية التكنولوجية ومفادها أن التطورات التكنولوجية تُشكل الهياكل والسلوكيات المجتمعية بطرق حتمية، فالحمية التكنولوجية غالبًا ما تتجاهل السياقات الاجتماعية والسياسية التي تُدمج فيها التقنيات، مُفترضةً أن حلول الذكاء الاصطناعي تؤدي بطبيعتها إلى التقدم. لكن في حالة التحقق من التزييف العميق، غالبًا ما يُقَدَّم الذكاء الاصطناعي كحل شامل للتضليل الإعلامي إلا أنه يجب دراسة حدوده وتحيزاته بدقة (Feenberg ٢٠١٧).

تُعد الاعتبارات الأخلاقية في التحقق المعتمد على الذكاء الاصطناعي بالغة الأهمية، لا سيما فيما يتعلق بقضايا الشفافية والمساءلة والتحيز الخوارزمي إذ يُسلط Bender وآخرون (٢٠٢١) الضوء على مخاطر تعزيز أنظمة الذكاء الاصطناعي للتحيزات القائمة، حيث غالبًا ما تعكس مجموعات بيانات التدريب أوجه عدم المساواة، فإذا أخطأت أدوات الكشف بالذكاء الاصطناعي في تحديد هوية أفراد معينين بشكل غير متناسب، أو فشلت في مراعاة الاختلافات الثقافية في محتوى الوسائط، فإنها تُخاطر بتفاقم التضليل الإعلامي بدلًا من تخفيفه، لذا يجب أن يتضمن أي إطار نظري متين

الدكتور محمد وهاب عبود

للتحقق من التزييف العميق مبادئ أخلاقية للذكاء الاصطناعي، بما يضمن توافق التدخلات التكنولوجية مع معايير الإنصاف والدقة الصحفية.

وتشير الحتمية التكنولوجية إلى أن التكنولوجيا قوة مستقلة تشكل المجتمع وتحدد مسار تطوره، ويواجه هذا المنظور الانتقاد لأنه يتجاهل السياقات الاجتماعية وتأثير الاختيارات البشرية. إذ أن أخلاقيات الذكاء الاصطناعي تركز على مبادئ مثل الشفافية والمساءلة والإنصاف وتجنب التحيز في تصميم واستخدام أنظمة الذكاء الاصطناعي.

الحذر من هذه المقاربة يتمثل في النظر إلى أدوات الذكاء الاصطناعي كحل سحري أو حتمي لمشكلة التزييف العميق. وشكل استفادة الدراسة من هذا المنظور النقدي لتجنب تبني رؤية حتمية، وبدلاً من ذلك، تقييم أدوات الذكاء الاصطناعي ضمن سياقها الاجتماعي والأخلاقي والعملي حيث يساعد الإطار الأخلاقي في تحليل قضايا مثل (عدم قابلية تفسير قرارات الذكاء الاصطناعي)، واحتمالية التحيز الخوارزمي في أدوات الكشف، وتأثير ذلك على العدالة والمصادقية. بالتالي، توجه الدراسة نحو تقييم ليس فقط فعالية الأدوات، بل أيضاً مدى مسؤولية استخدامها وتأثيرها الاجتماعي، فالفرضية الضمنية هنا هي أن تبني أدوات الذكاء الاصطناعي يجب أن يسترشد بالمبادئ الأخلاقية لضمان توافقها مع معايير الإنصاف والدقة الصحفية.

تقنية التزييف العميق

تشير تقنية التزييف العميق إلى استخدام تقنيات التعلم الآلي المتقدمة، وخاصةً شبكات التوليد التنافسية لإنشاء صور فائقة الواقعية. فالمضامين الإعلامية المتلاعب بها تعمل من خلال تدريب شبكتين عصبيتين، إحداها تولّد محتوى مُزيّفًا والأخرى تُقيّم صحته مما يؤدي إلى عمليات تلاعب مُتطورة بشكل متزايد ومعقد (Goodfellow et al., 2014). فقد أثارت قدرة التزييف العميق على محاكاة أشخاص حقيقيين بدقة شبه مثالية مخاوف كبيرة بشأن إمكانية خداعها، لا سيما في السياقات السياسية والمالية والصحفية.

