

طبيعة استخدام القائم بالاتصال في وسائل الإعلام الليبية لأساليب كشف التزييف العميق
بواسطة الذكاء الاصطناعي

دراسة وصفية تطبيقية في ظل نظرية دافع الحماية

د. خالد أبو القاسم علي خبريش

قسم الإعلام - كلية الآداب والتربية - جامعة صبراتة - ليبيا

Khebrech@gmail.com

The Nature of Communication Practitioners' Use of Artificial Intelligence-
Based Deep fake Detection Methods in Libyan Media

A Descriptive Applied Study within the Framework of the Protection Motivation
Theory

Dr. Khaled Abualqasim Ali Khebrish

Department of Media, Faculty of Arts and Education, Sabratha University, Libya

تاريخ الاستلام: 2025-06-15، تاريخ القبول: 2025-06-28، تاريخ النشر: 2025-07-22.

الملخص:

يواجه الإعلام المعاصر تحدياً غير مسبوق يتمثل في التزييف العميق، وهو أحد أخطر مظاهر التلاعب بالمعلومات المدعوم بتقنيات الذكاء الاصطناعي، وهذا البحث؛ يركز على طبيعة تعامل القائم بالاتصال في وسائل الإعلام الليبية مع هذه الظاهرة، ومدى استخدامه لتقنيات كشف التزييف العميق. وانطلق البحث في إطار نظرية دافع الحماية - التي تفسر استجابة الأفراد للتهديدات بناءً على إدراكهم للمخاطر وقدرتهم على مواجهتها - بهدف استكشاف طبيعة استخدام القائم بالاتصال لأساليب كشف التزييف، واعتمد البحث على المنهج الوصفي التحليلي وتطبيق نظرية دافع المعرفة، حيث تم تصميم وتوزيع استبيان إلكتروني على عينة عشوائية مكونة من 112 مفردة من القائمين بالاتصال في الإعلام الليبي، مع استخدام الاختبارات الإحصائية اللامعلمية لتحليل النتائج، التي كشفت عن وجود إدراك عالٍ لخطورة التزييف العميق لدى القائم بالاتصال، لكن معرفته العملية بالأدوات المتخصصة لكشفه محدودة، واستخدامها محدود، وتوافرها محدود أيضاً، فأكثر من نصف العينة يعتقدون على الخبرة الشخصية، وكلها عوامل أدت إلى ظهور فجوة معرفية وتطبيقية في التعامل مع هذه الظاهرة. كما أظهرت النتائج أن غياب الأدوات التقنية وضعف التدريب المؤسسي يشكلان عائقين رئيسيين أمام التطبيق الفعال لأساليب كشف التزييف العميق. ورغم أن غالبية القائمين بالاتصال أبدوا دافعاً قوياً للحماية والحفاظ على المصداقية، إلا أن هذه الدوافع لم تتحول إلى ممارسة فعالة بسبب ضعف البنية التقنية وغياب الدعم المؤسسي. وأكدت النتائج على الحاجة الملحة إلى تعزيز الموارد التقنية، وتوفير برامج تدريب متخصصة، وإنشاء أقسام مختصة بكشف التزييف داخل المؤسسات الإعلامية. ويؤكد البحث أن مواجهة التزييف العميق تتطلب استراتيجية متكاملة تشمل تدريب القائمين بالاتصال، وتعزيز الشراكات مع الجهات المتخصصة، وتبني حلول الذكاء الاصطناعي المتطورة لضمان المصداقية، مع تطوير سياسات إعلامية مستدامة تعزز التزام المؤسسات الإعلامية بممارسات كشف التزييف لبناء بيئة إعلامية أكثر موثوقية.

الكلمات المفتاحية: التزييف العميق، الذكاء الاصطناعي، وسائل الإعلام الليبية.

Abstract:

Contemporary media faces an unprecedented challenge in the form of deepfakes, one of the most dangerous manifestations of AI-driven information manipulation. This study examines how communication practitioners in Libyan media perceive and respond to this phenomenon, with a specific focus on their use of deepfake detection technologies. Grounded in the Protection Motivation Theory (PMT)—which explains individual responses to threats based on

perceived severity and self-efficacy—the research aims to explore the nature of practitioners' engagement with deepfake detection methods. A descriptive-analytical approach was employed, utilizing an electronic survey distributed to a random sample of 112 media practitioners in Libya. Data were analyzed using non-parametric statistical tests. The findings reveal a high level of awareness among practitioners regarding the risks posed by deepfakes. However, practical knowledge of specialized detection tools remains limited, as does their actual usage and availability. Over half of the respondents rely primarily on personal experience, indicating a significant gap between awareness and practical application. Key barriers to effective detection include the lack of technical tools and insufficient institutional training. Although the majority of practitioners exhibit a strong motivation to protect credibility and uphold journalistic integrity, this motivation has not translated into effective practice due to weak technical infrastructure and a lack of organizational support. The results underscore an urgent need to enhance technical resources, provide specialized training programs, and establish dedicated verification units within media institutions.

The study concludes that combating deepfakes requires a comprehensive strategy that integrates practitioner training, partnerships with technical experts, adoption of advanced AI-based detection tools, and the development of sustainable media policies that institutionalize verification practices—ultimately fostering a more trustworthy media environment.

Keywords: Deepfakes, Artificial Intelligence, Libyan Media.

المقدمة

يتزايد بشكل لافت استخدام تقنيات الذكاء الاصطناعي في تزيف المعلومات، وينمط متقن إلى درجة يصعب اكتشافه، بحيث أصبح من السهل نشر مقاطع صوت وصورة تتضمن معلومات مُظَلَّلَة وخاصة عبر شبكات التواصل الاجتماعي، ونتيجة لذلك؛ أصبح القائم بالاتصال عرضة للوقوع في مصيدة هذا التزيف، وحتى الذين يعتقدون أن لديهم الخبرة الكافية لكشف التزيف العميق، لا يمكنهم فعلياً كشف التزيف دون استخدام تقنيات متقدمة، وهو أثبتته دراسة تجريبية أجريت سنة 2024 على عينة من الصحفيين، حيث فشل جميعهم في كشف أي من مقاطع الفيديو المزيفة في اختبارين متتالين، وهو ما يؤكد صعوبة التعرف على "مخرجات" التزيف العميق بواسطة الذكاء الاصطناعي.⁽ⁱ⁾

وتشير بعض الدراسات والتقارير إلى أن العالم شهد - في السنوات الأخيرة - ارتفاعاً في استخدام تقنية التزيف العميق، وعلى سبيل المثال؛ ارتفع عدد حالات تزيف المحتوى العميق المنشور عبر الانترنت خلال عامي 2019 - 2020، إلى 900 %، فيما تشير الدراسات إلى أن تنامي تزيف المحتوى سيصل إلى نسبة 90% من المحتوى المنشور عبر الانترنت بحلول عام 2026. وفي جانب ذي صلة، تعرّض نحو 66% من المسؤولين على الأمن السيبراني في عام 2022 إلى هجمات التزيف العميق داخل مؤسساتهم.⁽ⁱⁱ⁾ وفي ذات السياق؛ أشارت دراسة (Hall, P., & Richards, T. 2021) ⁽ⁱⁱⁱ⁾ إلى أن أغلب العاملين في مجالات الإعلام يفتقرون إلى التدريب الكافي حول كيفية اكتشاف التزيف

العميق، ما جعلهم عرضة للانجرار وراءها في ظل نقص التدريب. وفي المقابل؛ أكدت دراسة (Nguyen, D., & Pham, Q. 2021)^(iv) أن العاملين في مجال الإعلام والمدربين بشكل متخصص على استخدام أدوات كشف التزييف، لديهم قدرة التمييز بين المحتوى الواقعي والمزيف.

الوضع المنشئ لمشكلة البحث:

مع تتطور تقنيات الذكاء الاصطناعي ازدادت مخاطر التزييف العميق، فقد أصبح اليوم بإمكان خوارزميات الذكاء الاصطناعي "بناء وإنشاء وتنفيذ" وسائط مزيفة يصعب تمييزها عن الحقيقة، في ظل سهولة الوصول إلى تلك الخوارزميات بشكل مجاني أو بتكلفة منخفضة^(v)، مما سهّل عمليات التزييف العميق سواءً لمجرد الدعاية والمرح، أو لأهداف تتنافى مع أخلاقيات ومبادئ العمل الإعلامي، وأحياناً الارتقاء إلى مستوى الجرائم الإلكترونية. حيث بينت دراسة (Mahmud, B. U., & Sharmin, A. 2020)^(vi) أن مقدرة الذكاء الاصطناعي على إنشاء مقاطع فيديو وصور أصبح واقعياً للغاية، وقد تُستخدم في نشر معلومات مضلّة عبر الإنترنت. وفي ذات الوقت، يشعر جُلّ العاملين في مجالات الإعلام المختلفة بالقلق، جراء تنامي انتشار التزييف العميق عبر مختلف المنصات، متضمنة معلومات مزيفة بشكل متقن، حيث أفادت دراسة (Liu, S., & Zhang, H. 2022)^(vii) بأن معظم "الصحفيين" يُعبرون عن قلقهم من استخدام هذه التقنية في نشر المعلومات المغلوطة، وأنهم يحتاجون إلى مزيد من التدريب.

وفي المقابل، ظهرت تقنيات عدة للكشف عن هذا التزييف، مثل منصة Sentinel، وكاشف Fake Catcher، وغيرها من الأدوات المتخصصة. ومع ذلك، ما يزال العديد من العاملين في وسائل الإعلام يفتقرون إلى الوعي الكافي بخطورة التزييف العميق أو بأساليب كشفه، نظراً لغياب التدريب والوعي المناسبين. يأتي هذا في ظل عدم اهتمام "بعض" وسائل الإعلام المعنية بتوظيف التقنيات والأدوات اللازمة للكشف عن التزييف أو في جوانب التوعية والتدريب، حيث أكدت دراسة (Simmons, L., & McDonald, F. 2022)^(viii) أن العاملين في مجالات الإعلام يواجهون صعوبة في توظيف تقنيات الكشف عن التزييف العميق في عملهم اليومي، بسبب نقص الأدوات والموارد. إلى جانب ذلك؛ أشارت دراسة (Harris, M., & Clark, J. 2022)^(ix) إلى ضرورة وجود سياسات واضحة وأدوات للكشف عن التزييف العميق في بيئات العمل الإعلامي.

المشكلة البحثية:

تعتبر ظاهرة التزييف العميق (Deepfake) من التحديات الحديثة التي تواجه وسائل الإعلام على مستوى العالم، إذ تهدد مصداقية الأخبار والمحتوى الإعلامي المنشور. وفي هذا السياق؛ يُعدّ القائم بالاتصال في وسائل الإعلام من الفئات الأساسية التي تتحمل المسؤولية في الكشف عن هذه الظواهر والتأكد من صحة المعلومات.

ويبرز تساؤل أساسي يتمحور حول "طبيعة استخدام القائم بالاتصال في وسائل الإعلام للبيبة لأساليب الكشف عن التزييف العميق بواسطة الذكاء الاصطناعي" في ظل الظروف المحلية التي تقتقر

إلى بعض الموارد والتدريب التقني المتقدم. أي أن هناك غموضاً حول طبيعة استخدام هذه الأساليب والتقنيات في وسائل الإعلام الليبية.

من هنا؛ تكمن المشكلة البحثية، المتمثلة في فحص طبيعة استخدام هذه الأساليب في بيئة العمل الإعلامية الليبية، من خلال محاولة فهم العوامل التي تؤثر في فعالية التطبيق ومدى توفر الأدوات والموارد التقنية اللازمة لذلك في ظل دوافع الحماية الذاتية.

الدراسات السابقة:

1. دراسة: أحمد، مصطفى محمود (2025)^(x) بعنوان: (التربية الإعلامية الرقمية وعلاقتها بمستوى تحصين وعي طلاب الإعلام التربوي بمخاطر تطبيقات التزييف العميق في إطار نظرية دافع الحماية)، سعت الدراسة إلى الكشف عن دور التربية الإعلامية الرقمية في تحصين مستوى وعي طلاب الإعلام التربوي بمخاطر تطبيقات التزييف العميق، ووظفت الدراسة المنهج الوصفي "الميداني" وطبقت أدوات الدراسة على عينة عشوائية قوامها (400) طالب من طلاب أقسام الإعلام التربوي بالجامعات المصرية، وقد خلصت إلى وجود تجانس تام في استجابات عينة الدراسة تجاه موافقتهم بدرجة "مرتفعة" على جميع اقتراحات تحصين الوعي، كما أثبتت صحة وجود ارتباط طردي ذي دلالة إحصائية بين مهارات التربية الإعلامية الرقمية، باستثناء مهارتي (المشاركة الرقمية المسؤولة - مواجهة المخاطر الرقمية).

