

تطوير نظام ذكاء اصطناعي قائم على التعلم العميق للكشف عن التضليل العلمي في المنصات الرقمية

م. شروق العواودة - جامعة الشرق الأوسط - عمان - الأردن
أ. آية الجعفري - جامعة الشرق الأوسط - عمان - الأردن

Developing a Deep Learning-Based Artificial Intelligence System for Detecting Scientific Misinformation on Digital Platforms

Eng. Shorouq Al-awawdeh - Middle East University, Amman - Jordan

Ms. Ayah Al-jafari - Middle East University, Amman - Jordan

تطوير نظام ذكاء اصطناعي قائم على التعلم العميق للكشف عن التضليل العلمي في المنصات الرقمية

م. شروق العواودة - جامعة الشرق الأوسط - عمان - الأردن

أ. آية الجعفري- جامعة الشرق الأوسط - عمان - الأردن

الملخص

الأهداف: تهدف الدراسة إلى تطوير وتقييم نظام للكشف عن التضليل العلمي باللغة العربية من خلال الضبط الدقيق لنموذج AraBERT-base-v2، وتحليل أثر الإيقاف المبكر وطول التسلسل النصي في الأداء، إلى جانب تفسير أبرز الخصائص اللغوية المرتبطة بالمحتوى المضلل وبيان حدود استخدام البيانات المترجمة.

المنهجية: اعتمدت الدراسة تصميمًا مختلطًا يوظف تكاملًا منهجيًا بين المنهجين الكمي والكيفي (النوعي)، مدعومًا بقراءة كيفية تفسيرية. تكوّنت القاعدة الأولية من 23,546 سجلًا، منها 123 مادة عربية جُمعت من منصات: أكيد، وشييك، وتيقن، و23,423 سجلًا أجنبيًا من مجموعتي GossipCop و PolitiFact ضمن FakeNewsNet. وبعد التنظيف، استُبقِي 73 مادة عربية للتحليل الوصفي والتفسيري، في حين تُرجمت البيانات الأجنبية آليًا إلى العربية باستخدام نموذج MarianMT (Helsinki-NLP/opus-mt-en-ar)، ثم نُظفت وطُبعت لغويًا. بلغ الحجم النهائي المستخدم في تجارب النموذج 21,595 سجلًا، قُسمت إلى 17,276 سجلًا للتدريب و4,319 سجلًا للتحقق. واعتمد التصنيف الثنائي إلى محتوى حقيقي ومضلل على الوسوم الأصلية للمصادر، ثم أُجري الضبط الدقيق لنموذج AraBERT-base-v2 عبر ثلاث تجارب: النموذج الأساسي، والإيقاف المبكر، وزيادة طول التسلسل إلى 512 رمزًا.

النتائج: أظهرت النتائج أن النموذج الأساسي حقق أعلى دقة (80.9%) في تصنيف الأخبار العلمية، إلا أنه عانى من فرط التكيف. في المقابل، أسهم تطبيق Early Stopping في تحسين التوازن بين الأداء والاستقرار بدقة بلغت 75.7% مع تعزيز القدرة على التعميم، أما زيادة طول التسلسل إلى 512 Token فقد أدت إلى انخفاض الدقة (73.1%) نتيجة إدخال معلومات غير ذات صلة، كما أظهرت النتائج تفوق النموذج في اكتشاف الأخبار العلمية المزيفة مقارنة بالحقيقية.

الخلاصة: توصلت هذه الدراسة إلى أن نموذج AraBERT يُعد أداة فعّالة في تصنيف الأخبار العلمية المضللة باللغة العربية، إلا أن تحقيق أداء مثالي يتطلب توازنًا دقيقًا بين

جودة البيانات واستراتيجيات التدريب المستخدمة. كما تؤكد النتائج أهمية توظيف تقنيات التنظيم، مثل Early Stopping؛ للحد من فرط التكيف وتعزيز قدرة النموذج على التعميم، في حين أن زيادة طول التسلسل النصي لا تضمن بالضرورة تحسناً في الأداء. وبناءً على ذلك، توصي الدراسة بالتركيز على تحسين جودة البيانات وتنوعها، إلى جانب الضبط الدقيق لمعلمات النموذج، بما يسهم في تطوير أنظمة أكثر كفاءة وموثوقية في البيئات التطبيقية الواقعية.

الكلمات المفتاحية: كشف الأخبار المضللة، AraBERT، التعلم العميق، معالجة اللغة الطبيعية، الإيقاف المبكر، نماذج Transformer، تصنيف النصوص.

Developing a Deep Learning-Based Artificial Intelligence System for Detecting Scientific Misinformation on Digital Platforms

Eng. Shorouq Al-awawdeh - Middle East University, Amman - Jordan

Ms. Ayah Al-jafari - Middle East University, Amman - Jordan

Abstract

Objectives: This study develops and evaluates an Arabic scientific misinformation detection system by fine-tuning AraBERT-base-v2. It examines the effects of early stopping and input sequence length on model performance, interprets selected linguistic characteristics associated with misleading content, and discusses the limitations of using machine-translated data.

Methodology: The study adopted a mixed-methods design, employing a systematic integration of quantitative and qualitative approaches, supported by an interpretive qualitative reading. The initial database consisted of 23,546 records, including 123 Arabic articles collected from the Akeed, Sheek, and Taqueen platforms, and 23,423 foreign-language records drawn from the GossipCop and PolitiFact collections within FakeNewsNet. After data cleaning, 73 Arabic articles were retained for descriptive and interpretive analysis, while the foreign-language data were machine-translated into Arabic using the MarianMT model (Helsinki-NLP/opus-mt-en-ar), and subsequently cleaned and linguistically normalized. The final dataset used in the model experiments comprised 21,595 records, divided into 17,276 records for training and 4,319 records for validation. The binary classification into genuine and misleading content was based on the sources' original labels. Fine-tuning of the AraBERT-base-v2 model was then carried out through three experiments: the baseline model, early stopping, and an increased sequence length of up to 512 tokens.

Results: The findings revealed that the baseline model achieved the highest classification accuracy (80.9%) in detecting scientific news;

however, it exhibited signs of overfitting. In contrast, applying the Early Stopping strategy improved the balance between performance and stability, achieving an accuracy of 75.7% while enhancing the model's generalization capability. Increasing the maximum sequence length to 512 tokens resulted in a lower accuracy (73.1%), likely due to the inclusion of irrelevant information. Furthermore, the model demonstrated superior performance in identifying fake scientific news compared with genuine scientific news.

Conclusion: The findings indicate that AraBERT is an effective tool for classifying Arabic scientific misinformation. However, achieving optimal performance requires a careful balance between data quality and the training strategies employed. The results also emphasize the importance of regularization techniques, such as Early Stopping, in mitigating overfitting and improving model generalization, whereas increasing the input sequence length does not necessarily lead to better performance. Accordingly, the study recommends prioritizing improvements in dataset quality and diversity, together with careful hyper parameters tuning, to support the development of more efficient and reliable misinformation detection systems for real-world applications.

Key words: Scientific misinformation detection; AraBERT; Deep Learning; Natural Language Processing (NLP); Early Stopping; Transformer Models; Text Classification.

المقدمة

يشهد العصر الرقمي تزايدًا كبيرًا في حجم المعلومات المتداولة عبر شبكة الإنترنت ووسائل الإعلام الرقمية، الأمر الذي أسهم في اتساع نطاق انتشار الأخبار المضللة والمعلومات غير الدقيقة، ولا سيما في المجالات العلمية والصحية، وقد أدى ذلك إلى بروز ما يُعرف بظاهرة التضليل المعلوماتي (Misinformation)، وهي ظاهرة قد تؤثر بشكل مباشر في فهم الجمهور للحقائق العلمية، كما قد تؤثر في القرارات المرتبطة بالصحة العامة والسياسات العلمية. وتشير العديد من الدراسات إلى أن انتشار المعلومات الخاطئة عبر الإنترنت أصبح تحديًا عالميًا متزايدًا، نظرًا لما يسببه من آثار سلبية في استقرار المجتمعات وفي مستوى الثقة بالمصادر العلمية والمؤسسات الإعلامية الموثوقة. (الصوالحة، والرجبي، 2025)

تُعد المؤسسات الإعلامية إحدى أهم القنوات التي يتم من خلالها نقل المعلومات العلمية إلى الجمهور، إلا أن سرعة نشر الأخبار والتنافس الإعلامي قد يؤديان أحيانًا إلى نشر معلومات غير دقيقة أو مُبالغ فيها، خاصة في المواضيع العلمية المعقدة. كما أسهمت منصات التواصل الاجتماعي في تسريع انتشار هذه الأخبار، حيث يمكن للمعلومات المضللة أن تنتشر بسرعة أكبر من المعلومات الصحيحة، مما يزيد من صعوبة اكتشافها أو تصحيحها في الوقت المناسب. (Tsikerdekis, & Zeadally, 2023)

في هذا السياق، برزت تقنيات الذكاء الاصطناعي، ولا سيما التعلم العميق (Deep Learning)، كأحد الحلول الواعدة لمواجهة مشكلة الأخبار المضللة، تعتمد هذه التقنيات على الشبكات العصبية العميقة القادرة على تحليل النصوص والأنماط اللغوية واستخراج السمات الدلالية المعقدة داخل البيانات، مما يسمح بتصنيف الأخبار وتمييز الأخبار الحقيقية من الأخبار المزيفة بدرجة دقة مرتفعة مقارنة بالأساليب التقليدية في التعلم الآلي. (Wang, 2023)

وقد أثبتت العديد من الدراسات أن نماذج معالجة اللغة الطبيعية المتقدمة (Natural Language Processing) مثل **Transformers** يمكنها تحقيق نتائج عالية في اكتشاف الأخبار المضللة، حيث تعتمد هذه النماذج على فهم السياق اللغوي للنصوص بدلاً من تحليل الكلمات بشكل منفصل، فعلى سبيل المثال، أظهرت بعض الأبحاث أن النماذج الهجينة التي تجمع بين BERT و LSTM تمكنت من تحقيق دقة تتجاوز 90% في الكشف عن المعلومات المضللة في النصوص الرقمية (Wang, 2025). كما أشارت دراسات أخرى إلى أن استخدام نماذج المحولات المتقدمة يمكن أن يحسن قدرة الأنظمة الذكية على تحليل الأخبار وتصنيفها ضمن عدة

فئات من المعلومات المضللة & (Hussna, Alam, Islam, Alkhamees, Hassan, & Uddin, 2024).

ومع ذلك، لا تزال هناك تحديات كبيرة تواجه تطوير أنظمة فعالة للكشف عن الأخبار العلمية المضللة، مثل تنوع مصادر المعلومات، وتعدد أشكال المحتوى الإعلامي، وصعوبة التحقق من صحة الادعاءات العلمية المعقدة. كما أن بعض النماذج قد تعاني من مشكلات مثل الإفراط في التكيف مع بيانات التدريب (Overfitting) أو الحاجة إلى مجموعات بيانات كبيرة ومتنوعة لتحقيق أداء موثوق.

انطلاقاً من ذلك، يهدف هذا البحث إلى تصميم نموذج ذكاء اصطناعي قائم على تقنيات الذكاء الاصطناعي لتحليل الأخبار العلمية المنشورة في المؤسسات الإعلامية ومواقع التواصل الاجتماعي، والكشف عن الأخبار المضللة فيها، ويسعى هذا النموذج مستقبلاً إلى تحسين دقة اكتشاف التضليل العلمي من خلال تحليل المحتوى النصي للأخبار واستخراج السمات الدلالية والسياقية المرتبطة بالمعلومات العلمية، مما يساهم في دعم المؤسسات الإعلامية والباحثين في الحد من انتشار المعلومات غير الدقيقة وتعزيز موثوقية المحتوى العلمي المنشور.

