

استخدام تقنيات الذكاء الاصطناعي في الكشف عن الأفكار الانتحارية لدى الأفراد عبر الإنترنت: مراجعة

سرديّة وتحليل SWOT

د. إيمان بنت علي المحمدي¹

أستاذ مشارك في قسم علم النفس بجامعة الملك عبدالعزيز

المملكة العربية السعودية

ealmhamdi@kau.edu.sa

الملخص

تهدف الدراسة إلى إجراء مراجعة سرديّة للدراسات التي تناولت استخدام تقنيات الذكاء الاصطناعي في الكشف عن الأفكار الانتحارية لدى الأفراد عبر الإنترنت، مع دمج تحليل SWOT. اتبعت المراجعة معايير SANRA. وقد تضمنت المراجعة 23 دراسة نُشرت خلال السنوات الثلاث الأخيرة، وتم تجميع النتائج باستخدام التحليل الموضوعي. وخلص تحليل SWOT إلى نقاط القوة: دقة تنبئية عالية في اكتشاف الحالات ذات التفكير الانتحاري، والقدرة على التعامل مع لغات عدّة، وفعالية عبر منصات متنوعة، دعم وتطوير البيانات الاصطناعية لملاءم الفجوات وتحسين الأداء، والكشف عن السلوكيات البصرية ذات الصلة بحالات التفكير الانتحاري، وتحليل المشاعر من رسائل المرضى لتعزيز فهم الحالة النفسية الفردية. أما نقاط الضعف، فهي: محدودية تمثيل الموضوعات الرئيسية والمتنوعة في مجموعات البيانات، ووجود اختلال في توازن البيانات، إلى جانب تحديات التعدد اللغوي، ونماذج غامضة أو غير شفافة يصعب تفسيرها وتحليلها، والاعتماد على أدوات وتقنيات لم يتم التحقق من صحتها بشكل كامل، وارتفاع معدلات النتائج الإيجابية والخاطئة؛ نتيجة صعوبة النقاط التراكمية اللغوية العامية والسخرية والسياق. الفرص: دمج أدوات الذكاء الاصطناعي في رسائل المرضى وتطبيقاتهم لتعزيز التواصل والدعم، وتطوير روبوتات الدردشة المعتمدة على الذكاء الاصطناعي وأدوات العلاج الرقمي، وتطبيق التحليلات التنبئية للاستفادة في تقييم الحالة النفسية للمراهقين، وإنشاء لوحات معلومات تفاعلية لرفع مستوى التوعية المجتمعية بأهمية الصحة النفسية، وتطوير نظم الذكاء الاصطناعي المتعددة الوسائط لتعزيز تقديم الخدمات الصحية النفسية المبتكرة. التهديدات: مخاطر تتعلق بالخصوصية، والموافقة، والمراقبة لحماية البيانات الشخصية، والتحيز الخوارزمي والتحديات المرتبطة بالعدالة بين الفئات المختلفة، ووجود ثغرات تنظيمية وغموض قانوني يعوق تطبيق هذه التقنيات بشكل مسؤول، والفجوة الرقمية وعدم المساواة في الوصول إلى الأدوات والخدمات الرقمية، خاصة في الفئات الضعيفة.

الكلمات المفتاحية: الذكاء الاصطناعي، الأفكار الانتحارية، مراجعة سرديّة، تحليل SWOT.

Using Artificial Intelligence Technologies to Detect Suicidal Ideation in Individuals Online: A Narrative Review and SWOT Analysis

Iman Aly Almohammadi¹

¹Associate professor, Department of Psychology, King Abdulaziz University
Kingdom of Saudi Arabia

ealmhamdi@kau.edu.sa

ABSTRACT

The study aimed to conduct a narrative review of studies that addressed the use of AI techniques to detect suicidal ideation among individuals online, incorporating a SWOT analysis. The review followed SANRA criteria. It included 23 studies published over the past three years. The findings were synthesized using thematic analysis. The SWOT analysis concluded: **Strengths:** high predictive accuracy in detecting suicidal ideation, multilingual and cross-platform efficacy, synthetic data augmentation to improve performance, visual behavior detection, and patient message sentiment analysis to enhance understanding of individual psychological states. **Weaknesses:** limited representation of key topics, data imbalance and language barriers, opaque models that are difficult to interpret and analyze, reliance on non-validated tools, and false positives/negatives due to difficulties capturing colloquial language structures, sarcasm, and context. **Opportunities:** integration into patient messaging and apps; ai chatbots and digital therapy tools; predictive analytics for adolescents; interactive dashboards to raise public awareness; multimedia ai systems to enhance the delivery of innovative mental health services. **Threats:** privacy risks, consent, and surveillance to protect personal data; algorithmic bias and fairness; regulatory gaps and legal ambiguity; and the digital divide and access inequality to digital tools and services especially for vulnerable groups. **Keywords:** Artificial intelligence, suicidal thoughts, narrative review, SWOT analysis

1- مقدمة الدراسة

يُعد الانتحار مشكلة صحية عامة ملحة على مستوى العالم، حيث يُقدّر أن نحو 720 ألف شخص يضعون حدًا لحياتهم سنويًا حول العالم، مع وجود عدد أكبر بكثير ممن يحاولون الانتحار. وقد يحدث الانتحار في جميع الفئات العمرية، إلا أنه يُمثل ثالث أسباب الوفاة شيوعًا بين الأشخاص الذين تتراوح أعمارهم بين 15 و 29 عامًا على مستوى العالم (WHO, 2025). كذلك جاء في تقرير منظمة الصحة العالمية (WHO) في عام 2019 أن حالة انتحار تحدث كل 40 ثانية (Kline & Sabri, 2023). وتزداد محاولات الانتحار بشكل كبير نتيجة لمشكلات الصحة العقلية، حيث أشار التحالف الوطني الأمريكي للأمراض العقلية (National Alliance on Mental Illness) إلى أن 46% من ضحايا الانتحار يعانون مشكلات في الصحة النفسية (Mohri, Rostamizadeh, & Talwalkar, 2018). ولا شك في أن كل حالة انتحار تمثل مأساة تؤثر في الأسر والمجتمعات والدول، وتترك آثارًا طويلة الأمد على الأشخاص الذين تركوا خلفهم؛ الأمر الذي يقتضي استجابة صحية عامة من خلال تدخلات سريعة ومستندة إلى الأدلة، بالإضافة إلى الحاجة الملحة إلى استراتيجيات وقائية فعالة للحد من هذه الظاهرة (O'Connor & Nock, 2014).

يلعب الكشف المبكر عن الأفكار الانتحارية دورًا حاسمًا في الوقاية من الانتحار، إلا أنه يواجه تحديات معقدة؛ فقد أكدت الدراسات (Blumenthal & Kupfer, 1988; Hawton & van Heeringen, 2009; Mann et al., 1999; Oquendo et al., 2014; Stanley & Brown, 2012) أن الكشف المبكر يُسهل التدخل في الوقت المناسب، بما في ذلك الوصول إلى موارد الصحة النفسية، وتطبيق خطط السلامة، وتقديم علاجات مصممة خصيصًا لمعالجة القضايا الكامنة مثل الاضطرابات النفسية، والسمات الشخصية، والمشكلات الاجتماعية، والتاريخ العائلي. إلا أن هناك عوامل عديدة تعيق اكتشاف الأفكار الانتحارية، أبرزها الوصمة المرتبطة بالانتحار التي تُثني الناس عن طلب الدعم. كذلك، فإن لأساليب الفحص التقليدية محدودية واضحة، فهي تتطلب تواصلًا مفتوحًا حول الميول الانتحارية، مثل أدوات الإبلاغ الذاتي Self-report Instruments التي تعتمد على المشاركة الطوعية، وقد تكون ذاتية وغير فعالة بشكل كافٍ في اكتشاف الأفراد البارعين في إخفاء نياتهم (Cohen et al., 2018)، وبالمثل، تعتمد المقابلات السريرية Clinical Interviews غالبًا على لجوء الأفراد أنفسهم إلى طلب المساعدة أو التواصل مع خدمات الأزمات، وهو ما قد لا يفعله العديد من الأفراد المعرضين للخطر أو حتى الاعتراف بمعاناتهم (Alambo et al., 2019). كذلك، يُعد الافتقار إلى أداة تنبؤية عالية الدقة (Runeson et al., 2017) من أهم التحديات التي تواجه الكشف المبكر للأفكار الانتحارية. تؤكد هذه القيود الحاجة إلى تطوير أساليب أكثر فعالية واستباقية للكشف عن الأفكار الانتحارية.

ومع انتشار منصات التواصل الرقمي، يعبر الأفراد بشكل متزايد عن أفكارهم ومشاعرهم عبر الإنترنت، وغالبًا ما يتشاركون أفكارهم الانتحارية بلغة خفية أو مُشفرة، أو يكشفون عن تلميحات خفية لمعاناتهم من خلال الأنشطة الإلكترونية

ومنشورات وسائل التواصل الاجتماعي. لذلك، ركزت العديد من الدراسات الحديثة (Aejaz, Owais, & Kalyani, 2023; Basyouni et al., 2024; Casu, Triscari, Battiato, Guarnera, & Caponnetto, 2024; Ghanadian et al., 2025; Yeskuatov, Chua, & Foo, 2025) على استكشاف فعالية دمج تقنيات الذكاء الاصطناعي والتعلم الآلي في مساعي الوقاية من الانتحار كوسيلة للكشف الاستباقي، في محاولة لتحديد عوامل الخطر المحتملة قبل وصول الفرد إلى نقطة الأزمة، وفحص قدرة هذه التقنيات على التنبؤ بالأفراد المعرضين للخطر بدقة واعدة.