دور الذكاء الاصطناعي في كشف التزييف العميق

يعتمد التحقق من محتوى التزييف العميق على نماذج كشف مُحركة بالذكاء الاصطناعي تُحلل مؤشرات مُختلفة للمحتوى الإعلامي، إذ تستخدم هذه النماذج العديد من التقنيات (Mukta et al., 2023)، بما في ذلك:

- تحليل الوجه والحركة: الكشف عن التناقضات في تعابير الوجه وأنماط الرمش وحركات الرأس غير الطبيعية.
- التناقضات السمعية والبصرية: تحديد عدم التطابق بين حركات الشفاه وأنماط الكلام.
- فحص البيانات الوصفية: تحليل خصائص الملفات وآثار الضغط التي قد تشير إلى تلاعب.

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام - الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام - آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م - (عدد خاص)

- تحليل الشبكات العصبية: استخدام خوارزميات التعلم العميق لتصنيف الوسائط إلى حقيقية أو اصطناعية بناءً على مجموعات بيانات مُدربة.

على الرغم من أن هذه الأساليب المُعتمدة على الذكاء الاصطناعي قد أظهرت إمكانات كبيرة، إلا أنها ليست معصومة من الخطأ. فمع تطور تقنية التزييف العميق أصبحت الهجمات ضد أنظمة الكشف أكثر تعقيداً، مما يستلزم تحسينات مستمرة في نماذج التحقق بالذكاء الاصطناعي.

تحديات تطبيق التحقق القائم على الذكاء الاصطناعي

تواجه تقنية الكشف عن التزييف العميق المُعتمدة على الذكاء الاصطناعي العديد من التحديات العملية والأخلاقية (Rana et al., 2022) :

١. النتائج الخاطئة: خوارزميات الكشف عُرضة للأخطاء حيث تُصنّف أحياناً المحتوى الأصلي على أنه تزييف عميق بشكل خاطئ أو تفشل في التعرف على عمليات التلاعب المُتقدمة للغاية.

٢. عدم القدرة على التفسير: تعمل العديد من نماذج الذكاء الاصطناعي بشكل غير مفهوم حيث تُقدم تصنيفات دون شرح الأساس المنطقي لقراراتها.

٣. إمكانية الوصول والتكلفة: غالباً ما تكون أدوات كشف التزييف العميق عالية الجودة مملوكة ومكلفة، مما يحد من وصول المؤسسات الإعلامية الصغيرة إليها.

٤. ثقة الجمهور وإدراكه: قد يُقابل الاعتماد على الذكاء الاصطناعي للتحقق بالتشكيك، خاصةً إذا كانت نتائج أدوات الكشف غير متسقة.

تتطلب معالجة هذه التحديات نهجاً متعدد الجوانب يدمج التحقق باستخدام الذكاء الاصطناعي مع أساليب التحقق الصحفي التقليدية من الحقائق، مما يضمن أن تُعزز الحلول التكنولوجية الخبرة البشرية بدلاً من أن تحل محلها.

دمج التحقق باستخدام الذكاء الاصطناعي في الصحافة

لكي يكون الكشف عن التزييف العميق باستخدام الذكاء الاصطناعي فعالاً في الممارسة الصحفية، يجب على المؤسسات الإعلامية اعتماد بروتوكولات تحقق شاملة (Al-Khazraji et al., 2023). قد تشمل هذه:

- نماذج التحقق الهجينة: الجمع بين الكشف باستخدام الذكاء الاصطناعي والتحقق اليدوي من الحقائق من قِبل صحفيين مُدربين.

- الشفافية في استخدام الذكاء الاصطناعي: توضيح كيفية عمل أدوات التحقق باستخدام الذكاء الاصطناعي والاعتراف بحدودها.

الدكتور محمد وهاب عبود

- المبادئ التوجيهية الأخلاقية: وضع معايير صناعية للاستخدام المسؤول للذكاء الاصطناعي في الصحافة.

- التعاون مع باحثي الذكاء الاصطناعي: الدخول في شراكات بين المؤسسات الإعلامية ومطوري الذكاء الاصطناعي لتحسين منهجيات الكشف.

من خلال دمج التحقق باستخدام الذكاء الاصطناعي في سير العمل الصحفي مع الحفاظ على الرقابة الأخلاقية، يمكن للمؤسسات الإعلامية تعزيز قدرتها على مكافحة المعلومات المضللة دون المساس بالمصداقية.