مشكلة الدراسة

تشهد البيئة الإعلامية المعاصرة تطوراً متسارعاً في وسائل نشر المعلومات، حيث أصبح تداول الأخبار عبر المنصات الرقمية (المواقع و الصحف الإلكترونية) ومواقع التواصل الاجتماعي يتم بسرعة كبيرة تفوق قدرة المؤسسات الإعلامية على التحقق الدقيق من المحتوى قبل نشره، وقد أدى هذا التسارع إلى زيادة انتشار الأخبار المضللة، بما في ذلك الأخبار العلمية التي تتعلق بقضايا حساسة مثل الصحة العامة، اللقاحات، التغير المناخي، والاكتشافات الطبية والسياسية؛ الأمر الذي قد يؤثر بشكل مباشر في وعي الجمهور وسلوكاته واتخاذ القرارات المبنية على معلومات غير دقيقة (Vosoughi, Roy, & Aral, 2018).

وتشير العديد من الدراسات إلى أن الأخبار المضللة تنتشر عبر الإنترنت بسرعة أكبر من الأخبار الصحيحة، خاصة عندما تتضمن موضوعات علمية أو صحية مثيرة للاهتمام العام، مما يزيد من احتمالية تضليل الجمهور وإضعاف الثقة بالمصادر العلمية والإعلامية. (Zhou & Zafarani, 2020) كما أن المؤسسات الإعلامية تواجه تحدياً كبيراً في التعامل مع الكم الهائل من المعلومات المتدفقة يومياً، حيث يصبح التدقيق اليدوي للأخبار عملية معقدة وتستهلك وقتاً وجهداً كبيرين، الأمر الذي يقلل من كفاءة عمليات التحقق التقليدية. (Shu et al., 2017)

وفي هذا السياق برزت تقنيات الذكاء الاصطناعي، وخاصة نماذج التعلم العميق، كأدوات واعدة لتحليل النصوص واكتشاف الأنماط اللغوية والدلالية المرتبطة بالأخبار المضللة، حيث أثبتت هذه النماذج قدرتها على معالجة كميات كبيرة من البيانات النصية والتعرف على الخصائص الخفية داخل المحتوى الإعلامي (Devlin et al., 2019). وعلى الرغم من تقدم نماذج التعلم العميق في تحليل النصوص، ما تزال تطبيقات الكشف عن التضليل العلمي باللغة العربية محدودة، كما تتأثر كفاءة النماذج بجودة البيانات، وعدم توازن الفئات، والترجمة الآلية، وطبيعة المجال الموضوعي، وفرط التكيف. ويستلزم ذلك تقييمًا واضحًا لأداء النموذج وحدوده بدلا من الاكتفاء بقيمة الدقة الكلية.

تحدد الدراسة موضوعيًا بالكشف النصي الثنائي عن التضليل المرتبط بالمحتوى العلمي والرقمي. ومصدرًا، تشمل منصات أكيد، وشييك، وتيقن، إضافة إلى مجموعتي GossipCop و PolitiFact المستخدمتين بوصفهما بيانات مساندة للتدريب. ولغويًا، تشمل محتوى عربيًا أصليًا ومحتوى إنجليزيًا مترجمًا آليًا إلى العربية. أما زمنيًا، فاعتمد جمع البيانات العربية على المواد المتاحة في المنصات حتى تاريخ 10 آذار 2026، في حين تتبع البيانات الأجنبية الفترات الزمنية المتاحة في مجموعات المصدر، وهو ما يمثل أحد حدود المقارنة.

وبناءً على ما سبق، تتمثل مشكلة الدراسة في الحاجة إلى تقييم فاعلية نموذج قائم على الضبط الدقيق لنموذج لغوي عربي مدرب مسبقًا في رصد الأخبار العلمية المضللة المنشورة عبر المنصات الرقمية والمؤسسات الإعلامية، وتصنيف المحتوى الإعلامي إلى محتوى حقيقي أو مضلل. كما تسعى الدراسة إلى بيان أثر كل من الإيقاف المبكر، وطول التسلسل النصي، وعدم توازن البيانات في أداء النموذج، إلى جانب تفسير الخصائص اللغوية والدلالية التي قد ترتبط بالمحتوى المضلل، بما يسهم في دعم عمليات التحقق من المعلومات العلمية وتعزيز موثوقية المحتوى الإعلامي.

وانطلاقًا من ذلك، تتمحور مشكلة الدراسة حول السؤال الرئيس الآتي:

ما مدى فاعلية النموذج القائم على الضبط الدقيق لنموذج AraBERT في رصد الأخبار العلمية المضللة وتصنيفها داخل المنصات الرقمية والمؤسسات الإعلامية، وما أثر إستراتيجيات تدريب النموذج وخصائص البيانات في أدائه؟

أهداف الدراسة

يتمثل الهدف الرئيس للدراسة في تطوير وتقييم نظام قائم على الضبط الدقيق لنموذج AraBERT-base-v2 للكشف عن المحتوى المضلل باللغة العربية، وبيان أثر خصائص البيانات واستراتيجيات التدريب في أدائه، وتحليل محتواها والكشف عن

المعلومات غير الدقيقة فيها، بما يسهم في تعزيز موثوقية المحتوى الإعلامي ودعم عمليات التحقق من المعلومات العلمية.

وينبثق عن هذا الهدف الرئيس مجموعة من الأهداف الفرعية تتمثل في:

1. تحديد السمات والخصائص اللغوية والدلالية التي تميز الأخبار العلمية المضللة عن الأخبار العلمية الصحيحة في المحتوى الإعلامي.
2. دراسة وتحديد تقنيات الذكاء الاصطناعي المناسبة لتحليل النصوص الإعلامية والكشف عن الأخبار العلمية المضللة.
3. تطوير نموذج قائم على الذكاء الاصطناعي قادر على تحليل الأخبار العلمية وتصنيفها إلى أخبار صحيحة وأخرى مضللة.
4. اختبار كفاءة النموذج المقترح وقياس مستوى أدائه في رصد الأخبار العلمية المضللة من خلال مؤشرات الدقة والموثوقية.
5. بيان دور النموذج المقترح في دعم المؤسسات الإعلامية لتحسين عمليات التحقق من المعلومات العلمية وتعزيز مصداقية المحتوى الإعلامي المنشور.

أهمية الدراسة

تنبع أهمية هذه الدراسة من تناولها قضية التضليل العلمي في المنصات الرقمية، في ظل الانتشار المتسارع للمعلومات غير الدقيقة وتأثيرها المحتمل في وعي الجمهور وقراراته. ويمكن تقسيم أهمية الدراسة إلى جانبين رئيسيين:

أولاً: الأهمية النظرية

- تسهم الدراسة في إثراء التراكم البحثي والأدبيات العربية التي تجمع بين الإعلام الرقمي والذكاء الاصطناعي ومعالجة اللغة الطبيعية والتعلم العميق، ولا سيما في مجال الكشف عن التضليل العلمي باللغة العربية.
- تسهم في توسيع نطاق البحوث المتعلقة باستخدام نماذج اللغة العربية المدربة مسبقاً في التحقق من المحتوى العلمي، في ظل محدودية الدراسات وقواعد البيانات العربية المتخصصة في هذا المجال.
- تقدم إطاراً علمياً يوضح أثر إعداد البيانات والمعالجة المسبقة واستراتيجيات التدريب، مثل الإيقاف المبكر وطول التسلسل النصي، في أداء نماذج تصنيف المحتوى المضلل.
- تسلط الضوء على عدد من القضايا المنهجية المؤثرة في نتائج نماذج الذكاء الاصطناعي، مثل عدم توازن فئات البيانات، ومحدودية المحتوى العلمي العربي الأصلي، والاعتماد على البيانات الأجنبية المترجمة، واحتمالية فرط التكيف.

- توفر أساسًا علميًا يمكن أن تستند إليه الدراسات المستقبلية في بناء قواعد بيانات عربية أكثر تنوعًا وتمثيلًا، وتطوير نماذج متعددة الفئات للكشف عن المحتوى المضلل كليًا أو جزئيًا والمحتوى غير المؤكد.

ثانيًا: الأهمية التطبيقية

- تقدم الدراسة نموذجًا أوليًا يمكن تطويره ليكون أداة مساندة للمؤسسات الإعلامية ومنصات التحقق في فرز المحتوى العلمي، وتحديد المواد التي تحتاج إلى مراجعة بشرية متخصصة.
- يمكن أن تسهم نتائج الدراسة في تحسين سرعة عمليات التحقق من المحتوى العلمي المنشور عبر المنصات الرقمية، والحد من تداول المعلومات غير الدقيقة، دون أن يحل النظام المقترح محل القرار المهني أو التقييم العلمي المتخصص.
- تساعد المؤسسات الإعلامية على توظيف تقنيات الذكاء الاصطناعي في دعم عمليات رصد المحتوى المضلل وتحليل النصوص، بما يعزز جودة المحتوى الإعلامي ومصداقيته.
- تقدم النتائج مؤشرات تطبيقية تتعلق بمخاطر فرط التكيف، وعدم توازن البيانات، وجودة الترجمة الآلية، وحدود زيادة طول المدخلات، بما يساعد الباحثين والمطورين على تصميم نماذج أكثر كفاءة وموثوقية.
- تسهم الدراسة في تعزيز الوعي العلمي لدى الجمهور ومساعدته على التمييز بين المعلومات العلمية الصحيحة والمضللة، من خلال دعم جهود المؤسسات الإعلامية ومنصات التحقق في تقديم محتوى أكثر دقة وموثوقية.
- تفتح المجال أمام تطوير تطبيقات مستقبلية للكشف المبكر عن التضليل العلمي، ودراسة تأثيره في الجمهور والمؤسسات الإعلامية، بما يدعم الاستخدام المسؤول للذكاء الاصطناعي في البيئة الإعلامية العربية.

الدراسات السابقة

1. دراسة (الروياتي والشوماني، ٢٠٢٦) بعنوان "الكشف عن الأخبار الزائفة باللغة العربية باستخدام نماذج التعلم العميق المحسنة بالخوارزمية الجينية". هدفت الدراسة إلى تصميم نموذج كشف آلي للأخبار الزائفة النصية باللغة العربية باستخدام التعلم العميق مدعوماً بالخوارزمية الجينية لتحسين أداء النموذج، إذ يعتمد النموذج على الشبكات العصبية والـ CNN ضمن إطار معالجة النصوص الطبيعية (NLP)، ويركز على استخراج الأنماط اللغوية الدالة على التضليل،

وأظهرت النتائج أن النموذج الهجين المحسّن بالخوارزمية الجينية تفوق على النماذج التقليدية مثل LSTM نموذج الذاكرة طويلة المدى، والنماذج الحديثة مثل نموذج المحول اللغوي العربي محققاً دقة كلية تجاوزت 96% ومعدل أداء متوازن F1-score مرتفع لفئة الأقلية (الأخبار الحقيقية).

2. دراسة (ميرزاه ورازمارا، ٢٠٢٥) بعنوان "نماذج التعلم العميق الهجينة للكشف عن الأخبار الكاذبة: دراسة حالة على اللغتين العربية والإنجليزية" هدفت هذه الدراسة إلى تطوير نموذج تعلم عميق هجين للكشف عن الأخبار المضللة في اللغتين العربية والإنجليزية، حيث اعتمدت الدراسة على دمج نموذج الشبكات الالتفافية (CNN) مع الشبكات المتكررة ثنائية الاتجاه (BiLSTM) لتحليل الخصائص الدلالية للنصوص الإخبارية، وأظهرت النتائج أن النموذج المقترح حقق دقة مرتفعة في تصنيف الأخبار المضللة، مما يدل على فعالية استخدام نماذج التعلم العميق الهجينة في تحليل الأخبار في بيئات لغوية مختلفة.