2. دراسة: Turner, M., & Brooks, L. (2023)^(xi) بعنوان: (The role of AI in identifying deepfake content in journalism) هدفت هذه الدراسة إلى فحص فعالية الذكاء الاصطناعي في الكشف عن التزييف العميق، وتقييم كيف يمكن للذكاء الاصطناعي أن يعزز مصداقية وسائل الإعلام، واعتمدت على منهج "التحليل التجريبي" وجمعت البيانات بواسطة تجارب حية على أدوات الكشف عن التزييف العميق، واستبيانات للصحفيين، وأكدت الدراسة أن الذكاء الاصطناعي يقدم حلاً عملياً للكشف عن التزييف العميق، وكشفت الدراسة أن 85% من الصحفيين محلّ البحث اعتقدوا أن الذكاء الاصطناعي يمكن أن يكون أداة مفيدة في مكافحة الأخبار المزيفة.

3. دراسة: Zhang, T. (2022)^(xii) بعنوان: (Deepfake generation and detection, a survey). وهدفت الدراسة إلى فحص قدرة البشر على التمييز بين مقاطع الفيديو الحقيقية والمزيفة (التزييف العميق) باستخدام تقنيات الذكاء الاصطناعي، بالإضافة إلى تحليل كيف يمكن للبشر إدراك مدى صحة المحتوى الذي يشاهدونه. وتم تنفيذ الدراسة باستخدام منصة ويب تعتمد على مفاهيم (gamification) والتي تم توفيرها لـ 110 مشارك منهم (55 لغتهم الأم الإنجليزية و55 غير ناطقين بها)، حيث قيّموا (40) فيديو منها (20 حقيقي و20 مزيف). وأجريت التجربة مرتين باستخدام نفس الـ 40 فيديو في ترتيب عشوائي مختلف. وتمت مقارنة أداء البشر مع خمسة نماذج متطورة لاكتشاف التزييف العميق السمعي البصري. وجد الباحثون أن جميع النماذج الذكية أدت أداءً

- أفضل من البشر في اكتشاف التزييف العميق على نفس الـ 40 فيديو. وأظهرت النتائج أن البشر يميلون إلى المبالغة في تقدير قدراتهم على الكشف عن التزييف العميق.
4. دراسة: شمس الدين، فتحي (2022)^(xiii)، هدفت إلى استشراف رؤية القائم بالاتصال في مصر لمستقبل الإعلام والإعلاميين في عصر الذكاء الاصطناعي، ضمت عينتها (50) إعلامياً من القائمين بالاتصال والقيادات الإعلامية بالمؤسسات المصرية والعربية، وتوصلت نتائج الدراسة إلى أن استخدام تقنيات الذكاء الاصطناعي سيكون واقعاً وضرورياً في كل وسائل الإعلام المستقبلية، وأشارت النتائج أيضاً إلى أن القائم بالاتصال يدرك أن هناك تأثيرات مستقبلية على مستقبل القائمين بالاتصال في ظل استخدام الروبوت الإعلامي.
5. دراسة: علي والنهباني (2022)^(xiv) بعنوان: (الذكاء الاصطناعي والتزييف العميق: التحديات والفرص في الإعلام العربي)، اعتمدت على المنهج الاستقصائي باستخدام استبيانات تم توزيعها على 100 صحفي في مختلف وسائل الإعلام العربية، بهدف دراسة تأثير الذكاء الاصطناعي على الكشف عن التزييف العميق في الإعلام العربي، وتقييم وعي الصحفيين بهذه التقنية. وتوصلت إلى أن الصحفيين في الإعلام العربي غير مدركين تماماً لكيفية استخدام الذكاء الاصطناعي في الكشف عن التزييف العميق، وهو ما يزيد من احتمالية نشر أخبار مزيفة.
6. دراسة: عبد الله، هشام (2022)^(xv)، بعنوان: (تأثير الذكاء الاصطناعي على الصحافة العربية في مكافحة الأخبار المزيفة)، وهي دراسة تحليلية هدفت إلى معرفة كيفية استخدام الذكاء الاصطناعي في الصحافة العربية لمكافحة الأخبار المزيفة، ودراسة أثر الذكاء الاصطناعي في تحسين جودة الأخبار ومكافحة الأخبار المزيفة. وأكدت نتائجها على أن الذكاء الاصطناعي يلعب دوراً مهماً في كشف الأخبار المزيفة، لكن هناك حاجة ملحة لتطوير أدوات الكشف في المنطقة العربية.
7. دراسة: Williams, S., & Green, R. (2021)^(xvi)، بعنوان: (الذكاء الاصطناعي والتزييف العميق: تهديد لمصداقية وسائل الإعلام)، وهي دراسة تحليلية اعتمدت أداة تحليل المحتوى ومراجعة الأدبيات ودراسة حالات لمحتويات تزييف عميق انتشرت عبر الإنترنت، بهدف تحليل تأثير التزييف العميق على مصداقية وسائل الإعلام، وفهم الطرق الفعالة لاكتشاف التزييف العميق. وتوصلت الدراسة إلى أن التزييف العميق يزيد من فقدان الثقة في وسائل الإعلام التقليدية، خاصة عند الجمهور غير المتخصص، وأشارت إلى أن الأساليب التقليدية للكشف عن التزييف العميق غير فعالة مع تطور التقنيات الحديثة.
8. دراسة: Mohammed, R., & Stevens, M. (2021)^(xvii) بعنوان: (Journalists' awareness of deepfakes and its impact on news integrity)، هدفت هذه الدراسة إلى قياس الوعي الإعلامي تجاه التزييف العميق، واستكشاف تأثير التزييف العميق على مصداقية الأخبار، واعتمدت على المنهج الاستقصائي واستخدمت أداة الاستبيان لجمع المعلومات من عينة

شملت 200 صحفي، وأظهرت الدراسة أن الوعي بتقنيات التزييف العميق كان مرتفعاً بين الصحفيين الذين يعملون في وسائل الإعلام الكبرى، بينما كان أقل بين الصحفيين في المؤسسات الصغيرة.

9. دراسة: النمس والحداد (2021)^(xviii) ، بعنوان: (التزييف العميق في الإعلام: التحديات وأساليب الكشف)، وهي دراسة تحليلية قائمة على مراجعة الأدبيات وتحليل حالات من التزييف العميق في الإعلام، تضمنت عينة من المقالات الإعلامية في الصحف الكبرى، بهدف دراسة مدى تأثير التزييف العميق على مصداقية الإعلام وطرق كشفه باستخدام تقنيات الذكاء الاصطناعي، وخلصت إلى أن التزييف العميق قد يؤثر بشكل كبير على مصداقية الأخبار ويصعب التمييز بين المحتوى المزيف والحقيقي.

10. دراسة: حسن والعسيري (2021)^(xix) ، بعنوان: (التأثيرات الاجتماعية والإعلامية للتزييف العميق في الإعلام العربي)، وهي دراسة ميدانية اعتمدت على المقابلات المعمقة مع 30 من الصحفيين والخبراء في من الصحف والمواقع الإلكترونية العربية، بهدف دراسة التأثيرات الاجتماعية للتزييف العميق في الإعلام العربي وكيفية تصدي الصحفيين له، وأظهرت النتائج أن هناك نقصاً كبيراً في الوعي بالتهديدات التي تمثلها أدوات التزييف العميق.

لقد أظهرت معظم الدراسات وجود وعي "جزئي وضعيف" بالتزييف العميق، مع نقص في التدريب على استخدام أدوات الذكاء الاصطناعي كما في دراسة (حسن والعسيري، 2021)، كما تفوقت تقنيات الذكاء الاصطناعي على البشر في الكشف عن التزييف، وتمثلت أوجه الاستفادة من هذه الدراسات وإضافة البحث الحالي في تصميم الأدوات واعتماد المنهج الوصفي.

التعريفات الإجرائية:

التزييف العميق (Deepfake): يشير المفهوم إلى تقنية تُستخدم لتحرير الصور والفيديوهات بشكل مزيف باستخدام الذكاء الاصطناعي (AI)، بحيث يتم - على سبيل المثال - إنشاء مقاطع صوتية أو مرئية مزيفة تتطابق أشخاص حقيقيين أو أحداث حقيقية. يتم ذلك عن طريق تطوير شبكات كبيرة على مجموعة ضخمة من البيانات لإنشاء محتوى يبدو حقيقياً، حيث يمكن استخدام هذه التقنية للتلاعب بالفيديوهات والصوتيات لأغراض متعددة، منها التلاعب الإعلامي أو الإعلانات الكاذبة، وتُعرف الهيئة السعودية للبيانات والذكاء الاصطناعي^(xx) التزييف العميق بأنه تقنية تعتمد على الذكاء الاصطناعي لإنشاء وسائط رقمية مزيفة، مثل مقاطع الفيديو والصور، تبدو واقعية للغاية لكنها في الواقع مزيفة. تستخدم هذه التقنية خوارزميات التعلم العميق لتعديل أو استبدال وجوه الأشخاص أو أصواتهم في المحتوى الرقمي، مما يجعل من الصعب تمييز المحتوى المزيف عن الحقيقي.

والتعريف الإجرائي للتزييف العميق: (Deepfake) هو ذلك الانتاج الإعلامي المزيف الناجم عن استخدام الذكاء الاصطناعي، ويُقصد به أي فيديو أو صورة تم إنشاؤها أو تعديلها باستخدام تقنيات الذكاء الاصطناعي بحيث تظهر محتويات تبدو حقيقية، ويمكن اختباره من خلال استخدام أدوات تقنية للكشف عن وجود تعديل باستخدام الذكاء الاصطناعي.

الذكاء الاصطناعي: (Artificial Intelligence - AI) وهو فرع من فروع علوم الكمبيوتر يركز على إنشاء برمجيات قادرة على محاكاة وظائف العقل البشري. وفي سياق التزييف العميق؛ يُستخدم الذكاء الاصطناعي في تقنيات مثل تعلم البرمجيات أو ما يعرف بالشبكات العصبية العميقة، لإنشاء وتوليد المحتوى المزيف، كما تساهم هذه التقنيات في الكشف عن التزييف العميق عن طريق تطوير خوارزميات قادرة على معرفة التلاعب في الصور والفيديوهات^(xxi) والتعريف الإجرائي: يُعرّف الذكاء الاصطناعي في هذا البحث بوصفه تقنيات أو أنظمة حاسوبية قادرة على محاكاة الذكاء البشري.

القائم بالاتصال: هو الشخص المسؤول عن إنتاج وتوزيع المحتوى الإعلامي^(xxii)، والتعريف الإجرائي يشمل ذلك (الصحفيين والمحررين والمراسلين والمخرجين والمنفذين والمنتجين) وغيرهم من المتخصصين مجالات إنتاج المواد الإعلامية.

أساليب كشف التزييف العميق: (Deepfake Detection Techniques) هي مجموعة إجراءات تتم عبر أساليب خاصة بكشف التزييف العميق، تشمل الأدوات التكنولوجية والبرمجيات التي تستخدم الذكاء الاصطناعي لتحليل البيانات المرئية والصوتية، بهدف كشف التلاعب، تتراوح هذه الأساليب من تحليل التفاصيل الدقيقة في الصور والفيديوهات إلى التعرف على أنماط غير متسقة أو غير قابلة للتفسير من خلال الخوارزميات^(xxiii).

فرضيات البحث:

الفرضية الأولى: يفترض البحث أن استخدام القائم بالاتصال في وسائل الإعلام الليبية لأساليب كشف التزييف العميق بواسطة الذكاء الاصطناعي محدود بسبب نقص المعرفة التقنية وغياب الدعم المؤسسي.

الفرضية الثانية: يفترض البحث أن دافع الحماية لدى القائم بالاتصال في وسائل الإعلام الليبية عالٍ، ولكن لم يتم استغلاله بشكل كافٍ بسبب نقص الوعي المؤسسي والتدريب المتخصص وقلة الموارد التقنية المتاحة.

الفرضية الثالثة: يفترض البحث أن هناك فروقاً ذات دلالة إحصائية بين طبيعة استخدام القائمين بالاتصال لأساليب كشف التزييف العميق وبين المتغيرات الديموغرافية والمهنية مثل: (النوع، نوع الوسيلة الإعلامية، طبيعة العمل، التخصص العلمي، مستوى الخبرة).

الفرضية الرابعة: يفترض البحث أن مستوى الدعم المؤسسي يؤثر بشكل سلبي على طبيعة استخدام تقنيات كشف التزييف العميق.

