



الكاشف الذكي عن المحتوى غير المرغوب في المواقع الإلكترونية

بسام عركوك¹، أكرم محمد زكي²

¹جامعة الحديدة - اليمن

²الجامعة الإسلامية العالمية بماليزيا

¹bassam.arkok2@gmail.com, ²akramzeki@iiu.edu.my

الخلاصة: أدى الانتشار الكبير للمواقع على الإنترنت إلى زيادة انتشار المحتوى الضار وغير المرغوب فيه بشكل مخيف، على سبيل المثال المحتوى الإباحي، والأفكار الهدامة مثل الإلحاد والتطرف. ولمعالجة هذه المشكلة، يقترح هذا البحث تقنية ذكية قادرة على تحديد المحتوى الضار أخلاقياً وفكرياً وتصفيته من المواقع الإلكترونية. يعمل نظامنا على الاستفادة من تقنيات التعلم الآلي المتقدمة، بما في ذلك معالجة اللغة الطبيعية وتقنيات الذكاء الاصطناعي، لتحليل النصوص والصور والمقاطع الصوتية والفيديو. يهدف هذا البحث إلى تحقيق الدقة والاستدعاء العالي في اكتشاف المحتوى غير المرغوب فيه، سواء كان نصاً أو صوتاً أو محتوى مرئياً وعناصر أخرى غير مرغوب فيها، من خلال تدريب نموذجنا على مجموعات بيانات مختلفة من المحتوى المعني. يمكن دمج الكاشف المقترح في منصات الإنترنت المختلفة، مثل وسائل التواصل الاجتماعي ومواقع التجارة الإلكترونية ومحركات البحث، لتعزيز تجربة المستخدم وبيئة رقمية أكثر أماناً.

الكلمات الجوهرية: الكاشف الذكي - المحتوى غير المرغوب - المواقع الإلكترونية - الذكاء الاصطناعي.

1. المقدمة:

مع الانتشار الواسع لاستخدام الإنترنت وتزايد عدد المواقع الإلكترونية، أصبح من الضروري تطوير تقنيات ذكية قادرة على كشف وتصنيف المحتوى غير المرغوب فيه. يشمل هذا المحتوى الضار المواد الإباحية، والأفكار المتطرفة، والمعلومات المضللة، وغيرها من العناصر التي قد تؤثر سلباً على المستخدمين. يهدف هذا البحث إلى إقترح نظام كاشف ذكي يعتمد على تقنيات التعلم الآلي والذكاء الاصطناعي لتحليل النصوص والصور والمقاطع الصوتية والفيديوهات، وتحديد المحتوى غير المرغوب فيه بدقة عالية. يسعى البحث إلى تحقيق أداء متميز في اكتشاف وتصنيف المحتوى الضار، من خلال تدريب نماذج على مجموعات بيانات متنوعة، مما يساهم في تحسين تجربة المستخدم وتعزيز بيئة رقمية أكثر أماناً.

الكاشف الذكي هو نظام يعتمد على الذكاء الاصطناعي للكشف عن المحتوى غير المرغوب فيه على المواقع الإلكترونية لحماية المستخدمين من المحتويات الضارة وغير المرغوب فيها فكرياً وأخلاقياً، وضمان سلامة المعلومات.

ستتناول هذه الدراسة منهجية النظام الكاشف الذكي وآلية تنفيذها وكذلك يستعرض هذا البحث الفوائد العملية التي يحققها النظام الذكي الفعال. سيتم تقديم توصيات لتحسين أداء النظام وتوسيع نطاق تطبيقاته ليشمل منصات الإنترنت المختلفة، مثل وسائل التواصل الاجتماعي ومواقع التجارة الإلكترونية ومحركات البحث.