3. دراسة (إسلام وآخرون، ٢٠٢٠) بعنوان "التعلم العميق للكشف عن المعلومات المضللة على الشبكات الاجتماعية عبر الإنترنت: دراسة استقصائية ووجهات نظر جديدة" ركزت هذه الدراسة على مراجعة وتحليل الأبحاث المتعلقة باستخدام تقنيات التعلم العميق في الكشف عن المعلومات المضللة في شبكات التواصل الاجتماعي، واستعرضت الدراسة عدداً من النماذج المعتمدة على الشبكات العصبية مثل CNN و RNN و LSTM، وناقشت كيفية استخدام هذه النماذج لتحليل المحتوى النصي والسياق الاجتماعي للأخبار المنشورة. وأشارت نتائج الدراسة إلى أن تقنيات التعلم العميق توفر قدرة عالية على اكتشاف الأنماط الخفية في البيانات، مما يجعلها أدوات فعالة في الكشف المبكر عن الأخبار المضللة.

4. دراسة (الشويكي، ٢٠٢٤) بعنوان "الكشف عن الأخبار والمعلومات الكاذبة والمضللة باستخدام خوارزميات الذكاء الاصطناعي (نموذج تحقق نصي عربي)" قدمت هذه الدراسة تصميم أداة ذكاء اصطناعي عربية باسم صياد (Sayyad) لتحليل النصوص العربية على شبكة الإنترنت وتقييم مدى مصداقيتها. الأداة تعتمد على تقنيات التعلم الآلي والذكاء الاصطناعي لتحليل المحتوى وفق قواعد تقييم ومقارنة منهجية تستلهم من ممارسات التحقق الصحفي، وتبين قدرات المنصة على اكتشاف الأخبار الكاذبة، مع التأكيد على التطوير.

5. دراسة (الغامدي وآخرون، ٢٠٢٢) بعنوان "دراسة مقارنة لتقنيات التعلم الآلي والتعلم العميق للكشف عن الأخبار الكاذبة"، هدفت هذه الدراسة إلى مقارنة أداء

تقنيات التعلم الآلي والتعلم العميق في الكشف عن الأخبار المضللة. استخدمت الدراسة عدة خوارزميات مثل (SVM) و (Logistic Regression) و (Random Forest)، بالإضافة إلى نماذج التعلم العميق مثل (CNN) و (BiLSTM) و (BERT). تم اختبار هذه النماذج على عدد من قواعد البيانات الخاصة بالأخبار المضللة. وأظهرت النتائج أن نماذج التعلم العميق، خصوصاً النماذج المعتمدة على المحولات مثل BERT، حققت أداءً أعلى في تحليل النصوص الإخبارية والكشف عن الأخبار الزائفة مقارنة بالخوارزميات التقليدية.

6. دراسة (اليحيى، ٢٠٢١) بعنوان "كشف الأخبار الزائفة باللغة العربية: دراسة مقارنة

بين الشبكات العصبية والنماذج القائمة على المحولات" سعت هذه الدراسة إلى مقارنة أداء نماذج التعلم العميق التقليدية، مثل الشبكات العصبية، مع النماذج الحديثة المعتمدة على المحولات في مجال كشف الأخبار الزائفة باللغة العربية، حيث اعتمدت الدراسة على عدة مجموعات بيانات، وطبقت نماذج مختلفة لتصنيف الأخبار، وأظهرت النتائج تفوق النماذج القائمة على المحولات من حيث الدقة، نظرًا لقدرتها على فهم السياق اللغوي بشكل أعمق.

7. دراسة (الحسين وآخرون، ٢٠٢١) بعنوان "مكافحة وباء المعلومات المضللة حول

كوفيد-19 في تويتر العربي باستخدام نماذج لغوية قائمة على المحولات" سعت هذه الدراسة إلى الحد من انتشار المعلومات المضللة المرتبطة بجائحة كوفيد-19 في تويتر العربي باستخدام نماذج لغوية متقدمة، تم تدريب نماذج على بيانات متخصصة بالمحتوى الصحي، وتم تقييم أدائها في تصنيف الأخبار، وأظهرت النتائج فعالية هذه النماذج في كشف الأخبار العلمية المضللة، مع التأكيد على أهمية المعالجة المسبقة وجودة البيانات.

8. دراسة (انطوني وآخرون، ٢٠٢٠) بعنوان "نموذج قائم على المحولات لفهم اللغة

العربية" سعت هذه الدراسة إلى تطوير نموذج لغوي عربي يعتمد على بنية المحولات (Transformers) لتحسين فهم النصوص العربية في مهام المعالجة اللغوية الطبيعية. اعتمدت الدراسة على تدريب نموذج باستخدام بيانات عربية ضخمة، ثم تقييم أدائه في مجموعة من المهام مثل التصنيف وتحليل المشاعر، وأظهرت النتائج تفوق النموذج على العديد من الأساليب التقليدية، مما يجعله مناسباً لتطبيقات متعددة، من بينها كشف الأخبار المضللة في البيئة الرقمية العربية.

التعليق على الدراسات السابقة

توضح مراجعة الدراسات السابقة نمو الاهتمام باستخدام التعلم الآلي والتعلم العميق للكشف عن الأخبار الزائفة والمعلومات المضللة، وتبين أن النماذج القائمة على المحاولات تتفوق غالبًا في فهم السياق مقارنة بالأساليب المعتمدة على السمات السطحية وحدها. يمكن ملاحظة ما يلي:

- تتقاطع الدراسات في استخدام التصنيف النصي ومقاييس الأداء المعيارية، لكنها تختلف في اللغة والمجال ونوع البيانات وبنية النموذج. وقد ركز جزء كبير منها على الأخبار العامة أو السياسية أو الصحية الطارئة، ولم يتناول التضليل العلمي العربي بوصفه مجالًا واسعًا بصورة كافية.
- تُظهر الدراسات السابقة اهتمامًا متزايدًا بتوظيف تقنيات الذكاء الاصطناعي، وخاصة التعلم العميق، في مجال الكشف عن الأخبار المضللة، مع تركيز واضح على تحليل النصوص الإخبارية واستخلاص الأنماط اللغوية والدلالية المرتبطة بالتضليل.
- بالرغم من اختلاف أنواع النماذج (تقليدية، عميقة، هجينة، محسّنة بخوارزميات) إلا أن جميع الدراسات تؤكد فاعلية التعلم العميق وخوارزميات الذكاء الاصطناعي في كشف الأخبار المضللة.
- تتمثل الفجوة البحثية في محدودية الدراسات التي تجمع بين تكييف نموذج لغوي عربي مدرب مسبقًا لمهمة الكشف عن التضليل العلمي، وتحليل أثر الإيقاف المبكر وطول التسلسل، ومناقشة تأثير البيانات المترجمة وعدم توازن الفئات في موثوقية النتائج.
- تستجيب الدراسة لهذه الفجوة من خلال إعداد قاعدة بيانات متعددة المصادر، وإجراء الضبط الدقيق لنموذج AraBERT-base-v2، وتنفيذ ثلاث تجارب تدريبية، وتقديم قراءة تفسيرية للخصائص والقيود التي ظهرت في البيانات والنتائج.
- تأتي هذه الدراسة لسد هذه الفجوة من خلال تطوير نموذج قائم على تقنيات الذكاء الاصطناعي يركز بشكل خاص على رصد الأخبار العلمية المضللة في المنصات الرقمية، مع الاستفادة من نقاط القوة في الدراسات السابقة وتجاوز محدودياتها.
- ومع ذلك، لا تدّعي الدراسة أن البيانات الأجنبية المترجمة تمثل المحتوى العلمي العربي تمثيلًا كاملاً، بل تستخدمها بوصفها بيانات تدريب مساندة، وتتعامل مع

هذا الاختلاف باعتباره قيداً منهجياً يستلزم التحقق الخارجي مستقبلاً عبر قاعدة عربية علمية أصلية.

- تكتسب المنصات الرقمية في الدراسات السابقة أهمية خاصة بوصفها مصدراً ديناميكياً للبيانات، إذ تمثل البيئة الرئيسة لانتشار الأخبار المضللة والكاذبة، نتيجة لسرعة تداول المعلومات وطبيعة التفاعل المفتوح الذي يسهل إعادة نشر المحتوى دون التحقق من مصداقيته.

مصطلحات ومفاهيم الدراسة

- **الذكاء الاصطناعي (اصطلاحاً):** فرع من علوم الحاسوب يهدف إلى تصميم وتطوير أنظمة قادرة على محاكاة القدرات الذهنية البشرية، مثل التعلم، والاستدلال، وحل المشكلات، واتخاذ القرار، وفهم اللغة الطبيعية (Luger, 2005, p. 17).
- **الذكاء الاصطناعي (إجرائياً):** تقنية حديثة تمكّن الحاسوب من التفكير والتعلم بطريقة تشبه الإنسان، حيث يستطيع تحليل المعلومات وفهمها واتخاذ قرارات من خلال تحليل النصوص والتمييز بين المعلومات الصحيحة والخاطئة.

- **التعلم العميق (اصطلاحاً):** أسلوب متقدم من أساليب تعلم الآلة يعتمد على الشبكات العصبية العميقة لمعالجة البيانات واستخلاص الأنماط المعقدة منها (الحسبان وآخرون، 2026، ص 22).

- **التعلم العميق (إجرائياً):** قدرة النموذج المقترح على معالجة البيانات النصية، والتصنيف الدقيق للمحتوى، واكتشاف مؤشرات التضليل العلمي اعتماداً على التدريب المسبق للنموذج على مجموعات بيانات موثوقة.

- **الأخبار العلمية المضللة (اصطلاحاً):** معلومات مُلَفَّقة تُنشر عمدًا بهدف خداع المتلقي ودفعه إلى تصديق معطيات غير صحيحة أو التشكيك في حقائق يمكن التحقق منها، وتُصاغ بأسلوب يحاكي شكل المحتوى الإخباري المعتقد لإكسابها مظهرًا من المصداقية (الصادقي والأشعري، 2024، ص 1).

- **الأخبار العلمية المضللة (إجرائياً):** كل محتوى إعلامي يتناول موضوعات علمية ويحتوي على معلومات غير دقيقة أو غير موثوقة أو مُحرَّفة، سواء بشكل متعمد أو غير متعمد، بما يؤدي إلى تكوين فهم خاطئ لدى المتلقي.

- **معالجات اللغات الطبيعية (اصطلاحاً):** حقل بحثي ينتمي إلى الذكاء الاصطناعي وعلوم الحاسوب واللسانيات الحاسوبية، يُعنى بتمكين الحاسوب من معالجة اللغة البشرية وفهمها وتوليدها، بما يخدم مهامًا تطبيقية متعددة كالترجمة الآلية،

وتصنيف النصوص، والإجابة عن الأسئلة، والتلخيص الآلي (Martin, 2026, p. 20 & Jurafsky).

معالجات اللغات الطبيعية (إجرائيًا): مجموعة من التقنيات والخوارزميات المستخدمة لتمكين النموذج القائم على الذكاء الاصطناعي من فهم وتحليل النصوص الإعلامية المكتوبة باللغة الطبيعية، واستخلاص المعاني والأنماط اللغوية منها، بهدف الكشف عن الأخبار العلمية المضللة.

- **المنصات الرقمية (اصطلاحًا):** بنى تحتية رقمية وسيطة تُسهّل التفاعل والتبادل بين أطراف متعددة — كالمستخدمين ومقدمي الخدمة والمنتجين والمعلنين — عبر الإنترنت، وتتيح التواصل المباشر بين عدد كبير من المستخدمين، وتزداد قيمتها كلما اتسعت قاعدة مستخدميها، وتشمل أنماطًا متنوعة كمنصات التواصل الاجتماعي، ومنصات مشاركة المحتوى، ومنصات نشر الأخبار (Srnicek, 2016, p. 45).

المنصات الرقمية (إجرائيًا): الوسائط الإلكترونية والتطبيقات القائمة على الإنترنت التي تُستخدم لنشر وتداول المحتوى الإعلامي والعلمي، مثل مواقع الأخبار الإلكترونية ووسائل التواصل الاجتماعي، والتي تُعد بيئة أساسية لانتشار الأخبار العلمية الصحيحة والمضللة. حيث تم جمع بيانات الدراسة من منصات أكيد، وشييك، وتيقن، إضافة إلى مجموعتي GossipCop وPolitiFact ضمن FakeNewsNet Dataset.

