

أبحاث منظمة كريست CREST Research: كيف الذكاء الاصطناعي يُغيّر مهنة اختبارات الاختراق Penetration Testing؟

لقد أصبح الذكاء الاصطناعي جزءاً لا يتجزأ من عملية تقديم خدمات الأمن السيبراني. ومع ازدياد فهمنا له ونضوج أدواته، ازدادت ثقة مزودي الخدمات في إمكانية استخدامه في ممارساتهم اليومية. وقد أدى ذلك إلى خلق فرص وكذلك مخاطر جديدة في قطاع قائم على ضمان جودة المستويات الأمنية. والسؤال المطروح أمام مزودي خدمات الأمن السيبراني هو: كيف يُمكن تسخير الذكاء الاصطناعي لتحسين القدرات مع الحفاظ على الثقة، لا سيما في الخدمات عالية الموثوقية مثل اختبارات الاختراق؟

تستكشف كريست هذا التحدي الإداري من خلال مبادرة عالمية حول التنبؤ المسؤول للذكاء الاصطناعي في الأمن السيبراني، والتي تشمل بحثاً مستقلاً وحواراً مع أصحاب المصلحة الرئيسيين. ونهدف إلى مساعدة القطاع على فهم الوضع الراهن، وتحديد ملامح سبل التطبيق المسؤول، وتطوير نماذج ضمان الموثوقية التي تلبي احتياجات السوق.

كيف يُستخدم الذكاء الاصطناعي في اختبارات الاختراق اليوم؟

يقدم تقرير CREST الأخير، ["داخل اختبارات الاختراق: كيف يُعيد الذكاء الاصطناعي تشكيل الممارسة"](#)، رؤى قيمة حول كيفية استخدام الذكاء الاصطناعي في اختبارات الاختراق.

استناداً إلى بحث أصلي شمل 62 مزوداً لخدمات الأمن السيبراني في 19 دولة، تكشف النتائج كيف أصبح استخدام الذكاء الاصطناعي من ضمن الممارسة المعيارية لدى غالبية المشاركين، حيث أن 69% منهم يستخدمونه في عمليات اختبارات الاختراق وذلك مع زيادة استخدامهم له بنسبة 76% خلال العام الماضي.

حالياً، يستخدم معظم المشاركين الذكاء الاصطناعي في أجزاء الاختبارات التي تتضمن كميات كبيرة من المعلومات، مثل صياغة التقارير وتلخيصها وتحليل البيانات. ويستطيع الممارسون توليد مخرجات هذا العمل والتحكم بها باستخدام نماذج التعلم الآلي العامة مثل Gemini و ChatGPT. وأما بالنسبة لمنصات اختبارات الاختراق المتخصصة والمُدعمة بالذكاء الاصطناعي، فقد تم اعتمادها على نطاق أضيق، وهو ما قد يعكس الطلب على مستويات أعلى من الدقة والأدلة لأنشطة الاختبارات الأساسية.

مع تحسن الأدوات ونضوج الممارسة، يخطط معظم مزودي الخدمات الذين شملهم الاستطلاع لزيادة استخدام الذكاء الاصطناعي. ومن المرجح أن يؤدي ذلك إلى مزيد من التدقيق في كيفية فحص الاختبارات وتوثيقها والتحقق من صحتها. لذلك ستكون هناك حاجة إلى أطر حوكمة لمساعدة مقدمي الخدمات على تلبية هذه التوقعات، مما سيمنح العملاء الثقة وسيدعم المختصين المسؤولين عن هذه الأعمال.

ما يقوله الممارسون عن الذكاء الاصطناعي في خدمات الأمن السيبراني

في فبراير 2026، ناقشنا نتائج بحثنا في اجتماع "طاولة مستديرة" مع خبراء من قطاع الأمن السيبراني في أبوظبي. شارك كبار الخبراء من الشركات الأعضاء وفريق كريست تجاربهم في استخدام الذكاء الاصطناعي في العمليات والخدمات. وتؤكد

ردودهم ما أظهره البحث بأن الذكاء الاصطناعي جزء لا يتجزأ من تقديم خدمات الأمن السيبراني، إلا أن الممارسين يطبقونه بشكل انتقائي.