وقد أحدثت التكنولوجيا والذكاء الاصطناعي ثورة في تقييم الصحة النفسية من خلال تمكين الوصول والكشف المبكر Early Detection والعلاج الشخصي Personalized Treatment. تُمكن المنصات الرقمية، مثل تطبيقات العلاج عن بعد (على سبيل المثال، BetterHelp، Talkspace) وأدوات التقييم الذاتي، الأفراد من مراقبة صحتهم النفسية عن بُعد، في حين أن تقنيات الذكاء الاصطناعي، مثل تحليل الكلام واللغة والوجه، يمكن أن تحدد علامات الاكتئاب أو القلق أو الميول الانتحارية، من خلال منشورات وسائل التواصل الاجتماعي أو مكالمات الفيديو أو التسجيلات (Fitzpatrick et al., 2017; Reece et al., 2017). كما توفر روبوتات الدردشة Chatbots مثل Woebot و Wysa دعماً فورياً وقابلاً للتطوير من خلال تقديم تقنيات العلاج السلوكي المعرفي، والتدخلات في الوقت الفعلي (Grygas et al., 2020). بشكل عام، تجعل هذه الابتكارات التكنولوجية الرعاية الصحية النفسية أكثر سهولة ودقة واستباقية؛ ما يؤدي في النهاية إلى تحسين التدخل المبكر ونتائج العلاج.

وعلى الرغم من إظهار النماذج القائمة على الذكاء الاصطناعي إمكانات واعدة في إحداث نقلة نوعية في خدمات الصحة النفسية من خلال تحسين دقة التشخيص وتوفير رعاية شخصية، فإن هناك قيوداً عديدة يجب معالجتها لضمان تطبيقها بفعالية وأخلاقية (Saeidnia et al., 2024; Sogancioglu et al., 2024; Warriar et al., 2023). ومن أبرز هذه القيود التحيز الخوارزمي، الذي قد ينشأ عن مجموعات بيانات متحيزة لا تمثل بشكل كافٍ فئات سكانية متنوعة. كما يُشكل غياب الشفافية في عمليات صنع القرار المتعلقة بالذكاء الاصطناعي تحديات؛ إذ إن طبيعة هذه النماذج التي تُعرف بـ"الصندوق الأسود" قد تُعيق قدرة كل من الأطباء والمعالجين النفسيين والمرضى على فهم التوصيات المقدمة والثقة بها. أيضاً، تُمثل خصوصية البيانات تحدياً أخلاقياً بالغ الأهمية، حيث تُشكل أنظمة الذكاء الاصطناعي مخاطر الوصول غير المصرح به، واستغلال البيانات. كذلك، يُمكن أن يحدث دمج تقنيات الذكاء الاصطناعي تغييراً جذرياً في ديناميكيات العلاقة التقليدية بين الطبيب والمريض، ويظل إيجاد التوازن الأمثل بين التدخلات القائمة على الذكاء الاصطناعي والحكم البشري مُعضلةً أخلاقية. تُبرز هذه الثغرات الحاجة إلى مراجعات شاملة تُقيم بشكل نقدي نقاط القوة والضعف في تقنيات الذكاء الاصطناعي الحالية في الكشف عن الأفكار الانتحارية لدى الأفراد عبر الإنترنت.

وتُعدّ المراجعة السردية Narrative Review (Braun & Clarke, 2019; Baethge, Goldbeck-Wood & Mertens, 2019) من الأدوات المنهجية التي تتيح تقييمًا نقدياً ومنهجياً للأساليب والتقنيات القائمة على الذكاء الاصطناعي

في الكشف المبكر عن الأفكار الانتحارية عبر الإنترنت. كذلك، يُمكن أن يُستخدم تحليل The Strengths, Weaknesses, Opportunities, & Threats analysis (SWOT) كأداة استراتيجية فعالة لتحديد نقاط القوة والضعف في المعالجات المستخدمة، والاستفادة من الفرص المتاحة، وتحديد التهديدات المحتملة؛ ما يسهم في وضع استراتيجيات مُحكمة لتحسين وتطوير الأساليب المستخدمة وتعزيز فعاليتها في الكشف المبكر عن الأفكار الانتحارية عبر الإنترنت. بناءً على ذلك، تهدف الدراسة الحالية إلى إجراء مراجعة سرديّة للدراسات التي تناولت استخدام تقنيات الذكاء الاصطناعي في الكشف عن الأفكار الانتحارية لدى الأفراد عبر الإنترنت، مع تسليط الضوء على نقاط قوتها ونقاط ضعفها ومجالات تحسينها من خلال دمج تحليل SWOT (نقاط القوة والضعف والفرص والتهديدات) لتوفير إطار استراتيجي شامل يُسهّل فهم الأبعاد التكنولوجية والأخلاقية والعملية لهذه التقنيات.

2- منهجية الدراسة

اتبعت المراجعة السردية الحالية معايير SANRA (Baethge et al., 2019)، من خلال تقييم ستة معايير جودة، هي: (1) تبرير أهمية المراجعة، (2) توضيح أهداف المراجعة، (3) وصف إجراءات البحث في الأدبيات، (4) التوثيق، (5) الاستدلال العلمي، (6) العرض المناسب للبيانات. تم استيفاء المعيارين الأول والثاني في مقدمة الدراسة، في حين تم تناول المعايير الثالث والرابع والخامس بشكل أكثر تفصيلاً في إجراءات ومنهجية البحث عن الدراسات. أما المعياران الخامس والسادس، فقد تمت الإشارة إليهما في تلخيص النتائج ومناقشتها.

2.1. إجراءات البحث في الأدبيات

قامت الباحثة بمراجعة الأدبيات السابقة التي تتناول استخدام تقنيات الذكاء الاصطناعي في الكشف عن الأفكار الانتحارية لدى الأفراد عبر الإنترنت، مع التركيز على نقاط القوة والضعف والفرص والتهديدات لهذه التقنيات، وذلك من خلال إجراء بحث شامل في الأدبيات عبر قواعد بيانات إلكترونية مثل PubMed و PsycINFO و Web of Science و Scopus و ScienceDirect و Google Scholar. وقد استخدم البحث مزيجاً من الكلمات المفتاحية مثل “Artificial Intelligence,” “Suicidal,” “Machine Learning,” “Deep Learning,” “Natural Language Processing,” “Suicide Risk,” “Online Mental Health,” “Ideation,” “Suicide Risk” .

2.2. معايير إدراج واستبعاد الدراسات السابقة

تم وضع معايير إدراج واستبعاد محددة لتوجيه اختيار الدراسات، حيث أُدرجت الدراسات ذات الصلة التي تناولت أساليب الكشف القائمة على الذكاء الاصطناعي للأفكار الانتحارية في السياقات عبر الإنترنت. وقد شملت معايير الإدراج المقالات التي راجعها النظراء، ونُشرت خلال السنوات الثلاث الأخيرة (2023 – 2025)، وكانت مكتوبة باللغة الإنجليزية، ومتوفرة للباحثة

بشكل كامل. واستُبعدت الأدبيات الأكاديمية غير المُحكّمة، مثل الرسائل ووقائع المؤتمرات، وكذلك الكتب وفصول الكتب وأوراق المؤتمرات. ولم تُفرض أي قيود على تصميم الدراسة لضمان نظرة عامة شاملة على الموضوع. في البداية، قامت الباحثة بفحص العناوين والملخصات والكلمات المفتاحية للتأكد من صلتها بالموضوع، ثم مراجعة الدراسات كاملة للتأكد من ملاءمتها للإدراج.

3- النتائج

تم تجميع النتائج باستخدام التحليل الموضوعي Thematic analysis (Braun & Clarke, 2006) لتحديد الموضوعات والنتائج الرئيسية المتعلقة بنقاط القوة والضعف، بالإضافة إلى التحديات والتهديدات المرتبطة باستخدام تقنيات الذكاء الاصطناعي في الكشف المبكر عن الأفكار الانتحارية بين الأفراد عبر الإنترنت. وقد تضمنت المراجعة 23 دراسة، ويستعرض الملحق (1) الدراسات المدرجة في المراجعة.

1.3. المراجعة السردية

أسفرت المراجعة عن المحاور الأربعة التالية:

المحور الأول: نقاط القوة - القدرات التكنولوجية للذكاء الاصطناعي في الكشف عن الأفكار الانتحارية عبر الإنترنت

سلطت العديد من الدراسات الضوء على الأداء القوي للذكاء الاصطناعي في تحديد الأفكار الانتحارية باستخدام خصائص نصية وسلوكية ومتعددة الوسائط.

