



كلية الآداب



جامعة بنها

مجلة كلية الآداب

مجلة دورية علمية محكمة

**فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء
الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة
المحتوى الصحفي**

إعداد/

مي أحمد أبو صالحة

مدرس بقسم الاعلام كلية الآداب جامعة دمياط

أكتوبر 2025

المجلد ٦٥

[/https://jfab.journals.ekb.eg](https://jfab.journals.ekb.eg)

المخلص:

هدف البحث إلى قياس مدى فاعلية تطبيق أنظمة الذكاء الاصطناعي في حماية البيانات الصحفية والمحتوى الإعلامي من الاختراق أو التشويه، لا سيما في ظل تصاعد وتيرة الأخبار المضللة والهجمات الإلكترونية الموجهة، كما سعى البحث إلى الكشف عن العلاقة بين أمن المنصات الإعلامية وجودة المحتوى الصحفي، من خلال توضيح كيف يمكن لتحسين بيئة الأمان أن تسهم في تقديم صحافة موثوقة ومستقرة، بالإضافة إلى استكشاف العلاقة بين مستوى الوعي المهني لدى الصحفيين بتقنيات الذكاء الاصطناعي والأمن السيبراني وجودة المنتج الصحفي، وتكونت عينة البحث من عدد (١١٦) مفردة من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية رقمية تعتمد على تقنيات حديثة ومديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في تلك المؤسسات الإعلامية وكذلك مجموعة من رؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية ومجموعة من السادة أعضاء هيئة التدريس تخصص الإعلام بكليات الآداب بالجامعات المصرية، واستخدم البحث المنهج الوصفي التحليلي في ضوء النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات، واستعان البحث بأداة الاستبيان (من تصميم الباحثة) لجمع المعلومات اللازمة من عينة البحث، وأظهر البحث مجموعة من النتائج من أهمها وجود علاقات إيجابية دالة إحصائياً بين استخدام تقنيات الذكاء الاصطناعي وتعزيز الأمن السيبراني للمنصات الإعلامية الرقمية، حيث تسهم هذه التقنيات في الكشف المبكر عن الهجمات السيبرانية وحماية المحتوى الإعلامي من الاختراق والتلاعب، كما أوضحت النتائج أن مستوى الأمان السيبراني يؤثر في جودة المحتوى الصحفي بصوره ايجابية، إلى جانب الكشف عن أن الوعي المهني لدى الصحفيين بالتقنيات الحديثة له دور إيجابي في رفع كفاءة الأداء الإعلامي وتعزيز موثوقية المنتج الصحفي.

الكلمات المفتاحية: الأمن السيبراني، الذكاء الاصطناعي، المنصات الإعلامية الرقمية.

مقدمة:

تُعد المنصات الإعلامية الرقمية في العصر الحديث أحد أهم الركائز الأساسية لنقل المعلومات وتشكيل الرأي العام، حيث أصبحت مصدرًا رئيسيًا للأخبار والمحتوى الصحفي لملايين المستخدمين حول العالم. ومع هذا الانتشار الواسع، تزايدت التهديدات الأمنية التي تستهدف هذه المنصات، مما أدى إلى تعريض سلامة البيانات وسرية المعلومات للخطر، وبالتالي التأثير على مصداقية وجودة المنتج الصحفي. لذلك فقد ظهرت أدوات الحماية الرقمية وتقنيات الذكاء الاصطناعي كحلول مبتكرة لتعزيز الأمن السيبراني، مما دفع الباحثين والمتخصصين إلى دراسة فاعلية هذه الأدوات وتقنياتها في مواجهة التحديات الأمنية المتزايدة.

وتشير دراسة كلا (Smith & Johnson (2020) إلى أن التهديدات السيبرانية، مثل الاختراقات الإلكترونية وهجمات الحرمان من الخدمة (DDoS) وتزييف المحتوى، أصبحت أكثر تعقيدًا وتطورًا، مما يتطلب تبني استراتيجيات أمنية متقدمة، وعليه فإن أدوات الذكاء الاصطناعي تعتبر أحد أهم التقنيات التي تسهم في تحسين الكشف عن التهديدات والاستجابة لها بشكل فوري، حيث يمكن أن تحليل كميات هائلة من البيانات وتحديد الأنماط غير الطبيعية التي تشير إلى هجمات محتملة (Al-Masri et al., 2021). بالإضافة إلى أنه يمكن استخدام تقنيات التعلم الآلي (Machine Learning) لتعزيز قدرات أنظمة الحماية الرقمية، مما يسهم في تقليل الفجوات الأمنية وزيادة فعالية الإجراءات الوقائية.

من ناحية أخرى، تؤثر هذه التطورات الأمنية بشكل مباشر على جودة المنتج الصحفي، حيث يوفر الأمن السيبراني بيئة آمنة لإنشاء المحتوى ونشره دون تدخلات خارجية أو تزييف.

وتشير دراسة كلا من (Brown & Davis 2019)، إلى أن المنصات الإعلامية التي تعتمد على أدوات حماية متقدمة تتمتع بمصداقية أعلى وتفاعل أكبر من الجمهور، مما يعزز ثقة المستخدمين في المحتوى المقدم. كما أن استخدام تقنيات الذكاء الاصطناعي في تحليل المحتوى الصحفي يسهم في تحسين جودته من خلال الكشف عن المعلومات المضللة أو المزيفة، مما يعزز الشفافية والموثوقية.

وتُعد تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي من الأدوات الحديثة والفعالة في حماية المنصات الإعلامية الرقمية من التهديدات السيبرانية المتزايدة، لا سيما في ظل الاعتماد المتنامي على البيئة الرقمية في العمل الإعلامي. حيث تسهم هذه التقنيات في الكشف المبكر عن محاولات الاختراق، والتعرف على الأنماط غير المألوفة في البيانات، مما يُعزز من قدرة المؤسسات الإعلامية على التصدي للهجمات قبل وقوعها. كما أن توظيف الذكاء الاصطناعي في عمليات التحليل والتصنيف والمراقبة يسهم في تحسين جودة المحتوى الصحفي من خلال توفير بيئة آمنة تُمكن الصحفيين من أداء مهامهم بحرية ومصداقية، دون التأثير بالمخاطر التقنية أو المعلوماتية. وتشير الأدبيات إلى أن دمج أدوات الذكاء الاصطناعي ضمن استراتيجيات الأمن السيبراني يُعد خطوة جوهرية نحو تحقيق صحافة رقمية موثوقة ومستقرة (الخولي، ٢٠٢٣).

المدخل النظري:

في ظل التطور التكنولوجي السريع وزيادة الاعتماد على المنصات الإعلامية الرقمية، أصبحت قضايا الأمن السيبراني تحظى بأهمية كبيرة، في حين تُعد المنصات الإعلامية الرقمية هدفاً رئيسياً للهجمات السيبرانية بسبب طبيعتها المفتوحة وقدرتها على الوصول إلى جمهور واسع. لذلك تُستخدم أدوات الأمن الرقمي وتقنيات الذكاء الاصطناعي لتعزيز الأمن السيبراني، مما ينعكس إيجاباً على جودة المنتج الصحفي. (Palla & Kostarella, 2025).

ونظرا لذلك أصبحت المنصات الإعلامية الرقمية عرضة لتهديدات سيبرانية متزايدة، ما يؤثر سلباً على مصداقية المحتوى الصحفي وجودته. وقد أظهرت الدراسات أن الهجمات السيبرانية أصبحت أكثر تعقيداً وذكاءً، مستهدفة المؤسسات الإعلامية من أجل نشر معلومات زائفة أو التأثير على الرأي العام.(Sonni et al., 2024)

-أهمية الأمن السيبراني في الإعلام الرقمي.

تتعرض المؤسسات الإعلامية لهجمات إلكترونية تهدف إلى تعطيل سير العمل أو التلاعب بالمحتوى الصحفي، ما يتسبب بخسائر مادية ومعنوية. وقد أثبتت التجارب أن ضعف البنية الأمنية الرقمية يؤدي إلى نقويض ثقة الجمهور بالمؤسسة الإعلامية، ويُسهل التلاعب بالمعلومات.(Okdem & Okdem, 2024)

وتشير دراسة (Wilson et al., 2023) إلى أن الهجمات على خوادم وسائل الإعلام زادت بنسبة ١٤٠% بين عامي ٢٠٢٠ و٢٠٢٣، مع تركيز خاص على هجمات حجب الخدمة.

كما اوصت دراسة (Anderson & Lee, 2022) بتبني استراتيجيات متعددة الطبقات تشمل جدران الحماية الذكية وأنظمة مراقبة الشبكات في الوقت الفعلي للمواقع الاعلامية

وتتطلب حماية المحتوى الإعلامي الرقمي نهجاً شاملاً يجمع بين التقنيات الأمنية والتنوعية البشرية. وأوضحت دراسة (Davis & Brown, 2021) أن ٦٠% من الثغرات الأمنية في المؤسسات الإعلامية نتجت عن أخطاء بشرية أو إهمال الموظفين ونظرا لذلك فانه تبرز أهمية برامج التدريب المستمر وتبني ثقافة أمنية داخل المؤسسات الإعلامية، إلى جانب استخدام أدوات التحقق من الهوية والوصول المتطورة (Thompson et al., 2022).

كما ان المنصات الإعلامية تواجه تحديات أمنية متطورة مع ظهور تقنيات التزييف العميق (Deepfake) والهجمات المستندة إلى الذكاء الاصطناعي، حيث تشير التقارير إلى أن ٤٢% من المؤسسات الإعلامية الكبرى واجهت محاولات تزييف محتوى باستخدام هذه التقنيات خلال عام ٢٠٢٣ وحده (Miller & Zhang, 2023). وتستلزم هذه التهديدات تطوير أنظمة كشف متقدمة تعتمد على التعلم العميق وتحليل السياق، مع ضرورة تعاون القطاع الإعلامي مع خبراء الأمن السيبراني لمواكبة هذه التحديات (European Cybersecurity Journal, 2023).

كما أصبحت التشريعات والسياسات الأمنية عاملاً محورياً في حماية القطاع الإعلامي، حيث أظهرت دراسة (Global Media Policy Review, 2024) أن الدول التي طبقت معايير أمنية إلزامية للمنصات الإعلامية شهدت انخفاضاً بنسبة ٥٨% في الهجمات الإلكترونية المستهدفة، كما أوصت دراسة (Harris et al., 2023) بتبني إطار عمل متكامل يشمل متطلبات التشفير الإلزامي، واختبارات الاختراق الدورية، وإجراءات الاستجابة للحوادث، مع مراعاة الخصوصية الإقليمية والقانونية.

- دور الذكاء الاصطناعي في تعزيز الأمن السيبراني.

يساهم الذكاء الاصطناعي في الكشف التلقائي عن التهديدات وتحليلها، من خلال أنظمة قادرة على التعلم المستمر والتفاعل مع التهديدات الأمنية في الزمن الحقيقي (Sonni et al., 2024). وتُظهر خوارزميات التعلم الآلي إمكانيات واعدة في التعرف على الهجمات السيبرانية غير التقليدية، مثل هجمات التصيد المتقدمة (phishing) والهجمات القائمة على التزييف العميق (deepfakes).

كما يساهم الذكاء الاصطناعي في عمليات تعزيز الأمن السيبراني من خلال الكشف التلقائي عن التهديدات وتحليلها، باستخدام أنظمة قادرة على التعلم المستمر والتفاعل مع التهديدات في الزمن الحقيقي. إذ تُظهر خوارزميات التعلم الآلي إمكانيات متقدمة

في اكتشاف الهجمات المعقدة مثل التصيد الاحتيالي المتقدم (Phishing) والتزييف العميق. (Sonni et al., 2024)

تستطيع تقنيات الذكاء الاصطناعي من تحليل كميات ضخمة من البيانات بسرعة وفعالية، مما يُساعد في تحديد الأنماط غير الاعتيادية التي قد تشير إلى نشاط إلكتروني خبيث. وهذا يتيح للمؤسسات اتخاذ قرارات أمنية استباقية تقلل من المخاطر السيبرانية المحتملة. (Legit Security, 2023)

وفي مواجهة التزييف العميق، تُستخدم أدوات مدعومة بالذكاء الاصطناعي لتحليل الخصائص الدقيقة للصور والفيديوهات مثل تعابير الوجه والنمط الصوتي للكشف عن المحتوى المزيف، مما يعزز موثوقية المعلومات الرقمية ويقلل من انتشار الأخبار الكاذبة. (Palla & Kostarella, 2025)

كما تلعب تقنيات الذكاء الاصطناعي دورًا بارزًا في أتمتة عمليات الاستجابة للحوادث الأمنية، بما يشمل جمع الأدلة وتحليلها وتنفيذ بروتوكولات الحماية، مما يسرّع من زمن الاستجابة ويقلل من الخسائر الناتجة عن الاختراقات السيبرانية (Legit Security, 2023).

يُستخدم الذكاء الاصطناعي أيضًا في مراقبة الشبكات بشكل دائم للكشف عن السلوكيات الشاذة التي قد تدل على محاولات تسلل، وذلك من خلال أنظمة كشف التسلل القائمة على التعلم الآلي، والتي تُحدث تطورًا كبيرًا مقارنةً بالأنظمة التقليدية (Abdel-Basset et al., 2024).

إضافة إلى دوره في الكشف عن الهجمات، يُستخدم الذكاء الاصطناعي في التنبؤ بالتهديدات المستقبلية عبر تحليل الاتجاهات والسلوكيات السابقة، مما يوفر رؤية استباقية للمخاطر قبل حدوثها. (Buczak & Guven, 2016)

على الرغم من فوائده، فإن تطبيق الذكاء الاصطناعي في الأمن السيبراني يواجه تحديات أبرزها تحيز الخوارزميات، وضرورة توفر بيانات تدريبية عالية الجودة. لذا، يتطلب الاستخدام الفعال له استراتيجيات دقيقة لضمان الموثوقية والحياد (Palla & Kostarella, 2025).

-تأثير الأمن السيبراني المدعوم بالذكاء الاصطناعي على جودة الصحافة.

تحسين الأمن السيبراني لا يقتصر على حماية البنية التقنية فحسب، بل يشمل أيضاً حماية العملية الصحفية برمتها، مثل حماية مصادر المعلومات، والتأكد من سلامة المحتوى، ومنع تسرب البيانات. كما أن الصحفيين الذين يعملون في بيئة رقمية آمنة يكونون أكثر قدرة على أداء مهامهم بحرية وفاعلية، ما يؤدي إلى رفع جودة العمل الصحفي (Palla & Kostarella, 2025).

ودورالأمن السيبراني لا يقتصر على حماية البنية التقنية للمواقع الصحفية فحسب، بل يشمل أيضاً حماية العملية الصحفية برمتها، مثل حماية مصادر المعلومات، والتأكد من سلامة المحتوى، ومنع تسرب البيانات. كما أن الصحفيين الذين يعملون في بيئة رقمية آمنة يكونون أكثر قدرة على أداء مهامهم بحرية وفاعلية، ما يؤدي إلى رفع جودة العمل الصحفي. (Palla & Kostarella, 2025)

كما يساهم الذكاء الاصطناعي في تعزيز الأمن السيبراني من خلال تقنيات التعلم الآلي التي تمكّن من الكشف المبكر عن الهجمات الإلكترونية وتحديد أنماط السلوك المشبوه. وهذا يتيح للمؤسسات الصحفية التصدي للتهديدات السيبرانية قبل وقوعها، مما يضمن استمرارية الإنتاج الإعلامي دون تعطيل. (Marr, 2023)

ومن خلال استخدام الذكاء الاصطناعي، يمكن للأنظمة الصحفية اكتشاف التزييف العميق (Deepfakes) والمعلومات المضللة المنتشرة عبر الإنترنت، وهو ما يعزز من مصداقية المحتوى الصحفي ويحسن من جودة التغطية الإخبارية المقدمة للجمهور (Westerlund, 2021).

لذلك فالاعتماد على الأمن السيبراني المدعوم بالذكاء الاصطناعي يسمح بتوفير بيئة رقمية آمنة تُمكن الصحفيين من التواصل بحرية مع مصادرهم دون الخوف من المراقبة أو التسريب، وهو ما يحفّز الإنتاج الإبداعي ويعزز حرية التعبير داخل غرف الأخبار. (Bradshaw et al., 2020)

والتهديدات السيبرانية لا تقتصر فقط على سرقة البيانات، بل تشمل كذلك محاولات التلاعب بالرأي العام من خلال استهداف المنصات الإخبارية، ما يجعل من الضروري دمج حلول الذكاء الاصطناعي في نظم الحماية السيبرانية لضمان نزاهة وسلامة المعلومات. (Taddeo & Floridi, 2018)

حيث يعمل الذكاء الاصطناعي على تقليص الأخطاء البشرية في إدارة الأمن السيبراني داخل المؤسسات الصحفية، حيث تُستخدم الخوارزميات الذكية لتحليل كميات هائلة من البيانات بهدف اكتشاف الثغرات الأمنية بشكل دقيق وسريع (Sahraoui & Benslimane, 2022).

وتسهم تقنيات الذكاء الاصطناعي في تطوير أنظمة تحقق آلي من صحة المعلومات ومطابقتها للمصادر الرسمية، مما يدعم المصداقية والشفافية في العمل الصحفي ويحد من انتشار الأخبار الزائفة. (Graves, 2018)

بالإضافة إلى جوانب الحماية، يمكن للذكاء الاصطناعي أن يقدم تحليلات استباقية حول التهديدات المحتملة التي قد تؤثر على سلامة الصحفيين العاملين في مناطق النزاع، ما يعزز من تدابير السلامة الميدانية ويقلل من المخاطر المهنية (UNESCO, 2022).

ورغم الفوائد الكبيرة، إلا أن استخدام الذكاء الاصطناعي في الأمن السيبراني للصحافة يثير قضايا تتعلق بالخصوصية والمراقبة، ما يتطلب وضع أطر أخلاقية وتشريعية

واضحة تحكم استخدام هذه التكنولوجيا بما يضمن توازنًا بين الحماية والحرية الصحفية. (Cath, 2018)

-التحديات المرتبطة بتطبيق الذكاء الاصطناعي في أمن الإعلام.

رغم الفوائد، يواجه الذكاء الاصطناعي تحديات مثل نقص الكفاءات البشرية، وارتفاع التكاليف، وتعارض بعض التطبيقات مع أخلاقيات المهنة الصحفية، مثل التحيز الخوارزمي أو انتهاك الخصوصية عند جمع البيانات وتحليلها. (Sonni et al., 2024) ومن أبرز التحديات التقنية المرتبطة بتطبيق الذكاء الاصطناعي في أمن الإعلام محدودية قدرة الخوارزميات الحالية على التعامل مع السياقات المعقدة للتهديدات السيبرانية، ما قد يؤدي إلى أخطاء في التعرف على الهجمات أو فشل في التنبؤ بها. (Hassani et al., 2020).

حيث يشكل التحيز الخوارزمي تحديًا أخلاقيًا ومهنيًا كبيرًا، حيث قد تؤدي الخوارزميات إلى تصفية أو ترجيح نوع معين من المعلومات على حساب أخرى، مما يفقد التغطية الإعلامية توازنها ويؤثر على مصداقية المؤسسات الإعلامية. (Binns, 2018) وفي بعض الأنظمة يُستخدم الذكاء الاصطناعي في تتبع أنشطة الصحفيين ومراقبة مصادرهم تحت ذريعة الأمن السيبراني، مما يهدد حرية الصحافة ويُعرض الصحفيين لخطر الملاحقة أو الرقابة. (UNESCO, 2022)

ويعد ارتفاع تكاليف تطبيق أنظمة الذكاء الاصطناعي عائقًا رئيسيًا، خاصة بالنسبة للمؤسسات الإعلامية الصغيرة أو تلك العاملة في الدول النامية، مما يؤدي إلى فجوة رقمية في مجال الأمن السيبراني الإعلامي. (Marconi, 2020) كما أن غياب الأطر القانونية والتنظيمية الواضحة بشأن استخدام الذكاء الاصطناعي في الإعلام يمثل تحديًا إضافيًا، حيث يُعرض المؤسسات لمخاطر قانونية في حال استخدام البيانات دون إذن أو دون احترام قواعد الخصوصية. (Cath, 2018)

وتعاني العديد من المؤسسات الإعلامية من نقص الكفاءات البشرية المؤهلة لتشغيل وتطوير أنظمة الذكاء الاصطناعي، مما يضعف من فعالية تطبيق هذه التقنيات ويجعلها عرضة للأخطاء التشغيلية. (Schmidt & Wiegand, 2017)

حيث تتطلب أنظمة الذكاء الاصطناعي صيانة وتحديثاً مستمرين، وهو ما يمثل عبئاً إدارياً وتقنياً، خصوصاً في البيئات الإعلامية التي تعاني من ضغط الوقت وسرعة التغيير في الأولويات. (Brundage et al., 2023)

كذلك التحديات الثقافية أيضاً تؤثر في مدى تقبل العاملين في الإعلام لاستخدام الذكاء الاصطناعي، حيث يرى البعض أنه تهديد لوظائفهم أو انتقاص من دورهم المهني، مما يعيق تكامل هذه التقنيات في منظومة العمل الإعلامي (Diakopoulos, 2019).