الفرضية الخامسة: يفترض البحث أن القائمين بالاتصال في وسائل الإعلام الليبية يعتمدون في الغالب على الخبرة الشخصية في الكشف عن التزييف العميق بسبب غياب الأدوات التقنية المتخصصة ونقص التدريب.

الفرضية السادسة: يفترض البحث أن القائم بالاتصال يواجه عقبات تطبيقية بسبب نقص الأدوات والتقنيات المتاحة.

أهداف البحث: يهدف البحث إلى تقديم تحليل شامل لطبيعة استخدام القائم بالاتصال في وسائل الإعلام الليبية لأدوات وأساليب كشف التزييف العميق. فضلاً عن الخروج بتوصيات لعلها تساهم في تحسين مستوى الوعي والمعرفة لدى القائم بالاتصال في وسائل الإعلام الليبية حول التزييف العميق وأساليب الكشف عنه والتعامل معه.

أهمية البحث: تكمن أهمية هذا البحث في كونه يعالج مشكلة ذات أبعاد خطيرة في وسائل الإعلام الحديثة، ويساهم في تعزيز ورفع مستوى الوعي العام بأخطار التزييف العميق، فضلاً عن التحفيز لاستخدام أدوات فعالة لمكافحة التزييف العميق باستخدام الذكاء الاصطناعي.

- تتبع أهمية البحث من كون التزييف العميق يمثل تهديداً متزايداً لقطاع الإعلام، حيث يُستخدم لتضليل الجماهير ونشر الأخبار المزيفة بطرق تبدو حقيقية بشكل كبير. في السياق الليبي، يمكن أن يكون لهذا التأثير أبعاد خطيرة على مختلف الصعد لا سيما السياسية والاجتماعية والدينية، مما يجعل من الضروري دراسة مدى اعتماد القائم بالاتصال على هذه تقنيات كشف التزييف.

- يهدف البحث إلى تحديد مدى استخدام القائم بالاتصال لهذه الأدوات، مما يساهم في تحديد الفجوات التقنية والتدريبية ويضع توصيات لتعزيز استخدام التكنولوجيا المتقدمة.

- يعد هذا البحث إضافة مهمة للأدبيات العلمية حول موضوع التزييف العميق في العالم العربي، وخاصة في السياق الليبي، حيث يركز على الجوانب المحلية ويقدم حلولاً موجهة لمواجهة هذه التحديات.

منهجية البحث وعينته:

يعد هذا البحث بحثاً ميدانياً تطبيقياً يعتمد على المنهج الوصفي في إطار تطبيق نظرية دافع الحماية، حيث يهدف إلى دراسة استخدام القائمين بالاتصال في وسائل الإعلام الليبية لأساليب كشف التزييف العميق، ويعتمد على البحث على المنهج الوصفي.

وتعد نظرية دافع الحماية (Protection Motivation Theory - PMT) واحدة من النظريات النفسية والاجتماعية التي تُستخدم لفهم كيفية استجابة الأفراد للمخاطر أو التهديدات وكيفية اتخاذهم للإجراءات الوقائية. قدمها (Rogers) في عام 1975، وتفسر النظرية الدوافع التي تحت الأفراد على حماية أنفسهم من تهديدات معينة بناءً على تصوراتهم للخطر وفعالية الاستجابة.^(xxiv)

توظيف النظرية:

في سياق هذا البحث الذي يسعى إلى تسلط الضوء على استخدام القائم بالاتصال في ليبيا لأساليب كشف التزييف العميق، تبيّن أن نظرية دافع الحماية تعد إطاراً نظرياً مناسباً يساعد على فهم كيفية استجابة القائم بالاتصال لمخاطر التزييف العميق والتفاعل مع أساليب الكشف، وذلك وفق عناصر النظرية التالية:

أولاً: عناصر نظرية دافع الحماية:

تقييم التهديد: يتعلق هذا الجانب بتصوير القائم بالاتصال في ليبيا لمخاطر التزييف العميق. ويشمل:

1. شدة التهديد (Perceived Severity) :

- مدى إدراك القائمين بالاتصال لخطر التزييف العميق على مصداقية الأخبار، والتأثير على الرأي العام، والثقة في وسائل الإعلام الليبية بشكل عام.
- مثال: تأثير فيديو مزيف على سمعة جهة إعلامية عند نشره أو تضليل الرأي العام.

2. قابلية التعرض (Perceived Vulnerability) :

- شعور القائم بالاتصال بمدى تعرضهم أو مؤسساتهم الإعلامية لخطر التزييف العميق. وهل دفعهم هذا الشعور إلى استخدام الأساليب المناسبة للكشف عن التزييف العميق أم لا؟
- تقييم الاستجابة: يركز هذا التقييم على الكيفية التي ينظر بها القائم بالاتصال إلى أساليب مواجهة التزييف العميق. ويشمل:

1. فاعلية الاستجابة (Response Efficacy) :

- استعمال القائم بالاتصال لأدوات الكشف وأساليب الحماية التقنية والمهنية المتاحة.
- مدى اعتماده على ما يُصنف بالخبرة المكتسبة دون أدوات أو برامج.

2. الكفاءة الذاتية (Self-Efficacy) :

- مدى حاجة القائم بالاتصال لتطوير مهارته في جانب استخدام الأدوات التقنية والتعامل مع التزييف العميق بشكل فعال.
- تكلفة الاستجابة: ويشير إلى التحديات أو العقبات التي قد تعيق القائم بالاتصال عن استخدام الأساليب الوقائية. وتشمل:

• نقص التدريب أو الموارد التقنية.

• غياب الدعم المؤسسي للكشف عن التزييف العميق.

ثانياً: تطبيق النظرية في البحث:

أ. تحديد طبيعة الاستخدام:

• تحديد وفهم طبيعة استخدام القائمين بالاتصال في ليبيا لأساليب التزييف العميق.

• تصميم استبيان تقيس طبيعة أساليب الكشف.

ب. دراسة استراتيجيات الاستجابة:

- استخدام عناصر "تقييم الاستجابة بالنسب المئوية وعلى مقياس ليكرت الخماسي والاختبارات الإحصائية" لتحديد مدى إلمامهم بالتقنيات والأدوات والأساليب الفعالة لكشف التزييف ومدى توفرها في بيئات العمل ومدى استخدامها.

ج. تحليل التحديات:

- تسليط الضوء على مدى قناعتهم بأن التزييف العميق يشكل خطراً على المؤسسات الإعلامية وعلى الرأي العام كعوامل تحدد مدى تبني القائم بالاتصال للأساليب الوقائية.
- اقتراح حلول لمعالجة نقص الموارد أو التدريب.

العينة وأداة جمع البيانات:

تم اختيار مفردات العينة بالطريقة العشوائية قوامها (112 مفردة)، من مجموعة وسائل إعلام ليبية مختلفة الجمهور والتخصص، وصناع المحتوى على شبكات التواصل الاجتماعي الحكومية "التابعة للدولة"، وقد صمّم الباحث أداة استبيان إلكترونية، مع الأخذ في الاعتبار عدداً من الإجراءات، منها صياغة أسئلة الاستبيان بشكل واضح ومباشر، والتأكد من موثوقية وصدق أدوات القياس المستخدمة من خلال عرضها على عدد من المحكمين من الأكاديميين والممارسين وقد تم الأخذ بمقترحاتهم، ثم طبقت الأداة باستخدام طريقة الاختبار وإعادة الاختبار (Test-ReTest) بفارق زمني بلغ ثلاثة أيام على اثني عشر مبحوثاً، واستُخدم معامل ارتباط بيرسون لإيجاد درجة الارتباط بين الاختبارين الذي بلغ (0.85) وهو معامل ثبات كافٍ لأغراض هذا البحث.

التحليل الإحصائي:

النسب والتكرارات: لحساب التوزيع التكراري للإجابات.

اختبار مان - وتني (Mann- Whitney) وهو من الاختبارات اللا معلمية استخدمه الباحث في هذا كبديل لاختبار "ت" لأن الشروط المتعلقة بالاختبارات المعلمية لم تتحقق، كما يستخدم غالباً في حالة العينات الصغيرة^(xxv)، وبالنسبة لتفسير نتائج هذا الاختبار فنقوم على عكس مبدأ تفسير نتائج اختبار "ت".

اختبار كروسكال - واليس (Kruskal- Wallis) لاختبار الفروق بين رتب أكثر من متوسطين في حالة الاستقلال، وقد تم استخدامه بدلاً عن اختبار التباين الأحادي في حالة مخالفة البيانات لشروط الاختبارات المعلمية، ويقوم هذا الاختبار على جداول التوزيع الطبيعي لاختبار $(\chi^2 \text{ Chi-square})$ ^(xxvi).

ويتم استخدام الاختبارات اللامعلمية في حالة البيانات غير الموزعة توزيعاً طبيعياً، بالإضافة إلى صغر حجم العينة، واحتواء البيانات على قيم متطرفة قد تؤثر على المتوسط والانحراف المعياري، فضلاً عن عدم تجانس التباين، كما هو الحال في هذا البحث^(xxvii).

اختبار ليفين **Levene's Test**: تم استخدامه للتحقق من مدى تجانس التباين وهو أحد شروط استخدام الاختبارات المعلمية، حيث تُقبل الفرضية الصفرية (تجانس التباين $H_0: \sigma^2_1 = \sigma^2_2$)، إذا كانت قيمة دلالة الاختبار أكبر من $(\alpha = 0.05)$ ، وتُقبل الفرضية البديلة (التباين غير متجانس $H_1: \sigma^2_1 \neq \sigma^2_2$) إذا كانت قيمة دلالة الاختبار تساوي أو أصغر من $(\alpha = 0.05)$. وقد وجد الباحث أن جميع البيانات محل التحليل لم تستوف شرط استخدام الاختبارات المعلمية (البارامترية)، حيث تم إجراء اختبار ليفين عند مستوى معنوية $(\alpha = 0.05)$ ، فتبين أن مستوى دلالة الاختبار أصغر من $(\alpha = 0.05)$ في جميع البيانات محل الدراسة، وبهذا نقبل الفرضية البديلة (التباين غير متجانس $H_1: \sigma^2_1 \neq \sigma^2_2$)، وعليه فإن البيانات محل التحليل مخالفة لشرط تجانس التباين وبهذا لا يمكن إجراء الاختبارات المعلمية (البارامترية) وتستبدل باختباري (مان - ويتني U وكروسكال - واليس H).

المتوسط الموزون "المرجح": في الحالات التي تختلف فيها القيم التي تأتي من عدد من الأفراد على هيئة درجات، يجب تمثيل كل قيمة بعدد أفرادها أو تكرارها أو وزنها، كما يفيد في ترتيب العبارات وفقاً لأعلى متوسط موزون^(xxviii) وذلك وفق التصنيف المستخدم في هذا البحث كالتالي:

جدول (1): تصنيف درجات مقياس ليكرت الخماسي والوزن النسبي المقابل لكل درجة

درجات مقياس ليكرت	الأولى	الثانية	الثالثة	الرابعة	الخامسة
المدى	من	1	1.8	2.6	3.4
	إلى	1.79	2.59	3.39	4.19
الوزن	من	20%	36%	52%	68%
النسبي	إلى	35.9%	51.9%	67.9%	83.9%
المقياس	لا أعرف عنه شيئاً	معرفة محدودة جداً	معرفة محدودة	معرفة عميقة	معرفة عميقة جداً
	غير متوفرة مطلقاً	غير متوفرة	ليست لدي فكرة	متوفرة	متوفرة وبكثافة

التزييف العميق: المفهوم والتقنيات وأثره على وسائل الإعلام.

بات التزييف العميق (Deepfake) - وهو من الظواهر التقنية الحديثة - يثير القلق في الأوساط الإعلامية والتكنولوجية، كونه يعتمد على تقنيات الذكاء الاصطناعي لإنشاء محتوى مرئي وصوتي مزيف إلى بدرجة أنه يبدو حقيقياً بشكل لافت. حيث يشكل التطور الكبير لتقنيات التزييف بواسطة الذكاء الاصطناعي، تهديداً على الأمن الشخصي والعام، فضلاً عن كونه يؤثر في مصداقية الإعلام بما في ذلك منصات التواصل الاجتماعي، خصوصاً في الأوضاع السياسية والاجتماعية مثل ليبيا.