- دقة تنبئية عالية في اكتشاف الحالات High Predictive Accuracy: أظهرت نماذج التحويل (Transformer-based models) قدرة في الكشف عن الأفكار الانتحارية، فقد أكدت دراسة (Yeskuatov et al., 2025) أن

نموذج BERT المخصص حقق أداءً متفوقاً، حيث سجل دقة عالية في اكتشاف الأفكار الانتحارية على منصة

Reddit، بمتوسط F1-score بلغ 99%. وبالمثل، أكدت دراسة (Basyouni et al., 2024) دقة نموذج الغابة

العشوائية (Random Forest classifiers) المخصص للميزات المختارة عبر الخوارزمية الجينية، حيث حقق معدل

دقة و F1-score بلغ 98.92% في الكشف عن الأفكار الانتحارية. كذلك، أظهرت دراسة (Buddhitha &

Inkpen, 2023) أن استخدام تقنيات التعلم المتعدد المهام وبيانات متعددة المنصات يحسّن بشكل كبير من دقة نماذج

الكشف المبكر عن الأفكار الانتحارية، خاصة عند دمج بيانات من مستخدمين يعانون اضطرابات نفسية متعددة، وأن

الميزات المشتركة بين المصابين تعزز من قدرة النماذج على التنبؤ بالمخاطر بشكل أكثر دقة وفعالية. وأخيراً، قام (Al-

Remawi et al., 2024) بإجراء مراجعة تهدف إلى استكشاف الدور التحويلي للذكاء الاصطناعي والتعلم الآلي في

تعزيز التنبؤ بمخاطر الانتحار وتطوير استراتيجيات وقائية فعالة، مع التركيز على كيفية دمج هذه التقنيات لعوامل

خطر معقدة تشمل العوامل النفسية، والاجتماعية، والاقتصادية، والنمط الغذائي. وقد أظهرت النتائج أن الذكاء

الاصطناعي والتعلم الآلي يُحدثان ثورة في استراتيجيات الوقاية من الانتحار، حيث تحقق النماذج دقة تنبئية تقترب من 90% عند دمج مصادر بيانات متنوعة.

- القدرة على التعامل مع لغات عدة وفعالية عبر منصات متنوعة **Multilingual and Cross-Platform Efficacy**: أكد (Mansoor & Ansari, 2024) أن أداء الذكاء الاصطناعي كان متسقًا عبر اللغات (F1: 0.827-0.872) وعبر المنصات (F1: 0.839-0.863)، حيث بلغ درجة F1 93.5% في الكشف عن الأفكار الانتحارية باللغات الإنجليزية والمندرين والعربية والإسبانية؛ ما يُظهر قدرة النموذج على التعميم عبر الثقافات.
- دعم وتطوير البيانات الاصطناعية لملء الفجوات وتحسين الأداء **Synthetic Data Augmentation**: إن النهج المبتكر الذي يعتمد على توليد بيانات اصطناعية باستخدام نماذج الذكاء الاصطناعي، مثل ChatGPT وFlan-T5 وLlama، يعزز بشكل ملحوظ أداء نماذج التصنيف اللغوي لاكتشاف الأفكار الانتحارية، حيث أظهرت نتائج دراسة (Ghanadian et al., 2024) أن النماذج المدربة على البيانات الاصطناعية أظهرت أداءً ثابتًا بمعدل F1 يبلغ 0.82، مقارنةً بنطاق 0.75 إلى 0.87 عند تدريبها على البيانات الحقيقية. كما أن دمج 30% فقط من البيانات الواقعية مع البيانات الاصطناعية أدى إلى تحسن في الأداء، حيث وصلت دقة النموذج إلى $F1 = 0.88$ ؛ ما يُشير إلى فاعلية هذه الطريقة في التغلب على تحديات نقص البيانات، وتوفير تنوع أكبر في التمثيل. كذلك، قامت (Ghanadian et al., 2025) بدمج البيانات المعززة باستخدام نماذج الذكاء الاصطناعي مثل GPT-3.5 Turbo في عملية ضبط النماذج؛ ما أدى إلى تحسين قدراتها على اكتشاف الأفكار الانتحارية، حيث زادت درجات F1 من 0.87 إلى 0.91 على مجموعة بيانات جامعة ميريلاند، ومن 0.70 إلى 0.90 على المجموعة الاصطناعية، وهو ما يُبرز فاعلية النهج في تحسين الدقة.
- الكشف عن السلوكيات البصرية ذات الصلة بحالات التفكير الانتحاري **Visual Behavior Detection**: كشفت أدوات الرؤية الحاسوبية بدقة عن سلوكيات الأزمات، مثل التمشي المتكرر والبقاء لفترة طويلة. حيث تمكنت خوارزمية التعرف التلقائي على السلوك من التحديد الصحيح لـ 80% من المقاطع التي تم تفعيلها كمقاطع لأزمات، مع استبعاد 90% من المقاطع غير المرتبطة بالأزمة بشكل صحيح (Onie et al., 2023).
- تحليل المشاعر من رسائل المرضى لتعزيز فهم الحالة النفسية الفردية **Patient Message Sentiment Analysis**: أظهرت نماذج البرمجة اللغوية العصبية (NLP models) باستخدام رسائل بوابة المريض (Patient Portal) إمكانية التنبؤ بالأحداث المتعلقة بالانتحار خلال مدة 30 يومًا ($AUC = 0.710$) مع حساسية بنسبة 56% وخصوصية بنسبة 69% (Bhandarkar et al., 2023).

المحور الثاني: نقاط الضعف - القيود التقنية والمنهجية

خلصت العديد من الدراسات إلى أن هناك نقاط ضعف في نماذج الذكاء الاصطناعي وتقييمها وتطبيقها.

- **محدودية تمثيل الموضوعات الرئيسية والمتنوعة في مجموعات البيانات Limited Representation of Key Topics:** تعاني البيانات الواقعية على وسائل التواصل الاجتماعي نقصًا واضحًا في تمثيل الموضوعات المتعلقة بالمجتمعات المهمشة. فقد أكدت دراسة (Ghanadian et al., 2025) وجود نقص في تغطية قضايا الأقليات، والصدمات النفسية، والضغوط المتعلقة بالهوية في مجموعات البيانات، الأمر الذي يحد من قدرة نماذج الذكاء الاصطناعي على التعميم. ومن خلال استخدام بيانات اصطناعية مولدة بواسطة GPT-3.5 Turbo، وتحسين مجموعة البيانات، تم تعزيز تنوع المواضيع بشكل كبير، مع بقاء مستوى القراءة والتعقيد مشابهًا للبيانات الحقيقية.
- **وجود اختلال في توازن البيانات والحواجز اللغوية Data Imbalance and Language Barriers:** تقع مجموعات البيانات الخاصة باللغة العربية في حالة نقص وركود من حيث التطوير، على الرغم من الجهود الأخيرة التي قام بها (Alatawi et al., 2024)، حيث أظهرت دراستهم أن استخدام نماذج اللغة المدربة مسبقًا باللغة العربية، مثل USE و MARBERT، فعال في الكشف عن الأفكار الانتحارية عبر وسائل التواصل الاجتماعي، حيث حققت نماذج USE والنموذج العميق الكثيف معدل F1 بلغ 82.6%، في حين أن نموذج CNN+LSTM المستند إلى MARBERT تفوق عليها، محققًا معدل F1 بلغ 80.8%.
- **نماذج غامضة أو غير شفافة يصعب تفسيرها وتحليلها Opaque Models:** وهي نماذج الذكاء الاصطناعي التي يصعب فهم كيفية اتخاذها لقراراتها، نظرًا لعدم وضوح العمليات الداخلية التي تتبعها. وتُعد هذه النماذج تحديًا في مجالات تتطلب الشفافية والثقة. وقد ناقش (Parsapoor et al., 2023) طبيعة التعلم العميق الغامضة (Deep learning's black-box) التي تعيق التفسير السريري.
- **الاعتماد على أدوات وتقنيات لم يتم التحقق من صحتها بشكل كامل Reliance on Non-Validated Tools:** لم يتم قياس أداء العديد من النماذج والتقنيات مقابل المقاييس السريرية أو اعتمادها من خلال دراسات طويلة للتحقق من صحتها عبر الزمن، حيث اقترحت دراسة (Parsapoor et al., 2023) إنشاء مجموعة بيانات قياسية تشمل البيانات النصية والصوتية، موحدة عبر إجراءات تقييم موحدة، لتمييز عوامل الخطر التي تؤدي إلى محاولة الانتحار بجانب الأفكار الانتحارية فقط. ويمكن أن تساعد هذه المجموعة على تطوير أنظمة ذكاء اصطناعي تعتمد على التعلم الآلي أو العميق، وتلعب دورًا مهمًا في الكشف المبكر والموثوق به عن مخاطر الانتحار.
- **ارتفاع معدلات النتائج الإيجابية والخاطئة نتيجة صعوبة التقاط التراكيب اللغوية العامية والسخرية والسياق False Positives/Negatives:** ذكرها (Basyouni et al., 2024) كإحدى الصعوبات الرئيسية التي تواجه نماذج الذكاء الاصطناعي عند محاولة التمييز بين المعاني الحقيقية وغير المباشرة، أو غير الواضحة. فالسخرية، والسياق الثقافي، والتعبيرات المبهمة تُعد من العوامل التي تعوق التصنيف الدقيق، حيث يصعب على النماذج فهم المعنى الحقيقي للنصوص؛ ما قد يؤثر في دقة أداء النماذج.

المحور الثالث: الفرص - التكامل الاستراتيجي في النظم الصحية

خلصت العديد من الدراسات إلى أن الذكاء الاصطناعي يُتيح إمكانيات تحويلية للمجالات السريرية، ومجالات الصحة العامة، والبحث.