يهدف هذا البحث إلى تقديم مساهمة فعالة في مجال كشف المحتوى غير المرغوب فيه، من خلال تطوير نظام ذكي يعتمد على أحدث التقنيات في مجال التعلم الآلي والذكاء الاصطناعي. هذا النظام يقوم باكتشاف المحتويات الضارة على الفرد المسلم فكرياً أو أخلاقياً مثل المواقع الإباحية والأفكار الإلحادية والمتطرفة. كما يهدف النظام إلى تحسين جودة محتوى المواقع الإلكترونية من خلال إزالة المحتوى غير المرغوب فيه. وأخيراً، يهدف النظام المقترح إلى حماية المستخدمين وأفكارهم من التعرض لمحتوى غير لائق في الموقع الإلكتروني وإزالته.

نأمل أن يسهم هذا النظام في تحسين تجربة المستخدمين على الإنترنت، وجعل البيئة الرقمية أكثر أماناً وموثوقية. القسم التالي سوف يتناول الدراسات الحديثة التي صنفت محتويات رقمية لاكتشاف محتوى غير مرغوب فيه ، بينما سنعرض تباعاً منهجية النظام المقترح . سيتم لاحقاً عرض الفوائد وأهمية النظام الذكي في عدة مجالات مختلفة.

2. الأدبيات:

نستعرض بعض الأبحاث الحديثة التي استخدمت تقنيات ذكية لتصنيف محتويات رقمية لاكتشاف محتويات غير المرغوب فيها.

هو وآخرون 2023، اقترحوا طريقة ذكية تُدعى TTEC اعتمدت على تقنية BERT الاخبار المزيفة من محتويات نصية فقط. الطريقة المقترحة اثبتت كفاءتها في اكتشاف الاخبار المزيفة والتي تفوقت على احدث الاساليب بنسبة 3,1 % في مقاييس Mac. F1 [1].

يانج وآخرون 2023، اقترحوا طريقة تُدعى TGA اعتمدت على تقنية Transformers و multi-modal fusion وذلك للتعرف واكتشاف الاخبار الزائفة من المحتويات النصية والصور ايضاً. حيث أظهرت الطريقة المقترحة تفوقها على الاساليب الحديثة والتي مقارنتها مع هذه TGA [2].

يافو وآخرون 2023، بنوا طريقة حديثة تُدعى Testbed لاكتشاف التزييف العميق في محتويات نصية والذين اشاروا إلى أن اكتشاف التزييف العميق ممكنة في سيناريوهات العالم الحقيقي على الرغم من وجود بعض التحديات [3].

شو وآخرون 2024، بنوا تقنية ذكية تُدعى GAMC اعتماداً على Graph Autoencoder لاكتشاف الاخبار المزيفة للمحتوي النصي غير خاضعة للإشراف. الطريقة المقترحة تفوقت بشكل فعال في اكتشاف الأخبار المزيفة، مما يلغي الحاجة إلى مجموعات بيانات موسومة واسعة النطاق. ومع ذلك، تتطلب هذه الطريقة مستوى معيناً من الانتشار للكشف عن الأخبار الكاذبة [4].

رينا وآخرون 2024، قدموا طريقة جديدة لاكتشاف الأخبار الزائفة التي تشتمل على المعرفة الخلفية لمقالات مضللة. المؤلفين صمموا طريقتهم بناءً على Supervised Contrastive Learning المحتويات النصية والصور المزيفة. الطريقة المقترحة تفوقت على الأساليب SOTA الحديثة [5].

الجدول التالي يحوي الأبحاث السابقة والتي اكتشفت محتوى غير مرغوب فيه كلاً حسب الهدف العام لهذه الأبحاث، والموضحة على النحو التالي:

رقم البحث	الهدف من البحث	الطريقة	نوع البيانات	التاريخ
1	Multimodal fake news detection	TTEC	Text	2023
2	Multi-modal transformer for fake news detection	TGA	Text, Images	2023
3	Deepfake Text Detection	Testbed	Text	2023
4	An Unsupervised Method for Fake News Detection	GAMC	Text	2024
5	Identifying multimodal misinformation leveraging	SCL	Text, Images	2024

كما في الجدول الموضح أعلاه، فأنا نلاحظ ان كل الأبحاث صنفنا واكتشفت محتويات نصية في المواقع الإلكترونية، بينما هناك بحثان صنفنا صوراً لاكتشاف محتوى غير مرغوب. بينما جوان وآخرون قدموا مقترح إمكانية تصنيف واكتشاف محتويات مرئية أيضاً اعتماداً على تقنية الشبكات التوليدية العميقة [6].