ويشير مفهوم التزييف العميق إلى استخدام تقنيات الذكاء الاصطناعي، لا سيما تقنيات التعلم العميق، لتعديل أو إنشاء محتوى مرئي أو صوتي يُحاكي الواقع بشكل كبير، ويصعب اكتشافه، ولكنه مزيف في حقيقته. حيث تعمل هذه التقنية على تعديل مقاطع الفيديو والصور أو إنتاج مقاطع جديدة بالكامل، وقد تظهر أشخاصاً في مواقف لم تحدث فعلياً. ومن الناحية التقنية، يعتمد التزييف العميق على الشبكات العصبية الاصطناعية (Neural Networks) و"تعلم الآلة" لتوليد محاكاة دقيقة للواقع، مما يفتح المجال لاستخدامات ضارة تشمل نشر الأخبار الزائفة والتلاعب بالرأي العام.^(xxix)

تقنيات التزييف العميق:

- الشبكات التوليدية التنافسية (GANs) تعد الشبكات التوليدية التنافسية من أبرز التقنيات المستخدمة في التزييف العميق، وتعتمد هذه الشبكات على وجود نموذجين يتنافسان: الأول يهدف إلى توليد محتوى مزيف، بينما الثاني يهدف إلى التمييز بين المحتوى الحقيقي والمزيف. من خلال هذا التفاعل التنافسي، تتحسن جودة التزييف تدريجياً، مما يجعل من الصعب اكتشاف التلاعب في المحتوى المزيف.^(xxx)

- تبديل الوجه (Face Swapping) هذه التقنية تستبدل وجه شخص بآخر في مقاطع الفيديو، مما يتيح إمكانات كبيرة في التلاعب بالصور. تُستخدم الخوارزميات المتقدمة لتحديد ملامح الوجه بدقة، ثم استبدالها بملامح شخص آخر، مما يعزز من الواقعية الزائفة (Korshunov & Marcel, 2018).

2018)

- إعادة تمثيل الوجه (Facial Reenactment) تقنية إعادة تمثيل الوجه تستخدم الذكاء الاصطناعي لمطابقة تعبيرات الوجه بين شخصين مختلفين، بحيث يظهر أحد الأشخاص وهو يؤدي حركات وتعابير شخص آخر، دون أن يكون قد فعل ذلك في الواقع. تعتمد هذه التقنية على أنظمة تعلم الآلة المدربة على قراءة وتحليل الحركات الدقيقة للوجه.^(xxxi)

التزييف العميق في وسائل الإعلام: يُشكل التزييف العميق تحدياً كبيراً في مجال الإعلام، حيث يمكن أن يُستخدم في تشويه الحقائق أو نشر الأخبار الزائفة، مما يؤثر في مصداقية وسائل الإعلام. **تأثيرات التزييف العميق:**

- تشويه الحقائق من خلال تعديل الفيديوهات والصور، يمكن إنشاء صور غير دقيقة للواقع، مما يؤدي إلى تأثير سلبي على مصداقية الأخبار، ما يُشكل تهديداً لاستقرار السياسي في الدول، وأيضاً يمكن استخدام التزييف العميق لتشويه سمعة الأفراد من خلال تقديمهم في مواقف مزيفة، مما يُسهم في انقسامات اجتماعية.

أساليب كشف التزييف العميق باستخدام الذكاء الاصطناعي

نظراً للتهديدات التي يُسببها التزييف العميق، تم تطوير العديد من الأساليب للكشف عن هذا النوع من التلاعب، ويقصد بالأساليب الطرق التي تعتمد عليها خوارزميات الذكاء الاصطناعي للكشف عن التزييف العميق الذي أُنتج أصلاً بالذكاء الاصطناعي، حيث تعتمد تقنيات الكشف على عدة أساليب مختلفة، منها:

- أسلوب تحليل الخصائص البصرية الدقيقة للكشف عن العيوب الدقيقة التي لا تكون مرئية بالعين البشرية، مثل عدم تطابق الإضاءة، وحواف مشوشة أو مشوهة، أو حركات غير طبيعية في الوجه أو العينين، وتعتمد هذا الأسلوب، أدوات (التعلم العميق مثل الشبكات العصبية التلافيفية (Convolutional Neural Networks - CNNs).^(xxxii)

- أسلوب الكشف عن التناقضات الصوتية والبصرية من خلال الجمع بين تحليل المكونات الصوتية والمرئية في مقاطع الفيديو لكشف الاختلافات، مثل عدم تطابق الصوت مع حركة الشفاه، وعدم التزامن بين الحركات الجسدية والإيقاعات الصوتية، ومن هذه الأدوات التي تنتهج هذا الأسلوب تقنيات دمج النماذج متعددة الوسائط (Multimodal Fusion Models).^(xxxiii) أسلوب تحليل الأنماط الزمنية للفيديو (Temporal Analysis) من خلال تحليل استمرارية الأنماط الزمنية في الفيديوهات المزيفة، إذ تظهر أحياناً انقطاعات غير طبيعية في الحركات، وتغيرات فجائية في الإطارات، ومن بين الأدوات التي تنتهج هذا الأسلوب: شبكات LSTM أو GRU التي تركز على معالجة البيانات المتسلسلة.^(xxxiv)

- أسلوب تحليل معدل الوميض والأنماط الحيوية، ويركز هذه الأسلوب على دراسة الأنماط الحيوية البشرية الطبيعية مثل: معدل وميض العين، والإيماءات الطبيعية، ومن بين الأدوات التي تعتمد على هذا الأسلوب خوارزميات التعلم الآلي التي تقارن الأنماط الطبيعية مع الأنماط المزيفة.^(xxxv)

- أسلوب الكشف بالاعتماد على بصمة الفيديو الرقمية (Digital Video Fingerprinting)، ويستخدم هذا الأسلوب تحليل البصمة الرقمية للمحتوى المرئي للتحقق من مصدره ومصادقيته، ومن الأدوات الشائعة تقنيات Blockchain لحفظ البيانات الأصلية ومقارنتها مع المحتوى المشكوك فيه. (xxxvi)
- أسلوب الكشف القائم على تحليل الشبكات العصبية المزيفة، حيث يعمل هذا الأسلوب على تحديد الأنماط التي تولدها الشبكات التوليدية العميقة (Generative Adversarial Networks - GANs)، مثل التشوهات في توزيع البيانات، والبصمات المميزة للشبكات التوليدية، ومن بين الأدوات التي تعتمد هذا الأسلوب خوارزميات GAN-specific Detectors. (xxxvii)
- وكل هذه الأساليب تستخدم تقنيات الذكاء الاصطناعي للكشف عن التلاعبات في المحتوى من خلال تدريب الأنظمة على التعرف على الأنماط التي تشير إلى التزييف، مثل الارتباك في الإضاءة أو الحركة غير الطبيعية. (xxxviii)

أمثلة لمنصات وبرامج الكشف عن التزييف العميق:

1. منصة Sentinel : تعد Sentinel منصة حماية متقدمة تعتمد على الذكاء الاصطناعي لتحليل الوسائط الرقمية وتحديد ما إذا كانت قد تعرضت للتلاعب. تستخدم خوارزميات متطورة للكشف عن التزييف العميق، مما يساعد في الحفاظ على مصداقية المعلومات المتداولة عبر الإنترنت.
2. كاشف التزييف العميق في الوقت الحقيقي (FakeCatcher) : طوّرت شركة إنتل تقنية FakeCatcher التي تتيح اكتشاف مقاطع الفيديو المزيفة بدقة تصل إلى 96% وفي أجزاء من الثانية. تعتمد هذه التقنية على تحليل الإشارات الحيوية الدقيقة في الفيديوهات، مثل تدفق الدم، للكشف عن التلاعب.
3. تقنية تدفق الدم: تستخدم هذه التقنية خوارزميات متقدمة لتحليل التغيرات في لون العروق نتيجة لتدفق الدم، مما يساعد في تمييز الأشخاص الحقيقيين عن المزيفين. تُترجم هذه التغيرات إلى خرائط تُستخدم لتحديد مصداقية الفيديوهات.
4. تقنية "نحن نتحقق" (We Verify) : تعتمد هذه التقنية على قاعدة بيانات عامة قائمة على تقنية البلوكشين للمزيفات المعروفة. تتيح للمستخدمين التحقق من صحة الوسائط الرقمية من خلال مقارنة المحتوى مع قاعدة البيانات المتاحة.
5. أداة مصادقة الفيديو من شركة Microsoft : تقدم الشركة أداة مصادقة الفيديو التي تمكن من تحليل عناصر التزييف العميق في الصور ومقاطع الفيديو. تُستخدم هذه الأداة للكشف عن التلاعبات الدقيقة التي قد لا تكون مرئية بالعين المجردة.
6. تقنية عدم التطابق (Phoneme-Viseme) : تعمل هذه التقنية على التحقق من مطابقة بصمات حركة الفم مع الصوت المنطوق، مما يساعد في الكشف عن التناقضات في الفيديوهات التي قد تشير إلى تزييف.

وتُظهر هذه التقنيات التقدم المستمر في مجال الكشف عن التزييف العميق، مما يسهم في تعزيز مصداقية المعلومات المتداولة عبر الوسائط الرقمية. (xxxix)

نتائج البحث:

توصيف مفردات العينة:

جدول(2): توزيع مفردات العينة حسب النوع "ذكور وإناث"

النوع	العدد	%
ذكور	85	75.9
إناث	27	24.1
المجموع	112	100.0

تشير النتيجة إلى وجود "قجوة" في توزيع العاملين في وسائل الإعلام الليبية حسب النوع، حيث يشكل الذكور ثلاثة أرباع العينة بنسبة بلغت 75.9%، مقابل 42.1% للإناث. هذا التفاوت ربما يعكس حالة التمثيل غير المتوازن بين الجنسين في المجال الإعلامي، ولعل النسبة المنخفضة للنساء في وسائل الإعلام الليبية تعكس تحديات اجتماعية وثقافية ومهنية متعددة.

جدول(3): توزيع مفردات العينة حسب التخصص العلمي

التخصص العلمي	العدد	%
إعلام	75	67.0
علوم إنسانية	15	13.4
علوم تطبيقية	22	19.6
المجموع	112	100

تظهر البيانات أن غالبية القائم بالاتصال في وسائل الإعلام الليبية ونسبة 67% هم متخصصون في مجالات الإعلام، مما يعكس التركيز الكبير على التخصص المهني في هذا القطاع. في المقابل، يشكل المتخصصون في العلوم التطبيقية نسبة 19.6%، يليهم المتخصصون في العلوم الإنسانية بنسبة 13.4%، وتعد هذه النتيجة طبيعية لأن سوق مخرجات المؤسسات التعليمية المتخصصة في الإعلام غالباً ما يكون وسائل الإعلام المختلفة.

جدول(4): توزيع مفردات العينة حسب المستوى الدراسي

المستوى الدراسي	العدد	%
متوسط	9	8.0
جامعي	82	73.2
ماجستير	17	15.2
دكتوراه	4	3.6
المجموع	112	100

تشير نتائج الجول السابق إلى هيمنة خريجي التعليم الجامعي على وسائل الإعلام الليبية 73.2%، مع تمثيل جيد لحملة درجة الماجستير 15.2%، ما يعكس المستوى الأكاديمي المرتفع بين أفراد العينة على رغم محدودية الفرص لحملة الدكتوراه 3.6%، وتراجع الاعتماد على المؤهلات المتوسطة إلى 8%.

جدول (5): توزيع عينة الدراسات على المؤسسات الإعلامية الليبية

جهة العمل	العدد	%
-----------	-------	---

29.5	33	قناة فضائية حكومية
17.9	20	راديو حكومي
22.3	25	منصة حكومية
8.0	9	صحيفة حكومية
2.7	3	وكالة أنباء حكومية
19.6	22	قناة فضائية خاصة
0.0	0	راديو خاص
0.0	0	منصة خاصة
0.0	0	صحيفة خاصة
100	112	المجموع

من خلال الجدول السابق يتبين أن أكثر من ربع العينة البالغ عددها 112 مفردة تمثل القنوات الفضائية الحكومية التابعة للدولة بنسبة 29.5%، تليها المنصات الحكومية بنسبة 22.3%، ثم القنوات الليبية الخاصة بنسبة 19.6%، ثم الإذاعات المسموعة الحكومية بنسبة 17.9%، فما حظيت الصحف الحكومية بنسبة 8% فقط.

جدول (6): توزيع عينة الدراسة حسب طبيعة العمل - الوظيفة

%	العدد	طبيعة العمل - الوظيفة
20.5	23	معد برامج
11.6	13	مخرج
19.6	22	صانع محتوى
8.9	10	محرر
3.6	4	مذيع أخبار
11.6	13	منفذ
8.9	10	مقدم برامج
2.7	3	منتج
2.7	3	مونتير
4.5	5	صحفي
5.4	6	مسؤول
100.0	112	المجموع

تعكس البيانات توزيعاً متنوعاً للأدوار الوظيفية لمفردات العينة من بين العاملين في وسائل الإعلام الليبية، حيث شكّل معدو البرامج بنسبة 20.5%، تليها نسبة 19.6% لصناع المحتوى في المنصات الحكومية، تليها نسبة جيدة للأدوار الداعمة للإنتاج، كالمخرجين ومنفذي البرامج بنسبة 11.6% لكل منهما، ثم يأتي المحررون ومقدمو البرامج بنسبة 8.9% لكل منهما، تليها بقية النسب.