- **دمج أدوات الذكاء الاصطناعي في رسائل المرضى وتطبيقاتهم لتعزيز التواصل والدعم Integration into Patient Messaging and Apps:** يمكن لنماذج البرمجة اللغوية العصبية (NLP) المُدرَّبة على رسائل بوابة المرضى أن تسهم في تعزيز جهود الوقاية من الانتحار داخل البيئات السريرية، من خلال الدمج في رسائل المرضى وتطبيقاتهم (Bhandarkar et al., 2023). كذلك، أظهرت دراسة (Aejaz et al., 2023) أن استخدام هذه النماذج مع المراهقين أدى إلى تحسين القدرة على الكشف المبكر عن الأفكار الانتحارية وتحليل المشاعر؛ ما يساعد في التدخل المبكر وتقليل معدلات الانتحار بين الفئة العمرية للمراهقين.

- **تطوير روبوتات الدردشة المعتمدة على الذكاء الاصطناعي وأدوات العلاج الرقمي AI Chatbots and Digital therapy Tools:** أثبتت روبوتات الدردشة (Chatbots)، مثل Woebot، ووكلاء الذكاء الاصطناعي (AI agents)، وهي تقنيات تقوم بمحاكاة المحادثات البشرية باستخدام تقنيات مثل معالجة اللغة الطبيعية (NLP) والتعلم العميق (Deep Learning)، فعاليتها في تقديم دعم فوري داخل مجال الصحة النفسية. فقد أشارت دراسة (Casuet al., 2024) إلى أن روبوتات الدردشة المدعومة بالذكاء الاصطناعي أظهرت قدرة محتملة على تحسين الرفاهية النفسية والعاطفية، ومعالجة حالات الصحة النفسية، وتيسير تغيير السلوك. كذلك، توصلت دراسة (Sambana et al., 2025) إلى أن إطار العمل القائم على العلاج السلوكي المعرفي، والمستخدم لتحليل المحتوى النصي والمرئي، يتيح تصنيف المشاعر إلى إيجابية أو سلبية باستخدام نماذج، مثل BERT وRoBERTa، بالإضافة إلى تلخيص النصوص عبر نظام T5 وPEGASUS وترجمة النصوص باستخدام mt5 للغات عدة. كل ذلك يسهم في الكشف المبكر عن المشاعر السلبية واضطرابات التفكير والتوجهات الانتحارية وغيرها من الأمراض النفسية الأخرى.

- **تطبيق التحليلات التنبؤية للاستفادة في تقييم الحالة النفسية للمراهقين Predictive Analytics for Adolescents:** إن استخدام تقنيات التعلم الآلي يمثل أداة واعدة لتوقع مخاطر الصحة النفسية في فئة المراهقين؛ حيث أظهرت نتائج دراسة (Su et al., 2023) أن نماذج التصنيف باستخدام طرق التعلم الآلي، خاصة مصنّفات الغابات العشوائية (Random Forest classifiers)، فعّالة في اختيار أكثر المتغيرات تأثيراً من بين مئات المتغيرات للتنبؤ بمخاطر الانتحار والإيذاء الذاتي لدى المراهقين، إذ حققت معدل دقة معتدلاً للتنبؤ بإيذاء النفس (AUC: 0.72) ومحاولات الانتحار (AUC: 0.74)، وهو أعلى من استراتيجيات التنبؤ المبنية فقط على التاريخ السابق للمحاولة (AUC: 0.60). كما أظهرت نتائج دراسة (Smith et al., 2023) أن استخدام تقنيات الذكاء الاصطناعي في تحليل مؤشرات الخطر الجيني والبيئي قد يُمكن من تحسين التنبؤ بالحالات النفسية الخطيرة بين الشباب. أيضاً، توصلت دراسة (Li

(et al., 2024) إلى أن الذكاء الاصطناعي يُحسن تقييم المخاطر ونماذج التنبؤ بالسلوك الانتحاري بين المراهقين، من خلال تحليل البيانات النصية باستخدام المعالجة اللغوية الطبيعية وتطوير استراتيجيات تدخل مبتكرة، مع تأكيد أن تقنيات الذكاء الاصطناعي، مثل الروبوتات الحوارية وأنظمة المراقبة، تمتلك إمكانات تحويلية لتعزيز جهود الوقاية من الانتحار في هذه الفئة العمرية.

- إنشاء لوحات معلومات تفاعلية لرفع مستوى التوعية المجتمعية بأهمية الصحة النفسية **Interactive Dashboards for Public Awareness**: يساهم استخدام أدوات الويب التفاعلية (Interactive web) في تعزيز مشاركة الجمهور وتتبعه حول موضوع الانتحار؛ فقد أظهرت (Hisham et al., 2023) أن استخدام تقنيات التصور التفاعلي عبر لوحة البيانات (Dashboard) المبنية باستخدام برنامج Tableau ساعد في جذب المستخدمين لقراءة معلومات حول الانتحار والتعرف على تاريخه بشكل أكثر فاعلية؛ الأمر الذي يعزز الوعي المجتمعي، ويقلل من معدل الانتحار على المدى الطويل. كذلك، اختبرت الدراسة قدرة نماذج التعلم الآلي على التنبؤ بمعدل الانتحار باستخدام ثلاث خوارزميات هي الغابة العشوائية (Random Forest Regressor)، وشجرة القرار (Decision Tree Regressor)، ومتجه الدعم (Support Vector Regressor)، حيث أظهرت خوارزمية الغابة العشوائية أداءً متفوقاً من حيث أعلى قيمة R^2 وأقل معدل خطأ في التنبؤ.

- تطوير نظم الذكاء الاصطناعي المتعددة الوسائط لتعزيز تقديم الخدمات الصحية النفسية المبتكرة **Multimodal AI Systems**: تمكن الأنظمة التي تدمج بين إشارات الصوت والسلوك واللغة من تقديم توقعات أكثر غنى وسياًفاً، بالإضافة إلى الكشف المبكر عن الأزمات بشكل أكثر فاعلية (Mansoor & Ansari, 2024). وقد أكدت دراسة (Qadeer et al., 2024) أن تقنيات الذكاء الاصطناعي، خاصة نماذج التعلم الآلي، يمكن أن تُستخدم بشكل فعال لتحسين جودة حياة الأشخاص المصابين بالاكتئاب الخفيف، من خلال التشخيص المبكر للأعراض عبر أجهزة الأندرويد، حيث تعتمد على تحليل البيانات الشخصية مثل الرسائل النصية، والبريد الإلكتروني، والمكالمات الصوتية وتاريخ التصفح. من خلال مراقبة سلوك المستخدم بشكل مستمر، يمكن تحديد الاتجاهات السلبية والتدخل سريعاً لتوفير الرعاية اللازمة قبل تفاقم الحالة؛ ما يساهم في إنقاذ الأرواح، خاصة في حالات الاكتئاب الشديد التي قد تؤدي إلى الانتحار. وقد تم تطوير واختبار تطبيق أندرويد قادر على أداء هذه المهام بشكل مباشر وبمتوسط دقة بلغ 95%.

المحور الرابع: التهديدات - المخاطر الأخلاقية والقانونية والاجتماعية

توجد العديد من المخاوف الأخلاقية والتشغيلية التي تُهدد الاستخدام المسؤول للذكاء الاصطناعي في الكشف المبكر والوقاية من الانتحار.

- مخاطر تتعلق بالخصوصية والموافقة والمراقبة لحماية البيانات الشخصية **Privacy, Consent, and Surveillance**: تؤكد الدراسات أهمية النظر في الاعتبارات الأخلاقية لضمان تطوير وتوظيف تقنيات الذكاء

الاصطناعي بأسلوب مسؤول وأخلاقي، يراعي حقوق الأفراد، ويعزز العدالة، ويحقق الفوائد القصوى مع الحد من الأضرار المحتملة. فقد خلصت دراسة (Solaiman et al., 2023) إلى أن تطبيق الذكاء الاصطناعي في المستشفيات الكبرى، مثل أنظمة المراقبة والتنبؤ بالانتحار وإدارة المستشفيات، يساهم في تحسين نتائج الرعاية الصحية، وتقليل التكاليف، وتوفير الوقت. إلا أنها شددت على أن استخدام الذكاء الاصطناعي في الأقسام النفسية يتطلب تقييمًا دقيقًا للمخاطر المرتبطة بخصوصية المرضى، والموافقة المستنيرة، ومعالجة البيانات، ويجب أن يتم بحذر شديد وتقييم مستمر لضمان الاستخدام المسؤول والأخلاقي. كما استعرضت دراسة (Saeidinia et al., 2024) 18 اعتبارًا أخلاقيًا رئيسيًا؛ منها: مسائل الخصوصية والسرية، والموافقة المستنيرة، والتحيز والعدالة، والشفافية والمساءلة، والاستقلالية، والوكالة الإنسانية، والسلامة والكفاءة. كما حددت المبادئ الأخلاقية الواجب اتباعها في أثناء تطوير وتنفيذ تقنيات الذكاء الاصطناعي في مجالات الصحة النفسية، مثل إطار العمل الأخلاقي، وانخراط الأطراف المعنية، والمراجعة الأخلاقية، وتقليل التحيز، والتقييم المستمر للتحسين. أيضًا، أوصت الدراسة بممارسات وإرشادات لتعزيز الاستخدام الأخلاقي للذكاء الاصطناعي؛ منها الالتزام بالمبادئ الأخلاقية، وضمان الشفافية، وحماية البيانات، وتقليل التحيز، وتضمين الأطراف المعنية، وإجراء مراجعات دورية، ومراقبة النتائج.