3. منهجية النظام الذكي:

لبناء نظام الكاشف الذكي لابد من إجراء عدة خطوات رئيسية لإكمال بناء هذا النظام:

- جمع البيانات:

سيتم جمع مجموعة متنوعة من البيانات التي تشمل النصوص والصور والمقاطع الصوتية والفيديوهات من مصادر متعددة من المواقع الإلكترونية المراد تصنيفها واكتشاف محتواها الضار. سيتم التركيز على جمع بيانات تحتوي على محتوى غير مرغوب فيه مثل المواد الإباحية، والأفكار المتطرفة، والمعلومات المضللة. سيتم استخدام تقنيات جمع البيانات الآلية واليدوية لضمان تنوع وشمولية البيانات.

- التحليل:

بعد جمع البيانات، سيتم معالجتها وتنقيتها لإزالة الضوضاء والمعلومات غير الضرورية. ثم بعد ذلك نقوم بتحليل البيانات باستخدام تقنيات معالجة اللغة الطبيعية (NLP) لتحليل النصوص. كما يتم استخدام تقنيات معالجة الصور والفيديو لتحليل الوسائط المتعددة للبيانات المعالجة سلفاً. كما إنه سيتم تقسيم البيانات المُعالجة والمحللة إلى مجموعات تدريب واختبار لضمان دقة التقييم.

- التصنيف والكشف:

سيتم تطوير نظام الكاشف الذكي والذي يعتمد على تقنيات التعلم الآلي والذكاء الاصطناعي لتصنيف والكشف عن المحتوى الضار . حيث سيتم استخدام خوارزميات التعلم العميق مثل الشبكات العصبية التلافيفية (CNN) والشبكات العصبية المتكررة (RNN) لتحليل البيانات واكتشاف المحتوى غير المرغوب فيه. كما سيتم تدريب النموذج باستخدام مجموعة البيانات التي تم جمعها ومعالجتها.

- التقييم والتحسين:

سيتم تقييم وتحليل النتائج التي تم الحصول عليها من تقييم النموذج وتفسيرها. سيتم استعراض النقاط القوية والضعف في النموذج المقترح وتقديم توصيات لتحسين الأداء. سيتم أيضًا مناقشة التحديات التي واجهت البحث وكيفية التغلب عليها. في الختام، سيتم تلخيص النتائج الرئيسية للبحث وتقديم توصيات للتطبيقات المستقبلية. سيتم استعراض الإمكانيات المستقبلية لتطوير نظام كاشف ذكي أكثر فعالية وشمولية.

- التنفيذ والتطبيق:

لتنفيذ وتطبيق نظام كاشف ذكي عن المحتوى غير المرغوب في المواقع الإلكترونية، يمكن اتباع الخطوات التالية:

1. تحديد المتطلبات: لأجل متطلبات النظام، ينبغي علينا تحديد أهداف النظام وكذلك تحديد المتطلبات الفنية لها، على النحو التالي:
 - تحديد الأهداف: تحديد الأهداف الرئيسية للنظام، مثل نوع المحتوى غير المرغوب فيه الذي سيتم اكتشافه (مثل المواد الإباحية، المعلومات المضللة، الأفكار المتطرفة).
 - تحديد المتطلبات الفنية: تحديد المتطلبات الفنية للنظام، مثل القدرة على معالجة النصوص والصور والفيديوهات، وسرعة الاستجابة، والدقة المطلوبة.
2. جمع البيانات: نحتاج الى معرفة مصادر البيانات وبعد ذلك نقوم بتتقيتها، كما يلي:
 - (1) مصادر البيانات: جمع بيانات متنوعة من مصادر متعددة تشمل النصوص والصور والفيديوهات التي تحتوي على محتوى غير مرغوب فيه.
 - (2) تنقية البيانات: معالجة وتنقية البيانات لإزالة الضوضاء والمعلومات غير الضرورية وضمان جودة البيانات المستخدمة في التدريب.
3. تطوير النموذج: يتم تطوير نموذج الكاشف الذكي عن طريق اختيار الخوارزميات المناسبة وكذلك بتدريب النموذج لتحسين فعاليته وتطويره، مثل الشبكات العصبية التلافيفية (CNN) لتحليل الصور والفيديوهات، والشبكات العصبية المتكررة (RNN) لتحليل النصوص.