جدول (7): توزيع عينة الدراسة حسب سنوات الخبرة

%	العدد	سنوات الخبرة
16.1	18	أقل من 5 سنوات
13.4	15	من 5 إلى 9 سنوات
28.6	32	من 10 إلى 14 سنة
8.9	10	من 15 إلى 19 سنة
33.0	37	20 سنة وأكثر

المجموع	112	100.0
---------	-----	-------

تظهر تركيبة عينة الدراسة أنها تجمع بين الخبرة والحدثة، في إشارة على استفادة وسائل الإعلام الليبية من الخبرات الطويلة 33% بالتزامن مع الخبرات القصيرة 16.1% و 13.4%، كما أن وجود نسبة كبيرة من ذوي الخبرة الطويلة يعزز استقرار المؤسسات الإعلامية، لكن التوسع في استقطاب الجيل الجديد ضروري لمواكبة التغيرات التقنية.

عرض وتحليل النتائج:

بسؤال مفردات العينة عن مدى معرفتهم بتقنيات التزييف، فجاءت الإجابات على النحو التالي:

جدول (8): معرفة مفردات العينة بتقنيات كشف التزييف العميق

الدرجة	معرفة عميقة	معرفة محدودة	لا أعرف	معرفة محدودة	معرفة عميقة	معرفة محدودة	معرفة عميقة	معرفة محدودة	معرفة عميقة	معرفة محدودة	معرفة عميقة	معرفة محدودة
التكرار	16	22	17	27	30	112	0.274	2.705	54.1%	0.273	محدودة	محدودة
%	14.3	19.6	15.2	24.1	26.8	100.0						

أظهرت النتائج أن الوزن النسبي الذي سجل 54.1% مستوى معرفة "محدود" وفق التصنيف المعتمد، حيث بلغ المتوسط الموزون 2.705، وهو يقع ضمن نطاق (2.6 - 3.39) على سلم ليكرت الخماسي، مما يضع تصنيف مستوى المعرفة ضمن الـ "محدود"، أما معامل الاختلاف (0.273) فيعكس وجود تباين متوسط بين أفراد العينة، مما يعني أن هناك بعض الفروقات في مستوى المعرفة بين المشاركين، لكنها ليست شديدة التطرف.

وبسؤال مفردات العينة عن مدى توافر تقنيات الكشف عن التزييف العميق في مؤسساتهم

الإعلامية) جاءت إجاباتهم على النحو التالي:

جدول (9): استجابات عينة البحث عن مدى توفر تقنيات كشف التزييف العميق في مؤسساتهم:

العدد	متوفرة	متوفرة بشكل	ليست لدي فكرة	غير متوفرة	مطلقة	مطلقة	مطلقة	مطلقة	مطلقة	مطلقة	مطلقة	مطلقة
العدد	9	15	32	30	26	112	0.2517	2.563	51.3%	0.445	غير متوفرة	غير متوفرة
%	8.0	13.4	28.6	26.8	23.2	100.0						

تشير نتائج الجدول إلى أن تقنيات كشف التزييف العميق غير متوفرة بشكل كافٍ في المؤسسات الإعلامية الليبية، حيث أفادت غالبية العينة بعدم توفر هذه التقنيات أو عدم امتلاكهم فكرة عنها، حيث سجل المتوسط الموزون (2.563) ويقع ضمن نطاق (1.8 - 2.59) على سلم ليكرت الخماسي، ما يشير إلى أن التقنيات "غير متوفرة" بشكل عام، وإيضاً قيمة الوزن النسبي 51.3% تؤكد هذه النتيجة، فيما أشار معامل الاختلاف (0.445) إلى وجد تباين مرتفع نسبياً بين إجابات العينة، مما يعني اختلاف وجهات النظر بين القائم بالاتصال.

وبسؤال العينة عن مدى معرفتهم واستعمالهم لتقنيات كشف التزييف العميق، جاءت إجاباتهم على النحو التالي:

جدول (10): استخدام عينة البحث بأساليب وأدوات الكشف عن التزييف العميق ومدى استخدامها

الخيارات	العدد	%
استخدمها	9	8.0
استخدمها نادراً	16	14.3
غير متاحة لي	56	50.0
لست مقتنعاً بها	1	0.9
لا أعرف عنها شيئاً	30	26.8
المجموع	112	100

أظهرت البيانات أن 59% من أفراد العينة أفادوا بأن أدوات كشف التزييف غير متاحة لهم، وهي النسبة الأكبر، مما يشير إلى ضعف البنية التحتية التقنية والدعم المؤسسي الذي يعاني منه القائمون بالاتصال في وسائل الإعلام الليبية، فيما أفاد 28.6% بأنهم لا يعرفون شيئاً عن هذه الأدوات، مما يعكس فجوة معرفية كبيرة وافتقاراً للتوعية اللازمة حول أساليب كشف التزييف العميق، واللافت أن نسبة من أشاروا إلى أنهم "يستخدمون هذه الأدوات" كانت 8% فقط، وهي نسبة ضئيلة تشير إلى قلة تطبيق هذه التقنيات في المؤسسات الإعلامية، أما النسبة التي أفادت باستخدام هذه الأدوات "نادراً" بلغت 14.3% ما يعني أن الاستخدام محدود وغير منهجي، فيما بلغت نسبة من أبدوا عدم اقتناعهم بالأدوات كانت منخفضة جداً عند 0.9%، وهو ما قد يكون دليلاً على أن المشكلة لا تتعلق بعدم القناعة بقدر ما ترتبط بعدم المعرفة أو التوفر.

وهذا يشير إلى أن دافع الحماية لدى القائمين بالاتصال مرتفع، ولكن لم يُترجم إلى أفعال بسبب غياب الدعم المؤسسي، ما وضوح الفجوة بين الوعي والتطبيق.

وبسؤال أفراد العينة عن التقنيات والأساليب المتبعة من قبلهم عن كيفية كشف التزييف العميق، فجاءت النتائج كالتالي:

جدول (11): التقنيات والأساليب المتبعة لكشف التزييف العميق لدى عينة البحث

التقنيات والأساليب	العدد	%
نحن نتحقق WeVerify	6	5.4
FakeCatcher	8	7.1
منصة Sentinel	7	6.3
استخدام عدم التطابق Phoneme-Viseme	5	4.5
أداة مصادقة الفيديو من Microsoft	11	9.8

الخبرة	66	58.9
أخرى تذكر:		
الرجوع إلى المسؤول	1	0.9
الرجوع إلي (منصات - مواقع - مصادر - جهات) رسمية وموثوقة	12	10.7
لا اعتمد على مصدر أشك في مصداقيته	1	0.9
لا يُعتمد بكل ما ينشر عبر شبكات التواصل الاجتماعي	1	0.9
مجموع الاستجابات	118	

تكشف بيانات الجدول السابق أن الاعتماد (المفرط) على الخبرة الشخصية هو الطريقة الأكثر شيوعاً بين عينة الدراسة، ويشير ذلك إلى نقص واضح في استخدام الأدوات التقنية أو غياب التدريب المناسب على تطبيقها، حيث أشارت نسبة 58.9% من مجموعهم إلى أنهم يعتمدون على خبرتهم، رغم أن الخبرة تُعد عنصراً مهماً، ولكن ليس في هذا الجانب، فهي لا تكفي للتعامل مع التزييف العميق الذي يتطلب حلولاً تقنية متقدمة، واللافت أن نسبة 10.7% من مجموعهم أشاروا إلى أنهم يرجعون إلى (منصات - مواقع - مصادر) رسمية وموثوقة، فيما يأتي اعتماد مفردات العينة على منصات وبرامج كشف التزييف متدنية بشكل متدنٍ، وهذا يرجع أساساً إلى عدم توافرها في مؤسساتهم " كما بينت نتائج الجدول رقم (9) السابق".

وهذه النتائج تستدعي تحليل استجابات العينة حيال (المعرفة، والتوفر، والاستخدام، والأساليب المتبعة) والمقارنة بينها على النحو التالي:

- الجدول 8 الذي يظهر مدى معرفة العينة بتقنيات كشف التزييف، والجدول 9 الذي يظهر مدى توفر التقنيات المعنية بالكشف، حيث أظهر الجدولان أن قلة المعرفة تتزامن مع نقص واضح في توفر التقنيات، هذا يعكس ضعفاً مؤسسياً مزدوجاً يتمثل في غياب التوعية وقلة الدعم الفني، ما يؤدي إلى فجوة في تطبيق أدوات كشف التزييف.
 - الجدول 9 الذي يظهر مدى توفر التقنيات المعنية بالكشف، والجدول 10 الذي يظهر مدى استخدام التقنيات للكشف. أوضح أن "قلة" التقنيات انعكست بشكل مباشر على الاستخدام، حيث إن 59% من العينة أشاروا إلى أن الأدوات غير متاحة لهم، مما يفسر ضعف النسبة التي تستخدم هذه الأدوات فعلياً (8%).
 - الجدول 8 الذي يظهر مدى معرفة العينة بتقنيات كشف التزييف والجدول 11 الذي يظهر مدى اعتماد عينة البحث على التقنيات الخاصة بالكشف، حيث اتضح أن قلة المعرفة أدت إلى اعتماد مفرط على الخبرة الشخصية (58.9%) بدلاً من استخدام أدوات وتقنيات متخصصة.
- ونخلص من هذه المقارنة أن النسبة الكبيرة التي تعتمد على الخبرة الشخصية (58.9%) تتناقض مع الطبيعة التقنية المعقدة للتزييف العميق، حيث إنَّ الخبرة البشرية ليست كافية للتعامل مع هذه التحديات، ورغم أن أكثر من نصف العينة أفادت أن الأدوات غير متوفرة (51.3%)، فإن جزءاً من

العينة أبدت معرفة "عميقة" أو "عميقة جداً" على التوالي (14.3% و 19.6%)، مما يشير إلى عدم استغلال المعرفة المكتسبة بسبب ضعف البنية التحتية.

وبالتالي يتضح أن هناك علاقة طردية بين عدم توفر التقنيات في (الجدول 9) وضعف الاستخدام (الجدول 10) تشير إلى أن المؤسسات الإعلامية لم توفر بيئة داعمة للاستفادة من أدوات الكشف عن التزييف.

كما وتشير المقارنة إلى وجود فجوة معرفية كبيرة كما في (الجدول 8) وهي مرتبطة أساساً بعدم اعتماد وسائل متقدمة في الكشف عن التزييف العميق كما في الجدول (الجدول 11).

وهذا ما يعني أن هناك ضعفاً في المعرفة التقنية وضعفاً في الدعم المؤسسي، ما أدى إلى الاعتماد غير الفعال على الخبرة الشخصية، ما نتج عنه فجوة تطبيقية، ورغم وجود معرفة جزئية لدى بعض أفراد العينة، إلا أن غياب الإمكانيات والدعم المؤسسي حال دون تحويل هذه المعرفة إلى ممارسات عملية.

وفي جانب التدريب على أساليب وأدوات الكشف عن التزييف العميق سألنا عينة الدراسة عن تلقيهم لدورات في هذا الشأن أم لا؟ فجاءت الإجابات على النحو التالي:

جدول (12): تدريب مفردات العينة حول كيفية التعرف على التزييف العميق والتحقق من صحة المعلومات.

العدد	%	
15	13.4	نعم
97	86.6	لا
112	100.0	المجموع

تكشف بيانات الجدول السابق أن التدريب المتخصص في الكشف عن التزييف العميق نادر أو محدود، حيث أفادت نسبة 13.4% فقط من عينة الدراسة إلى أنهم تلقوا تدريباً في هذا المجال، فيما أفادت نسبة 86.6% وتمثل الغالبية العظمى إلى أنها لم تتلق أي نوع من التدريب حول كيفية التعرف على التزييف العميق أو التحقق من صحة المعلومات.

ولمعرفة مدى حماس أفراد العينة للتدريب في مجال كشف التزييف العميق، سُئلوا عن رغبتهم في الحصول على تدريب في هذا المجال، فجاءت الإجابات كالتالي:

جدول (13): رغبات أفراد عينة البحث في تلقي التدريب في مجال كشف التزييف العميق.