- التحيز الخوارزمي والتحديات المرتبطة بالعدالة بين الفئات المختلفة **Algorithmic Bias and Fairness**:

أوضحت دراسة كل من (Bose et al., 2024; Ghanadian et al., 2025) أن النتائج التمييزية قد تظهر نتيجة لوجود بيانات تدريبية متحيزة؛ ما يؤدي إلى إصدار نماذج الذكاء الاصطناعي قرارات غير عادلة أو غير دقيقة.

- وجود ثغرات تنظيمية وغموض قانوني يعوقان تطبيق هذه التقنيات بشكل مسؤول **Regulatory Gaps and Legal Ambiguity**:

غالبًا ما تكون الأطر الأخلاقية غير موجودة أو غير متسقة؛ الأمر الذي يبرز الحاجة إلى وضع مبادئ واضحة لضمان استخدام مسؤول وأخلاقي لتقنيات الذكاء الاصطناعي في هذا السياق، حيث أكدت دراسة (Bose et al., 2024) الحاجة إلى وضع أطر أخلاقية قوية تتناول تعقيدات دمج الذكاء الاصطناعي في الرعاية النفسية، مع تأكيد أهمية تغيير السياسات الصحية وتطوير قواعد تنظيمية تضمن الاستخدام المسؤول لهذه التقنيات، من خلال التوعية وتقوية الرقابة؛ بهدف حماية حقوق المرضى، والحفاظ على العلاقات العلاجية في بيئة رقمية متزايدة. كذلك، شدد (Smith et al., 2023) على وجوب تقييم المخاطر الأخلاقية المرتبطة بتجاوز الخصوصية والتحيز، مع ضرورة وضع لوائح تنظيمية لضمان الاستخدام المسؤول والمستدام لتقنيات الذكاء الاصطناعي في الرعاية الصحية النفسية.

- الفجوة الرقمية وعدم المساواة في الوصول إلى الأدوات والخدمات الرقمية خاصة في الفئات الضعيفة **Digital Divide and Access Inequity**:

يمكن أن تؤدي مسألة عدم المساواة في الوصول إلى الأجهزة أو مستوى المهارات الرقمية إلى تفاقم الفوارق في الرعاية المدعومة بالذكاء الاصطناعي، وإعاقة العدالة في تقديم الخدمات الصحية (Samruddhi & Vaidya, 2024).

2.3. تحليل نقاط القوة والضعف والفرص والتهديدات (SWOT)

تم إجراء تحليل شامل لنقاط القوة والضعف والفرص والتهديدات (SWOT) لتقنيات الذكاء الاصطناعي المستخدمة في الكشف عن الأفكار الانتحارية لدى الأفراد عبر الإنترنت، والتي تناولتها الدراسات المتضمنة في المراجعة السردية. ويوضح جدول (1) نتائج تحليل SWOT، حيث يلخص جوانب القوة التي تعزز من فعالية هذه التقنيات، ويحدد نقاط الضعف التي تتطلب معالجة، ويوضح الفرص التي يمكن استثمارها لتطوير المجال، بالإضافة إلى التهديدات التي قد تؤثر في التطبيق المسؤول والمستدام لهذه التقنيات.

جدول 1

تحليل نقاط القوة والضعف والفرص والتهديدات (SWOT) لاستخدام تقنيات الذكاء الاصطناعي في الكشف عن الأفكار الانتحارية لدى الأفراد عبر الإنترنت

نقاط القوة	نقاط الضعف
- دقة تنبؤية عالية في اكتشاف الحالات ذات التفكير الانتحاري. - القدرة على التعامل مع لغات عدة وفعالية عبر منصات متنوعة. - دعم وتطوير البيانات الاصطناعية لملء الفجوات وتحسين الأداء. - الكشف عن السلوكيات البصرية ذات الصلة بحالات التفكير الانتحاري. - تحليل المشاعر من رسائل المرضى لتعزيز فهم الحالة النفسية الفردية.	- محدودية تمثيل الموضوعات الرئيسية والمتنوعة في مجموعات البيانات. - وجود اختلال في توازن البيانات، إلى جانب تحديات التعدد اللغوي. - نماذج غامضة أو غير شفافة يصعب تفسيرها وتحليلها. - الاعتماد على أدوات وتقنيات لم يتم التحقق من صحتها بشكل كامل. - ارتفاع معدلات النتائج الإيجابية والخطئة نتيجة صعوبة النقاط التراكيب اللغوية العامة والسخرية والسياق.
الفرص	التهديدات
- دمج أدوات الذكاء الاصطناعي في رسائل المرضى وتطبيقاتهم لتعزيز التواصل والدعم. - تطوير روبوتات الدردشة المعتمدة على الذكاء الاصطناعي وأدوات العلاج الرقمي. - تطبيق التحليلات التنبؤية للاستفادة في تقييم الحالة النفسية للمراهقين. - إنشاء لوحات معلومات تفاعلية لرفع مستوى التوعية المجتمعية بأهمية الصحة النفسية.	- مخاطر تتعلق بالخصوصية، والموافقة، والمراقبة لحماية البيانات الشخصية. - التحيز الخوارزمي والتحديات المرتبطة بالعدالة بين الفئات المختلفة. - وجود ثغرات تنظيمية وغموض قانوني يعوقان تطبيق هذه التقنيات بشكل مسؤول.

- تطوير نظم الذكاء الاصطناعي المتعددة الوسائط لتعزيز تقديم الخدمات الصحية النفسية المبتكرة. - الفجوة الرقمية وعدم المساواة في الوصول إلى الأدوات والخدمات الرقمية، خاصة في الفئات الضعيفة.

4- المناقشة

يعد استخدام تقنيات الذكاء الاصطناعي في الكشف عن الأفكار الانتحارية لدى الأفراد عبر الإنترنت من المجالات الواعدة التي تفتح آفاقًا جديدة في تحسين خدمات الرعاية النفسية؛ إذ يُقدم حلولاً مبتكرة للتحديات الطويلة الأمد التي يواجهها الأطباء والمعالجون النفسيون والباحثون وصانعو السياسات. لذلك، هدفت الدراسة الحالية إلى إجراء مراجعة سردية للدراسات التي تناولت تطبيقات الذكاء الاصطناعي في الكشف عن الأفكار الانتحارية على الإنترنت، مع التركيز على تحليل نقاط القوة والضعف، والفرص والتحديات باستخدام إطار تحليل SWOT، بهدف تقديم تصور شامل يساهم في فهم الأبعاد التكنولوجية، والأخلاقية، والعملية لهذه التقنيات؛ من أجل تحسين استدامتها وتوجيه استخدامها المسؤول في الرعاية النفسية.

وكما بينت النتائج الرئيسية المستخلصة من الدراسات السابقة، فإن أبرز نقاط قوة الذكاء الاصطناعي تكمن في دقته التنبؤية العالية في الكشف عن الأفكار الانتحارية، كما أوضحتها العديد من الدراسات التي استخدمت نماذج المحولات (مثل BERT) ومصنفات الغابات العشوائية، حيث تجاوزت درجة دقة بعض هذه النماذج 90% من درجات F. كذلك، تتيح القدرة على معالجة وتحليل تدفقات البيانات النصية والسلوكية والمتعددة الوسائط بفعالية الكشف المبكر عن الأزمات المحتملة، وهو ما يتماشى مع الأهداف الأوسع للرعاية الصحية النفسية الاستباقية. إلى جانب ذلك، تؤكد قدرة الذكاء الاصطناعي التعامل مع لغات متعددة والعمل عبر منصات متنوعة على مرونته في التكيف مع التطبيقات العالمية، وهو أمر بالغ الأهمية بالنظر إلى الطبيعة الواسعة لتحديات الصحة النفسية. كما يُحسن تطبيق تقنيات تعزيز البيانات الاصطناعية أداء الذكاء الاصطناعي من خلال معالجة مشكلات ندرة البيانات والتحيز، وتوليد بيانات تسد الثغرات، وتحسن أداء النموذج. وأخيرًا، أظهرت تقنيات الكشف عن السلوك البصري، باستخدام أدوات الرؤية الحاسوبية، دقة في الكشف عن علامات على وجود أزمة.

هذه النتائج تتوافق مع السرد الأوسع في الأدبيات التي تُسلط الضوء على الإمكانيات التحويلية للذكاء الاصطناعي في مجال الصحة النفسية، حيث إن أنظمة الذكاء الاصطناعي تُقدم رؤى واعدة لتحسين إمكانية الحصول على الرعاية. ومن خلال استخلاص الرؤى من معلومات مُعقدة ومتعددة العوامل، تُتيح تحليلًا وتقييمًا سريعين؛ ما يُساعد في اتخاذ القرارات. ومع ذلك، فإن هذا التقدم يصاحبه العديد من القيود والتحديات التي تلزم وضع إجراءات صارمة، لضمان استخدام تقنيات الذكاء الاصطناعي بشكل مسؤول وأخلاقي، يحفظ حقوق الأفراد، ويعزز مبدأ العدالة، ويحقق أقصى قدر من الفوائد مع الحد من الآثار السلبية المحتملة.