4. تدريب النموذج: تدريب النموذج باستخدام البيانات التي تم جمعها ومعالجتها، وضبط المعلمات لتحسين الأداء.

5. تقييم النموذج: لتقييم النظام المقترح، لابد من اختبار النموذج المبني وكذلك واستخدام بعض مقاييس التقييم، على النحو التالي:

- (1) اختبار النموذج: اختبار النموذج باستخدام مجموعة بيانات اختبارية لتقييم أدائه.
- (2) المقاييس: استخدام مقاييس مثل الدقة، والاسترجاع، والدقة التنبؤية لتقييم أداء النموذج.

6. نشر النظام: لابد من نشر النظام الذكي وذلك بعد إتباع الخطوات التالية:

- (1) البنية التحتية: إعداد البنية التحتية اللازمة لنشر النظام، مثل الخوادم وقواعد البيانات
- (2) التكامل: تكامل النظام مع المواقع الإلكترونية المستهدفة لضمان عمله بشكل سلس وفعال

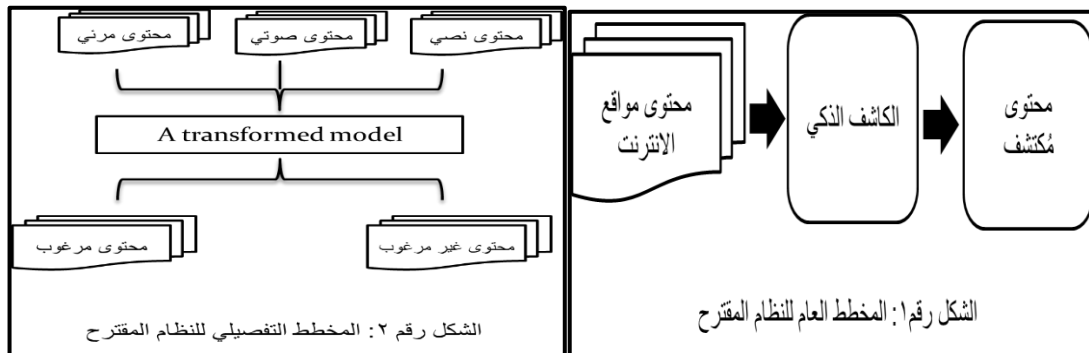
7. الصيانة والتحديث: ولصيانة النظام الذكي، علينا اتباع الاجراءات التالية:

- (1) مراقبة الأداء: مراقبة أداء النظام بشكل دوري لضمان استمرارية عمله بكفاءة
- (2) التحديثات: تحديث النظام بانتظام لتحسين الأداء والتكيف مع التغيرات في طبيعة المحتوى غير المرغوب فيه

8. التحديات والحلول: لكل نظام لابد من تحديات وحلول، للنظام الكاشف، كالتالي:

- (1) التعامل مع اللغات المختلفة: تطوير النظام ليكون قادرًا على التعامل مع محتوى بلغات مختلفة.
- (2) التكيف مع التغيرات: ضمان قدرة النظام على التكيف مع التغيرات المستمرة في طبيعة المحتوى غير المرغوب فيه.
- (3) خصوصية المستخدمين: ضمان حماية خصوصية المستخدمين وعدم انتهاكها أثناء عملية الكشف.