توفر الأطر النظرية والمفاهيمية الموضحة أعلاه أساساً شاملاً لفهم دور الذكاء الاصطناعي في كشف التزييف العميق في الصحافة. كما يبرز دمج نظرية ما بعد الحقيقة، وأطر مصداقية وسائل الإعلام، وأخلاقيات الذكاء الاصطناعي، إمكانات الحلول التكنولوجية وحدودها في معالجة المعلومات المضللة. فمن الناحية النظرية، تضع الدراسة التحقق من التزييف العميق ضمن المشهد الأوسع للصحافة المعتمدة على الذكاء الاصطناعي، مشددة على ضرورة وجود نماذج تحقق هجينة توازن بين الكشف الخوارزمي والإشراف البشري. ومع استمرار تطور تقنية التزييف العميق، يجب أن تتناول الأبحاث المستقبلية بشكل نقدي التحديات الأخلاقية والمعرفية والعملية للتحقق القائم على الذكاء الاصطناعي لضمان بقاء الصحافة ركيزة موثوقة للمجتمعات الديمقراطية.

المبحث الثالث: استخدام أدوات الذكاء الاصطناعي للتحقق من التزييف العميق في

المضامين الإعلامية

في ظل الانتشار الواسع للمحتوى الرقمي وتزايد استخدام تقنيات الذكاء الاصطناعي في إنشاء وتعديل الوسائط الإعلامية، أصبحت ظاهرة التزييف العميق (Deepfake) تهديداً حقيقياً للمصداقية الإعلامية. تتمثل هذه الظاهرة في إنشاء مقاطع فيديو أو صوتيات مزيفة تظهر واقعية عالية، مما يصعب التمييز بينها وبين المحتوى الأصلي. هذا التحدي يتطلب تطوير أدوات وتقنيات فعالة للكشف عن هذه التزييفات.

تسعى هذه الدراسة إلى استعراض وتحليل فعالية أدوات الذكاء الاصطناعي المستخدمة في الكشف عن التزييف العميق، مع التركيز على مدى دقتها، سرعتها، وقدرتها على التكيف مع مختلف أنواع المحتوى الإعلامي. من خلال هذه النتائج، نهدف إلى تقديم توصيات عملية للمؤسسات الإعلامية والباحثين في مجال الإعلام الرقمي لتعزيز قدراتهم في مواجهة هذه الظاهرة.

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام – الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام – آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م – (عدد خاص)

❖ نتائج البحث

نتائج تحليل أدوات الذكاء الاصطناعي

- دقة الكشف

كشف التحليل أن Sensity AI و Deepware Scanner أظهرتا دقة كشف عالية، لا سيما عند تحليل مقاطع الفيديو المتلاعب بها سياسياً. حددت هذه الأدوات المحتوى المُلقق بفعالية، مُظهرةً نسبة نجاح عالية في التعرف على عناصر التزييف العميق، مثل تشوهات الوجه ومزامنة الشفاه غير الطبيعية وعدم اتساق الإضاءة.

فيما تفاوتت الدقة بناءً على نوع المحتوى، فبينما تم اكتشاف مقاطع الفيديو السياسية بدرجة عالية من الموثوقية، كانت مقاطع الفيديو المُضللة والتزييف العميق الذي يتضمن تركيبيًا صوتيًا أكثر صعوبة. ويُشير هذا إلى أن أدوات الذكاء الاصطناعي الحالية قد تكون مُحسنة للكشف عن التلاعبات البصرية، ولكنها تواجه صعوبة في التعامل مع التعديلات الأكثر دقة، مثل الأصوات المُولدة بالذكاء الاصطناعي أو المعلومات المُضللة المُعدلة القائمة على النصوص.

- الإيجابيات والسلبيات

من القيود الأساسية لكل من Sensity AI و Deepware Scanner عدم قدرتهما على تأكيد صحة المحتوى الحقيقي، فبينما نجحت كلتا الأدوات في تحديد المحتوى المُزيف على أنه مُزيف، إلا أنهما لم تمتلكا آلية للتحقق من المحتوى الحقيقي على أنه حقيقي. وهذا يُنشئ حالة قد تتطلب فيها وسائل الإعلام عمليات تحقق يدوية لتجنب تصنيف المحتوى الموثوق بشكل خاطئ.