العدد	%	
110	98.2	نعم
2	1.8	لا
112	100.0	المجموع

تظهر بيانات الجدول السابق أن إدراك (عينة) البحث التي تمثل القائم بالاتصال في وسائل الإعلام الليبي بأهمية التدريب في مجالات كشف التزييف العميق كبير، فالغالبية الساحقة (98.2%) العينة أبدت اهتماماً كبيراً ورغبة واضحة في تلقي التدريب في هذا المجال كشف التزييف، مما يعكس احساسهم بالمسؤولية الكبيرة الملقاة على عواتقهم.

وبسؤال مفردات العينة عن الدور الذي يمكن أن تلعبه المؤسسات الإعلامية والجهات الحكومية في مكافحة التزييف العميق؟ جاءت الإجابات على النحو التالي:

جدول (14): الدور الذي ينبغي أن تلعبه المؤسسات الإعلامية حيال التزييف العميق

العدد	%	الدور
110	98.2	عقد دورات تدريبية وورش عمل
105	93.8	عدم السماح بنشر أي مقطع صوتي أو صورة أو مقطع فيديو إلا بعد التأكد من مصداقيته
103	92.0	الاعتماد على ما ينشر في المصادر الموثوقة
102	91.1	توفير برامج وأدوات الكشف في كل وسائل الإعلام الليبية
48	42.9	التوعية من خلال عقد مؤتمرات وندوات علمية للقائم بالاتصال
47	42.0	مكافحة المزيفين وسن قوانين تجرم التزييف
28	25.0	التوعية من خلال وسائل الإعلام
أخرى تُذكر:		
1	0.9	التركيز على تحليل الأخبار والتدقيق في مصدرها
1	0.9	وضع قانون الحماية والاختراق
1	0.9	البحث عن المعلومة الدقيقة مع المقارنة بالتفاصيل
1	0.9	توطيد العلاقات بين مختلف وسائل الإعلام والتشاور بينها حيال الأخبار المزيفة
1	0.9	تشجيع اقسام للتحقق من الاخبار قبل النشر في غرف التحرير
548		مجموع الاستجابات

تُظهر النتائج توافقاً كبيراً بين مفردات عينة الدراسة بشأن أهمية التدريب والتقنيات المتقدمة لكشف ومكافحة التزييف العميق، مع أهمية التحقق من مصداقية الأخبار والمحتوى المنشور. حيث ترى الغالبية 98.2% أن التدريب هو الحل الأمثل لمكافحة التزييف العميق، هذه النسبة العالية تعكس أهمية التدريب المستمر وتطوير المهارات الإعلامية لكشف التزييف من وجهة نظر القائم بالاتصال، كما أكد جُلّ أفراد العينة على عدم السماح بنشر محتوى غير موثوق بنسبة 93.8%، وأن نسبة 92% يرون ضرورة الاعتماد على المصادر الموثوقة، وتؤكد نسبة 91.1% على أهمية توفير برامج وأدوات كشف التزييف، إلى جانب ذلك، ترى نسبة 42.9% أن التوعية من خلال مؤتمرات وندوات علمية مهمة للقائم بالاتصال، وبالتوازي مع ما سبق، ترى نسبة 42% ضرورة العمل على "مكافحة المزيفين وسن قوانين تجرم التزييف"، فيما أشارت نسبة 25% إلى أهمية استخدام وسائل الإعلام نفسها كأداة للتوعية ضد التزييف العميق. ولعل من بين أبرز الاقتراحات "تشجيع اقسام للتحقق من الاخبار قبل النشر في غرف التحرير"، وهي فكرة قابلة للتطبيق بحيث يمكن تضمين وسائل الإعلام قسم للتحقق ومكافحة وكشف التزييف العميق، تحال إليه المواد محل الشك، للتحقق منها، ويعد هذا المقترح قابلاً للتنفيذ بشكل غير مكلف وفي وقت قياسي.

التحليل الإحصائي

أولاً: اختبار متغير سنوات الخبرة ومدى استخدام منصات وأدوات وأساليب التحقق من التزييف العميق:

تم استخدام اختبار Kruskal-Wallis ، وهو اختبار "لا معلمي" يستخدم للتحقق من وجود فروق دالة إحصائية بين مجموعات ذات بيانات لا تحقق شروط التوزيع الطبيعي أو تجانس التباين، كما هو الحال في بيانات هذا البحث، فكانت القيم الناتجة على النحو التالي:

جدول (15): نتائج اختبار (Kruskal-Wallis) لاختبار الفرضية القائلة: يوجد فروق ذات دلالة إحصائية

بين متغير سنوات الخبرة واستخدام تقنيات كشف التزييف العميق

الدالة	P	H	درجات الحرية	Rank Sum	Std Dev	العدد	سنوات الخبرة
دال	0.0000	79.6635	4	273.5	1.21133	18	أقل من 5 سنوات
				500	0.72375	15	من 5 إلى 9 سنوات
				1726.5	1.84915	32	من 10 إلى 14 سنة
				659.5	1.76698	10	من 15 إلى 19 سنة
				3168.5	0.88021	37	20 سنة وأكثر

أظهرت نتائج الاختبار أن القيمة الاحتمالية للاختبار تساوي 0.0000 وهي أصغر من مستوى الدلالة $\alpha=0.05$ ما يشير إلى وجود فروق دالة إحصائية حسب سنوات الخبرة، ما يفيد قبول فرضية البحث القائلة: (هناك اختلاف بين مستوى الخبرة في المجال الإعلامي واستخدام أدوات الكشف عن التزييف العميق)، ورفض الفرضية الصفرية (H_0) القائلة لا توجد فروق ذات دلالة إحصائية عند مستوى معنوية $\alpha=0.05$.

ويمثل مجموع الرتب "Rank Sum" ترتيب القيم في كل مجموعة بناءً على مقياس الاختبار، والفروق الواضحة في "مجموع الرتب" تشير إلى وجود اختلافات بين المجموعات.

وذلك استناداً إلى أن مجموع رتب المجموعة (20 سنة وأكثر) حقق أعلى قيمة (3168.5) ما يشير إلى أنها تختلف دالياً عن المجموعات الأخرى، خصوصاً المجموعة (أقل من 5 سنوات التي تساوي 273.5)، فيما تشير نتائج الاختبار إلى أن المجموعة (من 10 إلى 14 سنة) حققت مجموع رتب (1726.5) وهي أيضاً تختلف عن المجموعات الأقل خبرة مثل مجموعات (أقل من 5 سنوات، ومن 5 إلى 9 سنوات).

جدول (16): نتائج اختبار مان ويتني (Mann-Whitney U test) للتحقق من وجود فروق ذات دلالة إحصائية بين

متغيري "النوع" (ذكور وإناث) واستخدام أدوات الكشف عن التزييف العميق.

النوع	العدد	الانحراف المعياري	مجموع الرتب	قيمة u المحسوبة	قيمة z المحسوبة	p-value	قيمة z الجدولية
ذكور	85	1.403	4867.0	1083.000	-0.439	0.661	1.99
إناث	27	2.136	1461.0				

بناءً على نتائج الاختبار لا توجد فروق ذات دلالة إحصائية بين الذكور والإناث في استخدام أدوات الكشف عن التزييف العميق في عينة الدراسة. حيث إن قيمة (p-value 0.661) أكبر من $\alpha=0.05$ مما يعني قبول الفرضية الصفرية (H_0) التي تفترض أن هناك عدم وجود اختلافات ذات دلالة بين المجموعتين.

جدول (17): نتائج اختبار (Kruskal-Wallis) للتحقق من وجود فروق ذات دلالة إحصائية بين متغيري التخصص العلمي واستخدام أدوات الكشف عن التزييف العميق

التخصص العلمي	العدد	Std Dev	Rank Sum	درجات الحرية	H	p-value
إعلام	75	2.54034	4087.5	2	1.0381	0.59509
علوم انسانية	15	2.06559	867.5			
علوم تطبيقية	22	2.36908	1373			

أظهرت نتائج الجدول السابق أن القيمة الاحتمالية (0.59509) أكبر من $\alpha = 0.05$ ، فإن النتيجة غير دالة إحصائياً. مما يعني قبول الفرضية الصفرية (H_0) التي تفترض أن هناك عدم وجود اختلافات ذات دلالة بين المجموعتين.

جدول (18): نتائج اختبار (Kruskal-Wallis) للتحقق من وجود فروق ذات دلالة إحصائية بين متغيري المستوى التعليمي واستخدام أدوات الكشف عن التزييف العميق

المستوى التعليمي	العدد	Std Dev	Rank Sum	درجات الحرية	H	p-value
متوسط	9	2.34521	575	3	1.12373	0.77135
جامعي	82	2.2442	4677			
ماجستير	17	2.55095	888			
دكتوراه	4	4.04145	188			

أظهرت نتائج الجدول السابق أن القيمة الاحتمالية (0.77135) أكبر من $\alpha = 0.05$ ، فإن النتيجة غير دالة إحصائياً. مما يعني قبول الفرضية الصفرية (H_0) التي تفترض أن هناك عدم وجود اختلافات ذات دلالة بين المجموعتين.

جدول (19): نتائج اختبار (Kruskal-Wallis) للتحقق من وجود فروق ذات دلالة إحصائية بين متغيري جهات العمل واستخدام أدوات الكشف عن التزييف العميق

جهات العمل	العدد	Std Dev	Rank Sum	درجات الحرية	H	p-value
قناة فضائية حكومية	33	2.6039	1627.5	5	9.45547	0.09222
راديو حكومي	20	2.42574	1198.5			
منصة حكومية	25	2.36079	1207			
صحيفة حكومية	9	1.71594	609.5			
قناة فضائية خاصة	22	1.98315	1560			
وكالة أنباء حكومية	3	0.57735	125.5			

أظهرت نتائج الجدول السابق أن القيمة الاحتمالية (0.09222) أكبر من $\alpha = 0.05$ ، فإن النتيجة غير دالة إحصائياً. مما يعني قبول الفرضية الصفرية (H_0) التي تفترض أن هناك عدم وجود اختلافات ذات دلالة بين المجموعتين.

جدول (20): نتائج اختبار (Kruskal-Wallis) للتحقق من وجود فروق ذات دلالة إحصائية بين متغيري "طبيعة العمل" واستخدام أدوات الكشف عن التزييف العميق

طبيعة العمل	العدد	Std Dev	Rank Sum	درجات الحرية	H	p-value
مقدم برامج	23	2.0182	1796	9	15.3265	0.08235
مخرج + منتج	16	2.55278	806.5			
صانع محتوى	22	1.99036	896			
محرر	10	2.39444	289.5			
مذيع أخبار	4	1.89297	139			

			938	1.89128	13	منفذ
			655.5	2.50333	10	معد برامج
			226	1.1547	3	مونتير
			256.5	2.96648	5	صحفي
			325	3.444803	6	مسؤول
تم دمج فنتي (مخرج ومنتج)						

أظهرت نتائج الجدول السابق أن القيمة الاحتمالية (0.08235) أكبر من $\alpha 0.05$ ، فإن النتيجة غير دالة إحصائياً. مما يعني قبول الفرضية الصفرية (H_0) التي تقترض أن هناك عدم وجود اختلافات ذات دلالة بين المجموعتين.

أهم النتائج:

أظهر البحث عدداً من النتائج أهمها:

1. **طبيعة الاستخدام محدودة:** أظهرت الدراسة أن غالبية القائمين بالاتصال في وسائل الإعلام الليبية يمتلكون معرفة بتقنيات وأساليب كشف التزييف العميق تتراوح وبنسب متباينة بين معرفة عميقة ومعرفة محدودة، لكن استخدام لها ضيف جداً، حيث بلغ الوزن النسبي للاستخدام 45.1%، والتي يقع ضمن نطاق (52% - 67.9%) وفق التصنيف المتبع، والمتوسط الموزون بلغ 2.705، وهو يقع ضمن النطاق (2.6 - 3.39) على سلم ليكرت الخماسي، مما يعني أن درجة الاستخدام تقع في مستوى "المحدودة"، فيما أشار معامل الاختلاف (0.273) إلى أن هناك تبايناً متوسطاً في المعرفة بين الأفراد..

2. **نقص الدعم المؤسسي:** أشارت النتائج إلى ضعف كبير في الدعم المؤسسي لتقنيات كشف التزييف العميق، فغالبية القائمين بالاتصال (70%) تقريباً يعرفون الأدوات التقنية، لكن عدم توفرها - بسبب نقص التدريب المتخصص والموارد التقنية - شكّل عائقاً أساسياً أمام استخدامها.