قدم المشاركون في اجتماع الطاولة المستديرة رؤى واقعية حول التطبيقات العملية للذكاء الاصطناعي. ففي اختبارات الاختراق، يستخدم في مهام المراحل المبكرة مثل الاستطلاع والحصر ومراجعة الإعدادات، بالإضافة إلى تجهيز التقارير وضمان الجودة. تستخدمه بعض المؤسسات لتقليل المدة الزمنية المطلوبة للاستعداد لتقديم التقارير ومن ثم إعدادها. في بعض الحالات تم تقليل المدة من أسابيع إلى أيام. وفي بيئات مراكز عمليات الأمن، يساعد الذكاء الاصطناعي فرق العمل على كبح آثار الأحداث المتكررة، وتقليل الإنذارات الكاذبة، وإضافة سياق للتنبيهات، واتخاذ قرار توجيه الحالات إلى كبار المحللين بشكل أسرع.

تتحقق معظم مكاسب كفاءة الذكاء الاصطناعي من المهام التي يمكن للمختصين فيها مراجعة النتائج وتوثيقها. وقد أبدوا حذرًا أكبر بشأن استخدام الذكاء الاصطناعي في أنشطة الاختبارات الأساسية، وبيئات الإنتاج، وإعداد التقارير عالية الموثوقية. وقد وصفت مناقشة المائدة المستديرة هذا الموقف بأنه "واقعية ناضجة"، أي إقرار بأن الذكاء الاصطناعي يحمل فوائد واضحة، مع وجود مجالات مخاطر تتطلب إدارة دقيقة.

"أنا بمثابة حاجز للأمان للذكاء الاصطناعي".

رأي أحد المشاركين في "اجتماع المائدة المستديرة" في أبوظبي في فبراير 2026

يبرز هذا التركيز على الإشراف البشري أهمية الحوكمة مع تزايد استخدام الذكاء الاصطناعي. وتُشكل المخاوف بشأن ضمان الموثوقية الأمنية حاليًا عائقًا أمام التوسع في استخدامه، لذا يبقى نموذج الاستخدام في الوقت الراهن قائمًا على الاستفادة من على دعم الذكاء الاصطناعي تحت الإشراف البشري.

لا تمثل كفاءة الذكاء الاصطناعي سوى جزء من جدوى استخدامه

غالبًا ما يُشار إلى انخفاض التكاليف كأحدى مزايا الذكاء الاصطناعي. ومع ذلك، فبينما يُمكنه تقليل الجهد اليدوي، تشير أبحاثنا إلى أن هذا لا يُترجم مباشرةً إلى وفورات. عمليًا، يُحقق الذكاء الاصطناعي أفضل النتائج عندما يُدمج ضمن عمليات مُحكمة، مع معالجة بيانات واضحة، وخطوات مراجعة مُحَددة، وموافقة مسؤولة. يتطلب هذا استثمارًا في مهندسي الذكاء الاصطناعي، والمطورين، وعلماء البيانات، واختبارات التحقق، وأنظمة مراجعة قوية، مما قد يزيد التكاليف على المدى القصير.

بالنسبة للعديد من مُقدمي الخدمات، تُستخدم مكاسب الكفاءة لتحسين جودة الأداء بدلاً من خفض أسعار الخدمات. من خلال تقليل الوقت المُستغرق في المهام الروتينية، يُمكن للذكاء الاصطناعي أن يُوفر للمختصين قدرة أكبر على إجراء تحقيقات أعمق، وتغطية أوسع، وإنجاز أسرع. يُشير هذا إلى أن القيمة التجارية للذكاء الاصطناعي تكمن في تحسين القدرات وتقديم خدمات أكثر اتساقًا، وليس بالضرورة في خفض تكلفة الاختبار.

كما يجب أن يصمد العمل المدعوم بالذكاء الاصطناعي أمام التدقيق. أفاد المشاركون في الاستطلاعات ومناقشات المائدة المستديرة بوجود قيود مُستمرة تتعلق بجودة المُخرجات المُتغيرة، ومحدودية التفسير، والثقة المُفرطة، وأعمال التحقق اللازمة

لفحص المُخرجات وإعادة إنتاجها وتوثيقها. تُشكّل التصورات الخاطئة للذكاء الاصطناعي مخاطر تجارية وسمعة، لا سيما في القطاعات الخاضعة للتنظيم.

