

تطبيقات وخوارزميات الذكاء الاصطناعي و جرائم وإنتهاكات التزوير والتزييف العميق

د. ياسر الملك أحمد سليمان^١*

^١ أستاذ مشارك، قسم علوم الحاسوب نظم المعلومات، الجامعة التكنولوجية ، الخرطوم ، السودان

* dr.yaserking359@hotmail.com

Artificial Intelligence Applications and Algorithms in Crimes and Violations Involving Forgery, Fraud, and Deepfakes

Dr. Yasser Al-Malik Ahmed Suleiman¹

¹ Associate Professor, Department of Computer Science and Information Systems, Technological University, Khartoum, Sudan.

Received: 18/04/2026; Revised: 13/07/2026; Accepted: 01/09/2026; Published: 03/09/2026

الملخص العربي:

يشهد الذكاء الاصطناعي تطورًا متسارعًا أدى إلى توسيع نطاق استخداماته في مختلف المجالات، إلا أن هذا التطور صاحبه تصاعد في مخاطر إساءة استخدام تطبيقاته وخوارزمياته في ارتكاب الجرائم الإلكترونية وعمليات التزوير وانتحال الهوية والتضليل الرقمي. ويُعد التزييف العميق (Deepfake) من أبرز صور هذه المخاطر، لما تتيحه تقنيات الذكاء الاصطناعي والتعلم العميق من إمكانية إنشاء أو تعديل الصور ومقاطع الفيديو والتسجيلات الصوتية بصورة شديدة الواقعية، تجعل التمييز بين المحتوى الحقيقي والمحتوى المقلد أمرًا بالغ الصعوبة.

يهدف هذا البحث إلى تحليل الأسس التقنية لتطبيقات وخوارزميات الذكاء الاصطناعي المستخدمة في إنتاج التزييف العميق، ودراسة أبرز صور الجرائم والانتهاكات المرتبطة بها وتأثيراتها على الأفراد والمؤسسات والحكومات ووسائل الإعلام وثقة الجمهور بالمعلومات. كما يستعرض البحث مجموعة من أساليب وتقنيات الكشف عن المحتوى المزيف، بما في ذلك تحليل الوسائط الرقمية والبيانات الوصفية، واكتشاف التناقضات بين الصوت والصورة وحركات الوجه، والاستفادة من خوارزميات الذكاء الاصطناعي والتعلم العميق في التحقق من أصالة المحتوى.

اعتمد البحث على المنهج الوصفي والتحليلي في دراسة طبيعة التزييف العميق وآليات إنتاجه واستخداماته الإجرامية، وتحليل الأدوات والتقنيات المتاحة للكشف عنه والحد من مخاطره. وتوصل البحث إلى أن التطور المستمر في تقنيات التزييف العميق يمثل تحديًا متزايدًا للأمن السيبراني وسلامة المعلومات، وأن مواجهة هذه التهديدات تتطلب الجمع بين الحلول التقنية المتقدمة، وآليات التحقق الرقمي، والتشريعات والإجراءات التنظيمية، إلى جانب رفع مستوى الوعي والتثقيف الرقمي. ويوصي البحث بتعزيز التعاون بين الحكومات والمؤسسات والجامعات ومراكز الأبحاث، وتطوير قدرات متخصصة في كشف جرائم الذكاء الاصطناعي والتزييف العميق، بما يساهم في حماية البيانات والمعلومات وتعزيز الثقة في المحتوى الرقمي.

الكلمات المفتاحية: لذكاء الاصطناعي، خوارزميات الذكاء الاصطناعي، التزييف العميق، التزوير الرقمي، الجرائم الإلكترونية، الأمن السيبراني، كشف التزييف.

Abstract

Artificial intelligence (AI) has witnessed rapid technological advancement, expanding its applications across various fields. However, this development has also increased the risks associated with the misuse of AI applications and algorithms in cybercrime, forgery, identity impersonation, and digital misinformation. Deepfake technology represents one of the most significant manifestations of these risks, as AI and deep learning techniques can be used to create or manipulate images, videos, and audio recordings with a high degree of realism, making it increasingly difficult to distinguish authentic content from fabricated media.

This study aims to analyze the technical foundations of AI applications and algorithms used to generate deepfakes and to examine the major forms of crimes and violations associated with their misuse and their impact on individuals, institutions, governments, media organizations, and public trust in information. The study also reviews various methods and techniques for detecting manipulated content, including digital media and metadata analysis, the detection of

inconsistencies between audio and visual components and facial movements, and the use of AI and deep learning algorithms to verify the authenticity of digital content.

The study adopts a descriptive and analytical methodology to examine the nature of deepfake technology, its generation mechanisms, its criminal applications, and the available tools and techniques for detecting and mitigating its risks. The findings indicate that the continuous advancement of deepfake technologies poses an increasing challenge to cybersecurity and information integrity. Addressing these threats requires an integrated approach combining advanced technological solutions, digital verification mechanisms, appropriate legislation and regulatory measures, and increased public awareness and digital literacy. The study recommends strengthening cooperation among governments, institutions, universities, and research centers and developing specialized capabilities for detecting AI-related crimes and deepfakes to protect data and information and enhance trust in digital content.

Keywords: Artificial Intelligence, Artificial Intelligence Algorithms, Deepfakes, Digital Forgery, Cybercrime, Cybersecurity, Deepfake Detection..