الجدول ١: أداء اكتشاف الذكاء الاصطناعي لأنواع مختلفة من المحتوى

أداة الذكاء الاصطناعي	فيديوهات سياسية	معلومات علمية مضللة	التزييفات الصوتية العميقة	نتائج إيجابية	نتائج سلبية
Sensity AI	دقة عالية	دقة متوسطة	دقة منخفضة	دقة منخفضة	دقة متوسطة
Deepware Scanner	دقة عالية	دقة متوسطة	دقة منخفضة	دقة منخفضة	دقة متوسطة

يسلط معدل النتائج السلبية الضوء أيضًا على مصدر قلق كبير، فبعض محتوى التزييف العميق، وخاصة مقاطع الفيديو منخفضة الجودة أو المضغوطة، لم يتم تصنيفه على أنه مُتلاعب به ما يشير هذا إلى أن أدوات الذكاء الاصطناعي أكثر فعالية عند تحليل اللقطات عالية الدقة والواضحة، ولكنها قد تواجه صعوبة في التعامل مع المعلومات المضللة على وسائل التواصل الاجتماعي حيث غالبًا ما تكون مقاطع الفيديو مضغوطة.

وُجد أن Sensity AI أسرع بشكل طفيف من Deepware Scanner في تحليل المحتوى، وخاصةً لمقاطع الفيديو القصيرة التي تقل مدتها عن دقيقتين. ومع ذلك بالنسبة لملفات الوسائط الأطول، تطلبت كلتا الأداتين وقت معالجة إضافياً مع تأخيرات تصل إلى ٣٠ ثانية لكل تحليل.

على الرغم من أن هذا قد يبدو غير ذي أهمية، إلا أنه في الصحافة الآنية/العاجلة، يمكن أن تؤثر هذه التأخيرات على سرعة التحقق من الأخبار، وخاصةً في الأحداث السياسية سريعة الحركة. وهذا يُبرز الحاجة إلى نماذج ذكاء اصطناعي أكثر كفاءة يمكنها التحقق بسرعة من المحتوى دون التضحية بالدقة.

- الشفافية وسهولة التفسير

كان من أبرز المخاوف التي أشار إليها التحليل غياب الشفافية في كيفية توصيل أداتي Sensity AI و Deepware Scanner إلى استنتاجاتهما. فلم تقدم أيٌّ من الأداتين شرحاً مفصلاً لسبب تصنيف محتوى معين على أنه مُتلاعب به.

وبالنسبة للصحفيين يُؤدي هذا النقص في قابلية التفسير إلى فجوة ثقة، فبدون مبررات واضحة قد يتردد الإعلاميون في الاعتماد كلياً على أدوات الذكاء الاصطناعي للتحقق. وهذا يُعزز الحاجة إلى نماذج هجينة حيث تُكمل الخبرة البشرية قدرات الذكاء الاصطناعي في الكشف لضمان المصداقية.

- دقة الكشف المتفاوتة

كشفت نتائج تحليل أدوات الذكاء الاصطناعي Sensity AI و Deepware Scanner عن تباين ملحوظ في فعاليتها تبعاً لنوع المحتوى. فبينما أظهرت الأداتان دقة عالية في كشف مقاطع الفيديو ذات المحتوى السياسي المتلاعب به، وهو ما قد يعكس تركيز جهود التطوير على هذا النوع من التهديدات نظراً لأهميته وتأثيره المباشر على الرأي العام والعمليات الديمقراطية، إلا أن أدائها كان أقل إقناعاً في التعامل مع التزييف الصوتي والمعلومات العلمية المضللة، وهذا يشير، من وجهة نظر الباحث، إلى أن تقنيات كشف التزييف العميق الحالية قد لا تزال في مرحلة مبكرة نسبياً فيما يتعلق بالتعامل مع التلاعبات الصوتية الدقيقة أو التحقق من صحة الادعاءات العلمية المعقدة التي تتطلب فهماً للسياق يتجاوز مجرد التحليل البصري أو الصوتي.

ويشير هذا التفاوت في الأداء تساؤلات مهمة حول مدى جاهزية هذه الأدوات لتلبية احتياجات الصحفيين الشاملة. فالصحفي قد يواجه تزييفاً لا يقتصر على مقاطع الفيديو السياسية، بل يمتد ليشمل تسجيلات صوتية مفبركة لمسؤولين، أو مقاطع فيديو تزوج لمعلومات صحية أو علمية مغلوطة. فالاعتماد على أداة تتفوق في جانب وتضعف في جوانب أخرى قد يعطي إحساساً زائفاً بالأمان أو يؤدي إلى نشر محتوى مزيف دون قصد.