نوع الدراسة ومنهجها

تنتمي هذه الدراسة إلى البحوث المختلطة، حيث توظف تكاملاً منهجياً بين المنهجين الكمي والكيفي (النوعي) ضمن تصميم تتابعي تفسيري، يتمثل المنهج الكمي في المرحلة الأولى من الدراسة التي تم فيها تحليل البيانات النصية واسعة النطاق بهدف تطوير نموذج تصنيفي قائم على تقنيات الذكاء الاصطناعي لرصد الأخبار العلمية المضللة، إذ استُخدم هذا المنهج لقدرته على معالجة كميات كبيرة من البيانات بصورة منهجية، واستخلاص الأنماط العامة بشكل موضوعي، إضافة إلى قياس أداء النموذج باستخدام مؤشرات إحصائية معيارية، الأمر الذي يعزز دقة النتائج وقابليتها للتعميم.

في المقابل، يظهر المنهج الكيفي في المرحلة الثانية من الدراسة من خلال الاعتماد على تحليل المحتوى لعينة مختارة من الأخبار، للكشف عن أنماط الخطاب وأساليب التضليل، وتفسير النتائج التي توصل إليها التحليل الكمي.

يتيح هذا التكامل تحقيق فهم أكثر شمولاً للظاهرة، من خلال الجمع بين الدقة الكمية والعمق التفسيري، بما يخدم تطوير نموذج أكثر كفاءة في التعامل مع المحتوى الإعلامي المضلل.

مجتمع الدراسة وعينتها

يتمثل مجتمع الدراسة في جميع الأخبار العلمية المنشورة عبر المنصات الرقمية، بما يشمل المواقع الإخبارية، والمنصات الإعلامية الإلكترونية، ومواقع التواصل الاجتماعي، على المستويات المحلية والإقليمية والدولية، ونظرًا لاتساع هذا المجتمع وتنوعه فقد تكوّنت القاعدة الأولية من 23,546 سجلًا: 123 مادة عربية، جمعت قصديًا من منصات أكيد، وشييك، وتيقن، و23,423 سجلًا أجنبيًا من GossipCop وPolitiFact. وبعد التنظيف، استُبقى 73 مادة عربية للتحليل الوصفي والتفسيري، بينما بلغ الحجم النهائي للبيانات المترجمة المستخدمة في تجارب النموذج 21,595 سجلًا. وتُعد البيانات الأجنبية المترجمة corpus تدريبيًا مساندًا ولا تمثل بديلًا كاملاً عن المحتوى العربي العلمي الأصلي.

وهذا يضمن تمثيل مختلف لأنماط المحتوى ومصادره، حيث حرصت الدراسة على اختيار عينة متنوعة من حيث التوزيع الجغرافي، وطبيعة المنصة، ومستوى المصادقية. أما في الجانب الكيفي، فقد اعتمد اختيار البيانات العربية على المعاينة القصدية لارتباطها بمنصات التحقق وبالموضوعات الصحية والعلمية والبيئية والتقنية. أما البيانات الأجنبية، فاستُخدمت لزيادة حجم التدريب وفق وسومها الأصلية. ولم تُستخدم عينة اختبار خارجية مستقلة؛ إذ قُسمت البيانات النهائية إلى 17,276 سجلًا للتدريب و4,319 سجلًا للتحقق، وذلك بهدف إخضاعها لتحليل معمق يكشف عن الخصائص الدلالية والسياقية لخطاب التضليل العلمي.

أداة الدراسة وجمع البيانات

تعتمد هذه الدراسة في جمع البيانات على بناء قاعدة بيانات (Dataset) متكاملة تخدم هدف تطوير نموذج قائم على تقنيات الذكاء الاصطناعي لرصد الأخبار العلمية المضللة، تم جمع البيانات على مرحلتين رئيسيتين لضمان التنوع والشمولية في المحتوى، وبما يعزز من دقة النموذج وقدرته على التعميم.

في المرحلة الأولى، تكوّن مجتمع الدراسة من الأخبار العلمية المنشورة عبر المنصات الرقمية العربية التي تُعنى برصد وتدقيق المحتوى الإعلامي، وتمثلت عينة الدراسة في مجموعة من الأخبار الصحفية العلمية التي جمعت خلال المرحلة الأولى من البحث من عدد من المنصات المتخصصة، وهي منصة "أكيد"، وموقع "شييك"، وموقع "تيقن".

وقد جرى اختيار هذه العينة بشكل قصدي نظراً لارتباطها المباشر بموضوع الدراسة، ولما توفره من محتوى متنوع يتضمن نماذج من الأخبار العلمية المضللة والصحيحة، أما فيما يتعلق بآلية جمع البيانات ومعالجتها، فقد تمت على النحو الآتي:

- منصة أكيد

تم جمع البيانات الصحفية العلمية من المنصة وفق البنية المعتمدة في عرض المحتوى، حيث كانت البيانات مصنفة ضمن مجموعة من الحقول، كما هو موضح في الشكل رقم (1)، والتي شملت على: رقم الصفحة، نوع الخبر، التصنيف، الرابط، العنوان، نص المقدمة، التاريخ، المعلومات، الصورة، تاليا صورة للتوضيح:

الشكل (1): نموذج البيانات منصة أكيد

رقم الصفحة	النوع	التصنيف	الرابط	العنوان	نص المقدمة	التاريخ	المعلومات	الصورة	النص الكامل
------------	-------	---------	--------	---------	------------	---------	-----------	--------	-------------

وقد تم اختيار عينة من الأخبار العلمية بناءً على تصنيفات المنصة، بحيث شملت الأخبار المصنفة ضمن "شائعات" بعدد (12) خبراً، والأخبار المصنفة ضمن "تحقق" بعدد (5) أخبار، بالإضافة إلى فئة "خبر/تقرير" بعدد (22) خبراً، ليلبلغ إجمالي عدد الأخبار المجمعة من المنصة (39) خبراً، تالياً الجدول رقم (1) للتوضيح:

الجدول (1): إجمالي عدد البيانات وتصنيفاتها من منصة أكيد

العدد	التصنيفات
١٢	إشاعات (أخبار الطب والصحة، المناخ، التكنولوجيا والابتكار العلمي، وأخبار سياسية)
٥	تحقق (أخبار الطب والصحة، المناخ، التكنولوجيا والابتكار العلمي، وأخبار سياسية)
٢٢	خبر/ تقرير (أخبار الطب والصحة، المناخ، التكنولوجيا والابتكار العلمي، وأخبار سياسية)
٣٩	المجموع

بعد الانتهاء من جمع البيانات، تم إدخالها إلى النموذج مع التركيز على مجموعة محددة من الحقول المستخرجة من التصنيفات المعتمدة في المنصات. حيث تم الاعتماد على كل من: العنوان، ونص المقدمة، والنص الكامل للخبر، في حين تم استبعاد عدد من الحقول الأخرى التي لا تخدم هدف النموذج بشكل مباشر، وهي: رقم الصفحة، ونوع الخبر، والرابط، والتصنيف، والصورة، والتاريخ، والمعلومات.

وبعد تهيئة البيانات وإدخالها، خضعت لعمليات معالجة (تنظيف)، شملت إزالة التكرار، واستبعاد الأخبار غير المكتملة أو التي تحتوي على بيانات ناقصة، بالإضافة إلى

حذف الخلايا الفارغة. ونتيجة لهذه العمليات، انخفض عدد الأخبار من (39) خبراً إلى (27) خبراً صالحاً للاستخدام في النموذج.

- موقع شيبك

تم اختيار عينة من الأخبار الصحفية العلمية من موقع "شيبك"، وذلك بالاعتماد على التصنيفات والمحاور المعتمدة في الموقع. وقد جُمعت البيانات وفق البنية التنظيمية لعرض المحتوى في المنصة، حيث تضمنت مجموعة من الحقول، شملت: رقم الصفحة، نوع الخبر، التصنيف، الرابط، العنوان، نص المقدمة، التاريخ، المعلومات، الصورة، والنص الكامل. الشكل رقم (2) يبين عرض توضيحي لشكل البيانات كما وردت في الموقع:

الشكل (2): نموذج البيانات موقع شيبك

رقم الصفحة	النوع	التصنيف	الرابط	العنوان	نص المقدمة	التاريخ	المعلومات	الصورة	النص الكامل
------------	-------	---------	--------	---------	------------	---------	-----------	--------	-------------

وعند إدخال البيانات إلى النموذج، تم التركيز على الحقول الأكثر ارتباطاً بمحتوى الخبر، وهي: العنوان، ونص المقدمة، والمعلومات، والنص الكامل، في حين تم استبعاد بقية الحقول لعدم ضرورتها في عملية التحليل.

وقد شملت العينة أخباراً علمية متنوعة تدرج ضمن عدد من مجالات الصحافة العلمية، إذ بلغ عدد الأخبار المجمعة (45) خبراً، وبعد إجراء عمليات التنظيف والمعالجة، بقي عدد الأخبار ثابتاً دون تغيير، نظراً لخلو البيانات من التكرار أو النقص. والجدول رقم (2) يبين عرض تفصيلي لبيانات العينة بعد المعالجة:

الجدول (2): نموذج البيانات موقع شيبك

العدد	نوع الأخبار العلمية
٢٣	الأخبار الصحية والطبية
١٢	الأبحاث والدراسات العلمية
١٠	الغذاء والتغذية والبيئة والتغير المناخي
٤٥	المجموع

- موقع تيقن

تم اختيار عينة من الأخبار الصحفية العلمية من موقع "تيقن"، حيث بلغت (39) خبراً، وذلك بالاعتماد على التصنيفات والمحاور المعتمدة في الموقع. وقد جُمعت البيانات وفق البنية التنظيمية المعتمدة لعرض المحتوى في المنصة، والتي تضمنت مجموعة

من الحقول، من أبرزها: الرابط، التصنيف، النتيجة، العنوان، اسم الكاتب، التاريخ، الادعاء، نص الخبر، وفقرة "تحقق يقين"، كما هو موضح الشكل رقم (3).

الشكل (3): نموذج البيانات موقع تيقن

الرابط	التصنيف	النتيجة	العنوان	الكاتب	التاريخ	الادعاء	الخبر	تحقق يقين
--------	---------	---------	---------	--------	---------	---------	-------	-----------

وعقب إجراء عمليات تنظيف ومعالجة البيانات، تم الاكتفاء بخبر واحد فقط من العينة، وذلك نتيجة وجود عدد من المخالفات التي لا تتوافق مع متطلبات النموذج، والتي تم تحديدها استناداً إلى معايير المعالجة والتنظيف المعتمدة سابقاً. وقد شملت هذه المخالفات عناصر تتعلق ببنية النص، ووجود محتوى غير ملائم أو غير مكتمل، بالإضافة إلى عدم توافق بعض البيانات مع محاور وتصنيفات الدراسة، مما استدعى استبعاد بقية الأخبار.

فيما يلي عرض تفصيلي لخطوات معالجة البيانات المطبقة على جميع المنصات الثلاثة، بهدف تحسين جودة البيانات وجعلها أكثر ملاءمة لاستخدامها في تدريب النموذج:

أولاً: إزالة البيانات غير المفيدة

تم حذف التواريخ، والأرقام الوصفية، والمقدمات الإدارية المتكررة التي لا تضيف قيمة تحليلية للنموذج، نظراً لعدم ارتباطها المباشر بمحتوى الخبر العلمي.