١. المقدمة

شهد الذكاء الاصطناعي تطوراً كبيراً في تقنياته وبشكل سريع، مع ظهور الذكاء التوليدي أصبحت محتويات التقنية تبدو وكأنها من صنع البشر وأضحت تعوض الإنسان في بعض الأحيان، ما يثير المخاوف من إساءة استخدام الذكاء الاصطناعي لأغراض خبيثة مثل التضليل والإحتيال والخداع والتأثير على المجتمع. السنوات الأخيرة بدأت ظهور تقنيات التزييف العميق التي يمكن من خلالها تعديل الصور والفيديوهات بشكل يجعل من الصعب التمييز بين المحتوى الأصلي والمعدل وما يسمى بالتزييف العميق هو إحدى أبرز التقنيات التي تعتمد على الذكاء الاصطناعي لإنشاء أو تعديل الفيديوهات والصور بطريقة تجعلها تبدو حقيقية، تستخدم هذه التقنية بشكل رئيسي الشبكات العصبية التوليدية ومجموعه من الأدوات. أن تقنيات الأعمال المزيفة تكون متطورة ومعقدة للغاية وأصبحت متاحة بشكل متزايد، مما يجعل برامج الكشف عنها والقوانين الصارطة لها صعب ويحتاج الي توعية وتطوير تقنيات للحد من مثل هذه الجرائم والإنتهاكات. يعد الذكاء الاصطناعي بشكل عام من الأدوات التي يمكن إساءة استخدامها، مما يؤدي إلى مشاكل كثيرة، مثل التحيز، وإنتهاكات الخصوصية، ويؤدي ظهور تقنيات التزييف العميق إلى تضخيم هذه المخاطر، حيث تعتبر تقنيات التزييف العميق وسائط إصطناعية شديدة الواقعية تم إنشاؤها باستخدام تقنيات التعلم العميق، التي تتلاعب بوسائط الصوت أو الفيديو أو أي محتوى رقمي آخر، مما يجعل من الصعب التمييز بين المواد الأصلية والمزيفة.

٢. مشكلة البحث

تكمن مشكله البحث الأساسية في التحديات التي تواجه العالم وتتسبب في دمار ومشاكل لعدد من الدول والشعوب من خلال الإستخدام السلبي لتقنيات وخوارزميات الذكاء الإصطناعي فقد أصبح التزييف العميق أخطر أنواع هذه التقنيات التي تعتمد على خوارزميات التعلم العميق لإنشاء مقاطع فيديو أو تسجيلات صوتية تبدو وكأنها حقيقية و يمكن إستخدامها لتقليد أصوات وأشكال الأشخاص بدقة شديدة، مما يجعل المحتوى المزيف صعب التمييز عن الواقع، تعتمد التقنية على تدريب الشبكات العصبية على كميات هائلة من البيانات التي تتضمن الأصوات والصور والفيديوهات للأشخاص المستهدفين.

يتسبب التزييف في مجموعة من المشاكل يلخص الباحثون أهمها في النقاط التالية :

(١) إستخدام التزييف لانتحال صوت مدير في إحدى الشركات أو البنوك ، وتحويل أموال إلى حساب المهاجم، وهي مثال على الطريقة التي يمكن أن يستغل بها المهاجمون هذه التقنية للحصول على أموال أو معلومات حساسة.

(٢) إستخدام فيديوهات مزيفة لإقناع الأفراد بالقيام بمعاملات إحتيالية.

(٣) إستخدام فيديوهات مزيفة لزعماء و رؤساء دول او شخصيات يمكن أن تكون للتأثير على الشعوب والدول وإثارة حرب ونزاعات لإهداف معينة.

٣. تساؤلات البحث:

- (١) ما هي الخوارزميات والتقنيات المستخدمة في إكتشاف التزييف العميق والتحليل لبيانات التعريف الخاصة بالملف وضمان عدم التعديل والتأكد من المصادقية؟
- (٢) كيف تتم الانتهاكات التزوير والتزييف العميق وهل هناك تأثيرات للتزييف العميق؟
- (٣) هل توجد تدابير تقنية وإجرائية لإيجاد حلول للحد من الإنتهاكات؟
- (٤) هل تساهم التدابير التقنية في حماية الذكاء الاصطناعي؟

٤. الهدف من البحث:

يهدف للتعرف علي جرائم وإنتهاكات الذكاء الاصطناعي والطرق لمكافحة الجرائم والإنتهاكات وتوضيح التزييف العميق والتقنيات والأدوات التي يمكن إستخدامها للكشف عنه وتمكين مصادقة الوسائط الحقيقية وكيفية إستخدام التعلم الآلي كتقنية من تقنيات الذكاء الاصطناعي والخوارزميات الحديثة التي تهدف للكشف وتحديد الوسائط المزيفة دون الحاجة الي مقارنتها بالوسائط الأصلية غير المعدلة.

٥. أهمية البحث:

تتمثل في مجموعة نقاط هامة ظهرت حديثاً وهي وإنتهاكات الذكاء الاصطناعي من خلال التزوير والتزييف و إنتحال الهوية و نشر المعلومات مضلله و التزوير لفيدوهات ومكالمات وتسجيلات تضر بالأفراد والمؤسسات والدول وشهدت الشركات والبنوك وكالات الإعلام لمثل هذه الجرائم يتم توضيحها في البحث والسعي في معرفة الأسباب التي أدت لذلك والوقاية من الإنتهاكات ، ومثل هذه الجرائم ، وضمان أمن البيانات والمعلومات وتقاديا لها في المستقبل للتأثير العالمي الكبير.