لهذا السبب، يبقى الممارسون المهرة ركيزة أساسية في تقديم الخدمات. يمتلك الذكاء الاصطناعي القدرة على مساعدة الشركات على التوسع، ولكن بشرط أن تواكب الجودة القدرة، وأن تصبح أدوات الاختبار أكثر موثوقية.

تعتمد الثقة في الاختبارات المدعومة بالذكاء الاصطناعي على الحوكمة.

تُعَدّ الحوكمة مسألة تجارية وتقنية في آنٍ واحد. مع ازدياد استخدام الذكاء الاصطناعي في مجالات اختبارات الاختراق، سيرغب العملاء في الحصول على أدلة أوضح حول كيفية إدارة مقدمي الخدمات لاستخدامه والمخاطر المرتبطة به. وقد توقع المشاركون في الاستطلاع هذا التدقيق، حيث يتوقع 85% منهم أن يسأل العملاء عن كيفية استخدام الذكاء الاصطناعي في الاختبارات.

سلطت جلسة النقاش الضوء على مجالات المخاطر الرئيسية، مثل عدم كفاية التوثيق، واستخدام أغلفة بسيطة (طبقة برمجية رقيقة ومرنة) مبنية حول نماذج الذكاء الاصطناعي من الجهات الخارجية، وضعف سجلات التدقيق، وعدم وضوح المسؤولية. أوضح الحضور كيف تتزايد المخاطر عندما تكون أدوات الذكاء الاصطناعي غير موثقة بشكل جيد، أو عندما تكون مبنية دون هندسة برمجيات منضبطة، أو متصلة بنماذج خارجية، أو مسموح لها بتنظيم الأنشطة دون ضوابط واضحة. هذا يجعل من الصعب إثبات كيفية إنتاج المخرجات، وتحديد خطوط المسؤولية.

تُتيح الحوكمة الواضحة لمقدمي الخدمات طريقة عملية لمعالجة مخاوف العملاء، من خلال عمليات موثقة تُبين كيفية دمج استخدام الذكاء الاصطناعي في عملية التسليم، وكيفية سير خطوات المراجعة، وكيفية التحكم في البيانات، وكيفية اعتماد القرارات النهائية. ومع نمو الذكاء الاصطناعي، ستصبح القدرة على إثبات الحوكمة عنصراً أساسياً في نجاح الممارسات المتبعة.

ماذا يخبئ المستقبل للذكاء الاصطناعي في اختبارات الاختراق؟

بالنسبة لمقدمي خدمات اختبارات الاختراق، فإن الشفافية مطلب لاكتساب ثقة العملاء. وسيكون هنالك توقعات لتوضيح الغايات من استخدام الذكاء الاصطناعي، وكيفية التحقق منه من قبل البشر، ومدى دقة النتائج النهائية عند خضوعها للتدقيق.

كشف بحثنا حول الذكاء الاصطناعي في اختبارات الاختراق عن طلب كبير على دعم الحوكمة في القطاع. ولتلبية هذه الحاجة، نعمل على توثيق أفضل الممارسات الناشئة كأساس للمعايير. ونقطة انطلاق، اقترحت "مائدة أبوظبي المستديرة" تسعة مبادئ للأمن السيبراني المسؤول المدعوم بالذكاء الاصطناعي، بما في ذلك المساءلة والشفافية والتحقق والتحكم في البيانات والتطوير الآمن. وفي حال اعتماد هذه المبادئ على نطاق واسع كميثاق، فإنها ستُسهم في توحيد ممارسات الذكاء الاصطناعي، و ستُسهم في مساعدة القطاع على تطوير حوكمة قوية مع تزايد اعتماد الذكاء الاصطناعي.

وستواصل كريست دعم قطاع الأمن السيبراني من خلال الأبحاث والإرشادات العملية ونماذج ضمان الجودة التي تساعد مزودي الخدمات على إثبات تقديم حلول فعّالة مدعومة بالذكاء الاصطناعي.

حمّل تقرير CREST حول الذكاء الاصطناعي في اختبارات الاختراق للاطلاع على نتائج البحث كاملةً.