ثانياً: تنظيف النصوص (Text Cleaning)

تم إجراء عملية تنظيف مكثفة للنصوص، مع الحفاظ فقط على العناصر الأساسية، وهي: (الحروف العربية، الحروف الإنجليزية، الأرقام، المسافات) في المقابل، تم حذف الرموز الخاصة (مثل: #، %، @، -، ، ، "، /، وغيرها)، والعلامات الزخرفية، وجميع علامات الترقيم غير الضرورية (مثل: ، ، ، ؟، !، ، وغيرها)، بالإضافة إلى إزالة أرقام المراجع داخل الأقواس (مثل: [1]).

ثالثاً: إزالة المحتوى الثابت والمتكرر (Boilerplate Content)

تم التخلص من النصوص المتكررة والثابتة المرتبطة بالمنصات، والتي لا تعكس مضمون الخبر، وتشمل:

- عبارات الاشتراك والتنبيهات والإعلانات.
- عبارات التوجيه مثل: "اقرأ المزيد" و"تابعنا".
- أسماء المؤسسات والتواقيع الرسمية المتكررة.
- التواريخ الرقمية (مثل: 05-10-2023) والتواريخ المكتوبة نصياً.
- عبارات الإبلاغ عن المحتوى (مثل: أبلغ عن ادعاء).
- عبارات المتابعة على وسائل التواصل الاجتماعي.

- أسماء المواقع أو المنصات المكررة داخل النص.
- حقوق النشر وعبارات "جميع الحقوق محفوظة".
- المعلومات الجغرافية (مثل: الأردن – عمّان).
- عبارات إدخال البريد الإلكتروني والجمل التعريفية بالمؤسسات.

رابعاً: معالجة الاختصارات وتقليل الضوضاء

تم توسيع بعض الاختصارات الشائعة واستبدالها بصيغها الكاملة، مثل:

- "د." ← "دكتور"
- "أ.د." ← "أستاذ دكتور"
- "أ." ← "أستاذ"
- "م." ← "مهندس"

كما تم حذف الاختصارات الأجنبية غير المفيدة مثل RT, amp : ، والتخلص من الكلمات القصيرة جداً أو غير الدالة، بالإضافة إلى إزالة الضوضاء الناتجة عن الرموز أو الكلمات المفردة غير ذات المعنى.

خامساً: توحيد ومعالجة النصوص (Normalization)

بعد الانتهاء من مرحلة تنظيف البيانات، تأتي مرحلة التوطين، والتي تهدف إلى توحيد شكل النصوص ومعالجتها لغوياً، بما يساهم في تحسين جودتها وجعلها أكثر ملاءمة للتحليل باستخدام نموذج الذكاء الاصطناعي.

وقد ركزت هذه المرحلة على توحيد الحروف العربية، خصوصاً الحروف التي تأخذ أكثر من شكل، مثل الألف (ا) والواو (و) والتاء، وذلك من خلال تحويلها إلى صيغ موحدة، بما يقلل من التباين غير الضروري في النصوص.

كما شملت هذه المرحلة إزالة أقواس الجمل الاعتراضية، والرموز المصاحبة للنص، بهدف تبسيط البنية النصية وتحسين تمثيلها داخل النموذج، مما يساعد على رفع كفاءة المعالجة اللغوية. والجدول رقم (3) يوضح عمليات التوحيد والمعالجة التي تم تطبيقها على النصوص:

الجدول (3): عمليات تطبيع النص العربي باستخدام التعبيرات النمطية

#	العملية	النمط (Regex)	الاستبدال	الهدف من المعالجة
1	توحيد الألف	إ، أ، آ، ا	ا	إزالة اختلاف أشكال الألف لتوحيد الكلمات
2	توحيد الألف المقصورة	ى	ي	تحويل الألف المقصورة إلى ياء لتقليل التباين
3	توحيد الواو المهموزة	ؤ	و	إزالة الهمزة لتوحيد شكل الحرف
4	توحيد الياء المهموزة	ئ	ي	تبسيط الحرف وتحسين المطابقة النصية
5	تحويل التاء المربوطة	ة	ه	تقليل اختلافات نهاية الكلمات

أما في المرحلة الثانية، تم توسيع نطاق البيانات ليشمل مصادر دولية، حيث تم تجميع (23,423) خبر من الأخبار الأجنبية العلمية، بهدف إثراء قاعدة البيانات وتعزيز تنوعها من حيث اللغة والسياق والمصدر، ليسهم هذا التنوع في تحسين قدرة النموذج على التعامل مع أنماط مختلفة من الأخبار المضللة على المستوى العالمي، الصورة التالية تستعرض الشكل الأولي لقاعدة البيانات الأجنبية:

الشكل (4): نموذج أولى من بيانات FakeNews

id	news_url	title	tweet_ids
politifact14984	http://www.nfib-sbet.org/	National Federation of Independent Business	9671322598694871059671643687681966099672156186
politifact12944	http://www.cq.com/doc/newcomments	in Fayetteville NC	9429534598980098198162537173521668513250118459
politifact333	https://web.archive.org/web	Romney makes pitch, hoping to close deal : Elections : The Rocky Mountain News	
politifact4358	https://web.archive.org/web	Democratic Leaders Say House Democrats Are United Against GOP Default Act	
politifact779	https://web.archive.org/web	Budget of the United States Government, FY 2008	8980471037415424091270460595109888960396193003
politifact14064	http://www.politifact.com/t	Donald Trump exaggerates when he says China has 'total control' over North Korea	690248063990497286902540266638213126902789508
politifact14474	https://www.law.cornell.edu	25th Amendment	126260476210969740933118236439817507543393213
politifact5276	http://americanneedsmitt.com	هنگامی که می‌بینیم چه اتفاقی در آمریکا افتاده است،	
politifact1313	https://web.archive.org/web	Briefing by White House Press Secretary Robert Gibbs, 9/10/09	1351176226513512918230135138359001351424743113
politifact937	https://web.archive.org/web	A Solar Grand Plan: Scientific American	140962137332920320141057766704947200141669921
politifact8227	http://www.commonwealth	Covering Young Adults Under the Affordable Care Act: The Importance of Outreach and Expansion	

تم تجميع الأخبار الأجنبية من أربعة ملفات وبصيغة **CSV**. وبأعداد مختلفة، ولكي تصبح آلية التعامل مع البيانات أسهل وأدق تم دمج جميع البيانات بملف واحد وبنفس الصيغة ليتضمن الأعمدة التالية: `source`, `label`, `sector`, `date`, `sorce_lang`, `translation_type`, `origin_dataset`. يوضّح في الجدول رقم (4) عدد العينات التي

يحتوي عليها كل ملف، بالإضافة لصورة رقم (5) تبين الشكل المبدئي للعينة التي سيتم التعامل معها في المراحل التالية قبل استخدامها لتدريب النموذج.

الشكل (5): نموذج من بيانات FakeNews بعد دمج جميع بياناتها بملف واحد.csv.

source	text	label	sector	date	source_lang	origin_dataset	translation_type
Politifact	BREAKING	1			en	FakeNewsNet	
Politifact	Court Ord	1			en	FakeNewsNet	
Politifact	UPDATE: S	1			en	FakeNewsNet	
Politifact	Oscar Pist	1			en	FakeNewsNet	

الجدول (4): توزيع العينات تبعا لكل ملف من بيانات FakeNews

#	اسم الملف	عدد العينات
1	gossipcop_fake	5335
2	gossipcop_real	16818
3	politifact_fake	473
4	politifact_real	797

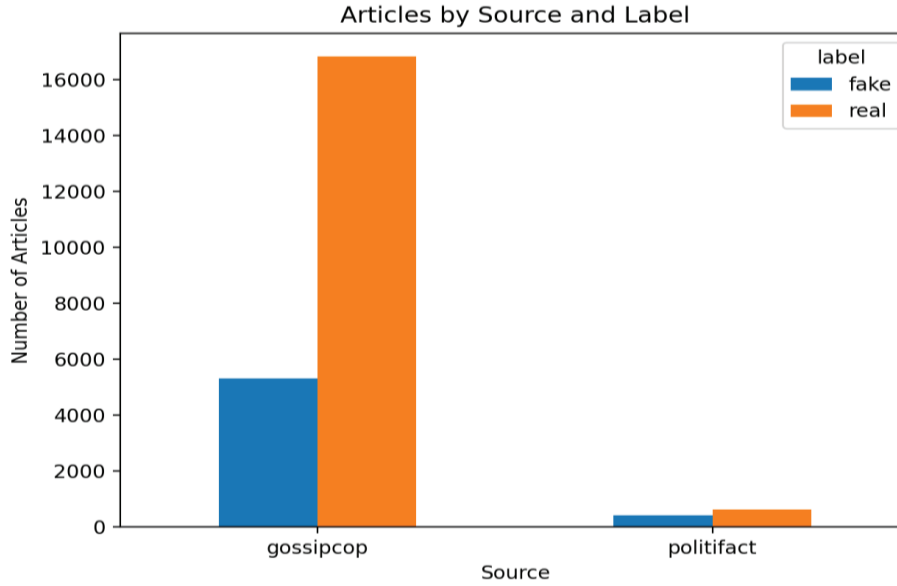
وحسب الإحصائية التي تمت لدراسة طبيعة البيانات الأجنبية، تبين أنها تتكوّن من مجموعة أخبار توزعت إلى 5,808 خبراً كاذباً و 17,615 خبراً حقيقياً، وأن معظم الأخبار تحتوي على روابط، حيث يصل عددها إلى 22,866 سجلاً، كما أن 21,695 خبراً منها يتضمن معرفات تغريدات (tweet_ids)، مما يشير إلى ارتباطها بمنصات التواصل الاجتماعي، أما من حيث خصائص العناوين، فيبلغ متوسط عدد كلمات العنوان 11.16 كلمة، مع وسيط قدره 11 كلمة. وبالنسبة للتفاعل، يبلغ متوسط عدد التغريدات المرتبطة بكل خبر 88.96 تغريدة، بينما يبلغ الوسيط 37 تغريدة، مما يعكس تفاوتاً في مستوى انتشار الأخبار وتداولها.

تتكوّن البيانات من مجموعة أخبار رقمية متنوعة تضم أخباراً حقيقية وأخرى كاذبة، كما تشمل هذه الأخبار أنماطاً مختلفة مثل الأخبار الصحفية العامة، والأخبار العلمية، بالإضافة إلى الأخبار السياسية، إلى جانب أخبار متداولة عبر منصات التواصل، ويعكس هذا التنوع استخدام البيانات لأغراض التحليل، وتصنيف الأخبار، ودراسة انتشارها، ومصداقيتها.

الشكل رقم (6) يعرض توزيع وخصائص البيانات بشكل بصري، وعليه فإن البيانات تنقسم إلى مجموعتين هما GossipCop و PolitiFact، حيث يتضح أن بيانات GossipCop أكبر بكثير من بيانات PolitiFact كما يوجد عدم توازن واضح بين الأخبار

الحقيقية والأخبار الكاذبة، وهو ما يُعد عاملاً مهماً يجب أخذه بعين الاعتبار قبل بناء أي نموذج تصنيف.