وقد حدد تحليل SWOT العديد من نقاط الضعف الرئيسية التي تحد من فائدة النماذج الحالية المستخدمة في الكشف عن الأفكار الانتحارية عبر الإنترنت، حيث تثير الممارسات الراهنة مخاوف أخلاقية عدة. ويتمثل التحدي الأبرز في التمثيل المحدود للموضوعات الرئيسية والمتنوعة في مجموعات البيانات، واختلال توازن البيانات الذي قد يؤدي إلى ضعف الأداء عبر مختلف الفئات الديمغرافية والفئات المهمشة، ولاسيما في بيانات اللغة العربية، التي لوحظ أنها محدودة. وتُشكل هذه القيود في التمثيل الثقافي عائقاً أمام تحقيق التنفيذ المسؤول. كذلك، فإن الاستخدام الشائع للنماذج الغامضة أو المُبهمة التي يصعب تفسيرها يعيق قبول الأطباء والمعالجين النفسيين لها، حيث إن التفسير السريري مهم في تحديد استخدام هذه النماذج. بالإضافة إلى ذلك، تعتمد النماذج الحالية على البيانات المُستخرجة من وسائل التواصل الاجتماعي، وهي معرضة لعوامل عدة قد تؤدي إلى كشف غير دقيق للبيانات، مثل صعوبة فهم تراكيب اللغة العامية، والسخرية، والسياق. بالإضافة إلى ذلك، فإن الاعتماد على أدوات لم يتم اختبارها والتحقق من موثوقيتها، من خلال دراسات علمية محكمة، قد يؤدي إلى مشكلات في تحديد الأفراد الذين يحتاجون إلى مساعدة عاجلة. وقد تتسبب هذه الأدوات غير الموثوق بها في تقديم نتائج غير دقيقة تؤثر في القرارات المتخذة بشأن التدخل والعلاج.

5- التوصيات

بناءً على نتائج المراجعة والتحليل الذي تم إجراؤه، تقترح الباحثة توصيات عدة لتعزيز الاستخدام المسؤول والفعال للذكاء الاصطناعي للكشف عن الأفكار الانتحارية:

- **إعطاء الأولوية لتنوع البيانات والحساسية الثقافية:** يجب أن تركز الأبحاث المستقبلية على جمع وتنظيم مجموعات بيانات متنوعة تمثل خلفيات لغوية وثقافية واجتماعية واقتصادية مختلفة. كما يجب تصميم توليد البيانات الاصطناعية بعناية لتجنب تعزيز التحيزات القائمة.
- **تعزيز شفافية النماذج وقابليتها للتفسير:** يجب على مطوري النماذج إعطاء الأولوية لابتكار تقنيات ذكاء اصطناعي قابلة للتفسير، توفر رؤية ثاقبة لعمليات صنع القرار. ويمكن للشفافية أن تعزز الثقة بين الأطباء والمعالجين والمرضى، وتمكّن من اتخاذ قرارات استخدام مدروسة.
- **تطوير أطر أخلاقية متينة:** يجب أن تتناول الإرشادات خصوصية البيانات، والموافقة المستنيرة، والعدالة الخوارزمية، وهناك حاجة إلى مسارات تنظيمية واضحة للإشراف على نشر أدوات الصحة النفسية المدعومة بالذكاء الاصطناعي، مع الموازنة بين الابتكار وسلامة المرضى والنزاهة الأخلاقية.
- **دمج الرقابة البشرية والخبرة السريرية:** يجب استخدام الذكاء الاصطناعي بوصفه أداة داعمة بدلاً من أن يكون بديلاً عن الحكم البشري. كما يجب إشراك الأطباء والمعالجين النفسيين في تطوير النماذج، والتحقق من صحتها، والمراقبة المستمرة؛ لضمان توافقها مع أفضل الممارسات الأخلاقية والسريرية.

- الاستثمار في دراسات التحقق الطولية: تُعدّ الدراسات الطولية ضرورية لتقييم فعالية أدوات الذكاء الاصطناعي وقابليتها للتوسع على المدى الطويل في بيئات العالم الحقيقي. ويجب أن تُقيم هذه الدراسات النتائج المتعلقة بالتفكير والسلوك الانتحاري والصحة النفسية بشكل عام.
- تسهيل المشاركة العامة والتثقيف: يمكن لحملات التوعية العامة أن تُعزز فهم دور الذكاء الاصطناعي في الصحة النفسية، وتُقلل من وصمة العار، وتُشجع على الحوار المفتوح حول الآثار الأخلاقية. كما أن إشراك أصحاب المصلحة في المجتمع في صنع القرار من شأنه أن يُعزز الثقة والقبول.

6- الخلاصة

تؤكد هذه المراجعة أن الذكاء الاصطناعي لديه القدرة على إحداث ثورة في استراتيجيات الوقاية من الانتحار، من خلال تسهيل الكشف المبكر والتدخل الشخصي. ومع ذلك، فلا تزال هناك تحديات كبيرة تتعلق بجودة البيانات وشفافية النماذج والنزاهة الأخلاقية، التي يجب معالجتها لتحقيق هذه الإمكانيات على أكمل وجه. وللمضي قدماً، فإن تضافر الجهود عبر التخصصات والقطاعات المختلفة يُعدّ أمراً حيوياً لتطوير أدوات الذكاء الاصطناعي المسؤولة والعادلة والفعالة، التي تدعم متخصصي الصحة العقلية وتحمي الفئات السكانية الضعيفة. ومع استمرار تطور الذكاء الاصطناعي، فإن دمجها في رعاية الصحة النفسية يُعدّ حجر الزاوية لمبادرات مبتكرة وأخلاقية تهدف إلى التوعية والوقاية من الانتحار، وبالتالي تحقيق انخفاض في معدلاته وإنقاذ الأرواح.

المراجع

- Aejaz, A., Owais, M. Z., & Kalyani, G. (2023). Content analysis of messages in social networks and identification of suicidal types. *International Research Journal of Modernization in Engineering Technology and Science*, 5(06), 1952–1955.
- Al-Remawi, M., Agha, A. A., Al-Akayleh, F., Aburub, F., & Abdel-Rahem, R. A. (2024). Artificial intelligence and machine learning techniques for suicide prediction: Integrating dietary patterns and environmental contaminants. *Heliyon*, 10(24), e40925. <https://doi.org/10.1016/j.heliyon.2024.e40925>
- Alatawi, H., Abudalfa, S., & Luqman, H. (2024). Empirical analysis for detecting Arabic online suicidal ideation. *Procedia Computer Science*, 244, 143–150.
- Alambo, A., Gaur, M., Lokala, U., Kursuncu, U., Thirunarayan, K., Gyrard, A., ... & Pathak, J. (2019). Question answering for suicide risk assessment using Reddit. In *2019 IEEE 13th International Conference on Semantic Computing (ICSC)* (pp. 468–473). IEEE. <https://doi.org/10.1109/ICOSC.2019.8665525>
- Baethge, C., Goldbeck-Wood, S., & Mertens, S. (2019). SANRA—a scale for the quality assessment of narrative review articles. *Research Integrity and Peer Review*, 4(5). <https://doi.org/10.1186/s41073-019-0064-8>
- Basyouni, A., Abdulkader, H., Elkilani, W. S., Alharbi, A., Xiao, Y., & Ali, A. H. (2024). A suicidal ideation detection framework on social media using machine learning and genetic algorithms. *IEEE Access*, 12, 124816–124833. <https://doi.org/10.1109/ACCESS.2024.3454796>

- Bhandarkar, A. R., Arya, N., Lin, K. K., North, F., Duvall, M. J., Miller, N. E., & Pecina, J. L. (2023). Building a natural language processing artificial intelligence to predict suicide-related events based on patient portal message data. *Mayo Clinic Proceedings: Digital Health*, 1(4), 510–518.
- Blumenthal, S. J., & Kupfer, D. J. (1988). Overview of early detection and treatment strategies for suicidal behavior in young people. *Journal of Youth and Adolescence*, 17(1), 1-23.
- Bose, A., Dutta, B., Taramol, K., Sarkar, M., Chatterjee, S., & Majumdar, A. (2024). Ethical considerations in the era of artificial intelligence: Implications for mental health, well-being, and digital health. *South Eastern European Journal of Public Health*, 1397–1419. <https://doi.org/10.70135/seejph.vi.2893>
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101. <https://doi.org/10.1191/1478088706qp063oa>
- Braun, V., & Clarke, V. (2019). Reflecting on reflexive thematic analysis. *Qualitative Research in Psychology*, 16(4), 585–599. <https://doi.org/10.1080/14780887.2019.1628806>
- Buddhitha, P., & Inkpen, D. (2023). Multi-task learning to detect suicide ideation and mental disorders among social media users. *Frontiers in Research Metrics and Analytics*, 8, 1152535.
- Casu, M., Triscari, S., Battiato, S., Guarnera, L., & Caponnetto, P. (2024). AI chatbots for mental health: A scoping review of effectiveness, feasibility, and applications. *Applied Sciences*, 14(13), 5889. <https://doi.org/10.3390/app14135889>
- Cohen, L. J., Ardalán, F., Yaseen, Z., Yaseen, A. S., & Galynker, I. I. (2018). Suicide crisis syndrome mediates the relationship between long-term risk factors and lifetime suicidal phenomena. *Suicide and Life-Threatening Behavior*, 48(6), 613-623.
- Fitzpatrick, K. K., Darcy, A., & Vierhile, M. (2017). Delivering cognitive behavior therapy to young adults with symptoms of depression and anxiety using a fully automated conversational agent (Woebot): A randomized controlled trial. *JMIR Mental Health*, 4(2), e19.
- Ghanadian, H., Nejadgholi, I., & Al Osman, H. (2024). Socially aware synthetic data generation for suicidal ideation detection using large language models. *IEEE Access*, 12, 14350–14363. <https://doi.org/10.1109/ACCESS.2024.3454796>
- Ghanadian, H., Nejadgholi, I., & Al Osman, H. (2025). Improving suicidal ideation detection in social media posts: Topic modeling and synthetic data augmentation approach. *JMIR Formative Research*, 9(1), e63272.
- Grygas, M., Filipe, J., & Martinho, C. (2020). AI and chatbots as mental health interventions—Real opportunities and challenges. *Behavioral Sciences*, 10(4), 58.
- Hawton, K., & van Heeringen, K. (2009). Suicide. *The Lancet*, 373(9672), 1372-1381.
- Hisham, N. A. Z., Paujah, N. R., Ismail, S. M., Razak, F. A., & Baharin, S. K. B. (2023). Suicide prediction using machine learning techniques and interactive visualization of suicide information. *International Journal of Academic Research in Business and Social Sciences*, 13(5), 1905–1917.
- Kline, S., & Sabri, T. (2023). The decriminalisation of suicide: A global imperative. *Lancet Psychiatry*, 10(4), 240–242.
- Li, X., Chen, F., & Ma, L. (2024). Exploring the potential of artificial intelligence in adolescent suicide prevention: Current applications, challenges, and future directions. *Psychiatry*, 87(1), 7–20.