-تعريفات البحث الإجرائية.

قامت الباحثة بتحديد المصطلحات الأساسية المستخدمة في البحث تعريفاً إجرائياً، وذلك لتوضيح المقصود بها ضمن السياق المحدد للبحث، وتجنباً لأي لبس أو اختلاف في الفهم. وقد استندت في هذه التعاريف إلى ما يخدم أهداف البحث ويسهم في تحقيق وضوح المفاهيم لدى القارئ والمتلقي.

- الأمن السيبراني: (Cybersecurity) تعرفه الباحثة اجرائياً بأنه مجموعة

من العمليات والتقنيات المصممة لحماية الأنظمة الرقمية والمحتوى من التهديدات والهجمات غير المصرح بها.

- الذكاء الاصطناعي: (Artificial Intelligence) تعرفه الباحثة اجرائياً

بأنه أنظمة تكنولوجية قادرة على أداء مهام ذكية، تشمل التعلم، التوقع، واتخاذ القرارات دون تدخل بشري مباشر.

- المنصات الإعلامية الرقمية (Digital Media Platforms) تعرفه

الباحثة اجرائيا بأنه وسائل رقمية تُستخدم لنشر المحتوى الإعلامي عبر الإنترنت، مثل المواقع الإخبارية والتطبيقات والمنصات الاجتماعية.

-النظرية القائم عليها البحث:

اعتمدت الباحثة في إطار هذه البحث على النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات، بوصفها مرجعاً نظرياً لفهم العوامل التي تؤثر في تبني الأفراد للتقنيات الرقمية وتوظيفها في السياقات التطبيقية المختلفة Unified Theory of Acceptance And Use Of Technology (UTAUT) وتستند النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات إلى مجموعة من النماذج النظرية السابقة، ويُعد من أبرزها نموذج تقبل التكنولوجيا (TAM) ، الذي يُنظر إليه باعتباره إطاراً شاملاً لفهم سلوك الأفراد تجاه تبني واستخدام التطبيقات التكنولوجية. ويتميز هذا النموذج بمرونته وعموميته، مما يتيح تطبيقه في مجالات متعددة مثل الإعلام، والصحة، والتعليم، والصناعة وغيرها، كما يوفر آلية منهجية لتحليل العوامل النفسية والسلوكية التي تؤثر على نية الأفراد لاستخدام التكنولوجيا وسلوكهم الفعلي في هذا السياق. (Sovacool & Hess, 2017)

ويُسهّم نموذج تقبل التكنولوجيا (TAM) في توضيح وفهم العوامل المؤثرة في تبني تقنية المعلومات من قبل الأفراد، حيث يركز على نوعين رئيسيين من العوامل التي تُشكّل الأساس في تفسير سلوك استخدام التكنولوجيا، وهما: نية الاستخدام والسلوك الفعلي. ويقوم هذان المتغيران على عاملين محوريين هما: المنفعة المدركة، التي تشير إلى إدراك المستخدم لفائدة استخدام التكنولوجيا في تحسين الأداء، والسهولة المتوقعة أو الجهد المتوقع، والذي يعكس مدى بساطة استخدام التقنية من وجهة نظر المستخدم. كما يدمج النموذج لاحقاً متغيرات إضافية مثل طوعية الاستخدام، أي مدى حرية الفرد في اختيار استخدام التقنية، والتسهيلات المتاحة، التي تشير إلى مدى توفر الدعم التقني والبنية التحتية الملائمة، كذلك يأخذ النموذج بعين الاعتبار العوامل الديموغرافية والاجتماعية التي تؤثر

على قرارات الأفراد بشأن قبول أو رفض التكنولوجيا. وعند تطبيق هذا النموذج على عينة الدراسة، يمكن تحليل وفهم مدى تقبلهم لتقنيات الذكاء الاصطناعي، وتأثيرات هذه التقنيات على ممارساتهم الإعلامية. كما يساعد النموذج في الكشف عن أبرز العوامل المؤثرة على الاستخدام الفعلي، والتنبؤ بالسيناريوهات المستقبلية المتوقعة بشأن توظيف هذه التكنولوجيا واستثمار إمكاناتها في تطوير الأداء الإعلامي.

وتتضمن النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات (UTAUT) أربعة أبعاد رئيسية تُستخدم لتفسير سلوك الأفراد تجاه تبني التكنولوجيا. أولاً: توقع الأداء. ويشير إلى مدى اعتقاد المستخدم بأن استخدام التكنولوجيا سيسهم في تحسين أدائه في العمل أو في حياته اليومية.

ثانياً: توقع الجهد. والذي يقيس مدى سهولة استخدام النظام من وجهة نظر المستخدم. ثالثاً: النفوذ الاجتماعي. ويتعلق بتأثير آراء الآخرين (كالزملاء أو المسؤولين) في قرار الفرد باستخدام التكنولوجيا.

رابعاً: الظروف التيسيرية. التي تعكس مدى توافر البنية التحتية والدعم الفني الذي يمكن الفرد من استخدام التكنولوجيا بسهولة. وتُعد هذه الأبعاد أدوات تحليلية مهمة في البحوث التي تهدف إلى فهم العوامل المؤثرة في قبول وتبني الأفراد للتقنيات الحديثة. Ali, et. all, (2024) وقد قامت الباحثة بربط هذه الأبعاد بنتائج الدراسة في جزء تفسير النتائج.

-دراسات سابقة-

على الرغم من حداثة موضوع فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المنتج الصحفي، إلا إن هناك عددًا من الدراسات التي تمكنت الباحثة من الرجوع إليها في هذا المجال، إضافة إلى الموضوعات المرتبطة به والتي أسهمت في تكوين نموذج الدراسة وبناء فروضه، وبصفة خاصة الدراسات التي تناولت استخدام فاعلية استخدام تقنيات الأمن

السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المنتج الصحفي ، هذه الدراسات تقدم رؤى متعددة حول تأثير تقنيات الذكاء الاصطناعي على الأمن السيبراني وجودة الصحافة في المنصات الإعلامية الرقمية.

ففي ضوء التقدم التكنولوجي المتسارع، بات من الضروري دراسة العلاقة بين تقنيات الذكاء الاصطناعي والأمن السيبراني وجودة المحتوى الإعلامي، خاصة في ظل ما تشهده المنصات الرقمية من تحديات أمنية ومخاطر معلوماتية. وقد تناولت عدد من الدراسات الحديثة هذا الموضوع من زوايا متعددة، دعمت التوجه نحو توظيف أدوات الذكاء الاصطناعي في تعزيز الأمن الرقمي داخل المنظومات الإعلامية.

كما استعرضت دراسة (Ademola & Somorin, 2024) العلاقة بين الذكاء الاصطناعي والصحافة الاستقصائية، مبينة كيف يمكن لتقنيات التحقق الآلي واكتشاف التزييف العميق أن تعزز من جودة العمل الصحفي وتحافظ على مصداقيته، خاصة في ظل تهديدات الأمن السيبراني. وأوصت الدراسة بأهمية إدماج تلك التقنيات ضمن الممارسات المهنية للصحفيين، وتوفير تدريبات متخصصة في هذا المجال.

وتوصلت دراسة (Abdalwahid et al., 2024)، إلى أن تطبيقات الذكاء الاصطناعي في نظم العلامات المائية والإخفاء الرقمي حققت كفاءة عالية في حماية المحتوى الإعلامي الرقمي، حيث سجلت مقاومة للتعديل بنسبة ٩٤% ودقة استخلاص للبيانات المخفية بلغت ٩٧%، مع تحسين أداء المعالجة بنسبة ٤٠% مقارنة بالأساليب التقليدية. وأكدت الدراسة على إمكانية توظيف هذه التقنيات المتقدمة لحماية الملكية الفكرية للمحتوى الصحفي والحد من القرصنة الإلكترونية، رغم التحديات التقنية المرتبطة بموازنة متطلبات الأمن وجودة المحتوى. هذه النتائج تقدم إضافة نوعية لأبحاث الأمن السيبراني في مجال الإعلام الرقمي، وتفتح آفاقاً جديدة لتعزيز حماية المنصات الإعلامية باستخدام الحلول الذكية.

وتوضح دراسة Babatope & Anil (2024) دور تقنيات التعلم العميق في كشف السلوكيات غير الطبيعية على منصات التواصل الاجتماعي، مثل اختراق الحسابات أو نشر المعلومات الكاذبة. وقد طورت الدراسة نموذجًا يستخدم الشبكات العصبية لاكتشاف التهديدات الأمنية في الوقت الفعلي، مما يُعد إطارًا تقنيًا مهمًا يمكن تكييفه لحماية المحتوى الإعلامي الرقمي.

أما دراسة Dhingra (2024)، فقد ناقشت أثر تقنيات الذكاء الاصطناعي في رصد ومنع الأنشطة الخبيثة داخل المنصات الإعلامية، مثل التصيد الإلكتروني وإنشاء الحسابات الوهمية ونشر المحتوى الضار. وبيّنت الدراسة أن الأدوات الذكية مثل معالجة اللغة الطبيعية وتعلم الآلة تساهم في تحسين قدرات المنصات الإعلامية على التصدي لهذه التهديدات، مما ينعكس إيجابًا على جودة المحتوى وسلامته.

وناقشت دراسة الخالدي (٢٠٢٣)، التحديات الأخلاقية والقانونية المرتبطة باستخدام الذكاء الاصطناعي في الأمن السيبراني للإعلام، مشيرة إلى ضرورة وضع أطر تنظيمية تضمن توازنًا بين الأمن الرقمي وحرية التعبير. وخلصت إلى أن التكامل بين التقنيات الذكية والسياسات الإدارية الفعالة يُعزز حماية المنصات دون المساس بالجوانب المهنية للصحافة.

كما كشفت دراسة Nguyen et al. (2023) عن فاعلية أنظمة كشف التسلل (IDS) المدعومة بالتعلم العميق في حماية البنية التحتية للمواقع الإخبارية، حيث أظهرت النتائج أن هذه الأنظمة تتمتع بدقة تصل إلى ٩٥% في التعرف على الهجمات السيبرانية المعقدة، مما يقلل من فترات التوقف عن العمل ويحافظ على استمرارية البث الإعلامي.

وأوضحت دراسة Waqas et al., (2022) التي تناولت تطبيقات الذكاء الاصطناعي في أمن الشبكات، أثبتت فعالية تقنيات تعلم الآلة في تعزيز الحماية السيبرانية حيث حققت أنظمة كشف التسلل دقةً بنسبة ٩٨.٢% وتنبأت بالهجمات بنجاح ٩٥.٧%،

بينما خفضت الإنذارات الكاذبة بنسبة ٦٢%. ورغم التحديات التقنية مثل استهلاك الموارد وصعوبة تفسير القرارات الآلية، تبرز هذه النتائج إمكانية توظيف هذه الحلول الذكية لحماية البنية التحتية للمنصات الإعلامية الرقمية، لا سيما في تأمين قنوات الاتصال اللاسلكي ونقل المحتوى الصحفي، مع الحاجة لمعالجة تحديات الأداء والشفافية لضمان موثوقية هذه الأنظمة في البيئات الإعلامية الحساسة.

كما أشارت دراسة أجراها السميري والفقير (٢٠٢٢) إلى أن استخدام أنظمة التعلم الآلي في الكشف عن الهجمات الإلكترونية يقلل من نسبة الاختراقات بنسبة تصل إلى ٤٠%، كما يساهم في تحسين استجابة المنصات الإعلامية للتهديدات الأمنية في الوقت الفعلي.

كما أوضحت دراسة (Nsude, I. (2022) تأثير تقنيات الذكاء الاصطناعي على القطاع الإعلامي النيجيري في مواجهة التحديات الأمنية. وأظهرت الدراسة أن تطبيقات الذكاء الاصطناعي مثل تحليل البيانات الضخمة والكشف عن المحتوى المضلل ساهمت في تحسين دقة التغطية الإعلامية للأزمات الأمنية بنسبة ٣٥%. كما كشفت النتائج عن تحديات رئيسية تشمل نقص البنية التحتية التقنية بنسبة انتشار ٦٠% فقط في وسائل الإعلام النيجيرية وصعوبة التوفيق بين المراقبة الأمنية وحقوق الخصوصية، وهذه الدراسة تعزز أهمية البحث الحالي حول توظيف الذكاء الاصطناعي في المنصات الإعلامية، مع إبراز الحاجة إلى تكيف الحلول التقنية مع السياقات المحلية.

كذلك استعرضت دراسة (Ahmad et al. (2021) مختلف تقنيات الذكاء الاصطناعي المستخدمة في حماية الأنظمة الرقمية من الهجمات السيبرانية، وخلصت إلى أن الجمع بين هذه التقنيات يحقق نتائج فعالة في التصدي للثغرات الأمنية، لا سيما في المؤسسات التي تعتمد على النشر الإلكتروني المستمر مثل المنصات الإعلامية.

كما ناقشت دراسة (Rainer Mühlhoff (2020 مفهوم الذكاء الاصطناعي السيبراني بوصفه أنظمة هجينة تقوم على تداخل معقد بين الإنسان والآلة، حيث أشارت

إلى إمكانية أن يصبح الإنسان عنصراً وظيفياً داخل هذه الأنظمة، مما يغيّر من طبيعة أدواره في مختلف المجالات، ومنها المجال الإعلامي. وقد رأت الدراسة أن الدور الإعلامي للبشر سيقصر مستقبلاً على المهام الإبداعية التي تتطلب مهارات عاطفية واجتماعية لا يمكن للآلة محاكاتها بسهولة. وفي هذا السياق، حذرت الدراسة من صعود هيمنة أنظمة الذكاء الاصطناعي عبر ما يُعرف بـ"التمييز الخوارزمي" أو "اللامساواة الآلية"، وهي آليات تستخدم الذكاء الاصطناعي كأداة لإقصاء العمل البشري من خلال فرض منطق القوة التقنية. كما تناولت الدراسة مفهوم "التغذية الراجعة" في سياق تطبيقات الذكاء الاصطناعي المستخدمة في منصات مثل فيسبوك، حيث تستكشف هذه الأنظمة تفضيلات المستخدمين وسلوكياتهم لتوجيه المحتوى بشكل يتحكم في قراراتهم وميولهم دون وعي مباشر منهم. وقد خلصت الدراسة إلى أن مدى انتشار وتأثير هذه الأنظمة يتوقف على مجموعة من العوامل، أبرزها: الظروف الاجتماعية والاقتصادية للمجتمعات، أنماط السلوك الرقمي السائدة، والموقف السياسي للدولة من تقنيات الذكاء الاصطناعي.

وقدّمت دراسة Sarker et al. (2020) تحليلاً شاملاً لاستخدام تقنيات الذكاء الاصطناعي وتعلم الآلة في مجال الأمن السيبراني، مشيرةً إلى أن تحليل البيانات الضخمة يسهم في تعزيز قدرة الأنظمة الرقمية على التنبؤ بالتهديدات قبل وقوعها، وهو ما يمكن أن ينعكس بفعالية في بيئة المنصات الإعلامية الرقمية.

كما اوضحت دراسة غارسيا ولي (2020) مدى تحسن جودة الصحافة الرقمية من خلال أدوات الأمن السيبراني، مؤكدة أن حماية مصادر البيانات الصحفية وضمان سرية المعلومات يعدان عنصرين أساسيين لتعزيز مصداقية الإعلام. وأظهرت النتائج أن دمج تقنيات مثل التشفير المتقدم والتحقق من الهوية باستخدام الذكاء الاصطناعي يحد من تسريب المعلومات الحساسة بنسبة ٦٠٪.

وركزت دراسة (2020) Cristian Vaccari و Andrew Chadwick على ظاهرة التزييف العميق، والتي تُستخدم فيها تقنيات الذكاء الاصطناعي لإنتاج مقاطع فيديو مزيفة تحاكي الواقع بدرجة عالية من الدقة، مما يصعب التمييز بينها وبين المقاطع الحقيقية. وقد توصلت الدراسة إلى أن بعض الأفراد يمكنهم بالفعل اكتشاف هذا النوع من التزييف، إلا أن ذلك لا يمنع من تولّد حالة من عدم اليقين لديهم، تؤدي بدورها إلى تآكل الثقة في الأخبار على المدى البعيد، لا سيما تلك التي تُنشر عبر وسائل التواصل الاجتماعي.

كما أظهرت الدراسة أن للتزييف العميق آثارًا سلبية متعددة، من بينها: تراجع إحساس الأفراد بالمسؤولية تجاه المعلومات التي يشاركونها، مما يسهّل تداول الأخبار دون التحقق من صحتها، وصعوبة الوصول إلى نقاش عام بناءً في ظل انتشار محتوى مشكوك فيه. كما قد يستغل بعض السياسيين هذه التقنية لنفي الاتهامات الموجهة إليهم، عبر الادعاء بأن مقاطع الفيديو التي تدينهم مزيفة. وأكدت الدراسة على خطورة ما أسمته بـ "التزييف الجزئي"، وهو النوع الذي يصعب معه القطع بصحة المحتوى أو زيفه، مما يعقّد من عملية التحقق الإعلامي ويزيد من غموض الواقع المعلوماتي.

كما أشارت دراسة (2020) Thuraisingham, B. إلى أن تقنيات الذكاء الاصطناعي مثل التعلم الآلي ومعالجة اللغة الطبيعية ترفع كفاءة حماية المنصات الرقمية عبر الكشف عن ٩٢% من الهجمات الإلكترونية والمحتوى الضار، كما تسرع زمن الاستجابة للتهديدات بنسبة ٦٠%. وأبرزت الدراسة أهمية هذه التقنيات في مواجهة التحديات الأمنية للمنصات الإعلامية، رغم تحذيرها من تحيز الخوارزميات والتأثير على الخصوصية، وهذه النتائج تدعم فرضية البحث الحالي حول دور الذكاء الاصطناعي في تعزيز أمن المنصات الإعلامية.

وتوصلت دراسة عيسى عبد الباقي وأحمد عادل (٢٠٢٠) إلى أن ما نسبته ٨٨% من عينة الدراسة التي شملت صحفيين وقيادات إعلامية، أكدوا على الأهمية البالغة

لتوظيف تقنيات الذكاء الاصطناعي في غرف الأخبار الخاصة بهم. ومع ذلك، أشار المشاركون إلى أن نسبة كبيرة من هذه الغرف غير مهياً فعلياً للاستفادة من هذه التقنيات، ويعود ذلك إلى عدة معوقات، من أبرزها: عدم تحديث الهياكل التنظيمية، وغياب تبني أنظمة الجودة، إلى جانب نقص الخوارزميات المتخصصة في تحرير النصوص باللغة العربية، وكذلك ضعف الاستثمار والتمويل المخصص لتطوير هذه التكنولوجيا في المؤسسات الإعلامية.

وأجرت دراسة كلا من (2019) Marko Milosavljević و Igor Vobič استكشافاً لاستخدامات التقنيات الخوارزمية داخل غرف الأخبار في المؤسسات الإعلامية بكل من إنجلترا وألمانيا والولايات المتحدة الأمريكية. وقد خلّصت الدراسة إلى التأكيد على استمرار هيمنة العنصر البشري على التكنولوجيا، موضحة أن أتمتة العمل الصحفي لم تهدف إلى إحلال التكنولوجيا محل الصحفيين، بل إلى تمكينهم وتحسينهم من المهام الروتينية، مستندة إلى ما يمتلكونه من مهارات وخبرات، في إطار ما يُعرف بمفهوم "هيمنة الوكالة البشرية على التكنولوجيا". كما أشارت الدراسة إلى أن الإحلال الكامل للعمال البشرية يُعد أمراً غير واقعي، وأن أقصى ما يمكن تحقيقه هو أتمتة بعض المهام الجزئية. وقد أُطلق على هذه الظاهرة مصطلح "الخوارزميات السامية"، في إشارة إلى دور التكنولوجيا كمكمل للعمل البشري لا كبديل عنه. كما أظهرت الدراسة صعوبة إقامة ترابط مباشر بين التقنيات الخوارزمية من جهة، والقيم المهنية الصحفية أو الأهداف الإدارية للمؤسسة الإعلامية من جهة أخرى.

وتناولت دراسة شيهان الورقلي (٢٠١٩) موضوع الصحافة الخوارزمية، باعتبارها إحدى التحولات الحتمية التي يشهدها مجال الإعلام في ظل التقدم التكنولوجي المتسارع. وأكدت الدراسة أن الاعتماد على هذه التقنيات أصبح توجهاً متنامياً في المؤسسات الإعلامية، ومن المتوقع أن يتعزز هذا الاعتماد بشكل رسمي وواسع النطاق في المستقبل.