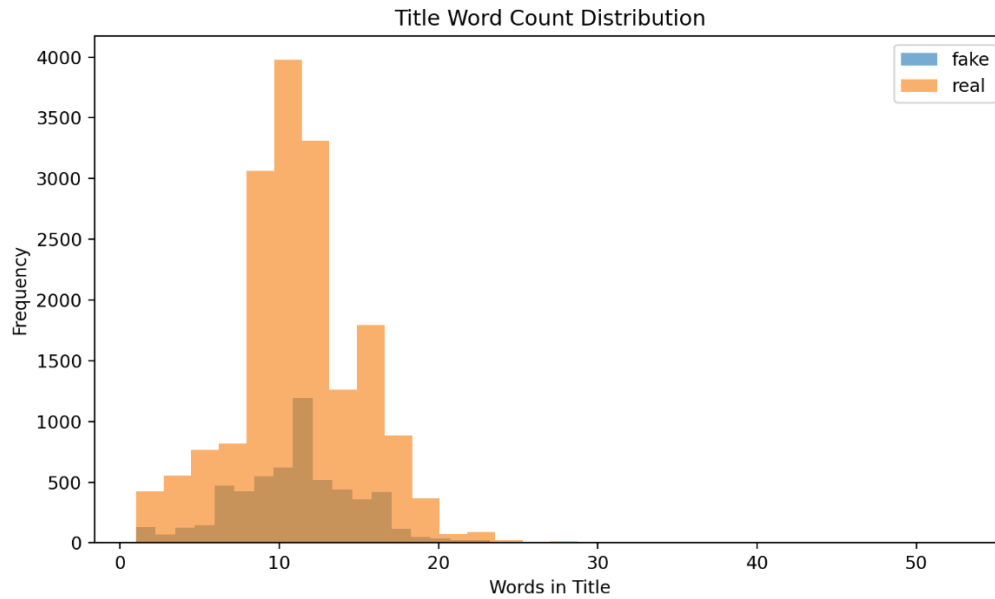
الشكل (6): نسب البيانات لكل صنف



وبعد إجراء تحليل شامل لخصائص البيانات، تبين وجود نقص في بعض السجلات، حيث تفتقر إلى حقول مهمة مثل `news_url` و `tweet_ids`، الأمر الذي استدعى تضمين خطوة ضمن الكود تهدف إلى تقييم جودة البيانات والتحقق من مدى اكتمالها. كما أظهرت النتائج وجود تباين ملحوظ في عدد التغريدات بين السجلات المختلفة والمجموعات، مما يشير إلى إمكانية اعتماد هذا المتغير كعامل مؤثر في عمليات التحليل أو بناء النماذج.

وتُظهر النتائج كذلك، من خلال الشكل رقم (7)، توزيع عدد الكلمات في عناوين الأخبار لكلٍ من الأخبار الحقيقية والمزيفة. حيث لوحظ أن عناوين الأخبار الحقيقية تميل إلى أن تكون أطول نسبياً، إذ يتركز معظمها ضمن نطاق يتراوح بين 8 و15 كلمة، مع انتشار أوسع وتكرار أعلى. في المقابل، تتسم عناوين الأخبار المزيفة بأنها أقصر وأكثر تركيزاً ضمن نطاق محدود من الكلمات، مع تكرار أقل. ويعكس هذا الاختلاف اعتماد الأخبار الحقيقية على عناوين أكثر تفصيلاً ووضوحاً، في حين تميل الأخبار المزيفة إلى استخدام عناوين مختصرة تستهدف جذب الانتباه بسرعة.

الشكل (7): تحليل توزيع طول العنوان (عدد الكلمات) لمقارنة الأخبار الحقيقية مقابل الأخبار المضللة



بما أن النموذج مخصص للغة العربية، ونظرًا لاعتماد بيانات المرحلة الثانية على مصادر إخبارية باللغة الإنجليزية، فقد جرى في هذه المرحلة إجراء عملية ترجمة لهذه البيانات إلى اللغة العربية وفق أسلوب علمي منظم، وذلك باستخدام نموذج الترجمة الآلية العصبية **MarianMT**، المتمثل في النسخة المدربة مسبقًا **Helsinki-NLP/opus-mt-en-ar** ضمن مشروع **OPUS-MT**، وهو مشروع مفتوح المصدر يهدف إلى إتاحة نماذج ترجمة آلية عصبية لعدد واسع من أزواج اللغات (Tiedemann & Thottingal, 2020). وتهدف هذه العملية إلى تهيئة البيانات للاستخدام في التجارب والتحليلات المرتبطة بتقييم وتشغيل النموذج، وفيما يلي عينة منها:

source	text	label
Politifact	أفريق كرة القدم الأولى يعلن الإفلاس على الركوع	1
Politifact	أوامر المحكمة الصادرة عن أوباما بدفع 400 مليون	1
Politifact	...ثانيًا، عميلة (روتر مور) تعمل لصالح (ميشيل أوب	1
Politifact	محاولتي الانتحار	1

بعد إجراء تحليل إحصائي لمحتوى البيانات، تم تحديد أكثر الكلمات تكرارًا في كل من الأخبار الحقيقية (Real) والمزيفة (Fake)، وذلك بهدف فهم خصائص النصوص والمساعدة في مواءمة البيانات مع متطلبات النموذج المستخدم في هذه التجربة.

وقد أظهرت النتائج أن بعض الكلمات الوظيفية مثل "في"، "من"، "على"، و"مع" تُعد من أكثر المفردات شيوعًا في كلا النوعين من الأخبار، مع ملاحظة ارتفاع معدل تكرارها في الأخبار المزيفة مقارنة بالحقيقية. كما تبين وجود كلمات أخرى تظهر

بشكل أكبر في الأخبار المزيفة مثل "إلى"، "عن"، "بعد"، و"أن"، في حين ظهرت بعض الرموز مثل (،) بتكرار متقارب في كلا النوعين. إضافة إلى ذلك، لوحظت مفردات مميزة لكل فئة، حيث ظهرت كلمات مثل "ما"، "هو"، و"التي" في الأخبار المزيفة فقط، بينما برزت كلمات مثل "هل"، "لم"، و"كارداشيان" في الأخبار الحقيقية.

كما كشفت النتائج عن الاستخدام المكثف للرموز وعلامات الترقيم في الأخبار المزيفة، مثل (#) والأقواس (،)، مقارنة بالأخبار الحقيقية التي احتوت على هذه العناصر بنسبة أقل، كما يوضح الجدول رقم (5). وبناءً على ذلك، تم تنفيذ مرحلة المعالجة المسبقة (Preprocessing)، والتي شملت إزالة الرموز الخاصة، وعلامات الترقيم، والروابط، والقيم الفارغة (NaN)، بالإضافة إلى تطبيق عملية توحيد الأحرف (Normalization) بعد الترجمة. وقد أسهمت هذه الإجراءات في تحسين جودة البيانات وجعلها أكثر ملاءمة لعملية النمذجة، مع ملاحظة انخفاض العدد الإجمالي للأخبار بعد التنظيف ليصل إلى 21,595 خبراً.

الجدول (5): تحليل تكرار الكلمات والرموز في الأخبار الحقيقية والمضللة

الكلمة	Fake التكرار	Real التكرار	ملاحظة
في	4882	1203	الأعلى في الاثنين
من	3121	1087	استخدام عام
على	2291	862	—
مع	1403	563	—
إلى	1197	311	أعلى في Fake
عن	1138	449	—
و	1061	374	—
بعد	1030	320	—
أن	957	414	—
لا	498	221	—
،	554	218	علامة ترقيم
-	736	282	علامة ترقيم
"	594	230	اقتباس
:	503	203	—

شبه متساوي	508	528	*
يظهر فقط في Fake	—	614	ما
يظهر فقط في Fake	—	537	هو
يظهر فقط في Fake	—	462	التي
هاشتاغ Fake أكثر	—	673	#
رموز Fake أكثر	—	2467	(
يظهر في Real	289	—	هل
يظهر في Real	188	—	لم
رموز في Real	~500	—	< >
اسم علم ((Real	227	—	كارداشيان

يوضح الشكل رقم (8) الشكل النهائي للبيانات بعد المعالجة، حيث يتضمن النصوص المترجمة إلى اللغة العربية بالإضافة إلى متغير الهدف (Label)، بحيث تشير القيمة (1) إلى الخبر المضلل، بينما تمثل القيمة (0) الخبر غير المضلل.

الشكل (8): نموذج من البيانات بعد عملية المعالجة

label	text_ar
1	فريق كرة القدم الأولى يعلن الإفلاس على الركوع
1	أوامر المحكمة الصادرة عن أولمبا بدفع 400 مليون
1	...ثلاثًا، صميلة (روجر مور) تعمل لصالح (ميشيل أوب
1	محاولتي الانتحار
1	الأمم المتحدة المصنوعة من أجل حقوق الإنسان للمثل

وبذلك، تمثل أداة الدراسة قاعدة بيانات متعددة المصادر والمراحل، تم إعدادها بطريقة منهجية تضمن التنوع والدقة، بما يدعم تطوير نموذج ذكاء اصطناعي قادر على رصد الأخبار المضللة بكفاءة عالية.

نتائج الدراسة

انقسمت نتائج الدراسة إلى أكثر من تجربة (مرحلة)، تم من خلالها تغيير بعض المعايير (العوامل) التي من الممكن أن تسهم في تحسين أداء النموذج وقدرته على تصنيف الأخبار العلمية المضللة (الحقيقية) أو غير المضللة (الكاذبة).

في التجربة الأولى من هذه الدراسة، تم الاعتماد على نموذج AraBERT base بوصفه أحد النماذج العميقة المتخصصة في معالجة اللغة العربية، وذلك بهدف تصنيف الأخبار إلى فئتين رئيسيتين: أخبار حقيقية وأخبار مزيفة. وقد سبقت عملية

التدريب مرحلة إعداد دقيقة للبيانات، حيث خضعت النصوص لسلسلة من خطوات المعالجة المسبقة لضمان توافقها مع متطلبات النموذج وتحسين جودة التمثيل النصي. شملت هذه الخطوات إزالة الرموز الخاصة، وعلامات الترقيم، والروابط، والقيم الفارغة، بالإضافة إلى تنفيذ عملية التطبيع اللغوي (Normalization) لتوحيد بعض الحروف العربية بعد الترجمة والمعالجة، الأمر الذي ساهم في تقليل التشويش النصي وتوحيد البنية اللغوية للمدخلات.

وبعد الانتهاء من عملية المعالجة، أصبح الحجم النهائي للبيانات 21,595 خبراً، ثم جرى تقسيمها إلى 17,276 عينة للتدريب و4,319 عينة للتحقق. وفي هذه التجربة تم ضبط الحد الأقصى لطول النص عند 128 رمزاً (Max_length = 128)، وهو اختيار يهدف إلى تحقيق توازن بين الاحتفاظ بالمعلومات الدلالية المهمة داخل النص وتقليل الكلفة الحسابية أثناء التدريب. بعد ذلك بدأ تدريب النموذج على مدى 6 عصور تدريبية (Epochs)، مع تتبع ثلاثة مؤشرات رئيسية للأداء هي: خسارة التدريب (Training Loss)، وخسارة التحقق (Validation Loss)، والدقة (Accuracy).

وقد أظهرت النتائج أن النموذج بدأ التعلم بكفاءة منذ العصر الأول Epoch، إذ انخفضت خسارة التدريب تدريجياً من قيمتها الابتدائية إلى قيم أقل في العصور اللاحقة، مما يدل على قدرة النموذج على التقاط الأنماط اللغوية داخل بيانات التدريب. كما ارتفعت الدقة تدريجياً من نحو 77.7% في البداية إلى ما يقارب 80.6% - 80.9% في أفضل مراحل التدريب، وهو ما يعكس تحسناً واضحاً في أداء النموذج على مهمة التصنيف. إلا أن تحليل منحنيات الأداء بين أيضاً أن خسارة التحقق بدأت بالارتفاع بعد العصور الأولى، على الرغم من استمرار انخفاض خسارة التدريب، وهي إشارة واضحة إلى بدء ظهور مشكلة فرط التكيف (Overfitting)، حيث أصبح النموذج أكثر قدرة على تمثيل بيانات التدريب نفسها، لكنه أقل كفاءة في التعميم على بيانات جديدة غير مرئية.