- Mann, J. J., Waternaux, C., Haas, G. L., & Malone, K. M. (1999). Toward a clinical model of suicidal behavior in psychiatric patients. *American Journal of Psychiatry*, 156(2), 181-189.
- Mansoor, M. A., & Ansari, K. H. (2024). Early detection of mental health crises through artificial intelligence-powered social media analysis: A prospective observational study. *Journal of Personalized Medicine*, 14(9), 958. <https://doi.org/10.3390/jpm14090958>
- Mohri, M., Rostamizadeh, A., & Talwalkar, A. (2018). *Foundations of Machine Learning*. Cambridge, MA, USA: MIT Press.
- O'Connor, R. C., & Nock, M. K. (2014). The psychology of suicidal behavior. *The Lancet Psychiatry*, 1(1), 73-85. [https://doi.org/10.1016/s2215-0366\(14\)70222-6](https://doi.org/10.1016/s2215-0366(14)70222-6)
- Onie, S., Li, X., Glastonbury, K., Hardy, R. C., Rakusin, D., Wong, I., ... & Larsen, M. E. (2023). Understanding and detecting behaviours prior to a suicide attempt: A mixed-methods study. *Australian & New Zealand Journal of Psychiatry*, 57(7), 1016-1022.
- Oquendo, M. A., Sullivan, G. M., Sudol, K., Baca-Garcia, E., Stanley, B. H., Sublette, M. E., ... & Mann, J. J. (2014). Toward a biosignature for suicide. *American Journal of Psychiatry*, 171(12), 1259-1277.
- Parsapoor, M., Koudys, J. W., & Ruocco, A. C. (2023). Suicide risk detection using artificial intelligence: The promise of creating a benchmark dataset for research on the detection of suicide risk. *Frontiers in Psychiatry*, 14, 1186569.
- Qadeer, S., Memon, K., & Palli, G. H. (2024). Predicting depression and suicidal tendencies by analyzing online activities using machine learning in android devices. *Mehran University Research Journal of Engineering & Technology*, 43(1), 213–224.
- Reece, A. G., Reagan, A. J., Lix, K. L., et al. (2017). Forecasting the onset and course of mental illness with social media data. *Scientific Reports*, 7, 45170.
- Runeson, B., Odeberg, J., Pettersson, A., Edman, G., Jansson, L., & Lindh, A. U. (2017). Instruments for the assessment of suicide risk: A systematic review evaluating the certainty of the evidence. *PLoS One*, 12(7), e0180292.
- Saeidnia, H. R., Ghasem Hashemi Fotami, S., Lund, B., & Ghiasi, N. (2024). Ethical considerations in artificial intelligence interventions for mental health and well-being: Ensuring responsible implementation and impact. *Social Sciences*, 13(3), 381. <https://doi.org/10.3390/socsci13070381>
- Saeidnia, H. R., Hashemi Fotami, S. G., Lund, B., & Ghiasi, N. (2024). Ethical considerations in artificial intelligence interventions for mental health and well-being: Ensuring responsible implementation and impact. *Social Sciences*, 13(7), 381. <https://doi.org/10.3390/socsci13070381>
- Samruddhi, N. C., & Vaidya, A. K. (2024). Artificial Intelligence in Mental Health: Challenges and Opportunities. *International Journal of Nursing Education and Research*, 12(4). Available from <https://www.proquest.com/scholarly-journals/artificial-intelligence-mental-health-challenges/docview/3158589326>
- Sambara, B. S., Archana, K., Reddy, S. I. S., Basha, S. M. J., & Karishma, S. (2025). Data augmentation for cognitive behavioral therapy: Leveraging ERNIE language models using artificial intelligence. In *2025 3rd International Conference on Intelligent Data Communication Technologies and Internet of Things (IDCIoT)* (pp. 204–209). IEEE.

- Smith, W. R., Appelbaum, P. S., Lebowitz, M. S., Gülöksüz, S., Calkins, M. E., Kohler, C. G., & Barzilay, R. (2023). The ethics of risk prediction for psychosis and suicide attempt in youth mental health. *The Journal of Pediatrics*, 263, 113583.
- Solaiman, B., Malik, A., & Ghuloum, S. (2023). Monitoring mental health: Legal and ethical considerations of using artificial intelligence in psychiatric wards. *American Journal of Law & Medicine*, 49(2-3), 250-266.
- Sogancioglu, G., Mosteiro, P., Salah, A. A., Scheepers, F., & Kaya, H. (2024). Fairness in AI-based mental health: Clinician perspectives and bias mitigation. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, 7(1), 1390-1400. <https://doi.org/10.1609/aies.v7i1.31732>
- Stanley, B., & Brown, G. K. (2012). Safety planning intervention: A brief intervention to mitigate suicide risk. *Cognitive and Behavioral Practice*, 19(2), 256-264.
- Su, R., John, J. R., & Lin, P. I. (2023). Machine learning-based prediction for self-harm and suicide attempts in adolescents. *Psychiatry Research*, 328, 115446. <https://doi.org/10.1016/j.psychres.2023.115446>
- Warrier, U., Warrier, A., & Khandelwal, K. (2023). Ethical considerations in the use of artificial intelligence in mental health. *Egyptian Journal of Neurology, Psychiatry and Neurosurgery*, 59, 139. <https://doi.org/10.1186/s41983-023-00735-2>
- World Health Organization: WHO. (2025, March 25). Suicide. Retrieved June 26, 2025, from <https://www.who.int/news-room/fact-sheets/detail/suicide>
- Yeskuatov, E., Chua, S.-L., & Foo, L. K. (2025). Detecting suicidal ideations on Reddit with transformer models. In *Artificial Intelligence and Human-Computer Interaction* (pp. 113–119). IOS Press.

ملحق 1

الدراسات المدرجة في المراجعة

#	Citation	Title	Link
1	Aejaz, A., Owais, M. Z., & Kalyani, G. (2023). CONTENT ANALYSIS OF MESSAGES IN SOCIAL NETWORKS AND IDENTIFICATION OF SUICIDAL TYPES. <i>International Research Journal of Modernization in Engineering Technology and Science</i> , 5(06), 1952-1955.	CONTENT ANALYSIS OF MESSAGES IN SOCIAL NETWORKS AND IDENTIFICATION OF SUICIDAL TYPES	https://www.suicideinfo.ca/wp-content/uploads/2024/04/Content-Analysis-of-Messages-in-Social-Networks-Identification-of-suicidal-types.pdf
2	Onie, S., Li, X., Glastonbury, K., Hardy, R. C., Rakusin, D., Wong, I., ... & Larsen, M. E. (2023). Understanding and detecting behaviours prior to a suicide attempt: A mixed-methods study. <i>Australian & New Zealand Journal of Psychiatry</i> , 57(7), 1016-1022.	Understanding and detecting behaviours prior to a suicide attempt: A mixed-methods study	https://journals.sagepub.com/doi/full/10.1177/00048674231152159
3	Bhandarkar, A. R., Arya, N., Lin, K. K., North, F., Duvall, M. J., Miller, N. E., & Pecina, J. L. (2023). Building a natural	Building a Natural Language Processing Artificial Intelligence to Predict	https://www.sciencedirect.com/science/article/pii/S2949761223000780