تتفق هذه النتيجة جزئياً مع ما أشارت إليه دراسة بلواضح وبن ابراهيم (2021) حول دقة استخدام التزييف العميق في فبركة محتوى مرتبط بالشخصيات البارزة (غالباً في سياقات سياسية)،

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام – الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام – آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م – (عدد خاص)

مما يجعل كشفه صعباً. لكن دراستنا تضيف بعداً آخر وهو ضعف الأدوات الحالية أمام أنواع أخرى من التزييف كالصوتي والعلمي. كما تتقاطع هذه النتيجة مع تحذيرات المتخصصين في الذكاء الاصطناعي الذين تمت مقابلتهم في هذه الدراسة، والذين أكدوا على أن تقنيات التزييف تتطور باستمرار وبوتيرة قد تسبق قدرات أدوات الكشف، مما يستدعي تحديثات مستمرة قد لا تغطي جميع أنواع التزييف بنفس الكفاءة. وهذا يختلف عن النتائج الواعدة التي توصلت إليها دراسة Rosca et al. (2025) في مجال كشف التزييف النصي، مما قد يشير إلى أن الكشف في الوسائط المختلفة (نص، صورة، صوت، فيديو) يواجه تحديات تقنية متفاوتة ويتطلب حلولاً متخصصة.

إن عدم قدرة الأداتين على تأكيد صحة المحتوى الحقيقي بشكل موثوق يمثل قيداً جوهرياً آخر، فمهمة الصحفي لا تقتصر على كشف الزيف، بل تشمل أيضاً تأكيد الحقيقة. إن فشل الأداة في تمييز المحتوى الأصلي قد يؤدي إلى إبطاء عملية النشر أو حتى التشكيك في محتوى صحيح، وهو ما عبر عنه الصحفيون الذين تمت مقابلتهم كأحد أسباب التردد في الاعتماد الكامل على هذه التقنيات. يرى الباحث أن هذا القصور يسلط الضوء على أهمية النموذج الهجين الذي لا يعتمد فقط على حكم الآلة، بل يدمجه مع التقييم والتحقق البشري كجزء أساسي من العملية. لذا يستوجب تطبيق نفس المنهجية في التعميق والتفسير والمقارنة لباقي النتائج، مثل تحديات الصحفيين وقيود الأدوات وانطباعات المتخصصين.

❖ نتائج المقابلات شبه المنظمة

• انطباعات الصحفيين

- تهديدات التزييف العميق للصحافة: أقرّ الغالبية- ١٥ صحفياً تمت مقابلتهم- بأن تقنية التزييف العميق تُشكّل تهديداً خطيراً للصحافة، لا سيما في المجالين السياسي والعلمي، كما ركزوا على دور التزييف العميق في التأثير على الانتخابات ونشر معلومات مضللة ومحاولات تشويه السمعة.

علق أحد الصحفيين بالقول:

"لا يكمن الخطر في مقاطع الفيديو المزيفة فحسب، بل في تآكل الثقة. فبمجرد أن يبدأ الجمهور بالاعتقاد بأن أي فيديو قد يكون مزيفاً سيشكك حتى في الأخبار الموثقة".

وتتوافق هذه الملاحظة مع مخاوف أوسع نطاقاً من أن تقنية التزييف العميق لا تخلق حقائق زائفة فحسب، بل تُقوّض الثقة في جميع المضامين حتى عندما تكون صادقة.

تحديات الكشف عن التزييف العميق: من التحديات الرئيسية التي يواجهها الصحفيون نقص الدعم المؤسسي لكشف التزييف العميق حيث أقرّ ١٠ صحفيين من أصل ١٥ صحفياً بأنهم تلقوا تدريباً محدوداً أو معدوماً على أدوات التحقق من التزييف العميق، وأن مؤسساتهم لم تُعط أولوية لكشف التزييف العميق.

الجدول ٢: تحديات الصحفيين في كشف التزييف العميق

التحديات	عدد الصحفيين
قلة الإلمام بأدوات كشف التزييف العميق	١٠ من أصل ١٥
قلة الاهتمام المؤسسي بالتحقق من التزييف العميق	١٠ من أصل ١٥
صعوبة التحقق من المعلومات العلمية المضللة	٧ من أصل ١٥
محدودية الوصول إلى خدمات الكشف المدفوعة	٥ من أصل ١٥

تشير هذه الاجابات إلى وجود فجوة خطيرة في جاهزية وسائل الإعلام، ما يشير إلى حاجة ملحة إلى برامج تدريبية واستثمار مؤسسي في استراتيجيات التحقق من صحة مقاطع الفيديو باستخدام الذكاء الاصطناعي.