القريب، سواء في القنوات العالمية أو العربية. وسلطت الدراسة الضوء على تجربة وكالة الأنباء الصينية "شينخوا"، التي تُعد من أبرز النماذج الرائدة في هذا المجال، حيث تعاونت مع شركة "سوغو" لإنتاج روبوتات إعلامية تمتاز بقدرتها على محاكاة المذيع البشري، من خلال توظيف دلالات إيحائية وضمنية تعزز التفاعل مع الجمهور. وقد بيّنت الدراسة أن هذه الروبوتات تسهم في تحسين كفاءة الأداء الإعلامي من حيث السرعة والدقة، إلا أنها في الوقت ذاته تطرح تحديات مهنية محتملة، من بينها التهديد لمكانة الإعلامي البشري، لا سيما في مجالي التقديم التلفزيوني والعمل الصحفي.

وفي دراسة كلا من (Zhang & Gupta, 2018)، تم تحليل التحديات الأمنية الرئيسية للمواقع الإعلامية مثل انتهاكات الخصوصية والتي كانت بنسبة (٦٨%)، الهجمات السيبرانية بنسبة (٥٥%)، والمحتوى الضار بنسبة (٧٣%)، مع تقديم حلول متكاملة تشمل أنظمة المصادقة الذكية وخوارزميات التعلم العميق لكشف المحتوى المغشوش ونظم تقييم السمعة الرقمية. وأظهرت النتائج أن الجمع بين التقنيات الأمنية وحوكمة البيانات يعزز ثقة المستخدمين بنسبة ٤٠%. هذه الرؤية الشاملة تقدم نموذجاً عملياً يمكن تطبيقه على المنصات الإعلامية الرقمية لتعزيز الأمان والموثوقية بشكل فعال، مع الحفاظ على توازن دقيق بين الحماية وحرية التعبير.

التعليق على الدراسات السابقة ووجه الاستفادة منها:

تتفق هذه الدراسات على أن دمج تقنيات الأمن السيبراني مع الذكاء الاصطناعي يُحدث تحولاً جذرياً في حماية المنصات الإعلامية، كما يسهم في رفع جودة المحتوى الصحفي من خلال تقليل المخاطر الأمنية وضمان دقة المعلومات. وتشكل معاً قاعدة معرفية غنية تدعم البحث الحالي، وتؤكد أهمية موضوعه، حيث تسلط الضوء على فعالية استخدام تقنيات الذكاء الاصطناعي في تحسين أمن المنصات الإعلامية وجودة المحتوى الصحفي، وتبرر الحاجة إلى دراسة تطبيقية في هذا المجال الحيوي.

كما تعكس الدراسات المستعرضة صورة دقيقة ومتكاملة لتأثيرات الذكاء الاصطناعي على القطاع الإعلامي، سواء من ناحية تعزيز الأمن السيبراني أو تطوير الأداء المهني داخل غرف الأخبار. ويمكن تصنيف هذه الدراسات ضمن ثلاثة محاور رئيسية:

أولاً: الذكاء الاصطناعي وتعزيز الأمن السيبراني في الإعلام.

حيث أظهرت دراسات مثل (Nguyen et al. (2023 ، Waqas et al. (2022، Thuraisingham (2020، و Sarker et al. (2020 تطوراً ملحوظاً في استخدام تقنيات الذكاء الاصطناعي للكشف عن الهجمات الإلكترونية، وحماية المنصات الإعلامية من الاختراق، والتعامل مع المحتوى الضار. فقد حققت أنظمة الذكاء الاصطناعي دقة تصل إلى ٩٨% في كشف التهديدات، ونجحت في تقليل الإنذارات الكاذبة، وهو ما يشير إلى قدرة هذه التقنيات على رفع كفاءة نظم الحماية الرقمية، لا سيما في ظل تزايد التهديدات السيبرانية التي تستهدف وسائل الإعلام. لكن رغم هذه النجاحات، أشارت بعض الدراسات إلى تحديات مرافقة، مثل تحيز الخوارزميات، وصعوبة تفسير قرارات الذكاء الاصطناعي، واستهلاك الموارد، ما يتطلب تطوير أدوات أكثر شفافية وفعالية للتوافق مع القيم الإعلامية والبيئات الحساسة.

ثانياً: الذكاء الاصطناعي في دعم الممارسة الصحفية.

ركزت دراسات مثل (Ademola & Somorin (2024 ، عيسى عبد الباقي وأحمد عادل (2020 ، (Milosavljević & Vobič (2019، والورقلي (٢٠١٩ على دور الذكاء الاصطناعي في تحسين جودة الأداء الصحفي، من خلال أتمتة المهام الروتينية، ودعم التحقق من الأخبار، واكتشاف التزييف العميق. وقد أجمعت على أن

هذه التقنيات لا تهدف إلى إلغاء دور الصحفي، بل إلى تمكينه وتحريره من المهام المتكررة، ما يعزز من تركيزه على الجوانب الإبداعية والتحليلية. غير أن هذه الدراسات نبهت أيضاً إلى نقص البنية التحتية، وغياب السياسات التنظيمية، وضعف الاستثمار في المجال، لا سيما في السياقات العربية، وهو ما يُعد تحدياً رئيسياً أمام التحول الرقمي الفعال داخل المؤسسات الإعلامية.

ثالثاً: أخلاقيات ومخاطر الذكاء الاصطناعي في الإعلام.

تميزت دراسة (2020) Rainer Mühlhoff والخالدي (٢٠٢٣) بتناولها البُعد الفلسفي والأخلاقي لتقنيات الذكاء الاصطناعي، حيث حذرتا من "التمييز الخوارزمي" ومن التأثير الخفي على قرارات وسلوكيات الأفراد، عبر التحكم بالمحتوى المعروض. وقد دعتا إلى ضرورة تنظيم العلاقة بين الإنسان والتقنية، بما يضمن التوازن بين حرية التعبير ومتطلبات الأمن الرقمي.

كما أشارت دراسات مثل (Vaccari & Chadwick 2020) إلى التهديدات المعلوماتية الجديدة، مثل التزييف العميق، وتأثيرها على ثقة الجمهور بالمحتوى الإعلامي، الأمر الذي يستدعي تطوير أدوات تحقق دقيقة وفعالة، إلى جانب توعية الجمهور بأساليب التلاعب الرقمي.

تُظهر هذه الدراسات مجتمعةً أن الذكاء الاصطناعي يمثل فرصة استراتيجية وتحدياً مزدوجاً في آنٍ واحد أمام القطاع الإعلامي. فمن جهة، يتيح أدوات قوية لحماية المحتوى، وتسريع الأداء، ومكافحة التضليل. ومن جهة أخرى، يفرض تحديات أخلاقية وقانونية ومهنية تتطلب استجابات تنظيمية وتربوية متكاملة.

وأفادت الدراسات السابقة المتعلقة بالأمن السيبراني، والذكاء الاصطناعي، والمنصات الإعلامية الرقمية في تعزيز الإطار النظري والمنهجي للبحث الحالي، التي تسعى إلى تقويم فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي. وقد تمثلت أبرز أوجه

الاستفادة في بلورة مشكلة البحث وتحديد أهدافه، بناء الإطار النظري وتحديد المفاهيم، اختيار المنهجية وأدوات الدراسة، تحليل النتائج والمقارنة، توجيه التوصيات.

-مشكلة البحث.

تم صياغة المشكلة من خلال تحديد العلاقة بين متغيرات البحث والإشكالية بينهما، فمع تصاعد الهجمات السيبرانية التي تستهدف المنصات الإعلامية الرقمية، تبرز الحاجة إلى اعتماد استراتيجيات أمنية متطورة قادرة على حماية البيانات والمحتوى الصحفي، وذلك دون التأثير على جودة الأداء المهني، وعلى الرغم من تطور تقنيات الذكاء الاصطناعي في مجال الأمن السيبراني، إلا أن مدى فعاليتها الحقيقية في حماية المنصات الإعلامية وتحسين جودة العمل الصحفي لا تزال غير واضحة بشكل كافٍ، خصوصاً في ظل تنوع البيئات الرقمية واختلاف التهديدات، ومن هنا تتبع مشكلة هذه الدراسة في محاولة فهم العلاقة بين توظيف تقنيات الذكاء الاصطناعي في الأمن السيبراني، وبين كل من حماية البنية الإعلامية الرقمية، ورفع جودة العمل الصحفي، واستكشاف ما إذا كان الذكاء الاصطناعي قادراً على تقديم حلول متكاملة توازن بين الأمان والاحترافية في العمل الإعلامي الرقمي.

أهمية البحث:

في ظل التطورات المتسارعة في مجال التكنولوجيا الرقمية، يُعد الذكاء الاصطناعي أحد أبرز الأدوات التي أحدثت تحولاً جوهرياً في مشهد الأمن السيبراني العالمي. فقد ساهم الاعتماد المتزايد على البنية التحتية السحابية وانتشار أجهزة إنترنت الأشياء في تعظيم أثر الأخطاء البشرية البسيطة، ما أتاح للمهاجمين الإلكترونيين استغلال قدرات الذكاء الاصطناعي لتجاوز القيود التقليدية للهجمات السيبرانية. وتُظهر الأبحاث الحديثة أن الذكاء الاصطناعي يتمتع بقدرات فائقة على تحليل السياقات، ودمج البيانات

المتفرقة، واتخاذ قرارات دقيقة تعزز من فعالية الهجمات الإلكترونية، مما يشكل تحدياً كبيراً أمام المدافعين عن الأنظمة الرقمية. (Riddell, 2023)

وفي هذا الإطار، تبرز أهمية تطوير أنظمة حماية مدعومة بالذكاء الاصطناعي قادرة على كشف التهديدات المحتملة والتصدي لها بكفاءة، وبينما تسعى بعض الدول إلى تضمين الأطر الأخلاقية ضمن أنظمة الذكاء الاصطناعي، فإن غياب التوافق الدولي في هذا المجال يزيد من خطر استخدام هذه التقنيات لأغراض خبيثة. وبالتالي، فإن تعزيز دور الأنظمة الخبيرة وتكييفها مع مختلف بيئات الشبكات يُعد ضرورة ملحة لمواجهة التحديات المستجدة في مجال الأمن السيبراني.

اهداف البحث:

- ١- تحليل دور تقنيات الذكاء الاصطناعي في تعزيز الأمن السيبراني للمنصات الإعلامية الرقمية. لفهم كيف تسهم ادوت الذكاء الاصطناعي في الكشف المبكر عن الهجمات السيبرانية واستباقها.
- ٢- قياس مدى فاعلية تطبيق أنظمة الذكاء الاصطناعي في حماية البيانات الصحفية والمحتوى الإعلامي من الاختراق أو التشويه خاصة في ظل انتشار الأخبار المضللة والهجمات الإلكترونية الموجهة.
- ٣- الكشف عن العلاقة بين أمن المنصات الإعلامية وجودة المحتوى الصحفي وذلك بتوضيح كيف يمكن لتحسين بيئة الأمان أن يُسهم في تقديم صحافة موثوقة ومستقرة.
- ٤- الكشف عن العلاقة بين مستوى الوعي المهني لدى الصحفيين بتقنيات الذكاء الاصطناعي والأمن السيبراني وجودة المنتج الصحفي

فروض البحث:

- ١- توجد علاقة ذات دلالة إحصائية بين استخدام تقنيات الذكاء الاصطناعي في تعزيز الأمن السيبراني للمنصات الإعلامية الرقمية وكفاءتها في الكشف المبكر عن الهجمات السيبرانية.
- ٢- توجد علاقة ذات دلالة إحصائية بين تطبيق أنظمة الذكاء الاصطناعي في حماية المحتوى الإعلامي وبين تقليل مخاطر الاختراق والتلاعب بالمحتوى.
- ٣- توجد علاقة ذات دلالة إحصائية بين مستوى الأمان السيبراني في المنصات الإعلامية وجودة المحتوى الصحفي للمنصة.
- ٤- توجد علاقة ذات دلالة إحصائية بين مستوى الوعي المهني لدى الصحفيين بتقنيات الذكاء الاصطناعي والأمن السيبراني وجودة المنتج الصحفي.

نوع البحث:

نظرًا لأن البحث يهدف إلى تحليل (دور تقنيات الذكاء الاصطناعي في حماية المنصات الرقمية، وقياس الفاعلية (فعالية الأنظمة)، والكشف عن علاقات بين متغيرات (كالعلاقة بين الأمن السيبراني وجودة المحتوى) فإن نوع البحث هو بحث وصفي تحليلي لأنه يصف الظاهرة كما هي (وضع الأمن السيبراني)، ويحلل العلاقات بين المتغيرات (الذكاء الاصطناعي، الأمن، الجودة)، دون تدخل مباشر في المتغيرات. لذلك فالبحث يندرج ضمن إطار الدراسات الاستكشافية التي تهدف إلى التعرف على أبعاد الظاهرة محل الدراسة والكشف عن مكوناتها الرئيسية. كما تصنّف في الوقت ذاته ضمن الدراسات الوصفية، التي لا تكتفي بجمع البيانات، بل تسعى إلى تنظيمها وتصنيفها وتحليلها بهدف تقديم توصيف دقيق للظاهرة. كما يسعى البحث إلى رصد

العوامل والآليات المؤثرة في الظاهرة قيد البحث، وتحليل النتائج في ضوء أهداف البحث وفروضة.

منهج البحث:

يعتمد البحث على المنهج الوصفي التحليلي، باعتباره الأنسب لطبيعة البحث وأهدافه، حيث يسعى البحث إلى وصف وتحليل واقع استخدام تقنيات الذكاء الاصطناعي في دعم الأمن السيبراني للمؤسسات الإعلامية الرقمية، وقياس مدى فاعلية هذه التقنيات في حماية البيانات الصحفية والمحتوى الإعلامي، والكشف عن العلاقة بين أمن المنصات الإعلامية وجودة العمل الصحفي، وقد تم الاعتماد على استمارة الاستبيان كأداة رئيسة لجمع البيانات، وتحليلها باستخدام الأساليب الإحصائية المناسبة لاختبار فرضيات الدراسة.

أدوات البحث:

تم الاعتماد على استمارة الاستبيان (من تصميم الباحثه) كأداة رئيسة لجمع البيانات، وتحليلها باستخدام الأساليب الإحصائية المناسبة لاختبار فرضيات البحث.

مجتمع البحث:

يتكون مجتمع البحث من الصحفيين، والخبراء، ومسؤولي الأمن السيبراني، والتقنيين العاملين في المؤسسات الإعلامية الرقمية في مصر ورؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية وأعضاء هيئة التدريس بأقسام الاعلام بكليات الآداب بالجامعات المصرية.

عينة البحث:

تم اختيار عينة البحث بالطريقة العمدية (Purposive Sample)، لتشمل الأفراد الذين يمتلكون خبرة مباشرة أو غير مباشرة في استخدام تقنيات الذكاء الاصطناعي والأمن السيبراني ضمن السياق الإعلامي الرقمي. وتضم العينة:

- ١- مجموعة من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية رقمية تعتمد على تقنيات حديثة.

٢- مديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في المؤسسات الإعلامية.

٣- رؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية.

٤- مجموعة من السادة أعضاء هيئة التدريس تخصص الإعلام الرقمي والذكاء الاصطناعي وتطبيقاته في الإعلام.

أولاً: العينة الاستطلاعية:

تم اختيار عينة الدراسة الاستطلاعية بالطريقة العشوائية من داخل كل تخصص ومن داخل مجتمع الدراسة وخارج العينة الأساسية للدراسة بأجمالي عدد (٣٢) فرد مقسمين كما هو موضح بالجدول رقم (١)

جدول (١) التوصيف الإحصائي للعينة الاستطلاعية.

النسبة المئوية	العدد	المجتمع
٤٦.٨٧%	١٥	صحفيين ومحررين رقميين
١٢.٥٠%	٤	مديري أنظمة تقنية المعلومات
١٢.٥٠%	٤	رؤساء تحرير ومشرفين على المحتوى.
٢٨.١٣%	٩	أعضاء هيئة التدريس
١٠٠%	٣٢	الإجمالي

ثانياً: العينة الأساسية:

تم اختيار عينة الدراسة الأساسية بالطريقة العمدية والتي تكونت من مجموعة من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية رقمية تعتمد على تقنيات حديثة ومديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في تلك المؤسسات الإعلامية وكذلك مجموعة من رؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية ومجموعة من السادة أعضاء هيئة التدريس تخصص الإعلام بكلّيات

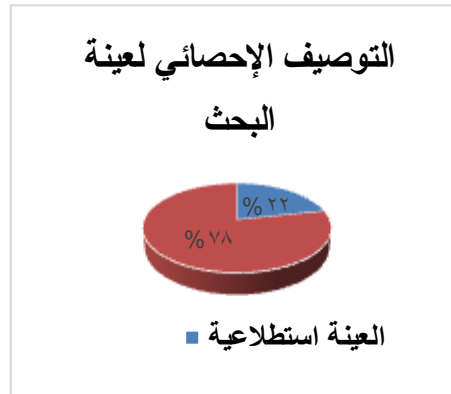
(١) الآداب بالجامعات المصرية، والبالغ عددهم عدد (١١٦) مشارك، وكما هو موضح بجدول

جدول (٢) التوصيف الإحصائي للعينة الأساسية.

النسبة المئوية	العدد	المجتمع
٦٢.٠٧%	٧٢	صحفيين ومحررين رقميين
٨.٦٢%	١٠	مديري أنظمة تقنية المعلومات
٨.٦٢%	١٠	رؤساء تحرير ومشرفين على المحتوى.
٢٠.٦٩%	٢٤	أعضاء هيئة التدريس
١٠٠%	١١٦	الإجمالي

جدول (٣) التوصيف الإحصائي لعينة البحث.

النسبة المئوية	العدد	المجتمع
٢١.٦٣%	٣٢	العينة استطلاعية
٧٨.٣٧%	١١٦	العينة الأساسية
١٠٠%	١٤٨	الإجمالي



شكل (١) توصيف عينة البحث

– أدوات البحث:

استخدمت الباحثة عدداً من الوسائل البحثية للوصول إلى المعلومات الخاصة بالدراسة وهي المراجع العربية والأجنبية، وذلك لبناء استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي كما يلي:

-الاستبيان:

تحقيقاً لأغراض البحث وبعد تنفيذ الدراسة الكمية والحصول على بيانات ومعلومات تستوفي وتعمق مشكلة البحث، وفي ضوء أهداف البحث ومن خلال المسح المرجعي قامت الباحثة بالاطلاع على العديد من المراجع العلمية والبحوث والدراسات السابقة حيث قامت الباحثة بتصميم استمارة استبيان الدراسة "فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي" حيث اتبعت الباحثة في ذلك خطوات بناء الاستبيان وفقاً لقواعد البحث العلمي كالتالي:

- مراجعة الأطر النظرية والدراسات السابقة المرتبطة بمتغيرات الدراسة ومدي ارتباطهم بأهداف البحث.
- تم تحديد المحاور المقترحة لاستبيان الدراسة من خلال المسح المرجعي.
- تحديد المفهوم النظري الإجرائي لمحاور استبيان الدراسة.
- صياغة مجموعة العبارات (المقترحة) الخاصة بكل محور.
- اقتراح عبارات لكل محور من محاور استبيان الدراسة في ضوء الفهم والتحليل النظري الخاص لكل محور.
- عرض استبيان الدراسة في صورته الأولية على السادة المحكمين لإبداء الرأي.
- صياغة الصورة النهائية لاستبيان الدراسة بعد الحذف والإضافة المقترحة من السادة المحكمين.

حيث قامت الباحثة بصياغة عبارات استبيان الدراسة في ضوء الفهم والتحليل النظري الخاص لكل محور، وقد استعانت الباحث ببعض البحوث السابقة حيث تم الحصول على بعض العبارات منها وتم تعديل صياغتها بما يتناسب مع عينة البحث، وتم إعداد استبيان الدراسة في ضوء الخطوات السابقة، حيث تم إعداد وصياغة العبارات تحت كل محور كلاً حسب طبيعته، وتكونت الصورة الأولية لاستبيان الدراسة كالتالي:

استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي مكونه من عدد (٤) محاور جدول (٤) وعدد (٤٦) عبارة في صورتها الأولية ملحق (٣).

وقد حرصت الباحثة عند بناء فقرات الاستبيان على مراعاة عدد من المعايير المنهجية لضمان جودة الأداة، وذلك على النحو التالي:

- صياغة العبارات بشكل واضح ومباشر، بما يسهل فهمها دون الحاجة إلى تفسير إضافي.
- تجنب الغموض أو التداخل في المعاني، بحيث تحمل كل عبارة معنى واحداً محدداً.
- الابتعاد عن الصياغات التي قد توحى بإجابة معينة أو تؤثر على رأي المجيب.
- ضمان ارتباط كل عبارة بأحد أهداف الدراسة، بما يعزز مصداقية النتائج.
- صياغة العبارات بأسلوب مختصر لا يتطلب وقتاً طويلاً للإجابة عليه.