3. **دوافع الحماية مرتفعة:** أظهرت النتائج أن القائمين بالاتصال يتمتعون بدوافع حماية عالية جداً، وذلك للحفاظ على مصداقية مؤسساتهم، إلا أن هذه الدوافع لم تُستغل الاستغلال الأمثل وبالشكل الفعال؛ بسبب نقص التدريب والتقنيات المتاحة، ما يُشير إلى أنّ ولائهم إلى مؤسساتهم وحرصهم على مصداقيتها كبير، إلا أنّها توليهم ذلك الاهتمام المطلوب في جوانب التدريب والدعم التقني.

4. **الاعتماد على الخبرة الشخصية بشكل مفرط:** تعتمد نسبة كبيرة من القائمين بالاتصال (58.9%) على خبراتهم الشخصية في الكشف عن التزييف العميق، مما يشير إلى غياب التدريب المتخصص والأدوات التقنية.

5. **قلة التدريب المتخصص:** بينت النتائج أن نسبة ضئيلة (13.4%) من القائمين بالاتصال تلقوا تدريباً متخصصاً في الكشف عن التزييف العميق، في المقابل أعربت الغالبية العظمى (98.2%) عن رغبتهم في تلقي هذا النوع من التدريب.

6. **فروق ذات دلالة إحصائية:** أظهرت التحليلات وجود فروق ذات دلالة إحصائية بناءً على سنوات الخبرة، حيث إن ذوي الخبرة الطويلة (20 سنة وأكثر) هم الأقل استخداماً لأساليب الكشف عن

التزييف، بينما الخبرات القصيرة (أقل من خمس سنوات) هم الأكثر استخداماً. بالمقابل؛ لم تُظهر النتائج فروقاً بين النوع أو المستوى التعليمي أو التخصص العلمي ونحوها وطبيعة الاستخدام. **التحقق من صحة فرضيات البحث:**

إن التحقق من صحة الفرضيات يمثل خطوة أساسية في أي دراسة بحثية، إذ يهدف إلى قياس مدى توافق النتائج التجريبية مع التوقعات التي تم صياغتها في الفرضيات. في هذه الدراسة، تم طرح مجموعة من الفرضيات التي ترتبط باستخدام القائمين بالاتصال في وسائل الإعلام الليبية لأساليب كشف التزييف العميق بواسطة الذكاء الاصطناعي في ظل نظرية دافع الحماية. وفيما يلي للتحقق من صحة هذه الفرضيات بناءً على النتائج المستخلصة:

1. **الفرضية الأولى:** "استخدام القائمين بالاتصال لأساليب كشف التزييف العميق محدود بسبب نقص التقنية وغياب الدعم المؤسسي".
 - **النتيجة:** تم قبول الفرضية؛ حيث أكدت النتائج محدودية الاستخدام وضعف الدعم المؤسسي.
 2. **الفرضية الثانية:** "يفترض البحث أن دافع الحماية لدى القائم بالاتصال في وسائل الإعلام الليبية عالٍ، ولكن لم يتم استغلاله بشكل كافٍ بسبب نقص التدريب المتخصص وقلة الموارد التقنية المتاحة".
 - **النتيجة:** تم قبول الفرضية؛ حيث تبين أن دافع الحماية المرتفع لا يُستغل بشكل فعال من قبل المؤسسات والمعنيين والجهات المختصة.
 3. **الفرضية الثالثة:** "توجد فروق ذات دلالة إحصائية بين استخدام القائمين بالاتصال لتقنيات كشف التزييف العميق وبين المتغيرات الديموغرافية والمهنية".
 - **النتيجة:** تم قبول الفرضية جزئياً؛ حيث ظهرت فروق بناءً على سنوات الخبرة فقط.
 4. **الفرضية الرابعة:** "مستوى الدعم المؤسسي يؤثر على فعالية استخدام تقنيات كشف التزييف العميق".
 - **النتيجة:** تم قبول الفرضية؛ حيث تبين أن الدعم المؤسسي يلعب دوراً حاسماً.
 5. **الفرضية الخامسة:** "يعتمد القائمون بالاتصال على الخبرة الشخصية بسبب غياب الأدوات التقنية والتدريب المؤسسي".
 - **النتيجة:** تم قبول الفرضية؛ حيث أثبتت النتائج الاعتماد الكبير على الخبرة الشخصية.
 6. **الفرضية السادسة:** "يواجه القائمون بالاتصال عقبات تطبيقية بسبب نقص الأدوات والتقنيات المتاحة".
 - **النتيجة:** تم قبول الفرضية؛ حيث أكدت النتائج وجود عقبات تعوق التطبيق العملي.
- مناقشة النتائج في إطار نظرية دافع الحماية.**
- استناداً إلى نظرية دافع الحماية، التي تقترض أن الأفراد يسعون إلى حماية مصالحهم من المخاطر التي تهددها، أمكن تفسير نتائج البحث، حيث إن دوافع الحماية لدى القائمين بالاتصال مرتفعة، ما يعكس

إدراكهم لأهمية الحفاظ على مصداقية الأخبار، إلا أن ضعف الدعم المؤسسي ونقص التدريب المتخصص حالاً دون الاستفادة الكاملة من هذه الدوافع.

ففي الوقت الذي اتضح فيه أن القائمين بالاتصال يدركون المخاطر المرتبطة بالتزيف العميق وأثره السلبي على مصداقية الأخبار، إلا أن نقص الأدوات التقنية والتدريب يمثل عقبة أساسية تعيق ترجمة دوافع الحماية إلى إجراءات عملية فعالة، ما يبرز الحاجة إلى تعزيز الدعم المؤسسي وتوفير الموارد اللازمة لضمان تفعيل دوافع الحماية. **طبيعة الاستخدام وفق نظرية دوافع الحماية:**

طبيعة استخدام القائم بالاتصال في وسائل الإعلام الليبية تعتمد على الجهود الفردية والخبرة الشخصية، مع قلة المعرفة بتقنيات الكشف عن التزيف، ومحدودية توفرها في المؤسسات الإعلامية، مع وجود إدراك واضح لأهمية تقنيات كشف التزيف العميق لدى عينة البحث. **وفيما يلي توضيح لطبيعة الاستخدام وفقاً لنتائج البحث:**

- أظهرت النتائج أن القائمين بالاتصال يمتلكون دافعاً عالياً لحماية مصداقية الأخبار والحفاظ على ثقة الجمهور، وهو ما يتماشى مع الأساس النظري لنظرية دوافع الحماية، التي تؤكد أن الأفراد يتخذون إجراءات وقائية لحماية أنفسهم أو مصالحهم من التهديدات المحتملة.
 - دافع الحماية هذا يعكس الوعي الكبير بالمخاطر التي يفرضها التزيف العميق على مصداقية المؤسسات الإعلامية.
 - على الرغم من الدافع العالي للحماية، تبين أن نقص الموارد التقنية والتدريب المتخصص يشكلان عائقين رئيسيين يحولان دون تفعيل هذا الدافع.
 - هذا يتماشى مع فرضيات النظرية التي تفترض أن الفعالية المدركة (Perceived Efficacy) للأدوات والقدرة على التعامل معها تلعب دوراً كبيراً في تحديد استجابة الأفراد للتهديدات.
 - بسبب غياب الدعم المؤسسي والتقني، يعتمد القائمون بالاتصال بشكل كبير على الخبرة الشخصية في التمييز بين المحتوى الحقيقي والمزيف.
 - يشير ذلك إلى محاولة تغطية العجز في الإمكانيات المؤسسية باستخدام البدائل الفردية، ولكنها تبقى محدودة وغير منهجية.
 - أظهرت النتائج أن القائمين بالاتصال الذين يعملون في مؤسسات تقدم دعماً تقنياً أكثر كانوا أكثر قدرة على استخدام أساليب كشف التزيف العميق.
 - الدعم المؤسسي يمكن أن يعزز فعالية دافع الحماية من خلال توفير الموارد والتدريب اللازمين.
 - القائمون بالاتصال يظهرون إدراكاً كبيراً لخطورة التزيف العميق وتأثيراته السلبية على العمل الإعلامي، مما يزيد من استعدادهم لاستخدام أدوات الكشف في حال توفرت لهم الموارد اللازمة.
- طبيعة استخدام القائم بالاتصال لأساليب الكشف عن التزيف في إطار نظرية دوافع الحماية:**
- نظرية دوافع الحماية تشير إلى أن الاستجابة للتهديد تعتمد على عاملين رئيسيين:**
1. الإدراك بالتهديد: القائمون بالاتصال يدركون حجم التهديد الذي يمثله التزيف العميق.

2. **الفعالية المدركة:** ضعف فعالية الأدوات المتاحة، ونقص التدريب أدى إلى الحد من الاستجابة الفعالة للتعامل مع التهديد.

وبناءً على ذلك، فإن طبيعة استخدام القائم بالاتصال لأساليب كشف التزييف العميق تعكس دافعاً قوياً للحماية؛ ولكنه غير مدعوم بالموارد المؤسسية والبنية التحتية اللازمة، مما يضعف قدرتهم على تطبيق هذه الأساليب بشكل فعال.

الخلاصة:

أظهرت النتائج أن طبيعة استخدام القائم بالاتصال وفق نظرية دوافع الحماية، تعتمد على الجهود الفردية والخبرة الشخصية، مع تدني مستويات معرفة واستخدام تقنيات كشف التزييف العميق، في المقابل هناك إدراك واضح لخطورة التزييف العميق لدى القائم بالاتصال، حيث يمتلك دافعاً قوياً للحفاظ على مصداقية الأخبار، ومع ذلك، تعيق محدودية الموارد التقنية والتدريب المؤسسي الفعالية التطبيقية لهذه الأساليب. وكشفت النتائج عن نقص واضح في المعرفة لدى عينة الدراسة بأساليب كشف التزييف العميق، مع محدودية توفرها في مؤسساتهم، مقترن بضعف استخدامها في المؤسسات الإعلامية الليبية، ويرتبط هذا القصور الحاد بعدم وجود تدريب مؤسسي كافٍ رغم ارتفاع دافع الحماية لدى القائمين بالاتصال، حيث أظهرت النتائج أن هناك محدودية في معرفة العينة بالتقنيات، تتزامن مع نقص كبير في توفر هذه الأدوات داخل المؤسسات الإعلامية، ويشير هذا التلازم إلى ضعف مؤسسي "مزدوج" يتمثل في غياب التوعية وقلة الدعم الفني، ما أدى إلى فجوة في التطبيق الفعلي لأساليب كشف التزييف العميق، كما أن قلة توفر التقنيات انعكس سلباً على الاستخدام، حيث أكد نسبة 59% من العينة عدم توفر الأدوات، ما فسر انخفاض نسبة من يستخدمونها فعلياً إلى 8% فقط، وهذا ذاته قد يكون الدافع إلى اعتماد قرابة 59% من العينة على الخبرة الشخصية بدلاً من التقنيات الخاصة. وعلى الرغم من أن نسبة (14.3%) من العينة تمتلك معرفة "عميقة" أو "عميقة جداً" (19.6%)، إلا أن هذه المعرفة غير مستثمرة عملياً بسبب محدودية التوفر وضعف البنية التحتية، ما يدل على أن حجم الفجوة بين المعرفة والتطبيق كبير جداً، حيث وضحت النتائج وجود علاقة بين ضعف توفر الأدوات كما وردت في (الجدول 9) وقلة استخدامها كما وردت في الجدول (الجدول 10) ما يشير إلى أن المؤسسات الإعلامية لا توفر بيئة داعمة للاستفادة من أدوات كشف التزييف، علاوة على ذلك؛ تبين أن الفجوة المعرفية (الجدول 8) مرتبطة بعدم تطبيق الأساليب المتقدمة (الجدول 11)، ما أدى إلى ضعف القدرة الفعلية على مواجهة التزييف العميق.

وبهذا خلصت النتائج إلى أن القائمين بالاتصال في المؤسسات الإعلامية الليبية يعانون من نقص في المعرفة التقنية، وضعف في الدعم المؤسسي، ما أدى إلى الاعتماد غير الفعال على الخبرة الشخصية، ورغم امتلاك بعض الأفراد معرفة نظرية، إلا أن غياب الإمكانيات "عدم توفرها" حال دون تحويلها إلى ممارسات عملية، ما يستدعي تعزيز التدريب المؤسسي، وزيادة الاستثمار في الأدوات التقنية، وتوفير بيئة داعمة لاستخدام تقنيات كشف التزييف العميق بفعالية.