الشكل رقم 1، يوضح المخطط العام للنظام الكاشف الذكي، بينما الشكل رقم 2، يوضح المخطط التفصيلي للنظام المقترح:



أ- المكونات الرئيسية للنظام:

لتنفيذ نظام الكاشف الذكي عن المحتوى غير المرغوب فيه، هناك عدة مكونات رئيسية يجب تضمينها لضمان فعاليته وكفاءته. إليك المكونات الرئيسية للنظام:

- 1- واجهة المستخدم: واجهة المستخدم يتكون من لوحة التحكم وعرض النتائج:
 - (1) لوحة التحكم: واجهة تفاعلية تسمح للمستخدمين بإدارة النظام ومراقبة الأداء.
 - (2) عرض النتائج: واجهة تعرض النتائج المكتشفة وتصنيفات المحتوى غير المرغوب فيه.
- 2- نظام جمع البيانات: يتكون نظام جمع البيانات من جامع البيانات وتنقية البيانات، على النحو التالي:
 - (1) جامع البيانات: مكون يجمع البيانات من مصادر متعددة مثل النصوص، الصور، الفيديوهات، والمقاطع الصوتية.
 - (2) تنقية البيانات: مكون يقوم بتنقية البيانات وإزالة الضوضاء والمعلومات غير الضرورية.
- 3- نظام معالجة البيانات: ولنظام معالجة البيانات مهمتين أساسيتين، هما:
 - (1) معالجة النصوص: استخدام تقنيات معالجة اللغة الطبيعية. (NLP) لتحليل النصوص
 - (2) معالجة الصور والفيديو: استخدام تقنيات معالجة الصور والفيديو لتحليل الوسائط المتعددة
- 4- نظام التعلم الآلي: لإكمال نظام التعلم الآلي، لا بد من اجراء الاجراءين التاليين:
 - (1) نموذج التعلم العميق: استخدام خوارزميات التعلم العميق مثل الشبكات العصبية التلافيفية
 - (2) لتدريب النموذج. (CNN) والشبكات العصبية المتكررة
 - (3) تدريب النموذج: تدريب النموذج باستخدام البيانات التي تم جمعها ومعالجتها.
- 5- نظام التقييم: يتم اختبار وتقييم النظام حسب العمليتين التاليتين:
 - اختبار النموذج: اختبار النموذج باستخدام مجموعة بيانات اختبارية لتقييم أدائه.
 - مقاييس الأداء: استخدام مقاييس مثل الدقة، والاسترجاع، والدقة التنبؤية لتقييم أداء النموذج.
- 6- نظام التكامل: ينبغي تحقيق نظام التكامل في النظام المقترح للوصول الى مصادر مواقع البيانات مع اعداد البنية التحتية له، كما موضح ادناه:
 - (1) التكامل مع المواقع: تكامل النظام مع المواقع الإلكترونية المستهدفة لضمان عمله بشكل سلس وفعال.
 - (2) البنية التحتية: إعداد البنية التحتية اللازمة لنشر النظام، مثل الخوادم وقواعد البيانات.
- 7- نظام الصيانة والتحديث: لنظام الصيانة والتحديث مكونين رئيسيين، وهما:
 - (1) مراقبة الأداء: مراقبة أداء النظام بشكل دوري لضمان استمرارية عمله بكفاءة.
 - (2) التحديثات: تحديث النظام بانتظام لتحسين الأداء والتكيف مع التغيرات في طبيعة المحتوى غير المرغوب فيه.

ب- أمثلة توضيحية للنظام المقترح

في هذه الفقرة يمكننا استعراض بعض الامثلة لمحتويات غير مرغوب فيها التي تحمل افكارا إحادية او مواقع إباحية.