#	Citation	Title	Link
	language processing artificial intelligence to predict suicide-related events based on patient portal message data. <i>Mayo Clinic Proceedings: Digital Health</i> , 1(4), 510-518.	Suicide-Related Events Based on Patient Portal Message Data	
4	Buddhitha, P., & Inkpen, D. (2023). Multi-task learning to detect suicide ideation and mental disorders among social media users. <i>Frontiers in research metrics and analytics</i> , 8, 1152535.	Multi-task learning to detect suicide ideation and mental disorders among social media users	https://www.frontiersin.org/journals/research-metrics-and-analytics/articles/10.3389/frma.2023.1152535/full
5	Smith, W. R., Appelbaum, P. S., Lebowitz, M. S., Gülöksüz, S., Calkins, M. E., Kohler, C. G., ... & Barzilay, R. (2023). The ethics of risk prediction for psychosis and suicide attempt in youth mental health. <i>The Journal of Pediatrics</i> , 263, 113583.	The Ethics of Risk Prediction for Psychosis and Suicide Attempt in Youth Mental Health	https://www.sciencedirect.com/science/article/abs/pii/S0022347623004468
6	Solaiman, B., Malik, A., & Ghuloum, S. (2023). Monitoring mental health: Legal and ethical considerations of using artificial intelligence in psychiatric wards. <i>American journal of law & medicine</i> , 49(2-3), 250-266.	Monitoring Mental Health: Legal and Ethical Considerations of Using Artificial Intelligence in Psychiatric Wards	https://www.cambridge.org/core/journals/american-journal-of-law-and-medicine/article/monitoring-mental-health-legal-and-ethical-considerations-of-using-artificial-intelligence-in-psychiatric-wards/C873AF5A4123E3C5224C8AB3E045F942
7	Su, R., John, J. R., & Lin, P. I. (2023). Machine learning-based prediction for self-harm and suicide attempts in adolescents. <i>Psychiatry research</i> , 328, 115446.	Machine learning-based prediction for self-harm and suicide attempts in adolescents	https://www.sciencedirect.com/science/article/pii/S0165178123003967
8	Hisham, N. A. Z., Paujah, N. R., Ismail, S. M., Razak, F. A., & Baharin, S. K. B. (2023). Suicide Prediction Using Machine Learning Techniques and Interactive Visualization of Suicide Information. <i>International Journal of Academic Research in Business and Social Sciences</i> , 13(5), 1905-1917.	Suicide Prediction Using Machine Learning Techniques and Interactive Visualization of Suicide Information	https://www.researchgate.net/profile/Nor-Rashidah-Paujah-Ismail/publication/373895131_Suicide_Prediction_Using_Machine_Learning_Techniques_and_Interactive_Visualization_of_Suicide_Information/links/6501cbe9a2e39316ce0838dc/Suicide-Prediction-Using-Machine-Learning-Techniques-and-Interactive-Visualization-of-Suicide-Information.pdf
9	Parsapoor, M., Koudys, J. W., & Ruocco, A. C. (2023). Suicide risk detection using artificial intelligence: the promise of creating a benchmark dataset for research on the detection of suicide risk. <i>Frontiers in psychiatry</i> , 14, 1186569.	Suicide risk detection using artificial intelligence: the promise of creating a benchmark dataset for research on the detection of suicide risk	https://www.frontiersin.org/journals/psychiatry/articles/10.3389/fpsy.2023.1186569/full
10	Saeidnia, Hamid Reza, Seyed Ghasem Hashemi Fotami, Brady Lund, and Nasrin Ghiasi. 2024. Ethical Considerations in Artificial Intelligence Interventions for Mental Health and Well-Being: Ensuring Responsible Implementation and Impact. <i>Social Sciences</i> 13: 381. https://doi.org/10.3390/socsci13070381	Ethical Considerations in Artificial Intelligence Interventions for Mental Health and Well-Being: Ensuring Responsible Implementation and Impact	https://www.mdpi.com/2076-0760/13/7/381
11	Bose, A., Dutta, B., Taramol, K., Sarkar, M., Chatterjee, S., & Majumdar, A. (2024). Ethical considerations in the era of artificial intelligence: Implications for mental health,	Ethical CONSIDERATIONS IN THE ERA OF ARTIFICIAL INTELLIGENCE: IMPLICATIONS FOR MENTAL HEALTH, WELL-	https://www.seejph.com/index.php/seejph/article/view/2893

#	Citation	Title	Link
	well-being, and digital health. South Eastern European Journal of Public Health, 1397–1419. https://doi.org/10.70135/seejph.vi.2893	BEING, AND DIGITAL HEALTH	
12	Al-Remawi, M., Ali Agha, A. S., Al-Akayleh, F., Aburub, F., & Abdel-Rahem, R. A. (2024). Artificial intelligence and machine learning techniques for suicide prediction: Integrating dietary patterns and environmental contaminants. <i>Heliyon</i> , 10(24). https://doi.org/10.1016/j.heliyon.2024.e40925	Artificial intelligence and machine learning techniques for suicide prediction: Integrating dietary patterns and environmental contaminants	https://www.cell.com/heliyon/fulltext/S2405-8440(24)16956-3?_returnURL=https%3A%2F%2Flinkinghub.elsevier.com%2Fretrieve%2Fpii%2FS2405844024169563%3Fshallow%3Dtrue
13	Alatawi, H., Abudalfa, S., & Luqman, H. (2024). Empirical Analysis for Detecting Arabic Online Suicidal Ideation. <i>Procedia Computer Science</i> , 244, 143-150.	Empirical Analysis for Detecting Arabic Online Suicidal Ideation	https://www.sciencedirect.com/science/article/pii/S1877050924029880
14	Samruddhi, N. C., & Vaidya, A. K. (2024). Artificial Intelligence in Mental Health: Challenges and Opportunities. <i>International Journal of Nursing Education and Research</i> , 12(4) https://www.proquest.com/scholarly-journals/artificial-intelligence-mental-health-challenges/docview/3158589326/se-2	Artificial Intelligence in Mental Health: Challenges and Opportunities	https://www.researchgate.net/publication/388807099_Artificial_Intelligence_in_Mental_Health_Challenges_and_Opportunities
15	Mansoor, M. A., & Ansari, K. H. (2024). Early Detection of Mental Health Crises through Artificial-Intelligence-Powered Social Media Analysis: A Prospective Observational Study. <i>Journal of personalized medicine</i> , 14(9), 958. https://doi.org/10.3390/jpm14090958	Early Detection of Mental Health Crises through Artificial-Intelligence-Powered Social Media Analysis: A Prospective Observational Study	https://www.mdpi.com/2075-4426/14/9/958
16	Casu, M., Triscari, S., Battiato, S., Guarnera, L., & Caponnetto, P. (2024). AI Chatbots for Mental Health: A Scoping Review of Effectiveness, Feasibility, and Applications. <i>Applied Sciences</i> , 14(13), 5889. https://doi.org/10.3390/app14135889	AI Chatbots for Mental Health: A Scoping Review of Effectiveness, Feasibility, and Applications	https://mirkocasu.github.io/assets/pdf/applsci-14-05889.pdf
17	Basyouni, A., Abdulkader, H., Elkilani, W. S., Alharbi, A., Xiao, Y., & Ali, A. H. (2024). A Suicidal Ideation Detection Framework on Social Media Using Machine Learning and Genetic Algorithms. <i>IEEE Access</i> , 12, 124816-124833. https://doi.org/10.1109/ACCESS.2024.3454796	A Suicidal Ideation Detection Framework on Social Media Using Machine Learning and Genetic Algorithms	https://ieeexplore.ieee.org/abstract/document/10666498
18	Li, X., Chen, F., & Ma, L. (2024). Exploring the potential of artificial intelligence in adolescent suicide prevention: current applications, challenges, and future	Exploring the Potential of Artificial Intelligence in Adolescent Suicide Prevention: Current Applications, Challenges, and Future Directions	https://www.tandfonline.com/doi/abs/10.1080/00332747.2023.2291945

#	Citation	Title	Link
	directions. <i>Psychiatry</i> , 87(1), 7-20.		
19	Ghanadian, H., Nejadgholi, I., & Al Osman, H. (2024). Socially aware synthetic data generation for suicidal ideation detection using large language models. <i>IEEE Access</i> , 12, 14350-14363.	Socially Aware Synthetic Data Generation for Suicidal Ideation Detection Using Large Language Models	https://ieeexplore.ieee.org/abstract/document/10413447
20	Qadeer, S., Memon, K., & Palli, G. H. (2024). Predicting depression and suicidal tendencies by analyzing online activities using machine learning in android devices. <i>Mehran University Research Journal Of Engineering & Technology</i> , 43(1), 213-224.	Predicting depression and suicidal tendencies by analyzing online activities using machine learning in android devices	https://search.informit.org/doi/abs/10.3316/informit.T2024041200008600682567662
21	Ghanadian H, Nejadgholi I, Al Osman H. (2025). Improving Suicidal Ideation Detection in Social Media Posts: Topic Modeling and Synthetic Data Augmentation Approach. <i>JMIR Formative Research</i> , 9(1):e63272.	Improving Suicidal Ideation Detection in Social Media Posts: Topic Modeling and Synthetic Data Augmentation Approach	https://formative.jmir.org/2025/1/e63272
22	Yeskuatov, E., Chua, S.-L., & Foo, L. K. (2025). Detecting Suicidal Ideations on Reddit with Transformer Models. In <i>Artificial Intelligence and Human-Computer Interaction</i> (pp. 113-119). IOS Press.	Detecting Suicidal Ideations on Reddit with Transformer Models	https://ebooks.iospress.nl/volumearticle/72078
23	B. Sambana, K. Archana, S. I. S. Reddy, S. M. J. Basha and S. Karishma, "Data Augmentation for Cognitive Behavioral Therapy: Leveraging ERNIE Language Models using Artificial Intelligence," 2025 3rd International Conference on Intelligent Data ommunication Technologies and Internet of Things (IDCIoT), Bengaluru, India, 2025, pp. 204-209, doi: 10.1109/IDCIoT64235.2025.10914873.	Data Augmentation for Cognitive Behavioral Therapy: Leveraging ERNIE Language Models using Artificial Intelligence	https://ieeexplore.ieee.org/abstract/document/10914873