



- (XAI) الدورة التدريبية: الذكاء الاصطناعي القابل للتفسير فهم سلوك نموذج الذكاء الاصطناعي

اغسطس ٢٠٢٦ ٢٨ - ٢٤

فيينا

(للشخص الواحد) € ٥٧٠٠

Ref: #AI9816_247077



مقدمة الدورة التدريبية / لمحة عامة





سلوك نموذج الذكاء التدريبية المتخصصة حول الذكاء الاصطناعي القابل يقدم BIG BEN Training Center هذه الدورة ثقة التعلم الآلي، والمطورين، والباحثين، ومديري الاصطناعي، وهي مصممة لعلماء البيانات، ومهندسي للتفسير (XAI) - فهم الآلي (Machine Learning) أكبر في أنظمة الذكاء الاصطناعي (AI)، وفهم كيفية المنتجات، وصناع القرار الذين يسعون لبناء الاصطناعي تعقيداً وانتشاراً في قطاعات حساسة مثل لقراراتها. في عالم تزداد فيه تطبيقات الذكاء اتخاذ نماذج التعلم وإمكانية التتبع أصبح فهم الشفافية (Transparency)، والقابلية الرعاية الصحية، والتمويل، والعدالة الجنائية، إلى تجاوز مفهوم "الصندوق الأسود" (Black Box) أمراً بالغ الأهمية. تهدف الدورة (Traceability) للتفسير (Interpretability)، ستغطي الدورة مفاهيم والتقنيات التي تمكننا من فهم وتفسير التنبؤات لنماذج الذكاء الاصطناعي، واستكشاف الأدوات التنظيمية والأخلاقية، والتقنيات المحلية والعالمية الذكاء الاصطناعي القابل للتفسير، وأهميته والقرارات الصادرة عنها. SHAP، و (Local Interpretable Model-agnostic Explanations) LIME للتفسير (Local and Global Interpretability)، Feature العصبية (Neural Network Interpretability)، وتفسير الشبكات InterpretML، و (Shapley Additive explanations) (Importance) آلي مختلفة، وتقييم مدى قابلية سيتعلم المشاركون كيفية تطبيق هذه التقنيات على وتحديد أهمية الميزات (Importance) إلى أصحاب المصلحة نماذج تعلم



اصطناعي أكثر شفافية ومسؤولية، غير التقنيين. تهدف الدورة إلى تمكين المختصين من التفسير، وتوصيل الرؤى المكتسبة
الدورة من أعمال والامتنال للمتطلبات التنظيمية المتعلقة بالذكاء وتعزيز الثقة في أنظمة الذكاء الاصطناعي، بناء نماذج ذكاء
عن بناء نماذج ذكاء Cynthia Rudin البروفيسور سينثيا رودين (Cynthia Rudin)، الاصطناعي المسؤول. نستلهم في هذه
نماذج الصندوق الأسود اصطناعي قابلة للتفسير بطبيعتها بدلاً من تفسير ، التي تدافع



لأ الفئات المستهدفة / هذه الدورة التدريبية مناسبة

- علماء البيانات^١
- مهندسي التعلم الآلي^١
- مهندسي الذكاء الاصطناعي^١
- محللي البيانات^١
- مديري المنتجات^١
- المطورين في مجال الذكاء الاصطناعي^١
- المهندسين المعماريين للذكاء الاصطناعي^١
- مسؤولي الامتثال (Compliance Officers)^١
- مديري المخاطر^١
- الباحثين في الذكاء الاصطناعي^١

القطاعات والصناعات المستهدفة^١:

- الخدمات المالية (تقييم الائتمان، كشف الاحتيال)^١
- الرعاية الصحية (التشخيص، خطط العلاج)^١
- القانون والعدالة (تطبيقات العدالة الجنائية)^١
- التأمين (تقييم المخاطر، تسعير البوليصات)^١
- القطاع الحكومي (صنع القرار، أنظمة الدعم)^١
- التصنيع (الصيانة التنبؤية، مراقبة الجودة)^١
- الخدمات الاستشارية^١
- التكنولوجيا والبرمجيات^١
- البحث والتطوير^١
- الخدمات اللوجستية^١



الأقسام المؤسسية المستهدفة:

- قسم علوم البيانات.
- قسم التعلم الآلي والذكاء الاصطناعي.
- قسم الامتثال القانوني.
- قسم إدارة المخاطر.
- قسم تطوير المنتجات (الخاصة بالذكاء الاصطناعي).
- قسم البحث والتطوير (R&D).
- قسم التدقيق الداخلي.
- قسم الحوكمة والسياسات.
- قسم هندسة البرمجيات.
- قسم التحليلات.