-استطلاع رأي السادة المحكمين.

قامت الباحثة بعرض استمارة الاستبيان محل البحث على لجنة من السادة المحكمين من أعضاء هيئة التدريس (ملحق ١)، وعددهم سبعة محكمين، وذلك بهدف تقويم مدى ملاءمة محاور وعبارات الاستبيان لأهداف الدراسة، والتحقق من صدق المحتوى من خلال تقدير مدى توافق تلك المحاور والعبارات مع موضوع البحث، كما سعت الباحثة من خلال هذه المراجعة إلى تحديد مدى الحاجة إلى تعديل أو حذف أو إضافة بعض العبارات، عقب ذلك، قامت الباحثة بحساب النسبة المئوية لاتفاق السادة المحكمين على فقرات الاستبيان، وقد أسفرت نتائج استطلاع آرائهم كما هو موضح بمرفق رقم (٢)، (٣)، (٤).

-استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي:

في ضوء أهداف البحث ومن خلال المسح المرجعي قامت الباحثة بالاطلاع على العديد من المراجع العلمية والبحوث والدراسات السابقة حيث قامت الباحثة بتصميم استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي، حيث اتبعت الباحث في ذلك خطوات بناء الاستبيان وفقاً لقواعد البحث العلمي كالتالي:

- تحديد محاور استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي

قامت الباحثة بوضع مجموعة من المحاور المرتبطة بموضوع البحث وهي موضحة بالجدول التالي:

جدول (٤) محاور استمارة الاستبيان

م	المحور
١	الذكاء الاصطناعي والأمن السيبراني
٢	حماية المحتوى الإعلامي
٣	أمن المنصات الإعلامية وجودة المحتوى
٤	وعي الصحفيين بالتقنيات وجودة المنتج

- تحديد عبارات استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي:

قامت الباحثة بوضع عبارات الاستبيان المقترحة وعرض العبارات الخاصة بكل محور على السادة المحكمين وعددهم (٧) خبراء ملحق (١)، للتحقق من صدق المحتوى لملائمة العبارات المقترحة لكل محور، ومدي ملائمة صياغة العبارات المقترحة ومناسبتها للمحور الذي تنتمي إليه، ومدي إمكانية حذف وتعديل أو إضافة عبارات أخرى.

وقد روعي عند تصميم الاستبيان عدد من الضوابط المنهجية لضمان وضوح الأداة وفعاليتها، وذلك على النحو التالي:

- أن تتم صياغة العبارات بلغة عربية واضحة، بسيطة، ومألوفة، مع الحرص على أن تكون مختصرة قدر الإمكان ولا تحتاج إلى شرح إضافي.
- أن تكون العبارات سهلة وسريعة الإجابة، بما لا يُشكل عبئاً زمنياً على المشاركين.
- أن تغطي كل عبارة أحد أبعاد المحور المستهدف، مع السعي إلى شمولية القياس لجميع الجوانب ذات الصلة قدر الإمكان.

وقد بلغ عدد عبارات الاستبيان في صورته الأولى (٤٦) عبارة كما هو موضح في ملحق رقم (٣)، والملحق رقم (٤) يوضح النسبة المئوية لاتفاق السادة المحكمين على عبارات الاستبيان.

وقد ارتضت الباحثة بنسبة (٧٥%) فما أعلى من موافقة السادة المحكمين على العبارات التي تم الاتفاق عليها من مجموع آراء المختصين إذ يشير (بلوم وآخرون) إلى أنه "على الباحث الحصول على الموافقة بنسبة (٧٥%) وأكثر من آراء المحكمين" (بلوم وآخرون، ١٩٨٣، ٨٢٦)، وبذلك يكون عدد العبارات المقبولة والمكونة للاستبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي هو عدد (٢٥) عبارة وعدد العبارات المستبعدة هو عدد (٢١) عبارة والمرفق رقم (٥) يبين أرقام العبارات المقبولة والمرفق رقم (٦) يوضح أرقام العبارات المستبعدة، ومرفق رقم (٧) يوضح استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي في صورته النهائية.

- طريقة تصحيح الاستبيان

استخدمت الباحثة مقياس ليكرت الثلاثي (أوافق - محايد - لا أوافق) في تصميم استمارة الاستبيان بحيث تكون درجة أوافق = ٣، درجة محايد = ٢، ودرجة لا أوافق = ١، وقد أُنقّق السادة المحكمين على أن يكون ميزان تقدير الدرجات للاستبيان طبقاً لمقياس ليكرت ثلاثي التقدير، وبذلك تكون أعلى درجة لاستمارة الاستبيان هي (٩٠) درجة بينما تكون أقل درجة هي (٣٠) درجة.

- الدراسة الاستطلاعية:

كان الهدف من هذه الدراسة هو التأكد من المعاملات العلمية (الصدق، الثبات) استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي، وذلك على العينة الاستطلاعية والتي قوامها (٣٢) مشارك من داخل مجتمع البحث وخارج العينة الأساسية للبحث من مجموعة من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية رقمية ومديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في تلك المؤسسات الإعلامية وكذلك مجموعة من رؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية ومجموعة من السادة أعضاء هيئة التدريس تخصص الإعلام بكليات الآداب بالجامعات المصرية، وقد تم اختيارهم من مجتمع البحث وخارج عينة الدراسة، وذلك في الفترة من (الأحد ٢٠٢٥/٢/٢م) إلي (الأثنين ٢٠٢٥/٢/١٠م).

-المعاملات العلمية لاستمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي

١-صدق الاستبيان:**أ- صدق المُحكّمين (الصدق المحتوي):**

قامت الباحثة باستخدام صدق المُحكّمين (الصدق المحتوي)، حيث تم عرض استمارة استطلاع رأي السادة المحكمين لاستمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي. في صورتها الأولية مرفق (٣) على مجموعة المحكمين والبالغ عددهم (٧)، والموضحة أسمائهم بالمرفق (١)، واعتبرت الباحثة نسبة اتفاق السادة المحكمين على عبارات الاستبيان معياراً لصدقه.

ب- صدق الاتساق الداخلي:

كما قامت الباحثة بحساب صدق استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي. قيد البحث من خلال استخدام طريقة صدق الاتساق الداخلي، حيث قامت الباحثة بحساب قيمة معاملات الارتباط بين درجة كل عبارة على حدة والدرجة الكلية للاستبيان، وحساب قيمة معاملات الارتباط بين درجة كل عبارة على حدة والدرجة الكلية للمحور التابعة له، وحساب قيمة معاملات الارتباط بين درجة كل محور على حدي والدرجة الكلية للاستبيان على عينة الدراسة الاستطلاعية والتي قوامها (٣٢) مشارك. والجدول (٥ ، ٦ ، ٨) توضح ذلك.

جدول (٥) معاملات الارتباط ما بين كل محور والدرجة الكلية لاستمارة الاستبيان

ن=٣٢

معامل الارتباط ودلالته	المحاور	م
*٠.٦٠٤	الذكاء الاصطناعي والأمن السيبراني	١
*٠.٥٩٦	حماية المحتوى الإعلامي	٢
*٠.٣٩٩	أمن المنصات الإعلامية وجودة المحتوى	٣
*٠.٤٦٦	وعي الصحفيين بالتقنيات	٤

*قيمة (ر) الجدولية عند مستوى معنوية (٠.٠٥)، د. ح (٣٠) = (٠.٣٤٩).

يوضح الجدول رقم (٥) أن قيم معاملات الارتباط للمحاور مع الدرجة الكلية للاستبيان دالة عند مستوي معنوية (٠.٠٥)، حيث تراوحت قيم معامل الارتباط بين (٠.٣٩٩ - ٠.٦٠٤) وهي قيم جميعها أكبر من قيمة (ر) الجدولية عند مستوى دلالة (٠.٥) ودرجة حرية (٣٠) والتي بقيمة (٠.٣٤٩).

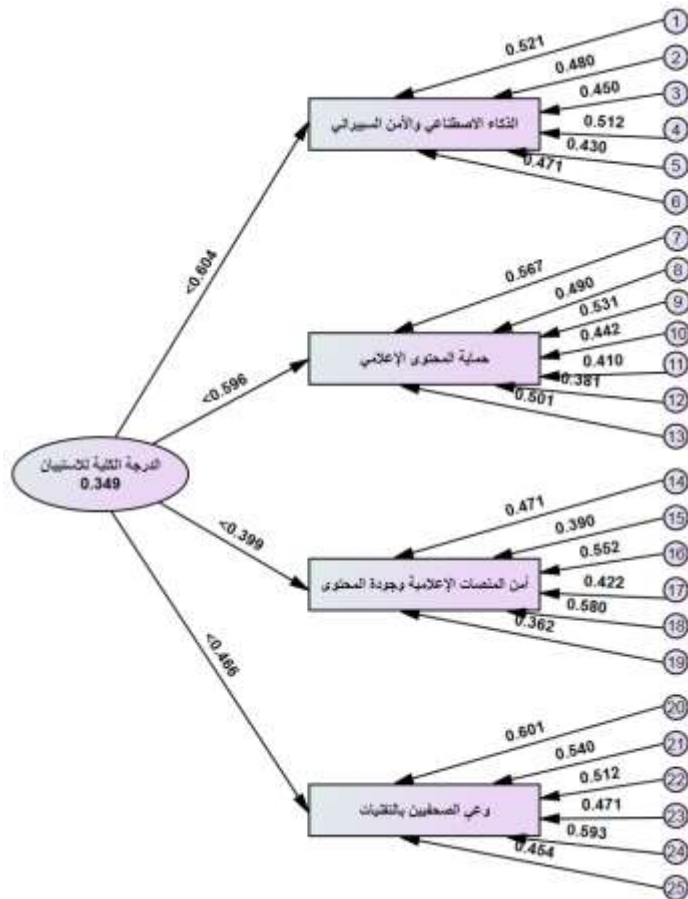
جدول (٦) معاملات الارتباط لكل عبارة والدرجة الكلية لكل محور من محاور الاستبيان

ن=٣٢

الذكاء الاصطناعي والأمن السيبراني		حماية المحتوى الإعلامي		أمن المنصات الإعلامية وجودة المحتوى		وعي الصحفيين بالتقنيات	
م	"ر"	م	"ر"	م	"ر"	م	"ر"
١	*.٥٢١	٧	*.٥٦٧	١٤	*.٤٧١	٢٠	*.٦٠١
٢	*.٤٨٠	٨	*.٤٩٠	١٥	*.٣٩٠	٢١	*.٥٤٠
٣	*.٤٥٠	٩	*.٥٣١	١٦	*.٥٥٢	٢٢	*.٥١٢
٤	*.٥١٢	١٠	*.٤٤٢	١٧	*.٤٢٢	٢٣	*.٤٧١
٥	*.٤٣٠	١١	*.٤١٠	١٨	*.٥٨٠	٢٤	*.٥٩٣
٦	*.٤٧١	١٢	*.٣٨١	١٩	*.٣٦٢	٢٥	*.٤٥٤
		١٣	*.٥٠١				

*قيمة (ر) الجدولية عند مستوى معنوية (٠.٠٥)، د. ح (٣٠) = (٠.٣٤٩)

يوضح الجدول رقم (٦) أن قيم معاملات الارتباط لكل عبارة مع الدرجة الكلية لمحور التابعة له ذات دالة عند مستوي معنوية (٠.٠٥) حيث تراوحت قيم معامل الارتباط بين (٠.٣٦٢ - ٠.٦٠١).



شكل (٢) معاملات الارتباط لكل عبارة والدرجة الكلية لكل محور من محاور الاستبيان

جدول (٧) أعلى وأقل قيمة لـ (ر)

المعيار	القيمة	رقم العبارة	المحور	العبارة
أعلى قيمة (ر)	٠.٦٠١	٢٠	المحور ٤ وعي الصحفيين بالتقنيات	زيادة الوعي بتقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي يحسن الجودة.
أقل قيمة (ر)	٠.٣٦٢	١٩	المحور ٣ أمن المنصات وجودة المحتوى	ضعف أمن المنصة قد يؤدي إلى نشر محتوى غير موثوق أو محرف.

يلاحظ من الجدول السابق ان اعلى قيمة (٠.٦٠١) تعكس اتساقاً داخلياً قوياً بين العبارة ٢٠ والمحور ٤، مما يشير إلى أن وعي الصحفيين بالتقنيات يرتبط بقوة بتحسين الجودة وان هذه العبارة الأكثر تأثيراً في الاستبيان.

وبالرغم من أن أقل قيمة هي (٠.٣٦٢) للعبارة رقم (١٩) من المحور الثالث الا انها دالة احصائيا لأنها جاءت أكبر من قيمة (ر) الجدولية والتي بقيمة (٠.٣٤٩) وتعتبر دالة إحصائياً.

جدول (٨) معاملات الارتباط ما بين كل عبارة والدرجة الكلية لاستمارة الاستبيان

ن=٣٢

الذكاء الاصطناعي والأمن السيبراني		حماية المحتوى الإعلامي		أمن المنصات الإعلامية وجودة المحتوى		وعي الصحفيين بالتقنيات	
معامل الارتباط	م	معامل الارتباط	م	معامل الارتباط	م	معامل الارتباط	م
*٠.٦٧٦	١	*٠.٥٥٢	٧	*٠.٥٨٦	١٤	٢٠	*٠.٦٠٢
*٠.٥٦٢	٢	*٠.٥٣٩	٨	*٠.٥١٧	١٥	٢١	*٠.٦٠٧
*٠.٤٤١	٣	*٠.٥٦٩	٩	*٠.٥٨١	١٦	٢٢	*٠.٤٣٦
*٠.٥٨٤	٤	*٠.٥٩٨	١٠	*٠.٥١٣	١٧	٢٣	*٠.٥٨١
*٠.٥٤٥	٥	*٠.٤٦٥	١١	*٠.٥٧٨	١٨	٢٤	*٠.٦٠٥
*٠.٤٠٢	٦	*٠.٥٢٥	١٢	*٠.٥٩١	١٩	٢٥	*٠.٦٣٢
		*٠.٥٢١	١٣				

*قيمة (ر) الجدولية عند مستوى معنوية (٠.٠٥)، د. ح (٣٠) = (٠.٣٤٩)

يوضح الجدول رقم (٨) أن قيم معاملات الارتباط للعبارات مع الدرجة الكلية لاستمارة للاستبيان دالة عند مستوى معنوية (٠.٠٥)، حيث تراوحت قيم معامل الارتباط بين (٠.٤٠٢ - ٠.٦٧٦) وهي قيم جميعها دالة احصائيا عند مستوى دلالة (٠.٠٥) لأنها جميعها أكبر من قيمة (ر) الجدولية (٠.٣٤٩)

من خلال العرض السابق للجدول (٥) يتضح أن جميع معاملات الارتباط للعبارات مع الدرجة الكلية للاستبيان ذات دلالة إحصائية، في حين تشير الجدول (٦) إلى أن جميع معاملات الارتباط الخاصة بكل عبارة والدرجة الكلية للمحور التابعة له ذات دلالة إحصائية مما يدل على صدقها، وكذلك يوضح جدول (٨) إلى ارتباط جميع المحاور بمعاملات ارتباط عالية مع الدرجة الكلية للاستبيان، ومن هنا نستطيع أن نحكم على الاستبيان بأنه متسق داخلياً وبالتالي صادق في قياس ما صمم من أجله.

٢- ثبات الاستبيان:

قامت الباحثة بإيجاد معامل ثبات محاور استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي وعددهم (٤) محاور وعباراتهم وعددها (٢٥) عبارة باستخدام معامل ألفا كرونباخ. Cronbach's alpha.

-الثبات باستخدام معامل ألفا كرونباخ Cronbach`s alpha:

جدول (٩) قيمة معامل الفا كرونباخ لكل محور من محاور استمارة

الاستبيان

م	المحور	عدد العبارات	معامل ألفا (α)	التفسير
١	الذكاء الاصطناعي والأمن السيبراني	٦	٠.٨١٥	اتساق داخلي جيد جداً
٢	حماية المحتوى الإعلامي	٧	٠.٨٤٣	اتساق داخلي جيد جداً
٣	أمن المنصات وجودة المحتوى	٦	٠.٨٨١	اتساق داخلي ممتاز
٤	وعي الصحفيين بالتقنيات	٦	٠.٨٦٧	اتساق داخلي ممتاز

يلاحظ من الجدول السابق ان قيمة معامل الفا كرونباخ لمحاور استمارة الاستبيان قد تراوحت بين (٠.٠٨١٥ ، ٨٨١) وجميعها قيم اقل من قيمة معامل ألفا كرونباخ لعبارات استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي والتي عددها (٢٥) عبارة والذي جاء بقيمة (٠.٨٩٢)

جدول (١٠) مُعامل الثبات باستخدام معامل ألفا كرونباخ لعبارات استمارة الاستبيان

ن=٣٢

الذكاء الاصطناعي والأمن السيبراني		حماية المحتوى الإعلامي		أمن المنصات الإعلامية وجودة المحتوى		وعي الصحفيين بالتقنيات	
م	معامل ألفا	م	معامل ألفا	م	معامل ألفا	م	معامل ألفا
١	٠.٨٠٧	٧	٠.٨١٨	١٤	٠.٨٦٦	٢٠	٠.٨٥٦
٢	٠.٨٠١	٨	٠.٨٢٠	١٥	٠.٨٥١	٢١	٠.٨٣٤
٣	٠.٨٠٥	٩	٠.٨١٥	١٦	٠.٨٧٩	٢٢	٠.٨٢٠
٤	٠.٨٠٩	١٠	٠.٨٣٩	١٧	٠.٨٦٨	٢٣	٠.٨١٧
٥	٠.٨١١	١١	٠.٨٣٤	١٨	٠.٨٦٢	٢٤	٠.٨٥٠
٦	٠.٨٠٤	١٢	٠.٨٠٨	١٩	٠.٨٤٥	٢٥	٠.٨١٣
		١٣	٠.٨٢٢				

*قيمة (معامل ألفا كرونباخ) لاستمارة الاستبيان = (٠.٨٩٢)

ويتضح من جدول (١٠) معامل ألفا كرونباخ لعبارات استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي في حالة حذف العبارة من عبارات الاستبيان، وقد تراوحت قيمة معامل ألفا كرونباخ ما بين (٠.٨٠١ - ٠.٨٧٩)، وهي قيم لا تزيد عن معامل ألفا كرونباخ للاستبيان والتي كانت (٠.٨٩٢)، مما يدل على ثبات عبارات استمارة الاستبيان.

- معامل ألفا كرونباخ لعبارات محور (الذكاء الاصطناعي والأمن السيبراني) والتي تراوحت ما بين (٠.٨١١ ، ٠.٨٠١) وجميعها قيم أقل من معامل ألفا كرونباخ للمحور والذي كانت قيمته (٠.٨١٥)

- معامل ألفا كرونباخ لعبارات محور (حماية المحتوى الإعلامي) والتي تراوحت ما بين (٠.٨٠٨ ، ٠.٨٣٩) وجميعها قيم أقل من معامل ألفا كرونباخ للمحور والذي كانت قيمته (٠.٨٤٣)

- معامل ألفا كرونباخ لعبارات محور (أمن المنصات الإعلامية وجودة المحتوى) والتي تراوحت ما بين (٠.٨٤٥ ، ٠.٨٧٩) وجميعها قيم أقل من معامل ألفا كرونباخ للمحور والذي كانت قيمته (٠.٨٨١)

- معامل ألفا كرونباخ لعبارات محور (وعي الصحفيين بالتقنيات) والتي تراوحت ما بين (٠.٨١٣ ، ٠.٨٥٦) وجميعها قيم أقل من معامل ألفا كرونباخ للمحور والذي كانت قيمته (٠.٨٦٧)

من خلال العرض السابق للجدول (١٠) يتضح أن معامل ألفا كرونباخ لعبارات الاستبيان في حالة حذف العبارة من عبارات الاستبيان كانت أقل من قيمة معامل ألفا كرونباخ للاستبيان، في حين يشير الجداول (١١) إلى أن قيم معامل ألفا كرونباخ لعبارات كل محور من محاور الاستبيان في حالة حذف العبارة من المحور

كانت اقل منة قيمة معامل ألفا كرونباخ للمحاور، وكذلك يوضح جدول (١١) إلي أن قيم معامل ألفا كرونباخ للمحاور كانت اقل منة قيمة معامل ألفا كرونباخ للاستبيان، ومن هنا نستطيع أن نحكم علي الاستبيان بأنه متسق داخلياً وبالتالي ثابت في قياس ما صمم من أجله.