٦. منهجية البحث:

إتبع الباحثون في منهجية البحث المنهج الوصفي والتحليلي المنهج الوصفي في الشرح والتوصيف للمشكلة ومن ثم التحليلي لشرح المشكلة والمساهمة في كيفية إيجاد طرق عديده تساعد في الحلول والتوصل لنتائج البحث الضرورية.

٧. حدود البحث:

أقتصر البحث على التعرف لتطبيقات وخوارزميات الذكاء الاصطناعي و جرائم وإنتهاكات التزوير والتزييف العميق وتأثيره على المؤسسات الحكومية الكبيرة في الدول المتقدمة والثغرات ونقاط الضعف وكيفية إجراء التدابير المناسبة.

- حدود زمانية: أجريت الدراسة في بداية مارس ٢٠٢٥.
- حدود مكانية: إهتم البحث بالمؤسسات الحكومية والوزارات المهمة بالدول العظمي.

أ. المحور الأول: الجانب النظري للبحث:

(١) مقدمة الذكاء الاصطناعي:

الذكاء الاصطناعي تم إستخدامه بشكل متزايد في عالم الجرائم الإلكترونية من قبل المهاجمين والأدوات والتقنيات التي تعتمد على الذكاء الاصطناعي توفر فرصاً لهجمات أكثر تعقيداً وصعوبة في الإكتشاف مقارنة بالهجمات

التقليدية، كما يمكن استخدام الذكاء الاصطناعي لتحليل الأنماط السلوكية، مما يساعد المهاجمين على معرفة أكثر النقاط ضعفًا في الأنظمة الأمنية" واستخدام تقنيات محاكاة أساليب تواصل محددة بناءً على تحليلات البيانات الشخصية للأفراد. (دهشان ، ٢٠٢٣) وباستخدام تقنيات التعلم العميق والشبكات العصبية، يمكن إنشاء مقاطع فيديو صوت تحاكي شخصيات مشهورة أو حتى أفراد عاديين في مواقف لم تحدث في الواقع في عصرنا الرقمي الحالي، أصبح الذكاء الاصطناعي (AI) أداة قوية تتسم بتعدد استخداماتها في مجالات عدة، بدءًا من تحسين الإنتاجية وزيادة الكفاءة وصولاً إلى تهديدات غير مرئية ترزع الأمن الرقمي. ومن بين أبرز المجالات التي أظهر فيها الذكاء الاصطناعي تأثيره هو الجرائم الإلكترونية والتزييف العميق (Deepfake)، والتي باتت تهدد الأفراد والمؤسسات والدول.

(٢) الذكاء الاصطناعي و الجرائم الإلكترونية:

أصبحت تقنيات و أدوات الذكاء الاصطناعي تؤثر بشكل كبير على الإقتصاد والطب والإعلام وجميع النواحي الحياتية ، و بشكل عام يمكن إساءة استخدام هذه الأدوات مما يؤدي إلى مشاكل كثيرة، مثل: التحيز، وانتهاكات الخصوصية، ويؤدي ظهور تقنيات التزييف العميق إلى تضخيم هذه المخاطر، حيث تعتبر تقنيات التزييف العميق وسائط إصطناعية شديدة الواقعية تم إنشاؤها باستخدام تقنيات التعلم العميق فمعالجة هذه القضايا أمر ضروري للحفاظ على سلامة المعلومات، وحماية سمعة الأفراد، وضمان السلامة العامة، إذ تؤدي صعوبة إكتشاف وتحديد أساليب التزييف العميق وتقنياته نظراً إلى أن تقنيات التزييف العميق أصبحت أكثر تقدماً ويمكن الوصول إليها، فقد تصاعدت المخاطر المرتبطة بها عالمياً، من خلال البحث يوضح الباحثون كيف يستخدم الذكاء الاصطناعي في ارتكاب الجرائم الإلكترونية، وكذلك تأثيره على إنتاج محتوى زائف وتوليد فيديوهات مزورة ومزيفة تشبه الواقع الحقيقي. وإستعراض يحي إبراهيم دهشان. (٢٠٢٣) سلسلة من الخطوات الآلية التي تساعد في تقديم التعليمات البرمجية بسرعة أكبر وأكثر أماناً للتعرف على التزييف العميق وكيفية الإكتشاف وإستراتيجيات التخفيف من المخاطر لكل مرحلة . للتخفيف من المخاطر المرتبطة بالتزييف العميق، يجب على مؤسسات الدولة تنفيذ أمثلة نموذجية حول كيفية إستفادة هذه القطاعات من التطبيقات المبتكرة والقيمة للتزييف العميق تستعرض في شكل مراحل تم التفصيل لها من خلال الباحثون كالاتي:

- تقنيات التزييف ويشتمل على كل الانواع (صور، صوت، فيديو).
- استخدام أدوات الذكاء الاصطناعي والتحقق من التزييف.
- جانب تطبيقي يوضح كيفية الإستخدام.