موقف متخصصي الذكاء الاصطناعي

قيود الكشف باستخدام الذكاء الاصطناعي: أكد المتخصصون الخمسة في الذكاء الاصطناعي الذين تمت مقابلتهم أنه لا يمكن الوثوق تماما بأي أداة ذكاء اصطناعي للكشف عن مقاطع الفيديو المزيفة بعمق. وأشاروا إلى أنه مع تطور تقنية التزييف العميق يجب على أدوات الكشف تحديث نماذجها باستمرار، وهي عملية تتأخر عن التطور السريع لتقنيات التزييف العميق الجديدة.

علق أحد خبراء الذكاء الاصطناعي قائلا:

"نُشئو مقاطع الفيديو المزيفة العميقة دائماً ما يكونون متقدمين على أدوات الكشف. ففي اللحظة التي ندرّب فيها الذكاء الاصطناعي على التعرف على أسلوب تلاعب واحد، تظهر أساليب جديدة". وهذا يؤكد على التنافس الشديد بين مُنشئي مقاطع الفيديو المزيفة العميقة وتقنيات الكشف ما يتطلب البحث وإجراء تحديثات مستمرة في مجال التحقق باستخدام الذكاء الاصطناعي.

- توصيات للمؤسسات الإعلامية

- ١- اقترح متخصصو الذكاء الاصطناعي أن تتبنى المؤسسات الإعلامية سياسات منظمة لمكافحة محتوى التزييف العميق إذ تضمنت توصياتهم ما يلي:
- ٢- تعيين متخصصين في الذكاء الاصطناعي داخل غرف الأخبار للإشراف على عمليات التحقق.
- ٣- عقد ورش عمل تدريبية للصحفيين حول كشف التزييف العميق باستخدام الذكاء الاصطناعي.

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام - الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام - آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م - (عدد خاص)

٤- دمج التحقق باستخدام الذكاء الاصطناعي مع أساليب التحقق التقليدية لتحقيق دقة أعلى.

٥- تطوير شراكات مع شركات الذكاء الاصطناعي لتحسين الوصول إلى أدوات الكشف المتقدمة.

وكان من أبرز المخاوف التي طرحها المتخصصون هي أن معظم أدوات الكشف باستخدام الذكاء الاصطناعي خدمات مدفوعة، ما يضع ضغوطاً مالية على المؤسسات الإخبارية. وهذا يؤكد الحاجة إلى نماذج ذكاء اصطناعي مفتوحة المصدر لجعل كشف التزييف العميق أسهل للصحفيين حول العالم.

❖ تلخيص النتائج والتوصيات

خلصت الدراسة الحالية إلى مجموعة من النتائج المحورية التي تلقي الضوء على واقع استخدام أدوات الذكاء الاصطناعي في التحقق من التزييف العميق في المضامين الإعلامية:

١- فعالية مشروطة وليست مطلقة لأدوات الكشف: أظهرت أدوات الذكاء الاصطناعي التي تم اختبارها Sensity AI و Deepware Scanner قدرة ملحوظة على كشف محتوى التزييف العميق في سياقات معينة، خاصة مقاطع الفيديو ذات الطابع السياسي. لكن فعاليتها تقل بشكل واضح عند التعامل مع أنواع أخرى من المحتوى مثل التلاعبات الصوتية أو المعلومات العلمية المضللة. وهذا يعني أن الاعتماد عليها كحل وحيد وشامل لا يزال محفوفاً بالمخاطر، حيث قد تفشل في كشف أنواع مهمة من التزييف أو تعطي نتائج غير دقيقة. بالإضافة إلى ذلك، فإن عدم قدرتها على تأكيد صحة المحتوى الحقيقي يضيف طبقة أخرى من التعقيد ويتطلب تدخلاً بشرياً.

٢- فجوة في الجاهزية لدى الصحفيين والمؤسسات: كشفت المقابلات عن نقص كبير في التدريب والدعم المؤسسي الموجه للصحفيين فيما يتعلق بتقنيات وأدوات كشف التزييف العميق. فغالبية الصحفيين المشاركين أقرّوا بقلّة إلمامهم بهذه الأدوات وعدم وجود أولوية مؤسسية واضحة لهذا الجانب، ويشير هذا بوضوح إلى وجود فجوة بين التهديد المتزايد للتزييف العميق وبين استعداد غرف الأخبار لمواجهته بفعالية باستخدام الأدوات المتاحة، مما يجعل الصحفيين عرضة للخطأ أو التردد.