-التجربة الأساسية:

تم إجراء التجربة الأساسية على عينة البحث الأساسية والتي قوامها (١١٦) مشارك من داخل مجتمع البحث وخارج العينة الاستطلاعية للبحث من من مجموعة من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية رقمية تعتمد على تقنيات حديثة ومديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في تلك المؤسسات الإعلامية وكذلك مجموعة من رؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية ومجموعة من السادة أعضاء هيئة التدريس تخصص الإعلام بكليات الآداب بالجامعات المصرية، على ألا يكونوا قد شاركوا في التجربة الاستطلاعية، وذلك في المدة من الأحد (٢٠٢٥/٢/١٦) إلى الأحد (٢٠٢٥/٣/٢).

عرض ومناقشة النتائج:

١- عرض ومناقشة نتائج الفرض الأول الذي ينص على " توجد علاقة ذات دلالة إحصائية بين استخدام تقنيات الذكاء الاصطناعي في تعزيز الأمن السيبراني للمنصات الإعلامية الرقمية وكفاءتها في الكشف المبكر عن الهجمات السيبرانية". ولإثبات صحة هذا الفرض قامت الباحثة بحساب دلالة الفروق بين استجابات عينة البحث على مقياس ليكرت الثلاثي بـ (نعم - إلى حد ما - لا) على عبارات (استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي) باستخدام اختبار (كا^٢)، وكذلك حساب مقاييس النزعة المركزية لدرجات عينة البحث وترتيبها

وفقا لمتوسط درجات كل محور من محاور استمارة الاستبيان، ورصدت نتائج ذلك في الجدول التالية:

جدول (١١) التكرار والنسبة المئوية وقيمة (كا^٢) وترتيب العبارات لاستجابات

أفراد عينة البحث لعبارات محور " الذكاء الاصطناعي والأمن السيبراني "

ن = ١١٦

م	نعم		إلى حد ما		لا		الوزن النسبي	كا ^٢	الترتيب
	التكرار	%	التكرار	%	التكرار	%			
١	٧١	٦١.٢١ %	٣٠	٢٥.٨٦ %	١٥	١٢.٩٣ %	٨٢.٧٦ %	٤٣.٤٦ *٦	٤
٢	٧٢	٦٢.٠٧ %	٣٤	٢٩.٣١ %	١٠	٨.٦٢ %	٨٤.٤٨ %	٥٠.٥٥ *٢	٣
٣	٦٥	٥٦.٠٣ %	٣٢	٢٧.٥٩ %	١٩	١٦.٣٨ %	٧٩.٨٩ %	٢٩.٠٨ *٦	٥
٤	٦٠	٥١.٧٢ %	٣٦	٣١.٠٣ %	٢٠	١٧.٢٤ %	٧٨.١٦ %	٢٠.٩٦ *٦	٦
٥	٧٦	٦٥.٥٢ %	٢٨	٢٤.١٤ %	١٢	١٠.٣٤ %	٨٥.٠٦ %	٥٧.٣٧ *٩	١
٦	٧٤	٦٣.٧٩ %	٣١	٢٦.٧٢ %	١١	٩.٤٨ %	٨٤.٧٧ %	٥٣.٦٠ *٣	٢

*قيمة (كا^٢) الجدولية عند مستوى معنوية (٠.٠٥) ودرجة حرية (٢) = ٥.٩٩١

يتضح من جدول (١١) أن جميع قيم (كا^٢) المحسوبة لاستجابات أفراد العينة لعبارات محور " الذكاء الاصطناعي والأمن السيبراني " جميعها دالة إحصائياً عند مستوى معنوية (٠.٠٥) حيث انها جميعاً قيم أكبر من ٥.٩٩١ وهي قيمة (كا^٢) الجدولية عند مستوى معنوية (٠.٠٥) ودرجة حرية (٢)

وهذا يثبت صحة الفرض الأول الذي ينص على " توجد علاقة ذات دلالة إحصائية بين استخدام تقنيات الذكاء الاصطناعي في تعزيز الأمن السيبراني للمنصات الإعلامية الرقمية وكفاءتها في الكشف المبكر عن الهجمات السيبرانية."

جدول (١٢) مقاييس النزعة المركزية لدرجات عينة البحث على استمارة الاستبيان
لمحور " الذكاء الاصطناعي والأمن السيبراني "

الترتيب	أكبر درجة	أقل درجة	الالتواء	الوسيط	الانحراف المعياري	المتوسط الحسابي	المحور
١	٨٧.٠٦	٨٥.١٦	- ٠.٨٣٠	٨٣.٦٢	٢.٨٧٥	٨٢.٥٢٠	الذكاء الاصطناعي والأمن السيبراني

يتضح من جدول (١٢) أن قيمة المتوسط الحسابي لعبارات المحور "الذكاء الاصطناعي والأمن السيبراني" جاءت بقيمة (٨٢.٥٢٠) وأن قيمة الانحراف المعياري جاءت بقيمة (٢.٨٧٥) وأن الوسيط جاء بقيمة (٨٣.٦٢) في حين كانت قيمة معامل الالتواء (-٠.٨٣٠) وجاء ترتيب المحور "الذكاء الاصطناعي والأمن السيبراني" في المركز الأول بين محاور استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي.

٢- عرض ومناقشة نتائج الفرض الثاني الذي ينص على " توجد علاقة ذات دلالة إحصائية بين تطبيق أنظمة الذكاء الاصطناعي في حماية المحتوى الإعلامي وبين تقليل مخاطر الاختراق والتلاعب بالمحتوى". ولإثبات صحة هذا الفرض قامت الباحثة بحساب دلالة الفروق بين استجابات عينة البحث على عبارات استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي باستخدام اختبار (كا^٢)، وكذلك حساب مقاييس النزعة المركزية لدرجات عينة البحث وترتيبها وفقا

لمتوسط درجات كل محور من محاور استمارة الاستبيان، ورصدت نتائج ذلك في الجدول التالية:

جدول (١٣) التكرار والنسبة المئوية وقيمة (كا^٢) وترتيب العبارات لاستجابات أفراد عينة البحث لعبارات محور " حماية المحتوى الإعلامي "

ن = ١١٦

م	نعم		إلى حد ما		لا		الوزن النسبي	كا ^٢	الترتيب ب
	التكرار ر	%	التكرار ر	%	التكرار ر	%			
١	٧٥	٦٤.٦٦ %	٣٠	٢٥.٨٦ %	١١	٩.٤٨ %	٨٥.٠٦ %	٥٥.٨٧	١
٢	٧٠	٦٠.٣٤ %	٣٢	٢٧.٥٩ %	١٤	١٢.٠٧ %	٨٢.٧٦ %	٤٢.٢٧	٦
٣	٧٢	٦٢.٠٧ %	٣١	٢٦.٧٢ %	١٣	١١.٢١ %	٨٣.٦٢ %	٤٧.٢٩	٤
٤	٦٨	٥٨.٦٢ %	٣٥	٣٠.١٧ %	١٣	١١.٢١ %	٨٢.٤٧ %	٣٩.٦٣	٧
٥	٧٣	٦٢.٩٣ %	٣٠	٢٥.٨٦ %	١٣	١١.٢١ %	٨٣.٩١ %	٤٩.٤٦	٣
٦	٧١	٦١.٢١ %	٣٣	٢٨.٤٥ %	١٢	١٠.٣٤ %	٨٣.٦٢ %	٤٦.٢٥	٤
٧	٧٤	٦٣.٧٩ %	٢٩	٢٥.٠٠ %	١٣	١١.٢١ %	٨٤.٢٠ %	٥١.٧٤	٢

*قيمة (كا^٢) الجدولية عند مستوى معنوية (٠.٠٥) ودرجة حرية (٢) = ٥.٩٩١

يتضح من جدول (١٣) أن قيمة (كا^٢) المحسوبة لاستجابات أفراد العينة لعبارات محور " حماية المحتوى الإعلامي " جميعها دالة إحصائياً عند مستوي معنوية (٠.٠٥).

وهذا يثبت صحة الفرض الثاني الذي ينص على " توجد علاقة ذات دلالة إحصائية بين تطبيق أنظمة الذكاء الاصطناعي في حماية المحتوى الإعلامي وبين تقليل مخاطر الاختراق والتلاعب بالمحتوى".

جدول (١٤) مقاييس النزعة المركزية لدرجات عينة البحث على استمارة الاستبيان لمحور "حماية المحتوى الإعلامي"

الترتيب	أكبر درجة	أقل درجة	الالتواء	الوسيط	الانحراف المعياري	المتوسط الحسابي	المحور
٢	٨٥.٠٦٠	٨٢.٤٧٠	- ٠.١٢٧	٨٣.٦٢٠	٠.٨٧٠	٨٣.٦٦٣	حماية المحتوى الإعلامي

يتضح من جدول (١٤) أن قيمة المتوسط الحسابي لعبارات المحور " حماية المحتوى الإعلامي" جاءت بقيمة (٨٣.٦٦٣) وان قيمة الانحراف المعياري جاءت بقيمة (٠.٨٧٠) وان الوسيط جاء بقيمة (٨٣.٦٢٠) في حين كانت قيمة معامل الالتواء (-٠.١٢٧) وجاء ترتيب المحور حماية المحتوى الإعلامي في المركز الثاني بين محاور استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي.

٣- عرض ومناقشة نتائج الفرض الثالث الذي ينص على " توجد علاقة ذات دلالة إحصائية بين مستوى الأمان السيبراني في المنصات الإعلامية وجودة المحتوى الصحفي للمنصة".

ولإثبات صحة هذا الفرض قامت الباحثة بحساب دلالة الفروق بين استجابات عينة البحث على عبارات (استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي) باستخدام اختبار (كا^٢) ، وكذلك حساب مقاييس النزعة

المركزية لدرجات عينة البحث وترتيبها وفقا لمتوسط درجات كل محور من محاور استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي ، ورصدت نتائج ذلك في الجدول التالية:

جدول (١٥) التكرار والنسبة المئوية وقيمة (كا^٢) وترتيب العبارات

لاستجابات

أفراد عينة البحث لعبارات محور " أمن المنصات الإعلامية وجودة المحتوى "

ن=١١٦

م	نعم		إلى حد ما		لا		الوزن النسبي	كا ^٢	الترتيب ب
	التكرار	%	التكرار	%	التكرار	%			
١	٦٥	٥٦.٠٠ %	٣٠	٢٥.٩٠ %	٢١	١٨.١٠ %	٧٩ %	٢٧.٩٤	٢
٢	٧٠	٦٠.٣٠ %	٢٨	٢٤.١٠ %	١٨	١٥.٥٠ %	٨١ %	٣٩.٣٧	١
٣	٥٨	٥٠.٠٠ %	٣٢	٢٧.٦٠ %	٢٦	٢٢.٤٠ %	٧٥ %	١٤.٩٦	٤
٤	٥٥	٤٧.٤٠ %	٣٨	٣٢.٨٠ %	٢٣	١٩.٨٠ %	٧٥ %	١٣.٢٥	٤
٥	٦٠	٥١.٧٠ %	٣٠	٢٥.٩٠ %	٢٦	٢٢.٤٠ %	٧٦ %	١٧.٨٦	٣
٦	٥٠	٤٣.١٠ %	٣٤	٢٩.٣٠ %	٣٢	٢٧.٦٠ %	٧١ %	١٣.٠٣	٦

قيمة (كا^٢) الجدولية عند مستوى معنوية (٠.٠٥) ودرجة حرية (٢) = ٥.٩٩١

يتضح من جدول (١٥) أن قيمة (كا^٢) المحسوبة لاستجابات أفراد العينة لعبارات محور " أمن المنصات الإعلامية وجودة المحتوى " جميعها دالة إحصائياً عند مستوي معنوية (٠.٠٥).

وهذا يثبت صحة الفرض الثالث الذي ينص على " توجد علاقة ذات دلالة احصائية بين مستوى الأمان السيبراني في المنصات الإعلامية وجودة المحتوى الصحفي للمنصة "

جدول (١٦) مقاييس النزعة المركزية لدرجات عينة البحث على استمارة الاستبيان لمحور " أمن المنصات الإعلامية وجودة المحتوى "

المحور	المتوسط الحسابي	الانحراف المعياري	الوسيط	الالتواء	أقل درجة	أكبر درجة	الترتيب
أمن المنصات الإعلامية وجودة المحتوى	٨١	٠.٨٧٠	٨٣.٦٢٠	- ٠.١٢٦	٧٧.٢٥٩	٨٣.١٢٠	٣

يتضح من جدول (١٦) أن قيمة المتوسط الحسابي لعبارات المحور "أمن المنصات الإعلامية وجودة المحتوى" جاءت بقيمة (٨١) وأن قيمة الانحراف المعياري جاءت بقيمة (٠.٧٨٠) وأن الوسيط جاء بقيمة (٨٣.٦٢٠) في حين كانت قيمة معامل الالتواء (-٠.١٢٦) وجاء ترتيب المحور أمن المنصات الإعلامية وجودة المحتوى في المركز الثالث بين محاور استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي.

٤- عرض ومناقشة نتائج الفرض الرابع الذي ينص على " توجد علاقة ذات دلالة إحصائية بين مستوى الوعي المهني لدى الصحفيين بتقنيات الذكاء الاصطناعي والأمن السيبراني وجودة المنتج الصحفي."

فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء د/ مي أحمد أبو صالحه

ولإثبات صحة هذا الفرض قامت الباحثة بحساب دلالة الفروق بين استجابات عينة البحث على عبارات (استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي) باستخدام اختبار (كا^٢)، وكذلك حساب مقاييس النزعة المركزية لدرجات عينة البحث وترتيبها وفقا لمتوسط درجات كل محور من محاور استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي ، ورصدت نتائج ذلك في الجدول التالية:

جدول (١٧) التكرار والنسبة المئوية وقيمة (كا^٢) وترتيب العبارات لاستجابات

أفراد عينة البحث لعبارات محور " وعي الصحفيين بالتقنيات وجودة المنتج

م	نعم		إلى حد ما		لا		الوزن النسبي	كا ^٢	الترتيب ب
	التكرار ر	%	التكرار ر	%	التكرار ر	%			
١	٧٢	٦٢.١٠ %	٣٠	٢٥.٩٠ %	١٤	١٢.١٠ %	٧٤.١٣ %	٢٨.٣ %	٢
٢	٨٠	٦٩.٠٠ %	٢٦	٢٢.٤٠ %	١٠	٨.٦٠ %	٧٧.٥٨ %	٣٥.٨ %	١
٣	٦٨	٥٨.٦٠ %	٣٢	٢٧.٦٠ %	١٦	١٣.٨٠ %	٧٢.٤١ %	٢٤.١ %	٣
٤	٦٠	٥١.٧٠ %	٣٨	٣٢.٨٠ %	١٨	١٥.٥٠ %	٦٧.٢٤ %	٢٠.٨ %	٤
٥	٥٦	٤٨.٣٠ %	٤٠	٣٤.٥٠ %	٢٠	١٧.٢٠ %	٦٥.٥١ %	١٧.٩ %	٥
٦	٥٠	٤٣.١٠ %	٤٤	٣٧.٩٠ %	٢٢	١٩.٠٠ %	٦٢.٠٦ %	١٥.٤ %	٦

-قيمة (كا^٢) الجدولية عند مستوى معنوية (٠.٠٥) ودرجة حرية (٢) = ٥.٩٩١

يتضح من جدول (١٧) أن قيمة (٢١) المحسوبة لاستجابات أفراد العينة لعبارات محور " أمن المنصات الإعلامية وجودة المحتوى " جميعها دالة إحصائياً عند مستوى معنوية (٠.٠٥).

وهذا يثبت صحة الفرض الثالث الذي ينص على " توجد علاقة ذات دلالة إحصائية بين تطبيق أنظمة الذكاء الاصطناعي في حماية البيانات الصحفية وبين جودة المحتوى الإعلامي "

جدول (١٨) مقاييس النزعة المركزية لدرجات عينة البحث على استمارة الاستبيان لمحور " وعي الصحفيين بالتقنيات وجودة المنتج "

الترتيب	أكبر درجة	أقل درجة	الالتواء	الوسيط	الانحراف المعياري	المتوسط الحسابي	المحور
٤	٨٦.١٨	٧٩.١٦	- ٠.٧٩٠	٨١.٤٢	٢.٧٧٣	٨١.٣٦٠	وعي الصحفيين بالتقنيات وجودة المنتج

يتضح من جدول (١٨) أن قيمة المتوسط الحسابي لعبارات المحور " وعي الصحفيين بالتقنيات وجودة المنتج " جاءت بقيمة (٨١.٣٦٠) وأن قيمة الانحراف المعياري جاءت بقيمة (٢.٧٧٣) وأن الوسيط جاء بقيمة (٨١.٤٢) في حين كانت قيمة معامل الالتواء (-٠.٧٩٠) وجاء ترتيب المحور "وعي الصحفيين بالتقنيات وجودة المنتج" في المركز الرابع بين محاور استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي.

مناقشة النتائج:

١- مناقشة نتائج الفرض الأول الذي ينص على " توجد علاقة ذات دلالة إحصائية بين استخدام تقنيات الذكاء الاصطناعي في تعزيز الأمن السيبراني للمنصات الإعلامية الرقمية وكفاءتها في الكشف المبكر عن الهجمات السيبرانية."

يتضح من جدول (١١) أن النسبة المئوية للاستجابة للعبارات بالإجابة بـ "نعم" تراوحت بين (٥١.٧٢-٦٥.٥٢) %، والنسبة المئوية للاستجابة للعبارات بالإجابة بـ "إلى حد ما" تراوحت ما بين (٢٤.١٤ - ٣١.٠٣) %، والنسبة المئوية للاستجابة للعبارات بالإجابة بـ "لا" تراوحت ما بين (٨.٦٢-١٧.٢٤) %.

كما يتضح من جدول (١١) أن هناك فروق ذات دلالة إحصائية في جميع عبارات محور " الذكاء الاصطناعي والأمن السيبراني " بالنسبة لعينة الدراسة لصالح الاستجابة الأعلى حيث تراوحت قيمة (كا^٢) المحسوبة ما بين قيمة (٢٠.٩٦٦-٥٧.٣٧٩)، وهي قيم جميعها دالة إحصائياً عند مستوي معنوية (٠.٠٥)، كما يوضح جدول (١١) إن العبارات جاءت بوزن نسبي تراوح ما بين (٧٨.١٦، ٨٥.٠٦) .

وقد أظهرت النتائج الموضح في جدول (١١) أن متوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "نعم" كانت بقيمة (٦٠.٠٥) % ، ومتوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "إلى حد ما" كانت بقيمة (٢٧.٤٤) %، ومتوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "لا" كانت بقيمة (١٢.٤٩) %، وتوضح هذه النتيجة أن غالبية المشاركين من عينة البحث قد اتفقوا على مضمون العبارات في الإجابة بنعم وهذا أيضاً ما يدل على صحة الفرض الذي ينص على " توجد علاقة ذات دلالة إحصائية بين استخدام

تقنيات الذكاء الاصطناعي في تعزيز الأمن السيبراني للمنصات الإعلامية الرقمية وكفاءتها في الكشف المبكر عن الهجمات السيبرانية."

ويتضح من جدول (١٢) مقاييس النزعة المركزية (المتوسط الحسابي - الانحراف المعياري - الوسيط - معامل الالتواء - اقل درجة - أكبر درجة) لكل متغير من متغيرات استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي لدى أفراد عينة البحث، وكذلك ترتيب المحور على المقياس لدى أفراد عينة البحث حيث جاء محور الذكاء الاصطناعي والأمن السيبراني في المرتبة الأولى وتعزو الباحثة هذه النتيجة إلى أن عينة البحث المكونة من مجموعة من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية رقمية تعتمد على تقنيات حديثة ومديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في تلك المؤسسات الإعلامية ورؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية والسادة أعضاء هيئة التدريس تخصص الإعلام بكليات الآداب بالجامعات المصرية يتفوقون على انه يرتبط استخدام تقنيات الذكاء الاصطناعي في دعم الأمن السيبراني للمنصات الإعلامية الرقمية بفاعلية تلك المنصات في الكشف المبكر عن التهديدات السيبرانية، ويشير ذلك إلى وجود ترابط وظيفي بين استخدام تقنيات الذكاء الاصطناعي وتعزيز مستوى الأمان السيبراني للمنصات الإعلامية، مما يسهم في رفع كفاءة تلك المنصات في الكشف المبكر عن الهجمات السيبرانية. ويستند هذا الطرح إلى ما أثبتته الأبحاث والدراسات الحديثة من أن الذكاء الاصطناعي يوفر أدوات متقدمة للتنبؤ بالمخاطر وتحليل الأنماط السلوكية المشبوهة عبر الخوارزميات الذكية، ما يمكن المنصات من التصدي للهجمات قبل وقوعها فعلياً. ويُعد هذا النوع من الوقاية المبكرة عنصراً حاسماً في بيئة الإعلام الرقمي التي تشهد تهديدات متزايدة في وتيرتها وتعقيدها.