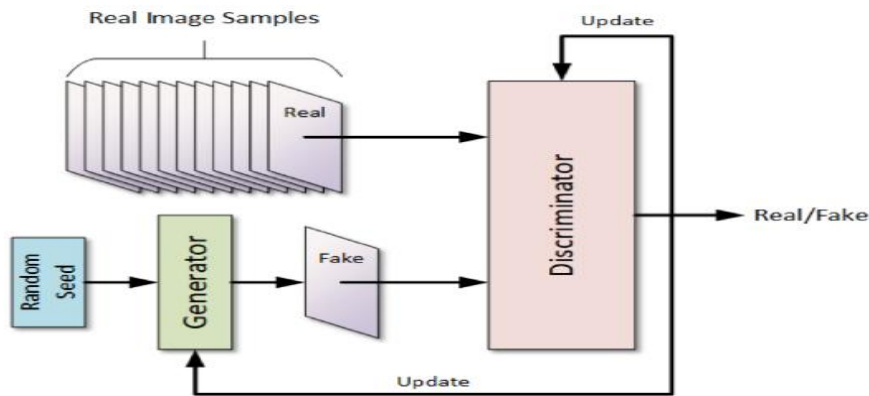
(٣) تقنيات التزييف الصوتي العميق:

إنشاء التزييفات الصوتية العميقة باستخدام شبكات الخصومة التوليدية (GANs) وخوارزميات التعلم العميق التي يمكنها تحليل وتكرار الفروق الدقيقة في أنماط الكلام البشري. ومن خلال تدريب هذه النماذج على مجموعات بيانات كبيرة من العينات الصوتية، يمكن لأنظمة الذكاء الاصطناعي تجميع تسجيلات صوتية واقعية للغاية تحاكي صوت فرد معين، مع التقدم في معالجة اللغة الطبيعية وتركيب الكلام، أصبح من الصعب بشكل متزايد التمييز بين المحتوى الصوتي الأصلي والمحتوى الصوتي الذي تم التلاعب به ، يشكل هذا تحدياً كبيراً لكشف ومكافحة إنتشار التزييف الصوتي العميق في المشهد الرقمي. ، إن انتشار الملفات الصوتية المزيفة يثير مخاوف جدية بشأن إدارة السمعة والخصوصية والأمن ، يمكن استهداف الزعماء والرؤساء بتصريح مزيف يؤدي إلى ردة فعل ، قد يقع المشاهير والشخصيات العامة ضحية التزييف الصوتي العميق أو إن إستغلال التزييف الصوتي العميق لأغراض الابتزاز أو الدعاية يؤكد بشكل أكبر على الحاجة لحماية المجتمع من التزييف الرقمي. (خليفة ، ٢٠١٩).

- توليد الصور والفيديوهات باستخدام الذكاء الاصطناعي (AI-Generated Content) : تم تطوير نماذج ذكاء إصطناعي قادرة على إنشاء محتوى مرئي بالكامل هذه النماذج تعتمد على تعلم الأنماط البصرية من مجموعات ضخمة من البيانات، وتمكّنها من توليد صور وفيديوهات واقعية بناءً على أوصاف نصية أو معطيات معينة و نماذج (Generative Models) مثل GANs و VAEs تستخدم هذه النماذج

في توليد محتوى مرئي جديد، مثل صور لأشخاص أو أماكن غير موجودة في الواقع على سبيل المثال، يمكن للنماذج مثل StyleGAN إنشاء صور لأشخاص خياليين تبدو حقيقية جداً رغم أنها غير موجودة في الواقع.

- **تحسين الجودة والتفاصيل:** يمكن استخدام الذكاء الاصطناعي لتحسين جودة الفيديوهات أو الصور المعدلة، مما يجعلها تبدو أكثر احترافية وأصيلة، حتى لو كانت قد خضعت للتعديل. يمكن إستبدال وجه شخص في فيديو بأخر باستخدام تقنيات التزييف العميق، كما يحدث في التطبيقات مثل "DeepFaceLab" أو "FaceSwap". يمكن تعديل تعابير الوجه: يمكن تعديل تعابير الوجه لجعل الشخص يظهر وكأنه يقول شيئاً لم يقله في الواقع، مما يتيح نشر معلومات مضللة أو التشهير. إعادة تشكيل الصوت: يمكن أيضاً استخدام الذكاء الاصطناعي لتقليد الأصوات البشرية، مما يجعل التزييف العميق أكثر إقناعاً.



الشكل (١) يوضح التزييف العميق للصور - (ايهاب خليفة ٢٠١٩)

تعتبر الخلايا العصبية بمثابة وحدات برمجية تقوم بوظيفة إستقبال وإرسال البيانات التي يتم تغذية النظام بها، ومن خلال وضع هذه الخلايا في صورة متضادة مع بعضها يساعدها ذلك في توليد نماذج إحصائية ومعلوماتية أكثر مكثت من رسم صور أكثر كفاءة ودقة، ويحدث ذلك من خلال وضع خلية عصبية في مواجهة خلية عصبية أخرى، أن الخليتين تم تدريبهم على نفس قواعد البيانات التي تم تغذية النظام بها، تقوم الأولى بمحاولة إبتكار وتصميم الصورة، وتقوم الثانية بدور المحقق أو المميز للاختلاف في الصورة، وذلك من خلال إكتشاف نقاط الخلل داخل الصورة، ويتم تسمية الخلية الأولى باسم المولد generator، والثانية باسم المميز discriminator. وعند قيام الخلية الثانية بإكتشاف الخلل تقوم الخلية الأولى تصحيحه حتى تعجز الخلية الثانية عن إكتشافه، ومن ثم يخلق النظام الشكل النهائي للصورة التي يصعب إكتشاف التزوير و التزييف. (خليفة ، ٢٠١٩).