أهداف الدورة التدريبية:

أتقن المهارات التالية: بنهاية هذه الدورة التدريبية، سيكون المتدرب قد



- وأهميته، فهم مفهوم الذكاء الاصطناعي القابل للتفسير
- التتبع، التمييز بين الشفافية، والقابلية للتفسير، وإمكانية
- نماذج التعلم الآلي، تطبيق تقنيات التفسير المحلية (LIME, SHAP) على
- (Interpretability methods) استخدام تقنيات التفسير العالمية (Global)
- تفسير سلوك الشبكات العصبية العميقة،
- نماذج الذكاء الاصطناعي، تحديد أهمية الميزات (Feature Importance) في
- نماذج الذكاء الاصطناعي، (Reliability) تقييم مدى عدالة (Fairness) وموثوقية
- المصلحة، توصيل الرؤى المكتسبة من تفسير النماذج لأصحاب
- الذكاء الاصطناعي، التعامل مع التحديات الأخلاقية والقانونية لنشر
- بناء نماذج ذكاء اصطناعي أكثر شفافية ومسؤولية.

منهجية الدورة التدريبية:



المشاركين من منهجية تدريبية تجمع بين المعرفة النظرية المتعمقة يعتمد BIG BEN Training Center في هذه الدورة على المنهجية محاضرات تفاعلية تستعرض فهم وتطبيق تقنيات الذكاء الاصطناعي القابل والتطبيق العملي المكثف، بهدف تمكين مكثفة حيث والأخلاقية، وأنواع التقنيات المختلفة للتفسير. المفاهيم الأساسية XAI، وأهميتها التنظيمية للتفسير. تشمل (LIME, SHAP, InterpretML) على سيقوم المشاركون باستخدام أدوات ومكتبات XAI تتبع هذه المحاضرات ورش عمل تطبيقية حساسة مثل حاسوبية). سيتم التركيز على دراسات حالة واقعية نماذج تعلم آلي مختلفة (تصنيف، انحدار، رؤية الرائدة (مثل) في بناء الثقة والامتثال. تتضمن XAI الرعاية الصحية والخدمات المالية، وكيف يمكن لابرز أهمية قابلية التفسير في قطاعات من نماذج الذكاء الاصطناعي وتقديم رؤى قابلة للتنفيذ. الدورة جلسات عمل جماعي لتحليل وتفسير مخرجات أن يساعد المجال الحيوي والمتنامي. المدربين الخبراء لضمان تطوير مهاراتهم في هذا يتلقى المشاركون تغذية راجعة مفصلة ومنظمة

خريطة المحتوى التدريبي (محاور الدورة التدريبية)

للتفسير (XAI) الوحدة الأولى: مقدمة إلى الذكاء الاصطناعي القابل



- مفهوم الذكاء الاصطناعي القابل للتفسير وأهميته^١
- تحدي "الصندوق الأسود" في نماذج الذكاء الاصطناعي^١
- تعريفات وفروقات^١ الشفافية، والقابلية للتفسير، وإمكانية التتبع:
- القانون^١ أهمية^١ XAI في القطاعات الحساسة (الصحة، التمويل،
- المسؤول^١ المتطلبات التنظيمية والأخلاقية للذكاء الاصطناعي
- أنواع قابلية التفسير: المحلية والعالمية^١
- تطور مجال XAI^١ وأبرز الباحثين فيه^١

١) (Interpretability الوحدة الثانية: تقنيات التفسير المحلية (Local)

- أهمية التفسير المحلي لفهم التنبؤات الفردية^١
- (LIME (Local Interpretable Model-agnostic Explanations مقدمة إلى^١
- تطبيق LIME^١ على نماذج التصنيف والانحدار^١
- (SHAP (Shapley Additive explanations مقدمة إلى^١
- تطبيق SHAP^١ لفهم مساهمة الميزات في التنبؤات^١
- مقارنة بين LIME^١ و SHAP^١
- الذكاء الاصطناعي أمثلة عملية لتفسير القرارات الفردية لنموذج