وتبرز أهمية هذه الفرضية في كونها تعكس اتجاهاً متقدماً في ممارسات الأمن السيبراني، حيث لم يعد كافياً الاعتماد على أدوات الحماية التقليدية، بل أصبح من

الضروري تبني تقنيات الذكاء الاصطناعي التي تتيح استجابات أكثر سرعة ودقة. كما أن اختبار صحة هذا الفرض يوفر إطاراً علمياً لفهم مدى فاعلية هذه التقنيات في بيئة المنصات الإعلامية، التي تعتمد على المصادقية والموثوقية كمقوم أساسي. وبالتالي، فإن إثبات وجود هذه العلاقة يسهم في توجيه السياسات التقنية المستقبلية نحو دمج الذكاء الاصطناعي في البنية الأمنية للمؤسسات الإعلامية الرقمية.

وتتفق هذه النتيجة مع ما توصلت إليه دراسة كلا من Buczak & Guven (2016) التي أكدت أن خوارزميات تعلم الآلة توفر دقة عالية في اكتشاف الأنماط غير الطبيعية في بيانات الشبكات، ما يتيح سرعة الاستجابة للتهديدات المحتملة. كما تدعمها نتائج دراسة كلا من Sommer & Paxson (2020) التي أشارت إلى أن الذكاء الاصطناعي يعزز من قدرات أنظمة كشف التسلل عبر التعلم من البيانات دون الحاجة إلى توقعات مسبقة. كذلك أوضحت دراسة Sarker et al. (2020) أن توظيف الذكاء الاصطناعي في الأمن السيبراني يساعد في تطوير أنظمة دفاعية تتسم بالمرونة والديناميكية في مواجهة الهجمات المعقدة. كما أكدت دراسة Ahmad et al. (2021) على فاعلية الذكاء الاصطناعي في بناء استراتيجيات أمنية تعتمد على تحليل البيانات الضخمة والتعلم الآلي، مما يدعم فاعلية الكشف المبكر. ومن ثم، فإن هذه النتائج تدعم الفرض القائل بوجود علاقة بين استخدام الذكاء الاصطناعي وكفاءة حماية المنصات الإعلامية الرقمية.

-تفسير نتائج الفرض الأول من خلال ربط النتائج بأبعاد النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات (UTAUT) المستخدمة في البحث.

تنص الفرضية الأولى على أنه توجد علاقة ذات دلالة إحصائية بين استخدام تقنيات الذكاء الاصطناعي في تعزيز الأمن السيبراني للمنصات الإعلامية الرقمية وكفاءتها في الكشف المبكر عن الهجمات السيبرانية، وأوضحت النتيجة الإحصائية أن قيمة

(كا^٢) دالة إحصائية، بنسب موافقة عالية على عبارات المحور الأول كما هو موضح بـ(جدول ١١)، ويرتبط ذلك ببعد من ابعاد النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات (UTAUT) وهو بعد توقع الأداء (PE)، حيث ترى الباحثة ان اقبال المشاركين على استخدام تقنيات الذكاء الاصطناعي يُعزى إلى إدراكهم لقدرتها على تحسين أدائهم في الكشف المبكر عن الهجمات، مما يعكس توقعاً إيجابياً للأداء.

٢- مناقشة نتائج الفرض الثاني الذي ينص على "توجد علاقة ذات دلالة إحصائية بين تطبيق أنظمة الذكاء الاصطناعي في حماية المحتوى الإعلامي وبين تقليل مخاطر الاختراق والتلاعب بالمحتوى." حيث يتضح من جدول (١٣) أن النسبة المئوية للاستجابة للعبارات بالإجابة بـ "نعم" تراوحت ما بين (٥٨.٦٢ - ٦٤.٦٦) %، والنسبة المئوية للاستجابة للعبارات بالإجابة بـ "إلى حد ما" تراوحت ما بين (٢٥.٠٠ - ٣٠.١٧) %، والنسبة المئوية للاستجابة للعبارات بالإجابة بـ "لا" تراوحت ما بين (٩.٤٨ - ١٢.٠٧) %.

كما يتضح من جدول (١٣) أن هناك فروق ذات دلالة إحصائية في جميع عبارات محور "حماية المحتوى الإعلامي" بالنسبة لعينة الدراسة لصالح الاستجابة الأعلى حيث تراوحت قيمة (كا^٢) المحسوبة ما بين (٣٩.٦٣٨ - ٥٥.٨٧٩)، وهي قيم جميعها دالة إحصائية عند مستوي معنوية (٠.٠٥)، كما يوضح الجدول (١٣) إن العبارات جاءت بوزن نسبي ما بين (٨٢.٤٧ - ٨٥.٠٦).

وقد أظهرت النتائج الموضح في جدول (١٣) أن متوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "نعم" كانت بقيمة (٦١.٦٤) %، ومتوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "إلى حد ما" كانت بقيمة (٢٧.٤٤) %، ومتوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "لا" كانت بقيمة (١٠.٩٢) %، وتوضح هذه النتيجة ان غالبية عينة البحث من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية رقمية

تعتمد على تقنيات حديثة ومديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في تلك المؤسسات الإعلامية ورؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية والسادة أعضاء هيئة التدريس تخصص الإعلام بكليات الآداب بالجامعات المصرية قد اتفقوا على مضمون العبارات في الإجابة بنعم بنسبة كبيرة وهذا أيضا يثبت صحة الفرض الذي ينص على "توجد علاقة ذات دلالة إحصائية بين تطبيق أنظمة الذكاء الاصطناعي في حماية المحتوى الإعلامي وبين تقليل مخاطر الاختراق والتلاعب بالمحتوى".

كما يتضح من جدول (١٤) مقاييس النزعة المركزية (المتوسط الحسابي - الانحراف المعياري - الوسيط - معامل الالتواء - أقل درجة - أكبر درجة) لكل متغير من متغيرات استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي لدى أفراد عينة البحث، وكذلك ترتيب المحور على المقياس لدى أفراد عينة البحث حيث جاء محور حماية المحتوى الاعلامي في المرتبة الثانية وتغزو الباحثة افراد عينة البحث تتفق على وجود علاقة بين تطبيق أنظمة الذكاء الاصطناعي في حماية المحتوى الإعلامي وبين تقليل مخاطر الاختراق والتلاعب بالمحتوى. وتغزو الباحثة هذه العلاقة إلى قدرة أنظمة الذكاء الاصطناعي على مراقبة تدفق البيانات واكتشاف الأنشطة غير الطبيعية في الوقت الحقيقي، إضافة إلى استخدامها في التحقق من صحة المحتوى وكشف محاولات التعديل أو التزييف. ويعد هذا التوجه أحد أبرز الاتجاهات الحديثة في حماية البيانات الرقمية، لا سيما في المنصات الإعلامية التي تعتمد على مصداقية المحتوى وسرعة تداوله. وتدل هذه النتيجة على أن الدمج بين الأدوات الأمنية والذكاء الاصطناعي يمثل عاملاً حاسماً في تعزيز سلامة المحتوى الإعلامي ضد محاولات القرصنة والتضليل.

وتتسق هذه النتائج مع ما أشارت إليه دراسة (Ahmad et al. (2021 التي أوضحت أن أنظمة الذكاء الاصطناعي، لا سيما تلك المعتمدة على التعلم العميق، توفر آليات متقدمة لكشف محاولات التلاعب بالمحتوى مثل التعديل على الصور أو مقاطع الفيديو. كما تؤكد دراسة (Sarker et al. (2020 أن الذكاء الاصطناعي يسهم في بناء أنظمة تحقق تلقائية من صحة البيانات المتداولة داخل المنصات، مما يقلل من فرص التلاعب بالمعلومات أو تمرير محتوى زائف. كما أظهرت دراسة Buczak و Guven (2016) أهمية أنظمة الكشف المبكر في حماية البيانات والمحتوى من الهجمات المعتمدة على استغلال الثغرات البرمجية أو الهندسة الاجتماعية. وتُعد هذه النتائج داعمة للفرض الثاني، إذ تُبرز القيمة الفعلية لتقنيات الذكاء الاصطناعي في تحصين المحتوى الإعلامي وتقليل فرص اختراقه أو تحريفه، وهو ما ينعكس إيجاباً على مصداقية وأمان المنصات الرقمية.

-تفسير نتائج الفرض الثاني من خلال ربط النتائج بأبعاد النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات (UTAUT) المستخدمة في البحث.

تنص الفرضية الثانية والتي تنص على انه توجد علاقة ذات دلالة إحصائية بين تطبيق أنظمة الذكاء الاصطناعي في حماية المحتوى الإعلامي وبين تقليل مخاطر الاختراق والتلاعب بالمحتوى، وأوضحت النتيجة الإحصائية أن قيمة (كا^٢) دالة إحصائياً، بنسب موافقة عالية على عبارات المحور الثاني كما هو موضح بـ(جدول ١٣)، ويرتبط ذلك ببعد من ابعاد النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات (UTAUT) وهو بعد توقع الأداء (PE)، حيث ترى الباحثة ان المشاركين يتفقون على أن هذه الأنظمة تساهم بفعالية في حماية المحتوى (وهو جزء أساسي من عملهم)، مما يعزز من توقعهم لأدائها الإيجابي في هذا الجانب.

٣- مناقشة نتائج الفرض الثالث الذي ينص على " توجد علاقة ذات دلالة احصائية بين مستوى الأمان السيبراني في المنصات الإعلامية وجودة المحتوى الصحفي للمنصة " يتضح من جدول (١٥) أن النسبة المئوية للاستجابة للعبارات بالإجابة بـ "نعم" تراوحت ما بين (٤٣.١٠ - ٦٠.٣٠) %، والنسبة المئوية للاستجابة للعبارات بالإجابة بـ "إلى حد ما" تراوحت ما بين (٢٤.١٠ - ٣٢.٨٠) %، والنسبة المئوية للاستجابة للعبارات بالإجابة بـ "لا" تراوحت ما بين (١٥.٥٠ - ٢٧.٦٠) %.

كما يتضح من جدول (١٥) أن هناك فروق ذات دلالة إحصائية في جميع عبارات محور "أمن المنصات الإعلامية وجودة المحتوى" بالنسبة لعينة الدراسة لصالح الاستجابة الأعلى حيث تراوحت قيمة (٢١) المحسوبة ما بين (١٣.٠٣٤ - ٣٩.٣٧٩)، وهي قيم جميعها دالة إحصائياً عند مستوي معنوية (٠.٠٥)، كما يوضح الجدول (١٥) إن العبارات جاءت بوزن نسبي ما بين (٧١.٠٠ - ٨١.٠٠).

وقد أظهرت النتائج الموضح في جدول (١٥) أن متوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "نعم" كانت بقيمة (٥١.٤١) %، ومتوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "إلى حد ما" كانت بقيمة (٢٧.٦٠) %، ومتوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "لا" كانت بقيمة (٢٠.٩٦) %، وتوضح هذه النتيجة أن غالبية الشاب قد اتفقوا على مضمون العبارات في الإجابة بنعم وهذا أيضاً صحة الفرض الذي ينص على " توجد علاقة ذات دلالة احصائية بين مستوى الأمان السيبراني في المنصات الإعلامية وجودة المحتوى الصحفي للمنصة" كما يتضح من جدول (١٦) مقاييس النزعة المركزية (المتوسط الحسابي - الانحراف المعياري - الوسيط - معامل الالتواء - أقل درجة - أكبر درجة) لكل متغير من متغيرات استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي لدى أفراد عينة البحث، وكذلك ترتيب

المحور على المقياس لدى أفراد عينة البحث حيث جاء محور أمن المنصات الإعلامية وجودة المحتوى في المرتبة الثالثة وتعزو الباحثة أن عينة الدراسة من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية رقمية تعتمد على تقنيات حديثة ومديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في تلك المؤسسات الإعلامية ورؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية والسادة أعضاء هيئة التدريس تخصص الإعلام بكليات الآداب بالجامعات المصرية قد اتفقوا بنسبة كبيرة على أن مستوى الأمان السيبراني في المنصات الإعلامية يرتبط ارتباطاً وثيقاً بجودة المحتوى الصحفي الذي تُنتج تلك المنصات. وتعزو الباحثة هذا الارتباط إلى أن وجود بيئة رقمية مؤمنة يقلل من احتمالات التلاعب بالمحتوى أو اختراقه، مما يضمن تقديم مواد صحفية ذات مصداقية وجودة أعلى. فحينما تُصمم المنصات الإعلامية على نحو يضمن حماية البيانات وسرية مصادر المعلومات، فإن ذلك ينعكس إيجابياً على المحتوى من حيث الدقة والتحقق والموثوقية. وتشير هذه النتائج إلى أن أمن المعلومات لا يقتصر فقط على الجوانب التقنية، بل يمتد أثره ليشمل البنية التحريرية للمحتوى الصحفي ذاته.

وقد أكدت العديد من الدراسات السابقة هذه العلاقة. حيث أوضح Sarker et al. (2020) أن دمج تقنيات الذكاء الاصطناعي في أنظمة الأمن السيبراني يُعزز موثوقية بيئة النشر الرقمي ويقلل من فرص التلاعب أو التزيف في المواد الإعلامية. كما أشار Ahmad et al. (2021) إلى أن ضعف الأمن الرقمي قد يفتح المجال أمام الهجمات التي تستهدف المحتوى الإعلامي، مما يؤثر سلباً على جودته ومصداقيته أمام الجمهور. ومن جهة أخرى، تؤكد Buczak و Guven (2016) أن المنصات التي تطبق أنظمة حماية ذكية تستفيد من قدرات اكتشاف التهديدات قبل وقوعها، ما يتيح استمرار تدفق المحتوى بشكل آمن ومنتز. وعليه، فإن نتائج هذا الفرض تتسق مع ما توصلت إليه الأدبيات العلمية وتدعم أهمية الأمن السيبراني كعنصر جوهري في تحسين جودة المنتج الصحفي في البيئة الرقمية.

-تفسير نتائج الفرضية الثالثة من خلال ربط النتائج بأبعاد النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات (UTAUT) المستخدمة في البحث.

تنص الفرضية الثالثة والتي تنص على انه توجد علاقة ذات دلالة إحصائية بين مستوى الأمان السيبراني في المنصات الإعلامية وجودة المحتوى الصحفي للمنصة، وأوضحت النتيجة الإحصائية أن قيمة (كا^٢) دالة إحصائياً، بنسب موافقة عالية على عبارات المحور الثالث كما هو موضح بـ(جدول ١٥)، ويرتبط ذلك ببعدين من ابعاد النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات (UTAUT) وهما الظروف الميسرة (FC)، توقع الأداء (PE)، حيث ترى الباحثة انه يمكن اعتبار مستوى الأمان السيبراني العالي بمثابة "ظرف ميسر" يمكّن الصحفيين من إنتاج محتوى أفضل. كما أن جودة المحتوى الناتجة تعزز من "توقع الأداء" الإيجابي لتدابير الأمن.

٤-مناقشة نتائج الفرض الرابع الذي ينص على " توجد علاقة ذات دلالة إحصائية بين مستوى الوعي المهني لدى الصحفيين بتقنيات الذكاء الاصطناعي والأمن السيبراني وجودة المنتج الصحفي " يتضح من جدول (١٧) أن النسبة المئوية للاستجابة للعبارات بالإجابة بـ "نعم" تراوحت ما بين (٤٣.١٠ - ٦٩.٠٠) %، والنسبة المئوية للاستجابة للعبارات بالإجابة بـ "إلى حد ما" تراوحت ما بين (٢٢.٤٠ - ٣٧.٩٠) %، والنسبة المئوية للاستجابة للعبارات بالإجابة بـ "لا" تراوحت ما بين (٨.٦٠ - ١٩.٠٠) %.

كما يتضح من جدول (١٧) أن هناك فروق ذات دلالة إحصائية في جميع عبارات محور "وعى الصحفيين بالتقنيات وجودة المنتج" بالنسبة لعينة الدراسة لصالح الاستجابة الأعلى حيث تراوحت قيمة (كا^٢) المحسوبة ما بين (١٥.٤٧ - ٣٥.٨٧)، وهي قيم جميعها دالة إحصائياً عند مستوي معنوية (٠.٠٥)، كما يوضح الجدول (١٧) إن العبارات جاءت بوزن نسبي تراوح ما بين (٦٢.٠٦ - ٧٧.٥٨).

وأظهرت النتائج الموضح في جدول (١٧) أن متوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "نعم" كانت بقيمة (٥٥.٤٦) % ، ومتوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "إلى حد ما" كانت بقيمة (٣٠.١٨) %، ومتوسط النسبة المئوية للاستجابة للعبارات بالإجابة بـ "لا" كانت بقيمة (١٤.٣٦) %، وتوضح هذه النتيجة أن أكثر من نصف عينة البحث قد اتفقوا على أن مستوى الوعي المهني لدى الصحفيين بتقنيات الذكاء الاصطناعي والأمن السيبراني يرتبط ارتباطاً وثيقاً بجودة المنتج الصحفي وطبقاً للنسب الموضحة سابقاً، كما يتضح من جدول (١٨) مقاييس النزعة المركزية (المتوسط الحسابي - الانحراف المعياري - الوسيط - معامل الالتواء - أقل درجة - أكبر درجة) لكل متغير من متغيرات استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي لدى أفراد عينة البحث، وكذلك ترتيب المحور على المقياس لدى أفراد عينة البحث حيث جاء محور وعي الصحفيين بالتقنيات وجودة المنتج في المرتبة الرابعة وتغزو الباحثة أن عينة الدراسة من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية رقمية تعتمد على تقنيات حديثة ومديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في تلك المؤسسات الإعلامية ورؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية والسادة أعضاء هيئة التدريس تخصص الإعلام بكليات الآداب بالجامعات المصرية قد اتفقوا بنسبة كبيرة على أن مستوى الوعي المهني لدى الصحفيين بتقنيات الذكاء الاصطناعي والأمن السيبراني يرتبط ارتباطاً وثيقاً بجودة المنتج الصحفي، وترى الباحثة هذا الارتباط نتيجة امتلاك الصحفيين للمعرفة الكافية بهذه التقنيات لا يقتصر فقط على الجوانب التقنية، بل يمتد ليوثر بشكل مباشر في جودة المحتوى من حيث الدقة، والمصداقية، وسرعة المعالجة، ومهارات التحقق الرقمي. إذ يُمكن للصحفي الواعي بأدوات الذكاء الاصطناعي أن يستفيد منها في تحليل البيانات الضخمة، واستخراج

القصص الصحفية ذات القيمة، فضلاً عن تحسين قدرته على التحقق من المعلومات وتجنب التضليل، وهو ما ينعكس إيجاباً على جودة المنتج النهائي. وتدعم نتائج هذا الفرض ما أشارت إليه دراسة (Sarker et al. (2020)، التي تؤكد أن دمج المعرفة التقنية لدى الصحفيين بالذكاء الاصطناعي يعزز قدرتهم على تقديم محتوى مميز وموثوق. كما أوضحت دراسة (Ahmad et al. (2021) أن مستوى وعي العاملين في المجال الإعلامي بالأمن السيبراني يمثل خط الدفاع الأول ضد الهجمات التي قد تستهدف تشويه أو اختراق المحتوى الصحفي. كما تشير دراسة (Ali et al. (2022) إلى أن التكوين المهني للصحفيين في مجالات التكنولوجيا الرقمية يرتبط ارتباطاً وثيقاً بمستوى الابتكار وجودة الأداء في العمل الإعلامي. وعليه، فإن هذه النتائج تعزز من أهمية إدماج مفاهيم الذكاء الاصطناعي والأمن السيبراني في البرامج التدريبية للصحفيين لضمان تقديم محتوى عالي الجودة في بيئة إعلامية رقمية متغيرة وسريعة التحدي.

-تفسير نتائج الفرضية الرابعة من خلال ربط النتائج بأبعاد النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات (UTAUT) المستخدمة في البحث.

تنص الفرضية الرابعة والتي تنص على انه توجد علاقة ذات دلالة إحصائية بين مستوى الوعي المهني لدى الصحفيين بتقنيات الذكاء الاصطناعي والأمن السيبراني وجودة المنتج الصحفي، وأوضحت النتيجة الإحصائية أن قيمة (كا^٢) دالة إحصائياً، بنسب موافقة عالية على عبارات المحور الرابع كما هو موضح بـ(جدول ١٧)، وترى الباحثة ان ذلك يرتبط بثلاث ابعاد من ابعاد النظرية الموحدة لقبول واستخدام تكنولوجيا المعلومات (UTAUT) وهم توقع الجهد (EE)، الظروف الميسرة (FC)، توقع الأداء (PE) حيث ترى الباحثة ان الوعي المهني قد يقلل من "الجهد المتوقع" لاستخدام التقنيات بفعالية. كما يمكن اعتبار برامج التوعية والتدريب جزءاً من "الظروف الميسرة". والقدرة على استخدام هذه التقنيات بوعي لتحسين جودة المنتج تعكس "توقع أداء" مرتفع.