وتُظهر هذه النتيجة أن أفضل أداء فعلي للنموذج تحقق في العصور الوسطى، وبخاصة حول العصر الثالث، قبل أن تبدأ فجوة واضحة بالظهور بين أداء التدريب وأداء التحقق. ويمكن تفسير هذه الظاهرة بأن نموذج AraBERT، نظراً لقوته التمثيلية العالية، استطاع تعلم الأنماط الأساسية في البيانات العربية بفعالية، لكنه في الوقت نفسه أصبح أكثر حساسية للتفاصيل الخاصة ببيانات التدريب مع استمرار التدريب لعدد أكبر من العصور. وبناءً على ذلك، تشير نتائج هذه التجربة الأولى إلى أن AraBERT يُعد نموذجاً واعدًا وفعالاً في كشف الأخبار المزيفة باللغة العربية، إلا أن استخدامه يتطلب ضبطاً دقيقاً لعدد عصور التدريب والاستفادة من تقنيات مثل الإيقاف المبكر (Early Stopping) أو أساليب التنظيم الأخرى، من أجل الوصول إلى أفضل توازن ممكن بين

التعلم والدقة والقدرة على التعميم. وتمثل هذه التجربة نقطة انطلاق أساسية لبقية تجارب الدراسة، إذ توفر خطأ مرجعيًا مهمًا يمكن مقارنة أداء النماذج الأخرى به لاحقًا.

2. التجربة الثانية

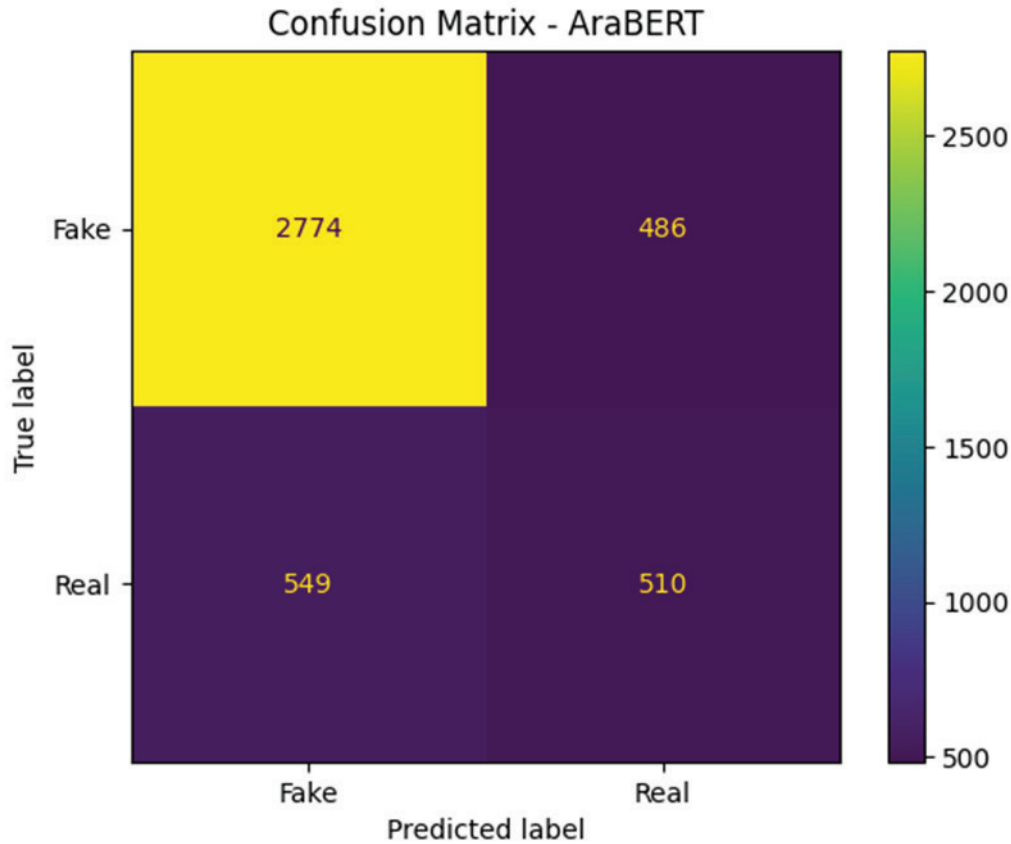
في هذه التجربة، تم تحسين أداء نموذج AraBERT من خلال تطبيق تقنية الإيقاف المبكر (Early Stopping)، والتي تُستخدم لإيقاف عملية التدريب عند النقطة التي يبدأ فيها النموذج بفقدان قدرته على التعميم على بيانات غير مرئية، وذلك بهدف الحد من مشكلة فرط التكيف (Overfitting).

في التجربة السابقة، استمر تدريب النموذج حتى العصر السادس (Epoch 6)، حيث لوحظ انخفاض تدريجي في قيمة دالة الخسارة على بيانات التدريب (Training Loss)، مما يشير إلى تحسن أداء النموذج على هذه البيانات. في المقابل، بدأت قيمة دالة الخسارة على بيانات التحقق (Validation Loss) بالارتفاع بعد عدد معين من العصور، وهو مؤشر واضح على تدهور قدرة النموذج على التعميم، واعتماده المتزايد على حفظ أنماط بيانات التدريب بدلاً من تعلم تمثيلات عامة.

لمعالجة هذه الإشكالية، تم تطبيق تقنية الإيقاف المبكر، حيث تم إنهاء التدريب عند العصر الثالث (Epoch 3)، وهي المرحلة التي حقق فيها النموذج أفضل توازن بين التعلم من بيانات التدريب والحفاظ على القدرة على التعميم. وقد انعكس ذلك على نتائج الأداء، حيث بلغت الدقة (Accuracy) نحو 75.7%، في حين سجلت مقاييس Precision و Recall و F1-score قيمًا تقريبية بلغت 50.4% و 58.1% و 54.0% على التوالي. أظهر تحليل النتائج أن النموذج أصبح أكثر استقرارًا بعد تطبيق الإيقاف المبكر، إذ لم يعد هناك ارتفاع ملحوظ في قيمة Validation Loss، مما يدل على تقليل ظاهرة فرط التكيف. كما تحسنت قدرة النموذج على التعميم، وأصبح أكثر كفاءة في التعامل مع بيانات غير مرئية مقارنة بالتجربة السابقة.

وبالاعتماد على مصفوفة الالتباس (Confusion Matrix)، كما في الشكل رقم (9)، تم إجراء تحليل تفصيلي لأداء النموذج في مهمة التصنيف، حيث تمكّن من تصنيف 2774 حالة من الأخبار المزيفة بشكل صحيح، و 510 حالات من الأخبار الحقيقية بشكل صحيح، في حين بلغ عدد حالات الخطأ 1035 حالة (486 + 549). وتشير هذه النتائج إلى أن النموذج يمتلك قدرة أعلى على اكتشاف الأخبار المزيفة مقارنة بالأخبار الحقيقية، وهو ما قد يُعزى إلى عدم توازن البيانات أو إلى انحياز النموذج نحو الفئة الأكثر تمثيلًا في مجموعة التدريب.

الشكل (9): مصفوفة الالتباس لتقييم أداء نموذج AraBERT في التمييز بين الأخبار الحقيقية والمضللة



بشكل عام، تُشير النتائج إلى أن تطبيق تقنية الإيقاف المبكر (Early Stopping) كان له تأثير إيجابي ملحوظ على أداء نموذج AraBERT، حيث أسهم في الحد من ظاهرة فرط التكيف (Overfitting) من خلال إيقاف عملية التدريب عند النقطة المثلى التي يتحقق فيها أفضل توازن بين أداء النموذج على بيانات التدريب وبيانات التحقق. وقد أدى ذلك إلى تحسين استقرار النموذج وزيادة موثوقيته عند تعميمه على بيانات غير مرئية.

في التجربة الثالثة، تم تحليل تأثير زيادة طول النصوص المدخلة من خلال رفع قيمة المتغير Max_length إلى 512، بهدف تمكين النموذج من الاستفادة من سياق نصي أطول قد يتضمن معلومات إضافية داعمة لعملية التصنيف. إلا أن النتائج أظهرت تراجعاً في الأداء، حيث انخفضت قيمة الدقة (Accuracy) إلى نحو 73.1% مقارنة بالتجارب السابقة، بالإضافة إلى انخفاض ملحوظ في مؤشري Precision و F1-score.

يمكن تفسير هذا التراجع بأن زيادة طول النص أدت إلى إدخال قدر أكبر من المعلومات غير ذات الصلة (Noise)، مما تسبب في تشتيت انتباه النموذج وإضعاف قدرته على التركيز على السمات الجوهرية للنصوص. وعليه، تُبرز هذه النتائج أن زيادة

طول المدخلات لا تُعد بالضرورة عاملاً محسناً للأداء، بل قد تؤدي إلى نتائج عكسية في حال عدم وجود حاجة فعلية لهذا الامتداد في السياق النصي. يوضح الجدول رقم (6) مقارنة بين ثلاث تجارب مختلفة باستخدام نموذج AraBERT، حيث تم تقييم تأثير كل من عدد العصور (Epochs) وطول النص (Max_length) على أداء النموذج.

الجدول (6): مقارنة التجارب وفق معايير معينة

المعيار	التجربة الأولى	التجربة الثانية	التجربة الثالثة
Max_length	128	128	512
Epochs	6	3	3
Accuracy	~80.9%	~75.7%	~73.1%
Precision	—	~50.4%	~45.9%
Recall	—	~58.1%	~55.5%
F1-score	—	~54.0%	~50.3%
Overfitting	مرتفع	منخفض	متوسط
التعميم	ضعيف	جيد	متوسط

بشكل عام، تُشير النتائج إلى أن أعلى قيمة للدقة (Accuracy) قد تحققت في التجربة الأولى، مما يعكس قدرة النموذج على تحقيق أداء مرتفع على مستوى التصنيف الكلي. في المقابل، تُعد التجربة الثانية الأكثر كفاءة من حيث التوازن بين المقاييس المختلفة، حيث أظهرت تحسناً في القدرة على التعميم نتيجة تطبيق تقنية الإيقاف المبكر (Early Stopping)، مما أسهم في تقليل ظاهرة فرط التكيف وتحقيق أداء أكثر استقراراً على بيانات غير مرئية.

أما في التجربة الثالثة، فقد تبين أن زيادة طول النص (Max_length = 512) لم تؤدي إلى تحسين الأداء، بل أسفرت عن تراجع في بعض المقاييس، الأمر الذي يشير إلى أن إدخال سياق نصي أطول قد يضيف معلومات غير ذات صلة تؤثر سلباً على كفاءة النموذج. وبناءً على ذلك، يتضح أن تحسين الأداء لا يعتمد بالضرورة على زيادة حجم المدخلات، بل على جودة المعلومات ومدى ملاءمتها لمهمة التصنيف.