التوصيات:

- تنظيم دورات تدريبية دورية حول تقنيات كشف التزييف العميق.
 - تضمين موضوعات التقييم النقدي وتحليل المصادر ضمن البرامج التدريبية.
 - توفير أدوات تقنية متطورة لتحليل المحتوى وكشف التزييف.
 - إنشاء شراكات مع شركات التقنية لتزويد المؤسسات الإعلامية بأحدث البرمجيات.
 - تنظيم ورش عمل ومؤتمرات توعوية حول مخاطر التزييف العميق.
 - إعداد برامج إعلامية تستهدف توعية الجمهور بمخاطر الأخبار المزيفة.
 - إنشاء مراكز بحثية مشتركة لدراسة التزييف العميق وتطوير تقنيات جديدة للكشف عنه.
 - إدخال موضوعات كشف التزييف ضمن المناهج الأكاديمية لطلاب الإعلام والصحافة.
 - الاستثمار في تقنيات الذكاء الاصطناعي لتطوير قدرات وسائل الإعلام.
 - تزويد غرف الأخبار بأدوات حديثة لتحليل المحتوى وكشف التزييف العميق.
 - تطوير مهارات القائمين بالاتصال في مجال الصحافة الاستقصائية.
 - إشراك الجمهور في تحقيقات التحقق من الأخبار عبر منصات تفاعلية.
 - وضع سياسات واضحة لضمان الشفافية والمصادقية في الأخبار.
 - إنشاء لجان مختصة بمراجعة مصداقية المحتوى المنشور.
 - سن قوانين تنظم استخدام أدوات كشف التزييف العميق وتفرض عقوبات على نشر الأخبار المزيفة.
 - تعزيز التعاون مع الهيئات التنظيمية لضمان التزام المؤسسات الإعلامية بالمعايير الأخلاقية.
- من خلال تطبيق هذه التوصيات، يمكن تعزيز قدرة القائمين بالاتصال في وسائل الإعلام الليبية على مواجهة تحديات التزييف العميق، مما يساهم في حماية مصداقية الأخبار وضمان تقديم محتوى موثوق للجمهور.

القيود:

1. حجم العينة المحدود: اقتضت العينة على 112 مفردة من القائمين بالاتصال في وسائل الإعلام الليبية، وهو عدد يعكس محدودية المجتمع الأصلي للعينة في ليبيا مقارنة بدول أخرى.
2. نقص الدراسات السابقة في السياق الليبي: لا توجد دراسات كافية حسب علم الباحث حول التزييف العميق والقائم بالاتصال في ليبيا، مما استدعى الاعتماد على أبحاث من سياقات أخرى لفهم الظاهرة. وتعكس هذه القيود طبيعة البحث لكنها لا تقلل من أهميته، بل تؤكد الحاجة إلى مزيد من الدراسات لتعميق الفهم وتطوير الحلول الممكنة.

الشكر والتقدير:

لا يسعنا في الختام إلا أن نتقدم بالشكر إلى كل من تجاوب مع الباحث من العاملين في المؤسسات الإعلامية الليبية العامة والخاصة.

وجزيل الشكر موصول إلى السيدة هبة الله خريش التي أسهمت في الترجمة من اللغة الإنجليزية إلى اللغة العربية، ما أثرى البحث وخاصة في ظل شح المصادر العربية.

- ⁱ - Hashmi, A., Shahzad, S. A., Lin, C.-W., Tsao, Y., & Wang, H.-M. (2024). Unmasking illusions: Understanding human perception of audiovisual deepfakes. *Computer Vision and Pattern Recognition*, 2, 123–145.
- ⁱⁱ - مركز المعلومات ودعم اتخاذ القرار. (2023، 25 مايو). الذكاء الاصطناعي ومخاطر التزييف العميق. مجلس - www.idsc.gov.eg: الوزراء المصري. متاح عبر
- ⁱⁱⁱ - Hall, P., & Richards, T. (2021). Artificial intelligence and deepfake technology: Challenges for journalists. *Journal of Media Studies*, 34 (4), 122–139.
- ^{iv} - Nguyen, D., & Pham, Q. (2021). Deepfake awareness among journalists and the role of media literacy. *Journal of Digital Journalism*, 29(1), 45–59.
- ^v - مركز المعلومات ودعم اتخاذ القرار. (2023، 25 مايو). مرجع سبق ذكره.
- ^{vi} - Mahmud, B. U., & Sharmin, A. (2020). Deep insights of deepfake technology: A review. *Dhaka University Journal of Applied Science & Engineering*, 5 (1&2), 13–23.
- ^{vii} - Liu, S., & Zhang, H. (2022). Exploring journalists' knowledge and attitudes toward AI-generated deepfake videos. *International Journal of Media and Communication Studies*, 42 (1), 23–35.
- ^{viii} - Simmons, L., & McDonald, F. (2022). The impact of deepfake technology on journalism ethics and integrity. *Journal of Ethics in Journalism*, 29 (3), 77–91.
- ^{ix} - Harris, M., & Clark, J. (2022). The role of media professionals in addressing the challenges of deepfake technology. *Media Ethics Review*, 38 (2), 59–75.
- ^x - محمود أحمد مصطفى (2025). التربية الإعلامية الرقمية وعلاقتها بمستوى تحصيل وعي طلاب الإعلام التربوي (1) مجلة البحوث الإعلامية، جامعة الأزهر، 73. بمخاطر تطبيقات التزييف العميق في إطار نظرية دافع الحماية 672..
- ^{xi} - Turner, M., & Brooks, L. (2023). The role of AI in identifying deepfake content in journalism. *Media Ethics and Technology Review*, 13 (2), 85–99.
- ^{xii} Zhang, T. (2022). Deepfake generation and detection, a survey. *Multimedia Tools and Applications*, 81(5), 6259–6276. <https://doi.org/10.1007/s11042-021-11733-y>
- ^{xiii} - شمس الدين، فتحي محمد (2022). رؤية القائم بالاتصال لمستقبل الإعلاميين في عصر الذكاء الاصطناعي (الجزء الثاني)، 1–26/24 العلمية لبحوث الإعلام، كلية الإعلام، جامعة القاهرة، 2022
- ^{xiv} - علي، خالد، والنبهاني، إيمان. (2022). الذكاء الاصطناعي والتزييف العميق: التحديات والفرص في الإعلام 210–229. (4) مجلة الدراسات الإعلامية، جامعة القاهرة، 30. العربي
- ^{xv} - عبد الله، هشام. (2022). تأثير الذكاء الاصطناعي على الصحافة العربية في مكافحة الأخبار المزيفة 30–144 الإعلامية، 32.
- ^{xvi} - Williams, S., & Green, R. (2021). Artificial Intelligence and Deepfakes: A Threat to Media Credibility. *International Journal of Communication*, 29 (1), 45–59.
- ^{xvii} - Mohammed, R., & Stevens, M. (2021). Journalists' awareness of deepfakes and its impact on news integrity. *Journal of Digital Journalism*, 19(4), 211–224.
- ^{xviii} - النمى، سامي، والحداد، فاطمة الزهراء. (2021). التزييف العميق في الإعلام: التحديات وأساليب الكشف. مجلة 123–146. (2) دراسات في الإعلام، 45
- ^{xix} - حسن، محمد، والعسيري، سارة. (2021). التأثيرات الاجتماعية والإعلامية للتزييف العميق في الإعلام العربي 87–106. (2) مجلة الدراسات الاجتماعية والإعلامية، 11
- ^{xx} منشورات الهيئة السعودية . - الهيئة السعودية للبيانات والذكاء الاصطناعي (سدايا). (2024)، مبادئ التزييف العميق للبيانات والذكاء الاصطناعي. متاح عبر <https://sdaia.gov.sa/ar/SDAIA/eParticipation/consulting/SDAIADeepfakesGuidelinesAr.pdf>
- ^{xxi} - Russell, S., & Norvig, P. (2016). *Artificial Intelligence: A modern approach* Pearson Education (p. 10).
- ^{xxii} - McQuail, D. (2010). *McQuail's Mass Communication Theory* (6th ed., p. 175). London: SAGE Publications.

- xxiii - Zhang, T. (2022). Deepfake generation and detection, a survey. *Multimedia Tools and Applications Journal*, 81 (5), 6259–6276. Retrieved from <https://doi.org/10.1007/s11042-021-11733-y>
- xxiv - للمزيد أنظر:
- Rogers, R. W. (1975). A protection motivation theory of fear appeals and attitude change. *The Journal of Psychology*, 91(1), 93–114.
- المجلة المصرية -. أبو الحسن، فاطمة شعبان. (2021). تأثير الآليات المعرفية ودافع الحماية على مدركات الشخص الثالث على العلاقة التنبؤية بإدمان مواقع التواصل (PMT) - سالم، انتصار محمد السيد. (2023). أثر متغيرات دافع الحماية مناح عبر: 361–472. (2) *المجلة العلمية لبحوث الصحافة*، 25. الاجتماعي والشعور بالاكتئاب لدى طالب الجامعات https://journals.ekb.eg/article_309023.html
- الشمراني، محمد موسى محمد. (2000). مشكلات استخدام تحليل التباين الأحادي والمقارنات البعدية وطرق جامعة أم القرى، كلية التربية، السعودية، 1421 هـ، ص. 35. رسالة ماجستير غير منشورة. علاجها
- خبريش، خالد أبو القاسم علي. (2015). مستقبل الإذاعات المحلية الليبية في ظل التغيرات السياسية والتقنية: دراسة كلية الآداب، جامعة المنيا، ص. 79. أطروحة دكتوراه غير منشورة تطبيقية)
- xxvii - Conover, W. J. (1980). *Practical Nonparametric Statistics* (2nd ed.). New York: John Wiley & Sons. P 423, 757, 974.
- xxviii مبادئ الإحصاء واستخداماتها في مجالات الخدمة الاجتماعية، ص 47. - كشك، محمد بهجت. (1996) الإسكندرية: دار الطباعة الحرة
- xxix - Chesney, R., & Citron, D. K. (2018). Deepfakes: A looming challenge for privacy, democracy, and national security. *California Law Review*, 107 (5), 1753–1808.
- xxx - Goodfellow, I., et al. (2014). Generative Adversarial Nets. *Advances in Neural Information Processing Systems*. 58–50, 27, Retrieved from: <https://arxiv.org/abs/1406.2661>
- xxxi - Thies, J., Zollhöfer, M., Stamminger, M., Theobalt, C., & Nießner, M. (2016). Face2Face: Real-time face capture and reenactment of RGB videos. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2387–2395.
- xxxii - Messer, B., et al. (2020). FaceForensics++: A Benchmark for Detecting Face Manipulations. *International Journal of Computer Vision*. 41–24, (1)128, <https://doi.org/10.1007/s11263.3-01204-019->
- xxxiii - Li, Y., Chang, M.-C., & Lyu, S. (2018). In icu oculi: Exposing AI-generated fake face videos by detecting eye blinking. *CVPR*, abs/1806.02877. Retrieved from <http://arxiv.org/abs/1806.02877>
- xxxiv - Xu, H., et al. (2021). Temporal Consistency and Detection of Deepfake Videos. *IEEE Transactions on Pattern Analysis and Machine Intelligence*. 1245–1233, (5)43, <https://doi.org/10.1109/TPAMI2020.3029237>.
- xxxv - Zhao, Y., et al. (2020). Biometric Video Deepfake Detection Techniques. *Journal of Biometrics*. 58–45, (4)8, <https://doi.org/10.1016/j.jbiom2020.06.007>.
- xxxvi - Shao, W., et al. (2021). Digital Fingerprinting for Video Authentication. *International Journal of Information Security*. 67–58, (2)20, <https://doi.org/10.1007/s102073-00512-021->.
- xxxvii - Goodfellow, I., et al. (2014). Generative Adversarial Nets. *Advances in Neural Information Processing Systems*. 58–50, 27, Retrieved from: <https://arxiv.org/abs/1406.2661>.
- xxxviii - Nguyen, T., Yamagishi, J., & Echizen, I. (2019). Capsule-forensics: Using capsule networks to detect forged images and videos. arXiv preprint arXiv:1905.06076. Retrieved from <https://arxiv.org/abs/1810.11215>.
- xxxix - للمزيد أنظر:
- Hernandez-Ortega, J., Tolosana, R., Fierrez, J., & Morales, A. (2020). DeepFakesON-Phys: DeepFakes detection based on heart rate estimation. *arXiv:2010.00400 [cs.CV]*. Retrieved from <https://doi.org/10.48550/arXiv.2010.00400>

- Agarwal, S., Farid, H., Fried, O., & Agrawala, M. (2020). Detecting deepfake videos from phoneme-viseme mismatches. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 1–10.
- Mercan, S., Cebe, M., Aygun, R. S., Akkaya, K., Toussaint, E., & Danko, D. (2021). Blockchain-based video forensics and integrity verification framework for wireless Internet-of-Things devices. *Computer Science Faculty Research and Publications*. Retrieved from https://epublications.marquette.edu/comp_fac/53