إهتمام البحث بمشكلة جرائم الذكاء الإصطناعي والتزييف العميق والتعرف على الكيفية التي تتم وطرق التخفيف والحد من الإنتهاكات، من خلال خطة علمية وعملية تشرح المشاكل الحقيقية و الواقعية ومن ثم توضيح إستغلال الذكاء الإصطناعي والتقنيات في هذه الجرائم التي يشكلها التزييف العميق ويتمثل أبرزها في التحايل على الشخصيات العامة، حيث يمكن للذكاء الإصطناعي أن يستخدم لإنشاء مقاطع فيديو أو تصريحات مزيفه لرؤساء دول، سياسيين، أو شخصيات عامة، مما يؤثر على سمعتهم ويزعزع الثقة العامة إعداد إستراتيجيات دقيقة وتفصيلية لجميع جوانب المخاطر الناجمة عن إستخدامات الذكاء الاصطناعي ، خاصة مايتعلق بالتزييف والتلاعب والسرقة والهجمات الإلكترونية وحماية البيانات (عبدالعزیز ، ٢٠٢٠) . يذكر الباحثون أمثلة للتزييف العميق:

- الشكوك حول فيديوهات لقائد مليشيات عسكرية في السودان كانت سبب الحرب الأخيرة وأنضح وفاته وتتواصل ظهور الفيديوهات وروبوت لإثارة ومواصلة الحرب وهي ستكون بمثابة جريمة العصر التكنولوجية.

- تزيف فيديو يظهر فيها مارك زوكربيرغ، المؤسس والمدير التنفيذي لموقع التواصل الاجتماعي (Facebook) ومدير شركة " (Meta) " وهو يعترف بتأمره في مشاركة بيانات المستخدمين.
- فيديو لرئيس الولايات المتحدة السابق باراك أوباما يتحدث عن رئيس آخر بصورة لا تليق بدولة عظمى، وأتضح أخيراً أن الفيديو كان مزيف وأستخدم للدعاية الانتخابية (شبكة قنوات CNN).
- في منطقة آسيوية، تم خداع موظف في شركة متعددة الجنسيات لتحويل مبلغ كبير من المال إلى المحتالين من خلال تقنية التزييف العميق لانتحال شخصية أحد كبار المسؤولين التنفيذيين خلال مكالمة عبر الفيديو.
- في الانتخابات الأمريكية تم استخدام تقنيات التزييف العميق لإنشاء مقاطع فيديو مزيفة للمرشحين، ونشر معلومات مضللة بهدف تغيير تصورات واصوات الناخبين، من خلال مقاطع الفيديو المزيفة أحد المرشحين البارزين وهو يُدلي بتصريحات مثيرة للجدل، مما أدى إلى انخفاض ملحوظ في شعبيته ومستوى دعمه، على الرغم من الجهود اللاحقة لدحض الفيديو وبيان زيفه.

(٤) استخدام أدوات الذكاء الاصطناعي والتحقق من التزييف:

- في العصر الرقمي الحديث ظهرت التزييفات العميقة باعتبارها تهديداً كبيراً لأصالة المحتوى الحقيقي، يمكن لمقاطع الفيديو المتطورة التي تم إنشاؤها بواسطة الذكاء الاصطناعي أن تحاكي الأشخاص الحقيقيين بشكل مقنع، مما يزيد من صعوبة التمييز بين الحقيقة والخيال، يستخدم أدوات الذكاء الاصطناعي المتقدمة لتحليل مقاطع الفيديو والصور بحثاً عن علامات التزييف، بالإضافة إلى ذلك فحص البيانات الوصفية للمحتوى الرقمي التي توفر تفاصيل حول تواريخ الإنشاء ومحفوظات التحرير والبرامج المستخدمة، مما يساعد في التحقق مما إذا كان المحتوى قد تم تغييره والتعديل وذلك يكون بفهم العناصر الرئيسية وعمليات التحقق (الصاوي، ٢٠٢٣).
- (٥) العناصر الرئيسية للتوصل للحلول مثل الجوانب التالية:

- فهم تقنية التزييف العميق: تقديم نظرة شاملة حول كيفية إنشاء حالات التزييف العميق، بما في ذلك خوارزميات وتقنيات الذكاء الاصطناعي الأساسية المعنية.
- آليات الكشف: تعليم الأساليب المتقدمة لتحديد التزييف العميق، مثل: تحليل التناقضات في الفيديو والصوت، باستخدام أدوات الكشف القائمة على الذكاء الاصطناعي، والتعرف على علامات الوسائط الاصطناعية.
- الاعتبارات الأخلاقية والنظامية: استكشاف الآثار الأخلاقية والأطر النظامية المتعلقة باستخدام التزييف العميق لضمان ممارسات مسؤولة وممتثلة للقانون.
- التدريب العملي: تقديم جلسات عملية حيث يمكن للمستخدمين النهائيين العمل مع أدوات الكشف ومجموعات البيانات لاكتساب خبرة عملية في تحديد التهديدات العميقة والتخفيف منها ولضمان بقاء المستخدمين النهائيين مطلعين على أحدث التطورات وأفضل الممارسات.

- (٦) عمليات التحقق تتم من خلال إتباع خطوات في عملية التحقق من خلالها يتم تحديد الآلية المتبعة في خطوات وهي:
- أولاً: التحقق اليدوي والإحترافي: خدمات التحقق من الحقائق : إستخدم الخدمات الإحترافية لتدقيق الحقائق للتحقق من صحة المحتوى المشبوه ويمكن من خلال تكوين فريق إشراك خبراء الذكاء الاصطناعي والمبرمجين وخبراء الطب الشرعي الرقمي لفحص المحتوى يدوياً بحثاً عن علامات بالتزييف العميق، مثل(ملاحم الوجه غير المنتظمة، والحركات غير الطبيعية، والتناقضات بين الصوت والمرئيات).