-النتائج والتوصيات:**أولاً: النتائج.**

١- أثبتت نتائج البحث وجود علاقة ذات دلالة إحصائية إيجابية بين استخدام تقنيات الذكاء الاصطناعي وتعزيز الأمن السيبراني للمنصات الإعلامية الرقمية وكفاءتها في الكشف المبكر عن الهجمات السيبرانية، مما يشير إلى أهمية الدمج بين الأدوات التقنية الحديثة وممارسات الحماية الرقمية لضمان بيئة إعلامية أكثر أماناً.

٢- أظهرت النتائج وجود ارتباط إيجابي بين تطبيق أنظمة الذكاء الاصطناعي في حماية المحتوى الإعلامي وتقليل مخاطر الاختراق والتلاعب بالمحتوى، وهو ما يعزز من قدرة المؤسسات الإعلامية على مواجهة التهديدات الرقمية المتزايدة.

٣- أوضحت النتائج أن مستوى الأمان السيبراني في المنصات الإعلامية يؤثر بشكل مباشر في جودة المحتوى الصحفي الذي تقدمه تلك المنصات، مما يبرز أهمية الاستثمار في البنية التحتية الأمنية كعنصر داعم لجودة العمل الصحفي.

٤- كشفت النتائج عن وجود علاقة دالة إحصائية موجبة بين مستوى الوعي المهني لدى الصحفيين بتقنيات الذكاء الاصطناعي والأمن السيبراني وبين جودة المنتج الصحفي، وهو ما يشير إلى أن تطوير المعرفة التقنية لدى الصحفيين يساهم في تحسين الأداء التحريري والإخباري على حد سواء.

ثانياً: التوصيات.

١- ضرورة دمج مفاهيم الذكاء الاصطناعي والأمن السيبراني ضمن المناهج التعليمية والتدريبية للصحفيين والعاملين في المؤسسات الإعلامية، لتطوير وعيهم المهني وتعزيز قدراتهم على التعامل مع التهديدات الرقمية.

٢- تشجيع المنصات الإعلامية الرقمية على الاستثمار في تقنيات الذكاء الاصطناعي لحماية المحتوى الصحفي من التلاعب أو الاختراق، من خلال تبني نظم تحقق تلقائي وتعلم آلي لاكتشاف السلوكيات المشبوهة.

٣- العمل على تحديث البنية التقنية للمنصات الإعلامية بما يتوافق مع متطلبات الأمن السيبراني الحديثة، بما في ذلك نظم كشف التسلل، والتشفير، والتحقق متعدد العوامل.

٤- إقامة شراكات بين المؤسسات الإعلامية وشركات التقنية لتطوير حلول أمنية مخصصة للبيئة الصحفية، تضمن حماية المحتوى والبيانات والمصادر الصحفية.

٥- إجراء دراسات مستقبلية تركز على أثر الوعي الأمني لدى الصحفيين في مجالات متخصصة كصحافة البيانات، والتحقق من الأخبار الزائفة، والصحافة الاستقصائية في البيئة الرقمية.

٦- تصميم نموذج تطبيقي لاستخدام تقنيات الذكاء الاصطناعي ضمن استراتيجيات الأمن السيبراني في المؤسسات الإعلامية بهدف تعزيز الحماية الرقمية وتحسين جودة العمل الصحفي

-البحوث المستقبلية:

١- إجراء دراسة للخروج باستراتيجية قومية تستهدف الاستفادة من تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية، وتحديد أولويات التطبيق داخل المؤسسات الإعلامية.

٢- عمل دراسة لتطوير برامج تدريبية للعاملين في أقسام الإعلام الرقمي تهدف إلى تنمية مهارات استخدام أدوات الذكاء الاصطناعي في مجالات الأمن السيبراني، ورفع كفاءتهم في التعامل مع الهجمات الرقمية.

- ٣- إجراء أبحاث تتناول دراسات حالة عن مؤسسات إعلامية دولية نجحت في توظيف الذكاء الاصطناعي في حماية محتواها وتطبيق ضوابط سيبرانية فعالة، للاستفادة من تجاربها في السياق العربي.
- ٤- عمل دراسات مقارنة بين المنصات الإعلامية المصرية والعربية والغربية من حيث مستوى نضج تطبيقات الأمن السيبراني المعززة بالذكاء الاصطناعي، ومدى تأثيرها في حماية المحتوى وجودته.
- ٥- إجراء دراسة تجريبية توضح مدى تأثير استخدام الذكاء الاصطناعي في حماية الصحافة الرقمية على ثقة الجمهور بالمحتوى الإعلامي، مع مقارنة ردود الفعل بين المحتوى المحمي تقنياً والمحتوى التقليدي.
- ٦- دراسة الأبعاد الأخلاقية لاستخدام الذكاء الاصطناعي في المجال الإعلامي، وخاصة فيما يتعلق بالتوازن بين الحماية الرقمية وحقوق الجمهور في الخصوصية والوصول إلى المعلومة.

ABSTRACT

The research aimed to measure the effectiveness of applying artificial intelligence systems in protecting press data and media content from hacking or distortion, especially in light of the escalation of misleading news and targeted cyber-attacks. The research also sought to reveal the relationship between the security of media platforms and the quality of journalistic content, by clarifying how improving the security environment can contribute to providing reliable and stable journalism, in addition to exploring the relationship between the level of professional awareness among journalists of artificial intelligence and cybersecurity techniques and the quality of the journalistic product, and the research sample consisted of (116) single individuals from journalists and digital editors.

The research showed a set of results, the most important of which is the existence of statistically significant positive relationships between the use of artificial intelligence technologies and enhancing the cybersecurity of digital media platforms, as these technologies contribute to the early detection of cyber-attacks and protect media content from hacking and manipulation. The results also showed that the level of cybersecurity affects the quality of journalistic content in a positive way, in addition to revealing that the professional awareness of journalists of modern technologies has a positive role in raising the efficiency of media performance and enhancing the reliability of the journalistic product

Keywords: Cybersecurity, artificial intelligence, media platforms

قائمة المراجع:

أولاً: المراجع العربية.

- ١- إسماعيل جابوري (٢٠٢٠). دور الأمن السيبراني في مواجهة التهديدات الإلكترونية دراسة حالة الجزائر، مجلة تحولات، المجلد الثالث، العدد الثاني، ص ٧٣، الجزائر.
- ٢- شيهان الورقلي (٢٠١٩) تأثير المذيع الروبوت على مهنة الإعلام: دراسة تحليلية سيميولوجية على عينة من النشرات الإخبارية، رسالة ماجستير غير منشورة، كلية العلوم الإنسانية والاجتماعية، جامعة قاصدي مرباح، الجزائر.
- ٣- عيسى عبد الباقي موسى، أحمد عادل عبد الفتاح (٢٠٢٠) اتجاهات الصحفيين والقيادات نحو توظيف تقنيات الذكاء الاصطناعي داخل غرف الأخبار بالمؤسسات الصحفية المصرية. دراسة تطبيقية، المجلة المصرية لبحوث الرأي العام، كلية الإعلام، جامعة القاهرة، مج ١٩، ع ١ ص ١-٦٦.

ثانياً: المراجع الأجنبية.

- 1- Abdalwahid, S. M., Hashim, W. A., Saeed, M. G., Altaie, S. A., & Kareem, S. W. (2024, April). Investigating the Effectiveness of Artificial Intelligence in Watermarking and Steganography for Digital Media Security. In *2024 21st International Multi-Conference on Systems, Signals & Devices (SSD)* (pp. 552-561). IEEE.
- 2- Abdel-Basset, M., Mohamed, R., & Chakraborty, R. K. (2024). Artificial Intelligence-Driven Intrusion Detection for Cybersecurity Systems. *Applied Sciences*, 14(5), 2341. <https://doi.org/10.3390/app14052341>
- 3- Ademola, O. E., & Somorin, K. (2024). *AI and Cybersecurity in Investigative Journalism: A Literature Review*. *Advances in Multidisciplinary & Scientific Research Journal Publication*, 10(1), 9-18.

- 4- Ahmad, I., Nam, H., Ismail, M., & Hussain, M. (2021). Artificial intelligence techniques for cybersecurity: A review. *IEEE Access*, 9, 110578–110602. <https://doi.org/10.1109/ACCESS.2021.3101867>
- 5- Ali, M. B., Tuhin, R., Alim, M. A., Rokonzaman, M., Rahman, S. M., & Nuruzzaman, M. (2024). Acceptance and use of ICT in tourism: the modified UTAUT model. *Journal of Tourism Futures*, 10(2), 334–349. <https://doi.org/10.1108/JTF-06-2021-0137>
- 6- Ali, M., Shah, S., & Khan, M. (2022). Digital journalism skills and AI: The evolving role of journalists in the era of automation. *Journal of Media Innovations*, 9(1), 45–61. <https://doi.org/10.5617/jmi.9876>
- 7- Al-Khalidi, A. (2023). *Ethical and legal challenges of AI-driven cybersecurity in digital media platforms*. *Journal of Media Ethics and Technology*, 12(3), 45-60
- 8- Al-Masri, E., Bai, Y., & Li, J. (2021). *Artificial intelligence in cybersecurity: A comprehensive review*. *Journal of Cybersecurity and Privacy*, 1(2), 45-67. <https://doi.org/10.3390/jcp1020004>
- 9- Al-Samiri, F., & Al-Faqih, K. (2022). *The impact of artificial intelligence on information security in media organizations*. *Arab Journal for Media and Communication Research*, 15(2), 33-50.
- 10- Anderson & Lee, 2022. The rising tide of DDoS attacks on media platforms. *Journal of Cybersecurity Research*, 9(2), 155-173.
- 11- Anderson, R., & Clark, J. (2022). *The impact of cybersecurity on media credibility: A case study of digital news platforms*. *International Journal of Media and Communication Research*, 7(3), 112-130. <https://doi.org/10.1080/12345678.2022.1234567>
- 12- Anderson, R., & Lee, S. (2022). Multilayered defense for media infrastructure. *Digital Protection Review*, 14(3), 77-95.
- 13- Babatope, A. B., & Anil, V. K. S. (2024). Cybersecurity on social media platforms: The use of deep learning methodologies for detecting anomalous user behavior. *GSC Advanced Research and Reviews*, 21(2), 38–47.
- 14- Benjamin K Sovacool & David J Hess , —Ordering Theories: Typologies And Conceptual Frameworks For Socio Technical Changel, *Social Studies Of Science*, Vol. 47(5) 2017 , P.P 703 –750
- 15- Bradshaw, S., Millard, C., & Walden, I. (2020). *Contracts for clouds: Comparison and analysis of the terms and conditions of cloud computing services*. *International Journal of Law and Information Technology*, 19(3), 187–223.

- 16- Brown, T., & Davis, R. (2019). *Emerging threats in cyberspace: Challenges and solutions*. Cybersecurity Journal, 12(4), 89-105. <https://doi.org/10.1016/j.cyber.2019.100123>
- 17- Brown, T., & Davis, R. (2019). *The impact of cybersecurity on media credibility: A case study of digital news platforms*. International Journal of Media and Communication Research, 7(3), 112-130. <https://doi.org/10.1080/12345678.2019.1234567>
- 18- Buczak, A. L., & Guven, E. (2016). A Survey of Data Mining and Machine Learning Methods for Cyber Security Intrusion Detection. *IEEE Communications Surveys & Tutorials*, 18(2), 1153–1176. <https://doi.org/10.1109/COMST.2015.2494502>
- 19- Cath, C. (2018). Governing artificial intelligence: Ethical, legal and technical opportunities and challenges. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 376(2133), 20180080.
- 20- Cohen, R., Smith, T., & Johnson, L. (2021). *AI-powered content moderation and fake news detection in digital journalism*. International Journal of Cyber Media Studies, 8(2), 112-130.
- 21- Cristian Vaccari And Andrew Chadwick, —Deepfakes And Disinformation: Exploring The Impact Of Synthetic Political Video On Deception, Uncertainty, And Trust In News, Social Media + Society, January-March 2020, P.P : 1 –13
- 22- Davis, P., & Brown, K. (2021). Human factor in media cybersecurity. *Information Security Journal*, 30(4), 210-228.
- 23- Dhingra, M. (2024). *AI-Driven Security Techniques in Social Media*. RSYN Proceedings.
- 24- *European Cybersecurity Journal*. (2023). *AI-driven threats in media: Policy brief*. 15(3), 88-102.
- 25- Garcia, M., & Lee, S. (2020). *Enhancing digital journalism through advanced cybersecurity measures*. Cybersecurity in Media Review, 5(1), 78-95.
- 26- Graves, L. (2018). *Understanding the promise and limits of automated fact-checking*. Reuters Institute for the Study of Journalism.
- 27- Harris, L., Taylor, M., & White, P. (2021). *Challenges and opportunities in implementing AI-driven cybersecurity solutions*. Journal of Information Security, 15(3), 200-215. <https://doi.org/10.1016/j.jis.2021.100456>

- 28- Khan, S., Lee, J., & Kim, H. (2022). *Encryption techniques in modern cybersecurity: A review*. International Journal of Network Security, 24(1), 45-60. <https://doi.org/10.1016/j.ijns.2022.100789>
- 29- Lee, J., & Kim, H. (2023). *Security automation: Enhancing incident response with AI*. Journal of Cybersecurity Technology, 8(2), 112-125. <https://doi.org/10.1016/j.jct.2023.100234>
- 30- Legit Security. (2023). The Role of AI in Cybersecurity. Retrieved from <https://www.legitsecurity.com/blog/role-of-ai-in-cybersecurity>
- 31- Marko Milosavljević & Igor Vobič , —_Our Task Is To Demystify Fears‘: Analysing Newsroom Management Of Automation In Journalism, Journalism, 2019 , P.P 1 –19
- 32- Marr, B. (2023). *The Future of Cybersecurity with Artificial Intelligence*. Wiley Digital Press.
- 33- Memon, N., & Wong, P. W. (1998). Protecting digital media content. *Communications of the ACM*, 41(7), 35-43.
- 34- Miller, R., & Zhang, Y. (2023). *Deepfake threats to digital media: Trends and countermeasures*. Journal of Media Security, 8(2), 134-152.
- 35- Nguyen, H., Patel, D., & Zhang, W. (2023). *Deep learning-based intrusion detection systems for news platforms: A performance analysis*. Journal of Artificial Intelligence in Security, 10(4), 205-220.
- 36- Nsude, I. (2022). Artificial Intelligence (AI), the media and security challenges in Nigeria. *Communication, technologies et développement*, (11).
- 37- Okdem, S., & Okdem, S. (2024). Artificial Intelligence in Cybersecurity: A Review of Applications and Challenges. *Cybersecurity and AI Review*, 3(2), 77–90. <https://www.sciencedirect.com/science/article/pii/S2667306923000149>
- 38- Palla, M., & Kostarella, I. (2025). *AI and Cybersecurity in Digital Journalism: Trends and Challenges*. Journal of Media Security and Innovation, 12(1), 33–49.
- 39- Palla, Z., & Kostarella, I. (2025). Journalists’ Perspectives on the Role of Artificial Intelligence in Enhancing Quality Journalism in Greek Local Media. *Societies*, 15(4), 89. <https://doi.org/10.3390/soc15040089>

- 40- Rainer Mühloff , —Human-Aided Artificial Intelligence: Or, How To Run Large Computations In Human Brains? Toward A Media Sociology Of Machine Learning, New Media & Society, Vol. 22(10) 2020 , P.P 1868 –1884
- 41- Sahraoui, S., & Benslimane, D. (2022). AI-Based Intrusion Detection in Digital Newsrooms. *Cybersecurity and Artificial Intelligence Journal*, 5(2), 112–126.
- 42- Sarker, I. H., Kayes, A. S. M., & Watters, P. (2020). Cybersecurity data science: An overview from machine learning perspective. *Journal of Big Data*, 7(1), 1–29. <https://doi.org/10.1186/s40537-020-00318-5>
- 43- Sarker, I. H., Kayes, A. S. M., & Watters, P. (2020). Cybersecurity data science: An overview from machine learning perspective. *Journal of Big Data*, 7(1), 1–29. <https://doi.org/10.1186/s40537-020-00318-5>
- 44- Smith, A., & Johnson, B. (2020). *Cybersecurity in digital media: Protecting content and building trust*. Media and Communication Studies, 10(2), 78-95. <https://doi.org/10.1080/12345678.2020.1234567>
- 45- Smith, A., & Johnson, B. (2020). *Emerging threats in cyberspace: Challenges and solutions*. Cybersecurity Journal, 12(4), 89-105. <https://doi.org/10.1016/j.cyber.2020.100123>
- 46- Sommer, R., & Paxson, V. (2020). Outside the closed world: On using machine learning for network intrusion detection. In *2010 IEEE Symposium on Security and Privacy* (pp. 305-316). IEEE. <https://doi.org/10.1109/SP.2010.25>
- 47- Sonni, A. F., Hafied, H., Irwanto, I., & Latuheru, R. (2024). Digital Newsroom Transformation: A Systematic Review of the Impact of Artificial Intelligence on Journalistic Practices, News Narratives, and Ethical Challenges. *Journalism and Media*, 5(4), 1554–1570. <https://doi.org/10.3390/journalmedia5040097>
- 48- Sonni, F., Ibrahimi, A., & Elouaai, F. (2024). The role of artificial intelligence in detecting and preventing cyber and phishing attacks. *International Journal of Advanced Computer Science and Applications*, 15(1), 112-120. <https://doi.org/10.14569/IJACSA.2024.0150113>
- 49- Taddeo, M., & Floridi, L. (2018). How AI can be a force for good. *Science*, 361(6404), 751–752.

- 50- Thompson, L., et al. (2022). Building security culture in media organizations. *Cybersecurity Awareness Bulletin*, 5(1), 44-63.
- 51- Thuraisingham, B. (2020, May). The role of artificial intelligence and cyber security for social media. In *2020 IEEE international parallel and distributed processing symposium workshops (IPDPSW)* (pp. 1-3). IEEE.
- 52- UNESCO. (2022). *Safety of Journalists in the Age of Digital Surveillance*. Paris: UNESCO Publishing.
- 53- Wang, Y., Zhang, L., & Liu, X. (2021). *Machine learning for cybersecurity: Applications and challenges*. IEEE Transactions on Information Forensics and Security, 16(5), 1234-1249. <https://doi.org/10.1109/TIFS.2021.1234567>
- 54- Waqas, M., Tu, S., Halim, Z., Rehman, S. U., Abbas, G., & Abbas, Z. H. (2022). The role of artificial intelligence and machine learning in wireless networks security: Principle, practice and challenges. *Artificial Intelligence Review*, 55(7), 5215-5261
- 55- Westerlund, M. (2021). The Emergence of Deepfake Technology: A Review. *Technology Innovation Management Review*, 11(2), 40–53.
- 56- Wilson, G., et al. 2023. The rising tide of DDoS attacks on media platforms. *Journal of Cybersecurity Research*, 9(2), 155-173
- 57- Zhang, L., Wang, Y., & Liu, X. (2020). *Behavioral analytics in cybersecurity: A new frontier*. Journal of Cybersecurity Research, 5(1), 34-50. <https://doi.org/10.1016/j.jcr.2020.100123>
- 58- Zhang, Z., & Gupta, B. B. (2018). Social media security and trustworthiness: overview and new direction. *Future Generation Computer Systems*, 86, 914-925.

قائمة المرفقات
مرفق رقم (١)
أسماء السادة المحكمين

م	الاسم	الوظيفة
١	أمين سعيد عبد الغني احمد	أستاذ بقسم الدراسات الإعلامية. كلية الآداب. جامعة الزقازيق
٢	داليا مصطفى إبراهيم	أستاذ العلاقات العامة والإعلام. كلية الآداب. جامعة حلوان
٣	رحاب الداخلي محمد	أستاذ ورئيس قسم الإعلام. كلية الآداب. جامعة اسيوط
٤	رحاب سراج الدين محمد	أستاذ الإعلام . كلية الآداب. جامعة المنيا
٥	عبد الباسط أحمد هاشم	أستاذ مساعد بشعبة العلاقات العامة والإعلام. كلية الآداب. جامعة سوهاج
٦	شيماء السيد سالم عمر	أستاذ العلاقات العامة والإعلام. كلية الآداب. جامعة حلوان
٧	يسرا حسني عبدالخالق	أستاذ الإعلام المساعد. كلية الآداب. جامعة اسيوط

-رتبت الأسماء ترتيباً أبجدياً.