الدكتور محمد وهاب عبود

٣- الحاجة إلى نماذج تحقق هجينة: أكدت نتائج تحليل الأدوات والمقابلات مع الصحفيين والمتخصصين على أن الاعتماد الحصري على الذكاء الاصطناعي للتحقق ليس كافيا ولا موثوقا بشكل كامل في الوقت الراهن. ف القيود المتعلقة بالدقة المتفاوتة، ونقص الشفافية وقابلية التفسير في عمل الخوارزميات وتطور تقنيات التزييف المستمر، كلها عوامل تجعل النموذج الهجين الذي يجمع بين قدرات الكشف الآلي والتحقق والتقييم البشري من قبل صحفيين مدربين هو النهج الأكثر فعالية وواقعية لضمان الدقة والمصداقية.

٤- سباق تقني مستمر: يتطور التزييف العميق بسرعة تفوق في كثير من الأحيان وتيرة تطوير وتحديث أدوات الكشف. وهذا يعني أن أي حل يعتمد على الذكاء الاصطناعي يتطلب التزاما مستمرا بالبحث والتطوير والتحديث لمواكبة التهديدات الجديدة، وهو ما يمثل تحديا تقنيا وماديا مستمرا للمطورين والمؤسسات الإعلامية على حد سواء.

❖ التوصيات

بناء على النتائج التي توصلت إليها الدراسة، تقدم الدراسة التوصيات التالية:

○ للمؤسسات الأكاديمية والتدريبية الإعلامية ونقابات الصحفيين: تطوير وتضمين برامج متخصصة في "محو الأمية الرقمية والذكاء الاصطناعي" ضمن المناهج الدراسية والدورات التدريبية للصحفيين الحاليين والجدد. كما يجب أن تركز هذه البرامج على فهم تقنيات التزييف العميق وكيفية عمل أدوات الكشف وتطوير المهارات النقدية اللازمة لتفسير نتائج التحقق الآلي وتقييم المحتوى الرقمي بشكل شامل.

○ لمراكز الأبحاث والمطورين في مجال الذكاء الاصطناعي والمؤسسات الداعمة للصحافة (مثل المنظمات غير الربحية والمبادرات الدولية) الدعوة إلى تكثيف الجهود نحو تطوير وإتاحة أدوات كشف للتزييف العميق تكون مفتوحة المصدر أو بتكلفة معقولة. وهذا من شأنه تقليل العوائق المالية أمام المؤسسات الإخبارية، خاصة الصغيرة والمستقلة وتمكينها من الوصول إلى تقنيات التحقق اللازمة.

○ للمؤسسات الإعلامية ومنظمات تقنين المعايير المهنية: العمل على وضع واعتماد إرشادات تحقق موحدة وواضحة تدمج استخدام أدوات الذكاء الاصطناعي ضمن عمليات التحقق الصحفي التقليدية (النموذج الهجين). ينبغي أن تحدد هذه الإرشادات متى وكيف يتم استخدام الأدوات وكيفية التعامل مع نتائجها ودور الصحفي في عملية التحقق النهائية.

وقائع المؤتمر العلمي التاسع (الدولي الثالث) لكلية الإعلام – الجامعة العراقية الموسوم: الذكاء الاصطناعي في

الإعلام – آفاق الابتكار وتحديات الحوار الثقافي للمدة من ٢٣-٢٤/٤/٢٠٢٥م – (عدد خاص)

- لشركات التكنولوجيا الكبرى (بما في ذلك مطورو الذكاء الاصطناعي ومنصات التواصل الاجتماعي) والمؤسسات الإعلامية: تعزيز التعاون والشراكات لتبادل الخبرات والبيانات وتطوير قدرات الكشف عن التزييف العميق بشكل أسرع وأكثر فعالية. ويمكن أن يشمل ذلك تطوير واجهات برمجة تطبيقات تسهل دمج أدوات الكشف في أنظمة إدارة المحتوى بغرف الأخبار.
- لصناع السياسات والهيئات التشريعية: النظر في وضع أطر تنظيمية وقانونية تعالج إنتاج ونشر التزييف العميق الضار مع مراعاة التوازن بين مكافحة التضليل وحماية حرية التعبير والابتكار التكنولوجي.