ثانياً: التحقق الآلي والمحوسب: إستخدم التقنيات والخوارزميات الحديثة وبعض البرامج لإجراء عمليات التدقيق للتأكد من الحقائق وهو جانب ومحور تطبيقي للبحث.

ب. المحور الثاني: محور تطبيقي:

مع تقدم التكنولوجيا الكامنة وراء التزييف العميق، تقدمت أيضاً الأدوات والتقنيات المصممة لإكتشافها، يناقش الباحثون في هذا المحور أفضل أدوات وتقنيات للكشف عن التزييف العميق المتاحة وسيتم مناقشة الأدوات من خلال البحث وأن أهم هذه التقنيات والأدوات الحديثة هي:

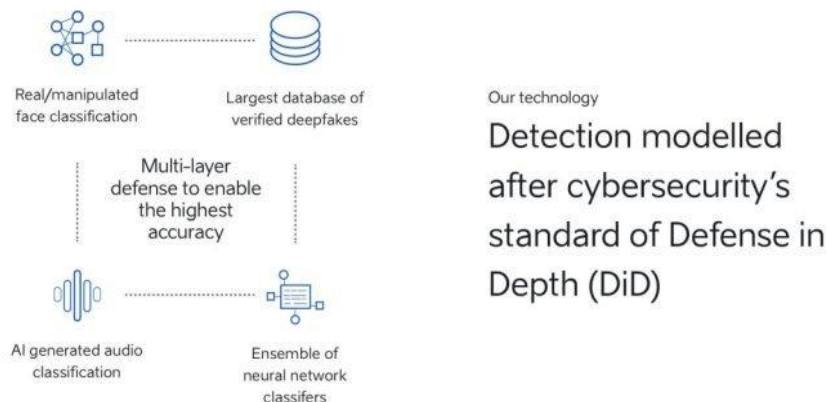
(١) مدافع الواقع Reality Defender: تتكيف Reality Defender بشكل مستمر مع تقنيات Deepfake المتطورة، وتحافظ على دفاع قوي ضد التهديدات في وسائل الإعلام والتمويل والحكومة.

المميزات الرئيسية للمدافع عن الواقع:

- يكتشف تطبيق Reality Defender عمليات التزييف العميق في الصور ومقاطع الفيديو والصوت والنصوص للمؤسسات والحكومات.
 - يوفر إكتشافاً في الوقت الفعلي وخالياً من العلامات المائية للتحقق السريع من المحتوى.
 - يمكن الوصول إليها عبر تطبيق الويب أو واجهة برمجة التطبيقات القابلة للتطوير لتحقيق التكامل المرن.
 - يقدم رؤية واضحة للتلاعب لتوجيه إجراءات الاستجابة.
 - يتم تحديثه باستمرار لمحاربة تهديدات الذكاء الاصطناعي المتطورة.
- (٢) حارس Sentinel: تستخدم خوارزميات الذكاء الاصطناعي المتقدمة لتحليل الوسائط التي تم تحميلها وتحديد ما إذا كان قد تم التلاعب بها، يقدم النظام تقريراً مفصلاً عن نتائجه، بما في ذلك تصور لمناطق وسائل الإعلام التي تم تغييرها، تم تصميم تقنية إكتشاف التزييف العميق من Sentinel لحماية سلامة الوسائط الرقمية، تساعد الحكومات ووكالات الدفاع والمؤسسات على وقف تهديد التزييف العميق. تعمل هذه التقنية من خلال السماح للمستخدمين بتحميل الوسائط الرقمية، و يتم تحليلها تلقائياً بعد ذلك من أجل التحديد من خلال النظام ما إذا كانت الوسائط مزيفة

المميزات الرئيسية - Sentinel:

- إكتشاف التزييف العميق المستند إلى الذكاء الاصطناعي.
- تستخدم من قبل المنظمات والحكومات الرائدة والمتقدمة في الدول المتقدمة.
- يسمح للمستخدمين بتحميل الوسائط الرقمية لتحليلها.



الشكل (٢) يوضح كاشف التزييف العميق - (يحي دهبان، ٢٠٢٣)

(٣) كاشف التزييف العميق FakeCatcher

من إنتاج شركة Intel ويستخدم أجهزة وبرامج Intel يمكن لهذه التقنية اكتشاف مقاطع الفيديو المزيفة بمعدل دقة يبلغ ٩٦٪ ، مما يؤدي إلى إرجاع النتائج في أجزاء من الثانية تم تصميم الكاشف بالتعاون مع Umur Ciftci من جامعة ولاية نيويورك في Binghamton ، ويستخدم أجهزة وبرامج Intel ، ويعمل على خادم ويتفاعل من خلال منصة على شبكة الإنترنت، يبحث FakeCatcher عن أدلة حقيقية في مقاطع فيديو حقيقية ، وتقييم ما يجعلنا بشراً - "تدفق الدم" الدقيق في وحدات البكسل في مقطع فيديو عندما تضخ قلوبنا الدم ، يتغير لون عروقنا. يتم جمع إشارات تدفق الدم هذه من جميع أنحاء الوجه وتقوم الخوارزميات بترجمة هذه الإشارات إلى خرائط زمانية مكانية. بعد ذلك ، باستخدام التعلم العميق ، يمكنه اكتشاف ما إذا كان مقطع الفيديو حقيقياً أم مزيفاً على الفور.