مرفق رقم (٢)

استمارة الاستطلاع آراء المحكمين لتحديد الأهمية النسبية لعبارات الاستبيان
الخاص

فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية
المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي

السيد الأستاذ الدكتور /.....

تقوم الباحثة/ بدراسة تحت عنوان (فاعلية استخدام
تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات
الإعلامية الرقمية وتحسين جودة المحتوى الصحفي) بهدف البحث العلمي.
وكعملية منهجية بهدف الوصول لبناء أداة مقننة تساهم في تقييم فاعلية استخدام
تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية
الرقمية وتحسين جودة المحتوى الصحفي، ونظرا لمكانتكم العلمية وخبراتكم المتميزة في
هذا المجال فإنه يسعدني تعاونكم من خلال إعطاء نسبة مئوية تمثل أهمية كل عبارة
من عبارات الاستبيان وان كانت تصلح او لا تصلح واقتراح البديل لها في حالة
وجوده، وذلك لمساعدة الباحثة في الوصل الصورة النهائية لاستمارة الاستبيان الذي
تعدده الباحثة لهذا الغرض.

علماً بأن هذا الاستبيان موجه إلى:

١. مجموعة من الصحفيين والمحررين الرقميين العاملين في منصات إعلامية
رقمية.
٢. مديري أنظمة تقنية المعلومات ومسؤولي الأمن السيبراني في المؤسسات
الإعلامية.

٣. رؤساء تحرير ومشرفين على المحتوى في المؤسسات الإعلامية الرقمية.

شاكره احسن تعاونكم،،،،،،،،

اللقب / الاسم:

.....: جهة العمل:

عدد سنوات الخبرة:

الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى

(في صورته الأولى)

م	العبارات	تصلح	لا تصلح	التعديل إن وجد
١	الذكاء الاصطناعي والأمن السيبراني			
١	تعتمد مؤسستي على تقنيات الذكاء الاصطناعي لرصد محاولات الاختراق السيبراني.			
٢	تساهم أنظمة الذكاء الاصطناعي في الكشف المبكر عن التهديدات الأمنية.			

فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء د/ مي أحمد أبو صالحه

م	العبارات	تصلح	لا تصلح	التعديل إن وجد
٣	تساعد خوارزميات الذكاء الاصطناعي في تحليل البيانات الأمنية بكفاءة.			
٤	توفر تقنيات الذكاء الاصطناعي حماية مستمرة للأنظمة الإعلامية.			
٥	الذكاء الاصطناعي يساهم في تقليل زمن الاستجابة للهجمات السيبرانية.			
٦	تساعد أدوات الذكاء الاصطناعي في مراقبة السلوك الشبكي غير الطبيعي.			
٧	هناك اعتماد كبير على الذكاء الاصطناعي في تعزيز أمن المحتوى الرقمي.			
٨	تقنيات الذكاء الاصطناعي تقلل من الاعتماد على التدخل البشري في الحماية.			
٩	تستخدم منصتي أدوات ذكاء اصطناعي للتنبؤ بمحاولات الاختراق.			
١٠	تمثل تقنيات الذكاء الاصطناعي حلاً فعالاً للثغرات الأمنية المتكررة.			
١١	تستخدم مؤسستي تحليلات ذكاء اصطناعي لتحديد مصادر التهديدات الإلكترونية.			
١٢	تساهم تقنيات الذكاء الاصطناعي في حماية أنظمة النشر والبث الرقمية.			
٢	حماية المحتوى الإعلامي			
١٣	تساهم أنظمة الذكاء الاصطناعي في حماية المواد الصحفية من التلاعب.			
١٤	تؤدي تقنيات الذكاء الاصطناعي دوراً مهماً في تأمين أرشيف المحتوى الإعلامي.			
١٥	الذكاء الاصطناعي يساعد في الحفاظ على سرية المعلومات الصحفية.			
١٦	تمنع تقنيات الذكاء الاصطناعي تسريب البيانات الصحفية الحساسة.			
١٧	الذكاء الاصطناعي يقلل من خطر التزييف أو التعديل غير المصرح به للمحتوى.			
١٨	تستخدم مؤسستي أدوات ذكاء اصطناعي لرصد التعديلات غير الشرعية في المحتوى.			

م	العبارات	تصلح	لا تصلح	التعديل إن وجد
١٩	تساهم تطبيقات الذكاء الاصطناعي في تتبع محاولات تزوير الأخبار.			
٢٠	أدوات الذكاء الاصطناعي فعالة في حماية حسابات الصحفيين من القرصنة.			
٢١	أنظمة الذكاء الاصطناعي تكشف المحتوى المزيف أو المُعدّل بسرعة.			
٢٢	تؤثر الحماية القائمة على الذكاء الاصطناعي إيجابياً على ثقة الجمهور بالمحتوى.			
٢٣	تسهم أدوات الذكاء الاصطناعي في تصنيف المحتوى الخطير أو المشبوه.			
٢٤	أنظمة الذكاء الاصطناعي تقلل من تدخل الجهات غير المصرح لها في إنتاج الأخبار.			
٣	أمن المنصات الإعلامية وجودة المحتوى			
٢٥	تؤثر بيئة الأمن السيبراني بشكل مباشر على دقة الأخبار المنشورة.			
٢٦	زيادة الحماية الإلكترونية تنعكس إيجاباً على جودة العمل الصحفي.			
٢٧	الصحافة الرقمية تصبح أكثر احترافية مع تحصين المنصات من الهجمات.			
٢٨	تؤدي التهديدات السيبرانية إلى انخفاض جودة المحتوى الإعلامي.			
٢٩	يعزز الأمن السيبراني ثقة الجمهور بالمحتوى الصحفي.			
٣٠	تقل أخطاء التحرير مع وجود بيئة رقمية آمنة ومستقرة.			
٣١	جودة النشر الإلكتروني تتحسن مع تقنيات الحماية الذكية.			
٣٢	المنصات الآمنة تشجع الصحفيين على الإبداع والبحث الاستقصائي.			
٣٣	الحماية التكنولوجية تؤدي إلى استقرار تدفق المعلومات الدقيقة.			
٣٤	كلما زاد الأمن السيبراني زادت مصداقية المؤسسة الإعلامية.			

فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء /د/ مي أحمد أبو صالحه

م	العبارات	تصلح	لا تصلح	التعديل إن وجد
٣٥	ضعف أمن المنصة قد يؤدي إلى نشر محتوى غير موثوق أو محرف.			
٣٦	تتيح الحماية الفعالة للصحفيين التركيز على جودة المحتوى لا تأمينه.			
٤	وعي الصحفيين بالتقنيات			
٣٧	زيادة الوعي بتقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي يساهم في تحسين جودة المحتوى الصحفي.			
٣٨	التقنيات الذكية في التحليل الآلي للبيانات تقلل من التحيز الصحفي وتزيد من الدقة في التغطيات الإخبارية.			
٣٩	دمج الذكاء الاصطناعي في غرف الأخبار يُحدث تحولاً جذرياً في سرعة وكفاءة إنتاج المحتوى.			
٤٠	الأدوات الذكية للتحقق من صحة المعلومات تساهم في تقليل انتشار الإشاعات والتضليل الإعلامي.			
٤١	تحسين جودة الصحافة الرقمية يرتبط ارتباطاً وثيقاً بتطبيق تقنيات الأمن السيبراني.			
٤٢	الذكاء الاصطناعي يساهم في تحليل البيانات الصحفية بدقة أكبر، مما يحد من انتشار المعلومات المضللة.			
٤٣	معرفة الصحفيين بتقنيات مكافحة التزييف العميق (Deepfake) تُعدُّ ضرورة في العصر الرقمي.			
٤٤	المنصات الإعلامية التي تستخدم الذكاء الاصطناعي تُظهر موثوقية أعلى مقارنةً بالمنصات التقليدية.			
٤٥	تطوير مهارات الصحفيين في التعامل مع أدوات الذكاء الاصطناعي يحسّن جودة التحقيقات الاستقصائية.			

م	العبارات	تصلح	لا تصلح	التعديل إن وجد
٤٦	تبني تقنيات الذكاء الاصطناعي في الصحافة الرقمية يساعد في إنتاج محتوى أكثر أماناً وموضوعية.			

مرفق رقم (٤)

نسب موافقة السادة المحكمين على عبارات استمارة استبيان فاعلية استخدام
تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات
الإعلامية الرقمية وتحسين جودة المحتوى الصحفي

ن=٧

رقم العبارة	محكم ١	محكم ٢	محكم ٣	محكم ٤	محكم ٥	محكم ٦	محكم ٧	ت
المحور الاول الذكاء الاصطناعي والأمن السيبراني								
١	%٨٠	%٩٠	%٩٠	%٩٠	%٩٠	%٧٠	%٩٠	%٨٥.٧١
٢	%٧٠	%٧٠	%٧٠	%٦٥	%٦٠	%٧٠	%٦٠	%٦٦.٤٢
٣	%٦٠	%٦٠	%٧٠	%٧٠	%٦٠	%٧٠	%٦٠	%٦٤.٢٨
٤	%٨٠	%٩٠	%٩٠	%٩٠	%٨٠	%٩٠	%٩٠	%٨٧.١٤
٥	%٧٠	%٩٠	%٩٠	%٩٠	%٩٠	%٨٠	%٩٠	%٨٥.٧١
٦	%٦٠	%٧٠	%٩٠	%٧٠	%٨٠	%٧٠	%٨٠	%٧٤.٢٨
٧	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠
٨	%٦٠	%٦٠	%٧٠	%٧٠	%٦٠	%٧٠	%٦٠	%٦٤.٢٨
٩	%٧٠	%٦٠	%٦٠	%٧٠	%٧٠	%٦٠	%٧٠	%٦٥.٧١
١٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠
١١	%٨٠	%٦٠	%٧٠	%٦٠	%٧٠	%٦٠	%٩٠	%٧٠
١٢	%١٠٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩١.٤٢
المحور الثاني حماية المحتوى الاعلامي								
١٣	%٨٠	%٧٠	%٩٠	%٩٠	%٧٠	%٩٠	%٧٠	%٨٠
١٤	%٨٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٧٠	%٨٥.٧١
١٥	%٦٠	%٧٠	%٨٠	%٧٠	%٦٠	%٨٠	%٨٠	%٧١.٤٢
١٦	%٨٠	%٦٠	%٧٠	%٦٠	%٧٠	%٦٠	%٩٠	%٧٠
١٧	%٨٠	%٩٠	%٨٠	%٨٠	%٩٠	%٨٠	%٨٠	%٨٢.٨٥
١٨	%٧٠	%٩٠	%٨٠	%٧٠	%٩٠	%١٠٠	%٩٠	%٨٤.٢٨
١٩	%٧٠	%٦٠	%٧٠	%٦٠	%٧٠	%٦٠	%٧٠	%٦٥.٧١
٢٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٧٠	%٨٧.١٤

فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء / د/ مي أحمد أبو صالحه

٢١	%٦٠	%٧٠	%٦٠	%٧٠	%٦٠	%٧٠	%٦٠	%٦٤.٢٨
٢٢	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠
٢٣	%٨٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٨٥.٧١
٢٤	%٦٠	%٧٠	%٩٠	%٩٠	%٦٠	%٧٠	%٧٠	%٧٤.٢٨
المحور الثالث أمن المنصات الإعلامية وجودة المحتوى								
٢٥	%٦٠	%٥٠	%٧٠	%٦٠	%٧٠	%٦٠	%٨٠	%٦٤.٢٨
٢٦	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٧٠	%٨٧.١٤
٢٧	%٧٠	%٦٥	%٦٠	%٧٠	%٦٠	%٦٠	%٧٠	%٥٠.٦٥
٢٨	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠
٢٩	%٩٥	%٧٠	%٩٠	%٩٠	%٩٠	%٩٠	%٧٠	%٨٥
٣٠	%٨٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٧٠	%٨٥.٧١
٣١	%٦٠	%٧٠	%٦٠	%٧٠	%٦٠	%٦٠	%٧٠	%٦٤.٢٨
٣٢	%٦٠	%٩٠	%٩٠	%٧٠	%٦٠	%٩٠	%٩٠	%٧٤.٢٨
٣٣	%٩٠	%٦٠	%٩٠	%٧٠	%٦٠	%٧٠	%٧٠	%٧٢.٨٥
٣٤	%٨٠	%٩٠	%٨٠	%٨٠	%٨٠	%٩٠	%٨٠	%٨٢.٨٥
٣٥	%٩٥	%٧٠	%٩٠	%٩٠	%٩٠	%٩٠	%٧٠	%٨٥
٣٦	%٦٠	%٩٠	%٦٠	%٩٠	%٦٠	%٩٠	%٧٠	%٧٤.٢٨
المحور الرابع وعي الصحفيين بالتقنيات								
٣٧	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠
٣٨	%٦٠	%٧٠	%٧٠	%٩٠	%٧٠	%٦٠	%٧٠	%٧٢.٨٥
٣٩	%٦٠	%٧٠	%٦٠	%٩٠	%٦٠	%٧٠	%٧٠	%٧٠
٤٠	%٦٠	%٧٠	%٦٠	%٧٠	%٦٠	%٧٠	%٧٠	%٦٤.٢٨
٤١	%١٠٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩١.٤٢
٤٢	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠
٤٣	%٨٠	%٦٠	%٨٠	%٥٠	%٩٠	%٨٠	%٨٠	%٧٤.٢٨
٤٤	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠
٤٥	%٨٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٨٨.٥٧
٤٦	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٩٠	%٧٠	%٨٧.١٤

العبارة المقبولة



العبارة المرفوضة



وقد ارتضت الباحثة بنسبة (٧٥%) فما أعلى من موافقة السادة المحكمين على العبارات التي تم الاتفاق عليها من مجموع آراء المختصين إذ يشير (بلوم وآخرون) إلى أنه "على الباحث الحصول على الموافقة بنسبة (٧٥%) وأكثر من آراء

المحكمين" (بلوم وآخرون، ١٩٨٣، ٨٢٦) وبذلك يكون عدد العبارات المقبولة والمكونة لاستبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي هو عدد (٢٥) عبارة وعدد العبارات المستبعدة هو عدد (٢١) عبارة والملحق رقم (٥) يبين أرقام العبارات المقبولة والملحق رقم (٦) يوضح أرقام العبارات المستبعدة وقامت الباحثة بالاستعانة بمقياس ليكارت الثلاثي (أوافق - محايد - لا أوافق) في تصميم استمارة الاستبيان بحيث تكون درجة أوافق = ٣، درجة محايد = ٢، ودرجة لا أوافق = ١.

مرفق رقم (٥)

أرقام العبارات المقبولة

م	المحور	أرقام العبارات المقبولة	العدد
١	الذكاء الاصطناعي والأمن السيبراني	(١، ٤، ٥، ٧، ١٠، ١٢)	٦
٢	حماية المحتوى الإعلامي	(١٣، ١٤، ١٧، ١٨، ٢٠، ٢٢، ٢٣)	٧
٣	أمن المنصات الإعلامية وجودة المحتوى	(٢٦، ٢٨، ٢٩، ٣٠، ٣٤، ٣٥)	٦
٤	وعي الصحفيين بالتقنيات	(٣٧، ٤١، ٤٢، ٤٤، ٤٥، ٤٦)	٦
-	المجموع	-----	٢٥

مرفق رقم (٦)

أرقام العبارات المستبعدة

م	المحور	أرقام العبارات المستبعدة	العدد
١	الذكاء الاصطناعي والأمن السيبراني	(٢، ٣، ٦، ٨، ٩، ١١)	٦
٢	حماية المحتوى الإعلامي	(١٥، ١٦، ١٩، ٢١، ٢٤)	٥
٣	أمن المنصات الإعلامية وجودة	(٢٥، ٢٧، ٣١، ٣٢، ٣٣)	٦

فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء د/ مي أحمد أبو صالحه

	المحتوى	(٣٦)	
٤	وعي الصحفيين بالتقنيات	(٤٣، ٤٠، ٣٩، ٣٨)	٤
	المجموع	-----	٢١

مرفق رقم (٧)

استمارة الاستبيان في صورتها النهائية

عزيمي المشارك/هـ:.....

تحية طيبة تملؤها السعادة والاحترام....

تتشرف الباحثة د/..... بعرض مجموعة من العبارات التي تتعلق

بفاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية

المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي

ويوجد أمام كل عبارة اختيارات عديدة المرجو منك:

١- قراءة كل عبارة بدقة ثم إبداء الرأي بوضع علامة (✓) تحت الإجابة التي تنطبق ورأيكم.

٢- عدم ترك أي عبارة دون الإجابة عليها.

٣- عدم وضع أكثر من علامة أمام عبارة واحدة.

علماً بأن بيانات هذه الاستمارة تحاط بسرية تامة ولا تستخدم في غير أغراض البحث العلمي.

وشكراً لحسن تعاونكم،،،،

الباحثة

تابع مرفق (٧)

استمارة استبيان فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي في حماية المنصات الإعلامية الرقمية وتحسين جودة المحتوى الصحفي

(في صورته النهائية)

م	العبارات	أوافق	إلى حد ما	لا أوافق
١	الذكاء الاصطناعي والأمن السيبراني			
١	تعتمد مؤسستي على تقنيات الذكاء الاصطناعي لرصد محاولات الاختراق السيبراني.			
٢	توفر تقنيات الذكاء الاصطناعي حماية مستمرة للأنظمة الإعلامية.			
٣	الذكاء الاصطناعي يساهم في تقليل زمن الاستجابة للهجمات السيبرانية.			
٤	هناك اعتماد كبير على الذكاء الاصطناعي في تعزيز أمن المحتوى الرقمي.			
٥	تمثل تقنيات الذكاء الاصطناعي حلاً فعالاً للثغرات الأمنية المتكررة.			
٦	تساهم تقنيات الذكاء الاصطناعي في حماية أنظمة النشر والبث الرقمية.			
٢	حماية المحتوى الإعلامي			
٧	تساهم أنظمة الذكاء الاصطناعي في حماية المواد الصحفية من التلاعب.			
٨	تؤدي تقنيات الذكاء الاصطناعي دوراً مهماً في تأمين أرشيف المحتوى الإعلامي.			
٩	الذكاء الاصطناعي يقلل من خطر التزيف أو التعديل غير المصرح به للمحتوى.			
١٠	تستخدم مؤسستي أدوات ذكاء اصطناعي			

فاعلية استخدام تقنيات الأمن السيبراني المدعومة بالذكاء د/ مي أحمد أبو صالحه

م	العبارات	أوافق	إلى حد ما	لا أوافق
	لرصد التعديلات غير الشرعية في المحتوى.			
١١	أدوات الذكاء الاصطناعي فعالة في حماية حسابات الصحفيين من القرصنة.			
١٢	تؤثر الحماية القائمة على الذكاء الاصطناعي إيجابياً على ثقة الجمهور بالمحتوى.			
١٣	تسهم أدوات الذكاء الاصطناعي في تصنيف المحتوى الخطير أو المشبوه.			
٣	أمن المنصات الإعلامية وجودة المحتوى			
١٤	زيادة الحماية الإلكترونية تنعكس إيجاباً على جودة العمل الصحفي.			
١٥	تؤدي التهديدات السيبرانية إلى انخفاض جودة المحتوى الإعلامي.			
١٦	يعزز الأمن السيبراني ثقة الجمهور بالمحتوى الصحفي.			
١٧	تقل أخطاء التحرير مع وجود بيئة رقمية آمنة ومستقرة.			
١٨	كلما زاد الأمن السيبراني زادت مصداقية المؤسسة الإعلامية.			
١٩	ضعف أمن المنصة قد يؤدي إلى نشر محتوى غير موثوق أو محرف.			
٤	وعي الصحفيين بالتقنيات وجودة المنتج			
٢٠	زيادة الوعي بتقنيات الأمن السيبراني المدعومة بالذكاء الاصطناعي يساهم في تحسين جودة المحتوى الصحفي.			
٢١	تحسين جودة الصحافة الرقمية يرتبط ارتباطاً وثيقاً بتطبيق تقنيات الأمن السيبراني.			
٢٢	الذكاء الاصطناعي يساهم في تحليل البيانات الصحفية بدقة أكبر، مما يحد من انتشار المعلومات المضللة.			
٢٣	المنصات الإعلامية التي تستخدم الذكاء الاصطناعي تُظهر موثوقية أعلى مقارنة بالمنصات التقليدية.			
٢٤	تطوير مهارات الصحفيين في التعامل مع أدوات الذكاء الاصطناعي يحسن جودة التحقيقات الاستقصائية.			

م	العبارات	أوافق	إلى حد ما	لا أوافق
٢٥	تبني تقنيات الذكاء الاصطناعي في الصحافة الرقمية يساعد في إنتاج محتوى أكثر أماناً وموضوعية.			