١) (Interpretability الوحدة الثالثة: تقنيات التفسير العالمية (Global)

- فهم السلوك العام لنموذج الذكاء الاصطناعي^١
- (Model-agnostic طرق التفسير العالمية المستقلة عن النموذج
- (Plots مخططات أهمية الميزات (Feature Importance)
- بطبيعتها^١ تفسير نماذج الشجرة ((Tree-based Models)
- (Plots - PDPs تحليل اعتماد الميزات الجزئي (Partial Dependence)
- (Expectation - ICE plots تأثيرات الميزات المعزولة (Individual Conditional
- (Interpretable-by-design models) بناء نماذج قابلة للتفسير بطبيعتها



الاصطناعي المرئي^١: الوحدة الرابعة: تفسير الشبكات العصبية والذكاء

- (Networks) تفسير الشبكات العصبية العميقة (Deep Neural)
- (CAMs) خرائط التنشيط الفئوية (- Class Activation Maps)
- التدرجات الموجهة (Guided Backpropagation)
- أهمية الميزات في الرؤية الحاسوبية^١.
- الصور^١: تفسير قرارات نماذج الرؤية الحاسوبية (التعرف على
- تطبيقات XAI^١ في تحليل الصور الطبية^١.
- (models) تحديات تفسير النماذج المعقدة (Transformer)

المستقبلية^١: الوحدة الخامسة: تقييم XAI^١، الأخلاقيات، والآفاق

- مقاييس لتقييم جودة التفسيرات^١.
- الذكاء الاصطناعي وكيفية كشفها بـ XAI^١ العدالة (Fairness) والتحيز (Bias) في نماذج
- وخصوصية البيانات. XAI^١
- المصلحة^١: كيفية توصيل تفسيرات الذكاء الاصطناعي لأصحاب
- XAI^١ التحديات القانونية والأخلاقية المتقدمة لنشر
- الاصطناعي المسؤول^١: مستقبل الذكاء الاصطناعي القابل للتفسير والذكاء
- الشفافية^١: بناء الثقة في أنظمة الذكاء الاصطناعي عبر

الأسئلة المتكررة^١:

التسجيل في الدورة؟ ما هي المؤهلات أو المتطلبات اللازمة للمشاركين قبل

لا توجد شروط مسبقة^١.

الإجمالي لساعات الدورة التدريبية؟ كم تستغرق مدة الجلسة اليومية، وما هو العدد



المدة إلى ٢٥٢٠- بمعدل يومي يتراوح بين ٤ إلى ٥ ساعات، تشمل فترات تمتد هذه الدورة التدريبية على مدار خمسة أيام، ساعة تدريبية، راحة وأنشطة تفاعلية، ليصل إجمالي

سؤال للتأمل:

بعض من كفاءة أكثر قابلية للتفسير، هل يمكن أن يؤدي هذا التركيز في ظل سعيها المستمر لجعل نماذج الذكاء الاصطناعي أمثل بين التفسيرية والأداء؟ ودقة هذه النماذج في مهام معينة، وهل هناك نقطة الشدida على الشفافية إلى التوضيح توازن

ما الذي يميز هذه الدورة عن غيرها من الدورات؟



اتخاذ القرار في نماذج مجال الذكاء الاصطناعي القابل للتفسير (XAI)، مع تتميز هذه الدورة بتقديمها نهجاً فريداً وشاملاً في النظرية المتينة لقابلية التفسير مع التطبيق الذكاء الاصطناعي المعقدة. ما يميزنا هو دمج الأسس التركيز على فهم آليات تقنيات التفسير، من التقنيات للمشاركين تحليل وتفسير مخرجات النماذج بفاعلية. العملي لأدوات وتقنيات XAI الرائدة، مما يتيح العالمية التي تكشف عن السلوك العام للنموذج. المحلية التي تشرح التنبؤات الفردية إلى التقنيات نغطي مجموعة واسعة من والامتثال للمتطلبات التنظيمية، اللازمة لبناء ثقة أكبر في أنظمة الذكاء الدورة تركز على تزويد المشاركين بالمهارات مما يجعلها ضرورية لأي محترف يتعامل مع نماذج والمساهمة في تطوير ذكاء اصطناعي مسؤول وأخلاقي، الاصطناعي، الذكاء الاصطناعي في بيئات حساسة تتطلب الشفافية.