المميزات الرئيسية لكاشف التزييف العميق في الوقت الحقيقي:

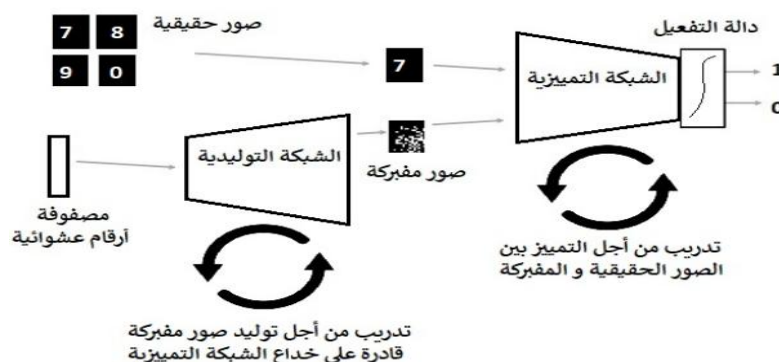
- يمكنه اكتشاف مقاطع الفيديو المزيفة بمعدل دقة ٩٦٪.
- إرجاع النتائج بالملي ثانية.
- يستخدم تقنية "تدفق الدم" الدقيقة في بكسلات الفيديو لاكتشاف التزييف العميق.

(٤) اكتشاف التزييف العميق باستخدام عدم تطابق Phoneme-Viseme

هذه التقنية المبتكرة ، تستغل حقيقة أن البصمات ، التي تشير إلى ديناميكيات شكل الفم ، تكون أحياناً مختلفة أو غير متوافقة مع الصوت المنطوق، هذا التناقض هو عيب شائع في التزييف العميق ، حيث يكافح الذكاء الاصطناعي غالباً لمطابقة حركة الفم مع الكلمات المنطوقة، تستخدم تقنية عدم التطابق Phoneme-Viseme خوارزميات متقدمة للذكاء الاصطناعي لتحليل الفيديو واكتشاف التناقضات، يقارن حركة الفم (البصمات) بالكلمات المنطوقة (الصوتيات) ويبحث عن أي عدم تطابق، إذا تم اكتشاف عدم تطابق ، فهذا مؤشر قوي على أن الفيديو مزيف عميق.

المميزات الرئيسية لاكتشاف التزييف العميق باستخدام عدم تطابق Phoneme-Viseme:

- يستغل التناقضات بين البصمات والصوتيات في التزييف العميق.
- يستخدم خوارزميات الذكاء الاصطناعي المتقدمة لاكتشاف حالات عدم التطابق.
- يعطي مؤشراً قوياً على التزييف العميق إذا تم الكشف عن عدم تطابق.



الشكل (٣) يوضح طريقة الالية كشف التزييف العميق - (يحي دهبان). (٢٠٢٣)

الشكل يوضح الطريق التي تتم بها عملية الكشف علي التزييف العميق من خلال الشبكة التوليدية التنافسية من خلال التمييز بين الصور الحقيقية والصور المزيفة والتدريب علي الاختلافات التي يتعرف عليها المولد .
٨. خاتمة البحث:

التزييف العميق (Deepfake) هو تقنية ذكاء إصطناعي تُستخدم لصنع صور، أو مقاطع فيديو، أو أصوات مزيفة. تعتمد هذه التقنية على دمج بيانات حقيقية لتركيب محتوى جديد. يبدو هذا المحتوى واقعياً جداً لدرجة يصعب معها تصديق أنه غير حقيقي. تناول البحث موضوع الجرائم و الإنتهاكات الإلكترونية من خلال إستخدام تطبيقات وخوارزميات الذكاء الاصطناعي باعتبارها ظاهرة تفتت في المجتمعات الحديثة ونالت من المجتمعات النامية مثلما نالت من المتقدمة، نظراً إلى أن تقنيات التزييف العميق أصبحت أكثر تقدماً ويمكن الوصول إليها، فقد تصاعدت المخاطر المرتبطة بها عالمياً، مما زاد من إمكانية تعرض الأفراد والمنظمات والدول لهذه المخاطر وتسبب مجموعة من الاحتمالات وتتمثل المشكلة في كيفية تحديد التزييف والتزوير وإجراء التدابير التقنية والتعرف على الحلول والمعالجات تساعد في الحماية.

ونظراً إلى التأثير المجتمعي العميق للتزييف العميق، فمن الأهمية من وجه نظر الباحثون بماكان من توجيه وتطوير تطبيقات إيجابية وبناءة مع تخفيف المخاطر المرتبطة بالمشاكل . يمكن تلخيص أهم النتائج التي توصل إليها الباحثون في مجموعة نقاط وهي:

(١) تمثل أدوات وتقنيات إكتشاف التزييف العميق التي تم شرحها من خلال الباحثون وما توصل له البحث في إستخدام خوارزميات الذكاء الاصطناعي المتقدمة لتحليل وإكتشاف التزييف العميق بدقة، تقدم كل أداة وتقنية نهجاً فريداً لاكتشاف التزييف العميق ، بدءاً من تحليل العناصر الرمادية الدقيقة لمقطع فيديو إلى تتبع تعابير الوجه وحركات الأشخاص و درجة ثقة في الوقت الفعلي تشير إلى ما إذا كان قد تم التلاعب بالصورة الثابتة أو الفيديو هذه الأدوات ، إلى جانب الأدوات الأخرى التي تم الوصول لها من خلال الباحثون ، تقود الي التوصل الي نتائج تساهم في التعرف علي التزييف العميق ومواكبة تقدم التكنولوجيا الكامنة وراء تقنية التزييف العميق ، مما يساعد على ضمان أصالة وثقة المعلومة.