النتائج العامة للدراسة

- أظهر الضبط الدقيق لنموذج AraBERT-base-v2 قدرة على تعلم أنماط التصنيف في البيانات محل الدراسة، ووصلت أعلى دقة تحقق إلى نحو 80.9%. غير أن هذه

- النتيجة ترتبط ببيانات يغلب عليها المحتوى الأجنبي المترجم ومتعدد الموضوعات، ولذلك لا تكفي وحدها لإثبات فاعلية مماثلة على المحتوى العلمي العربي الأصلي.
- أيضاً، يعكس كفاءة النموذج في تحليل المحتوى الإعلامي واكتشاف المعلومات غير الدقيقة، إذ تتفق النتائج مع دراسة (اليحيى، ٢٠٢١) التي أكدت تفوق النماذج القائمة على المحولات في كشف الأخبار المضللة، وهو ما انعكس في الأداء الجيد لنموذج AraBERT في هذه الدراسة.
 - أسهمت عملية تحليل البيانات في تحقيق الهدف المتعلق بتحديد الخصائص اللغوية والدلالية، حيث تبين أن الأخبار العلمية المضللة تميل إلى استخدام عناوين أقصر، مع تكرار أعلى لبعض الكلمات الوظيفية والرموز، في حين تعتمد الأخبار العلمية الحقيقية على عناوين أكثر تفصيلاً وتنظيماً. وتتفق هذه النتيجة مع ما أشارت إليه دراسة (إسلام وآخرون، ٢٠٢٠) التي أكدت أن نماذج التعلم العميق قادرة على اكتشاف الأنماط اللغوية الخفية في النصوص، كما تنسجم مع دراسة (ميرزاه ورازمارا، 2025) التي أوضحت أهمية الخصائص الدلالية في تحسين دقة تصنيف الأخبار، مما يعزز دور التحليل اللغوي في التمييز بين الأخبار العلمية المضللة والصحيحة.
 - أكدت النتائج تحقق الهدف المرتبط بتحديد تقنيات الذكاء الاصطناعي المناسبة، إذ أظهرت النماذج القائمة على المحولات، وخاصة AraBERT، قدرة عالية على فهم السياق اللغوي العربي وتحليل النصوص العلمية بدقة، كما تنسجم مع دراسة (الغامدي وآخرون، ٢٠٢٢) التي أظهرت أن نماذج التعلم العميق، وخاصة نماذج BERT، تحقق أداءً أعلى مقارنة بالخوارزميات التقليدية في تحليل النصوص الإخبارية.
 - تحقق الهدف المتعلق بتطوير نموذج تصنيفي، حيث تم بناء نموذج قادر على التمييز بين الأخبار العلمية المضللة والأخبار الصحيحة، مع قدرة واضحة على التعلم التدريجي واكتشاف الأنماط داخل البيانات، تدعم هذه النتائج الأساس النظري لدراسة (انطوني وآخرون، ٢٠٢٠) التي أثبتت كفاءة نموذج AraBERT في فهم اللغة العربية، وهو ما يفسر نجاحه في تحليل وتصنيف الأخبار العلمية المضللة في هذه الدراسة.
 - أظهرت نتائج التقييم تحقيق الهدف المرتبط بقياس كفاءة النموذج، حيث تم استخدام مؤشرات مثل الدقة (Accuracy) وخسارة التدريب والتحقق، مع ملاحظة تحسن الأداء خلال مراحل التدريب الأولى كما تتقاطع مع دراسة (ميرزاه ورازمارا، ٢٠٢٥) التي أكدت فعالية النماذج الهجينة، مما يفتح المجال لتطوير النموذج الحالي مستقبلاً.
 - كشفت النتائج عن ظهور مشكلة فرط التكيف (Overfitting)، مما يبرز أهمية استخدام تقنيات تنظيم مثل Early Stopping وضرورة الضبط الدقيق لمعاملات النموذج لتحقيق أفضل أداء. وتتفق هذه النتيجة مع ما أشارت إليه دراسة (الغامدي

وآخرون، ٢٠٢٢) ودراسة (إسلام وآخرون، ٢٠٢٠) حول أهمية ضبط عملية التدريب واستخدام أساليب تنظيم مناسبة لضمان قدرة النماذج على التعميم وعدم الاكتفاء بحفظ بيانات التدريب.

- أكدت النتائج إمكانية توظيف النموذج في دعم الهدف التطبيقي للدراسة، والمتمثل في تعزيز دور المؤسسات الإعلامية في التحقق من الأخبار العلمية والحد من انتشار المعلومات المضللة على المنصات الرقمية. ويتوافق ذلك مع ما توصلت إليه دراسة (إسلام وآخرون، ٢٠٢٠) التي أثبتت فعالية النماذج القائمة على المحولات في مواجهة المعلومات العلمية المضللة، وكذلك دراسة (الشويكي، ٢٠٢٤) التي أبرزت أهمية توظيف أدوات الذكاء الاصطناعي في دعم عمليات التحقق وتحسين مصداقية المحتوى الإعلامي.

التوصيات

- توصي الدراسة بتحسين جودة البيانات المستخدمة، من خلال جمع بيانات أكثر تنوعاً ودقة، ومعالجة مشكلة عدم التوازن بين فئات الأخبار، إلى جانب تحسين جودة الترجمة لضمان الحفاظ على المعنى وتقليل الضوضاء اللغوية، خاصة في مجال الأخبار العلمية.
- التأكيد على أهمية تطبيق المعالجة المسبقة (Preprocessing) بشكل منهجي، بما يشمل تنظيف النصوص وتوحيدها، لما لذلك من دور في تحسين تمثيل البيانات ورفع كفاءة النموذج.
- ضرورة استخدام تقنيات تنظيم مثل Early Stopping، إلى جانب الضبط الدقيق لمعاملات التدريب (Hyperparameters)، للحد من مشكلة فرط التكيف وتحسين قدرة النموذج على التعميم.
- يوصى بإجراء تجارب إضافية على أطوال مختلفة للنصوص (Max_length)، حيث لا تؤدي زيادة طول التسلسل بالضرورة إلى تحسين الأداء.
- التوسع في استخدام النماذج الهجينة، وإجراء مقارنات مع نماذج لغوية أخرى مثل CAMELBERT و mBERT، بهدف تحسين دقة التصنيف وتعزيز موثوقية النتائج.
- دمج خصائص إضافية غير نصية، مثل بيانات التفاعل على المنصات الرقمية، لتعزيز قدرة النموذج على كشف الأخبار العلمية المضللة بشكل أكثر دقة.
- توظيف النموذج المطور في هذه الدراسة ضمن المؤسسات الإعلامية ومنصات التواصل الاجتماعي، بحيث يُستخدم كأداة مساندة في عمليات التحقق

- من صحة المعلومات، مما يسهم في الحد من انتشار الأخبار المضللة وتعزيز مصداقية المحتوى المتداول لدى الجمهور.
- تطوير قواعد بيانات عربية علمية أصلية ومراجعة وسومها من متخصصين، ثم توسيع التصنيف مستقبلاً ليشمل: مضلل كلياً، ومضلل جزئياً، وصحيح جزئياً، وصحيح، وغير مؤكد، مع إجراء اختبار خارجي مستقل ومقارنة AraBERT بنماذج مثل mBERT و CAMELBERT.
 - تشجيع تطوير نماذج متخصصة في كشف الأخبار العلمية المضللة، مع التوسع مستقبلاً في دراسات تربط بين الخصائص اللغوية للمحتوى وسلوك المستخدمين في البيئة الرقمية، بما يسهم في فهم أعمق لآليات انتشار المعلومات المضللة وتفاعل الجمهور معها.
 - توسيع نطاق الدراسة ليشمل أنواعاً أخرى من البيانات، وفي مقدمتها البيانات المرئية (الصور)، نظراً لانتشارها الواسع بوصفها أحد أبرز أشكال المحتوى المُستخدم في نشر الأخبار المضللة عبر منصات التواصل الاجتماعي حالياً. ويقتضي ذلك توظيف نماذج متخصصة في معالجة الصور وتحليلها (كنماذج الرؤية الحاسوبية أو النماذج متعددة الوسائط)، بما يسهم في تطوير أنظمة كشف أكثر شمولاً وقدرة على التعامل مع المحتوى المضلل بمختلف صوره النصية والمرئية.

المراجع

- الشويكي، ل. (2024). *الكشف عن الأخبار والمعلومات الكاذبة والمُضللة باستخدام خوارزميات الذكاء الاصطناعي* [رسالة ماجستير، جامعة القدس]. جامعة القدس.
- الروياتي، ع. ع.، والشوماني، م. م. (2026). الكشف عن الأخبار الزائفة باللغة العربية باستخدام نماذج التعلم العميق المحسنة بالخوارزمية الجينية. *مجلة البحوث التقنية، عدد خاص لمؤتمر الدولي السابع للعلوم الهندسية والتقنية - (CEST-2025) الجزء الثاني، 865-879*.
<https://doi.org/10.26629/>
- الصادقي، ع. ع.، والأشعري، س. (2024). الأخبار الزائفة (The fake news) في وسائل التواصل الاجتماعي: تأصيل في المفهوم وبحث في الدوافع والأسباب. *مجلة البحث في العلوم الإنسانية والمعرفية* <https://crshc.com/3392/>.
- الحسبان، م. ع. (2026). *تطبيقات الذكاء الاصطناعي* (الطبعة الأولى). مكتبة المجتمع العربي للنشر والتوزيع.
- الصوالحة، س. د. م. م.، والرجبي، م. أ. م. (2025). اتجاهات القائم بالاتصال في وكالة الأنباء الأردنية (بترا) نحو تطبيق برمجيات الذكاء الاصطناعي في التحرير الصحفي: دراسة مسحية. *دراسات: العلوم الإنسانية والاجتماعية، 52(4)، 6437*.
- منصة شيبك. (د.ت). *شيبك: منصة للتحقق من الأخبار*. تم الاسترجاع في 10 مارس 2026، من <https://www.chayyek.com/>
- مرصد مصداقية الإعلام الأردني (أكيد). (د.ت). *أكيد: مرصد مصداقية الإعلام*. تم الاسترجاع في 10 مارس 2026، من <https://www.akeed.jo/>
- منصة تيقن. (د.ت). *تيقن: منصة للتحقق من الأخبار والمعلومات*. تم الاسترجاع في 10 مارس 2026، من <https://tayqan.net/>

References

- Wang, J., Wang, X., & Yu, A. (2025). Tackling misinformation in mobile social networks: A BERT-LSTM approach for enhancing digital literacy. *Scientific Reports*.
- Tsikerdekis, M., & Zeadally, S. (2023). *Misinformation detection using deep learning*. University of Kentucky Research Publications.
- Wang, R. (2023). Fake news detection based on deep learning. *Frontiers in Computing and Intelligent Systems*.
- Jurafsky, D., & Martin, J. H. (2026). *Speech and language processing: An introduction to natural language processing, computational linguistics, and speech recognition* (3rd ed., draft). Stanford University & University of Colorado Boulder.
- Hussna, A. U., Alam, M. G. R., Islam, R., Alkhamees, B. F., Hassan, M. M., & Uddin, M. Z. (2024). Dissecting the infodemic: Analysis of COVID-19 misinformation detection using machine learning and deep learning techniques. *Heliyon*.
- Vosoughi, S., Roy, D., & Aral, S. (2018). The spread of true and false news online. *Science*, 359(6380), 1146–1151.
- Zhou, X., & Zafarani, R. (2020). A survey of fake news: Fundamental theories, detection methods, and opportunities. *ACM Computing Surveys*, 53(5).
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT)*.
- Shu, K., Sliva, A., Wang, S., Tang, J., & Liu, H. (2017). Fake news detection on social media: A data mining perspective. *ACM SIGKDD Explorations Newsletter*, 19(1), 22–36.

- Merzah, B. M., & Razmara, J. (2025). Hybrid deep learning models for fake news detection: Case study on Arabic and English languages. *Frontiers in Big Data*.
- Islam, M. R., Liu, S., Wang, X., & Xu, G. (2020). Deep learning for misinformation detection on online social networks: A survey and new perspectives. *Social Network Analysis and Mining*.
- Srnicek, N. (2016). *Platform capitalism*. Polity Press.
- Alghamdi, J., Lin, Y., & Luo, S. (2022). A comparative study of machine learning and deep learning techniques for fake news detection. *Information*, 13(12), Article 576.
- Luger, G. F. (2005). *Artificial intelligence: Structures and strategies for complex problem solving* (6th ed.). Pearson.
- Tiedemann, J., & Thottingal, S. (2020). OPUS-MT — Building open translation services for the world. In *Proceedings of the 22nd Annual Conference of the European Association for Machine Translation* (pp. 479–480). European Association for Machine Translation.
- Al-Yahya, M. A. (2021). Arabic fake news detection: Comparative study of neural networks and transformer-based approaches. *Complexity*, 2021, 1–15.
- Antoun, W., Baly, F., & Hajj, H. (2020). AraBERT: Transformer-based model for Arabic language understanding. In *Proceedings of the 4th Workshop on Open-Source Arabic Corpora and Processing Tools*.
- Hussein, A., Hasanain, M., & Elsayed, T. (2021). Fighting the Arabic COVID-19 infodemic on Twitter using transformer-based language models. In *Proceedings of NLP4IF 2021*.

م. شروق العواودة - جامعة الشرق الأوسط - عمان - الأردن S.awawdeh@meu.edu.jo

أ. آية الجعفري- جامعة الشرق الأوسط - عمان - الأردن A.aljafari@meu.edu.jo