ختاماً، تُبرز الدراسة التهديد المتزايد للتزييف العميق والحاجة الملحة لاستراتيجيات تحقق تعتمد على الذكاء الاصطناعي. وعلى الرغم من أن أدوات الذكاء الاصطناعي تؤدي دوراً مهماً وتقدم مساعدة قيمة، إلا أنها لا تغني عن الخبرة البشرية، فالنهج المشترك يحقق كامل الاستفادة من تقنية الذكاء الاصطناعي، كما أن الحفاظ على التدقيق الصحفي التقليدي يُعد الحل الأكثر فعالية لمكافحة التضليل الإعلامي المرتبط بالتزييف العميق.

❖ المراجع العربية

- ابراهيم، سامح محمد محمد السيد. (٢٠٢٥). المخاطر الأمنية والمجتمعية للتزييف العميق وآليات المواجهة. المجلة القانونية، ٢٣(٧)، ٣٩٥٧-٤٠١٠.
- أبو الحمام، عزام. (٢٠١٧). نظرية حارس البوابة الإعلامية في ظل البيئة الجديدة لتكنولوجيا الاتصال. مجلة الرسالة للدراسات والبحوث الإنسانية، ٢(٣)، ٢٦٥-٢٩٢.
- بلواضح، بن ابراهيم. استخدام تقنية الذكاء الاصطناعي (التزييف العميق) في الفبركة الإعلامية دراسة تحليلية لعينة من الفيديوهات المنشورة على منصة تويتر الانتخابات الرئاسية الأمريكية لسنة ٢٠٢٠ نموذجاً (Doctoral dissertation, univ-ouargla).
- قادم، جميلة لصوان. (٢٠٢٤). التأثير السلبي لتقنية التزييف العميق على سمعة الشخصيات البارزة على منصات التواصل الاجتماعي دراسة تحليلية على عينة من الفيديوهات المفبركة The Negative Impact of Deepfake Technology on the Reputation of Prominent Figures on Social Media Platforms: An Analytical Study on a Sample of Fabricated Videos. مجلة العلوم وفاق المعارف، ٤(١)، ٥١٠-٥٣٢.
- مصطفى، محمود أحمد. (٢٠٢٥). التربية الإعلامية الرقمية وعلاقتها بمستوى تحصين وعي طلاب الإعلام التربوي بمخاطر تطبيقات التزييف العميق في إطار نظرية دافع الحماية. مجلة البحوث الإعلامية، ٧٣(١)، ٥٨١-٦٧٢.

- Abbood, M. (2024). Otherness of individual existence in the context of fake reality. *Izvestiya of Saratov University. Philology. Journalism*, 24(2), 212–219. <https://doi.org/10.18500/1817-7115-2024-24-2-212-219>
- Al-Khazraji, S. H., Saleh, H. H., Khalid, A. I., & Mishkhal, I. A. (2023). Impact of deepfake technology on social media: Detection, misinformation and societal implications. *The Eurasia Proceedings of Science Technology Engineering and Mathematics*, 23, 429-441.
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610-623.
- Berger, A. A. (2018). *Media and communication research methods: An introduction to qualitative and quantitative approaches*. Sage Publications.
- Cazzamatta, R., & Sarisakaloğlu, A. (2025). Mapping Global Emerging Scholarly Research and Practices of AI-supported Fact-Checking Tools in Journalism. *Journalism Practice*, 1-23.
- Feenberg, A. (2017). *Technosystem: The social life of reason*. Harvard University Press.
- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., & Bengio, Y. (2014). Generative adversarial nets. *Advances in Neural Information Processing Systems*, 27, 2672-2680.
- McIntyre, L. (2018). *Post-truth*. MIT Press.
- Mukta, M. S. H., Ahmad, J., Raiaan, M. A. K., Islam, S., Azam, S., Ali, M. E., & Jonkman, M. (2023). An investigation of the effectiveness of deepfake models and tools. *Journal of sensor and actuator networks*, 12(4), 61.
- Rana, M. S., Nobi, M. N., Murali, B., & Sung, A. H. (2022). Deepfake detection: A systematic literature review. *IEEE access*, 10, 25494-25513.
- Rosca, C. M., Stancu, A., & Iovanovici, E. M. (2025). The New Paradigm of Deepfake Detection at the Text Level. *Applied Sciences*, 15(5), 2560.
- Tandoc, E. C., Jenkins, J., & Craft, S. (2019). Fake news as a critical incident in journalism. *Journalism Studies*, 22(7), 927-944.