(٢) توصل الباحثون الي أن مشكلة التزييف العميق و إستخدام خوارزميات متقدمة للذكاء الاصطناعي أصبحت مهدد ومشكلة حقيقة تواجه عالمنا اليوم واصبحت الحاجة إلى أدوات وتقنيات فعالة للكشف عن التزييف العميق مع إستمرار تقدم التكنولوجيا الكامنة وراء تقنية التزييف العميق ، يجب أن تتقدم أيضاً أساليبنا في الكشف .
(٣) إن التكنولوجيا فقط لا تكون حل في مشكلة التزييف العميق، ويجب التركيز علي التعليم والوعي والتثقيف للمجتمعات والإطلاع بأحدث التطورات في تكنولوجيا التزييف العميق والكشف ، بحيث يمكن لعب دور في مكافحة هذا التهديد.

(٤) كما يجب التوعية بمجموعة من النقاط الهامة المتمثلة في الآتي:

- يجب علي الحكومات والشركات التوعية و الإستفادة من تقنيات الذكاء الاصطناعي للكشف عن البرمجيات الخبيثة وفيديوهات التزييف العميق، و تطوير أدوات متقدمة للكشف عن التلاعب في المحتوى الرقمي والتعاون الدولي.

- آليات التصدي لهذه العمليات ينجم من خلال تعزيز وعي الأفراد، حيث يجب على الأفراد أن يكونوا على دراية بمخاطر الذكاء الاصطناعي في الجرائم الإلكترونية والتزييف العميق، وأن يتبعوا ممارسات أمان رقمية مثل التحقق من مصادر الأخبار، والاستثمار في التقنيات الدفاعية.

- تتلخص توصيات الباحثون في مجموعة من النقاط يتم توضيحها في الآتي:

- إقامة مؤسسات بحثية داخل وحدات مكافحة جرائم الذكاء الاصطناعي تهتم بالأمن الدولي الإلكتروني، والتعامل مع التطورات التقنية التي تؤدي إلى تطور وسائل التزييف والتزوير العميق.

- بعد تحديد تقنيات التزيف العميق وتأثيرها على المجتمع يجب رفع مستوى الوعي بالتطبيقات الخبيثة وغير الخبيثة لتقنيات التزيف العميق.
- بعد تحديد إستراتيجيات تخفيف المخاطر وتمكين الجهات التنظيمية والجهات الحكومية, يجب تقديم تدابير وقائية لتحديد ومكافحة تقنيات التزيف العميق.
- تعزيز التعليم الفعّال والقدرة على مواجهة إساءة استخدام تقنية التزيف العميق.
- المشاركة الجماعية للجهات الحكومية، والمنظمات ومراكز الأبحاث والجامعات، في رسم السياسات لمنع التزيف العميق والجرائم المستترة والظاهرة، ووضع قوانين وتدابير إجرائية.

المراجع:

ولاً: المراجع العربية

البالول، نور سليمان يوسف يعقوب. (2021). *الأحكام الموضوعية لجرائم المعلوماتية* [أطروحة دكتوراه غير منشورة]. جامعة عين شمس، كلية الحقوق، قسم القانون الخاص.

الحبسي، عيسى عبد الله. (2020). *جرائم البريد الإلكتروني: دراسة مقارنة* [أطروحة دكتوراه غير منشورة]. جامعة المنصورة، كلية الحقوق، قسم القانون الجنائي.

خليفة، إيهاب. (٢٠١٩، ٣١ ديسمبر). *التهديد المتصاعد للخداع العميق عبر نظم الذكاء الاصطناعي*. مركز المستقبل للأبحاث والدراسات المتقدمة.

دهشان، يحيى إبراهيم. (٢٠٢٠). *المسؤولية الجنائية عن جرائم الذكاء الاصطناعي*. مجلة *الشرعية والقانون*، ٣٤ (82)، ١٠٠-144. <https://doi.org/10.65088/2020.34.2.2>

عبد العزيز، سارة. (٢٠٢٠، ١٧ فبراير). *تأثيرات الذكاء الاصطناعي على الأمن المعلوماتي للدول*. مركز المستقبل للأبحاث والدراسات المتقدمة.

فاضل، علي مولود، وعباس، سيف عدنان. (2021). *التزيف العميق: لغة الذكاء الاصطناعي في حروب السيبران الإعلامية*. دار أمجد للنشر والتوزيع.

الصاوي، محمد كرم كمال الدين. (٢٠٢٥). *تكنولوجيا التزيف العميق: دراسة بحثية حول الجوانب المظلمة للذكاء الاصطناعي*. مجلة *العمارة والفنون والعلوم الإنسانية*، ١٠ (51)، 694.٦٧٤-

مركز الإمارات للدراسات والبحوث الاستراتيجية. (2023). *تحديات التزيف العميق لمصادقية المعلومات وسبل معالجتها* (تقرير).

بوفاس، الشريف، وطلحي، فاطمة الزهراء. (٢٠١٤). *نحو بناء نظم لإدارة حماية المعلومات ISO 27001 في المؤسسات الجزائرية*. في *الملتقى الدولي الثاني حول البقطة الاستراتيجية ونظم المعلومات في المؤسسة الاقتصادية*. جامعة باجي مختار عنابة.

ثانياً: المصادر الإلكترونية

Microsoft. (2026, March, 19). *Zero Trust assessment overview*. Microsoft Learn.

منصة أريد العلمية. (د.ت). *منصة أريد العلمية*.