المثال الاول:

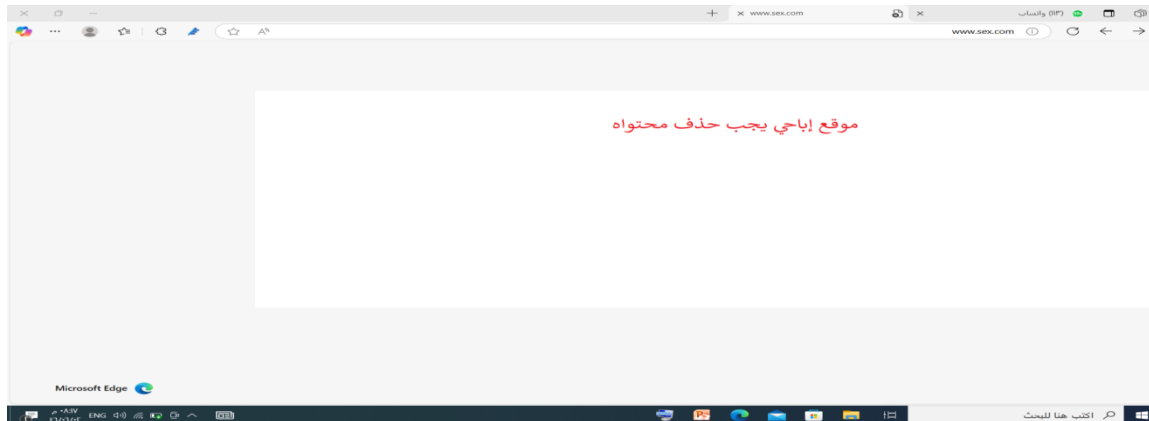
حيث نلاحظ في المثال نص إحدادي والتي تم اقتباسه من شبكة الإلحاد، والعياذ بالله.



الشكل (1) إكتشاف نص إحدادي

المثال الثاني:

وفي هذا المثال، موقع إباحي يحتوي على مقاطع وصور إباحية والنظام بدور يقوم بإغلاق وإخفاء محتواه تلقائياً بعد التعرف على فيديوهات وصور الموقع.



الشكل (2) إغلاق وإخفاء موقع إباحي

4. الفوائد العملية للنظام:

تطوير نظام كاشف ذكي عن المحتوى غير المرغوب فيه يمكن أن يقدم العديد من الفوائد العملية، منها:

1) تحسين تجربة المستخدم: يتم تحسين تجربة المستخدم من خلال الإجراءات التالية:

(2) تصفية المحتوى الضار: يساعد النظام في تصفية المحتوى الضار وغير المرغوب فيه، مما يجعل تجربة التصفح أكثر أمانًا وراحة للمستخدمين.

(3) تقديم محتوى ملائم: يساهم في تقديم محتوى ملائم وذو جودة عالية، مما يعزز من رضا المستخدمين.

(4) حماية الأطفال والمراهقين: يقوم النظام المقترح بحماية الأطفال والمراهقين وذلك من خلال الامور التالية:

- a. منع الوصول إلى المحتوى غير المناسب: يساعد النظام في منع الأطفال والمراهقين من الوصول إلى المحتوى غير المناسب، مثل المواد الإباحية والأفكار المتطرفة.
- b. تعزيز الأمان الرقمي: يساهم في تعزيز الأمان الرقمي للأطفال والمراهقين، مما يقلل من تعرضهم للمخاطر الإلكترونية.

(5) دعم الشركات والمؤسسات: كما ان النظام يحمي المستخدمين من الوصول الى محتوى ضار اخلاقياً وفكرياً، فإنه يحمي ايضا الشركات والمؤسسات وذلك عبر:

- a. تحسين سمعة العلامة التجارية: يساعد النظام الشركات والمؤسسات في الحفاظ على سمعة علامتها التجارية من خلال منع نشر المحتوى الضار على منصاتها.
- b. زيادة الثقة: يعزز من ثقة المستخدمين في المنصات الإلكترونية التي تستخدم النظام، مما يزيد من قاعدة العملاء والمستخدمين.

4. تحسين الأداء التقني: النظام الذكي يقوم بتحسين الأداء التقني للمواقع الإلكترونية وذلك باتباع الإجراءات التالية:

- a. تقليل الحمل على الخوادم: يساعد النظام في تقليل الحمل على الخوادم من خلال تصفية المحتوى غير المرغوب فيه قبل وصوله إلى المستخدمين.
- b. زيادة سرعة الاستجابة: يساهم في زيادة سرعة استجابة المواقع الإلكترونية من خلال تحسين إدارة المحتوى.

5. دعم الأبحاث والدراسات: النظام الكاشف الذكي يدعم الأبحاث والدراسات العلمية وذلك على النحو الآتي:

- a. توفير بيانات دقيقة: يساعد النظام الباحثين في الحصول على بيانات دقيقة وموثوقة حول المحتوى غير المرغوب فيه
- b. تطوير تقنيات جديدة: يساهم في تطوير تقنيات جديدة في مجال الذكاء الاصطناعي والتعلم الآلي لتحسين أداء النظام.

6- تعزيز الأمان والخصوصية: النظام الذكي يعزز الأمان والخصوصية باتباع الاجراءات الاتية:

- a. حماية البيانات الشخصية: يساعد النظام في حماية البيانات الشخصية للمستخدمين من خلال منع الوصول إلى المحتوى الضار.
- b. تقليل التهديدات الأمنية: يساهم في تقليل التهديدات الأمنية مثل البرمجيات الخبيثة والهجمات الإلكترونية.

5. النظرة المستقبلية للنظام:

- تتضمن النظرة المستقبلية لنظام الكاشف الذكي عن المحتوى غير المرغوب فيه العديد من التطورات والتحسينات التي يمكن أن تعزز من فعاليته وكفاءته. إليك بعض النقاط الرئيسية:
- (1) تحسين دقة الكاشف الذكي وذلك بتطبيق تقنيات أكثر ذكاءً.
 - (2) تطوير خوارزميات جديدة: يمكن تطوير خوارزميات أكثر تطوراً ودقة لتحليل المحتوى واكتشاف العناصر غير المرغوب فيها بشكل أفضل.
 - (3) التعلم المستمر: يمكن للنظام أن يتعلم باستمرار من البيانات الجديدة والتحديثات، مما يعزز من دقته في الكشف عن المحتوى الضار.
 - (4) التكيف مع اللغات والثقافات المختلفة.
 - (5) دعم لغات متعددة: يمكن توسيع نطاق النظام ليشمل دعم لغات متعددة، مما يجعله أكثر فعالية في الكشف عن المحتوى غير المرغوب فيه في مختلف الثقافات.
 - (6) التكيف مع السياقات الثقافية: يمكن تحسين النظام ليكون قادراً على فهم السياقات الثقافية المختلفة وتحديد المحتوى الضار بناءً على ذلك.
 - (7) التكامل مع تقنيات الذكاء الاصطناعي الأخرى.
 - (8) التكامل مع تقنيات التعلم العميق: يمكن دمج النظام مع تقنيات التعلم العميق لتحليل الصور والفيديوهات بشكل أكثر دقة.
 - (9) التكامل مع تقنيات معالجة اللغة الطبيعية (NLP): يمكن تحسين قدرة النظام على تحليل النصوص وفهمها بشكل أفضل من خلال استخدام تقنيات معالجة اللغة الطبيعية المتقدمة.
 - (10) تعزيز الأمان والخصوصية.
 - (11) حماية البيانات الشخصية: يمكن تطوير تقنيات جديدة لحماية البيانات الشخصية للمستخدمين وضمان عدم انتهاك خصوصيتهم أثناء عملية الكشف.
 - (12) التعامل مع التهديدات الأمنية: يمكن تحسين النظام ليكون قادراً على التعامل مع التهديدات الأمنية الجديدة والمتطورة.
 - (13) توسيع نطاق التطبيقات.
 - (14) استخدامات جديدة: يمكن توسيع نطاق استخدام النظام ليشمل تطبيقات جديدة مثل وسائل التواصل الاجتماعي، منصات التجارة الإلكترونية، ومحركات البحث.
 - (15) التكامل مع أنظمة الأخرى: يمكن دمج النظام مع أنظمة أخرى لتحسين أدائه وتوسيع نطاق تطبيقاته.

6. الختام

في هذا البحث، تم تطوير نظام كاشف ذكي يعتمد على تقنيات التعلم الآلي والذكاء الاصطناعي لكشف وتصنيف المحتوى غير المرغوب فيه على المواقع الإلكترونية. من خلال جمع ومعالجة بيانات متنوعة تشمل النصوص والصور والفيديوهات، تم تدريب النموذج على اكتشاف المحتوى الضار بدقة عالية. أظهرت النتائج أن النظام قادر على تحسين تجربة المستخدمين من خلال تصفية المحتوى غير المرغوب فيه وتقديم محتوى ملائم وآمن. على الرغم من التحديات التي سيواجهها تطوير النظام، مثل التعامل مع اللغات المختلفة والتكيف مع التغيرات المستمرة في طبيعة المحتوى غير المرغوب فيه، إلا أن النظام سيظهر قدرة عالية على التكيف والتحسين المستمر. تم تقديم توصيات لتحسين أداء النظام وتوسيع نطاق تطبيقاته ليشمل منصات الإنترنت المختلفة، مثل وسائل التواصل الاجتماعي ومواقع التجارة الإلكترونية ومحركات البحث.

في الختام، يمكن القول أن نظام الكاشف الذكي عن المحتوى غير المرغوب فيه يمثل خطوة هامة نحو تحسين الأمان الرقمي وتجربة المستخدمين على الإنترنت. من خلال تطوير تقنيات جديدة وتحسين الأداء، يمكن للنظام أن يساهم بشكل كبير في خلق بيئة رقمية أكثر أمانًا وموثوقية.

7. المراجع

- [1] Hua, J., Cui, X., Li, X., Tang, K., & Zhu, P. (2023). Multimodal fake news detection through data augmentation-based contrastive learning. *Applied Soft Computing*, 136, 110125.
- [2] Yang, P., Ma, J., Liu, Y., & Liu, M. (2023). Multi-modal transformer for fake news detection. *Mathematical Biosciences and Engineering*, 20(8), 14699-14717.
- [3] Li, Y., Li, Q., Cui, L., Bi, W., Wang, L., Yang, L., ... & Zhang, Y. (2023). Deepfake text detection in the wild. *arXiv preprint arXiv:2305.13242*.
- [4] Yin, S., Zhu, P., Wu, L., Gao, C., & Wang, Z. (2024, March). GAMC: an unsupervised method for fake news detection using graph autoencoder with masking. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 38, No. 1, pp. 347-355).
- [5] Kumari, R., Ashok, N., Agrawal, P. K., Ghosal, T., & Ekbal, A. (2023). Identifying multimodal misinformation leveraging novelty detection and emotion recognition. *Journal of Intelligent Information Systems*, 61(3), 673-694.
- [6] Cao, J., Qi, P., Sheng, Q., Yang, T., Guo, J., & Li, J. (2020). Exploring the role of visual content in fake news detection. *Disinformation, Misinformation, and Fake News in Social Media: Emerging Research Challenges and Opportunities*, 141-161.

8. الخلاصة باللغة الانجليزية

Intelligent detector of unwanted content on websites

Bassam Arkok¹, and Akram M. Zeki²

¹Hodeidah University, Yemen

² International Islamic University Malaysia, Malaysia

¹bassam_arkok@yahoo.com, ² akramzeki@iium.edu.my

Abstract

The widespread proliferation of websites on the internet has led to a frightening increase in the dissemination of harmful and unwanted content, such as pornography and destructive ideas like atheism and extremism. To address this issue, this research proposes an intelligent technique capable of identifying morally and intellectually harmful content and filtering it from websites. Our system leverages advanced machine learning techniques, including natural language processing and artificial intelligence, to analyze texts, images, audio clips, and videos. This research aims to achieve high accuracy and recall in detecting unwanted content, whether it be text, audio, visual content, or other undesirable elements, by training our model on various datasets of the relevant content. The proposed detector can be integrated into different online platforms, such as social media, e-commerce sites, and search engines, to enhance user experience and create a safer digital environment.

Keywords: Smart detector, unwanted content, websites, artificial intelligence