

هل يمكن تفويض الذكاء الاصطناعي بمهمة التوظيف؟

كتبه بيتر كابيلي | 5 نوفمبر 2020



ترجمة وتحرير نون بوست

يعدنا الذكاء الاصطناعي بشكل أساسي - كأداة ذات قاعدة واسعة - بحل المشكلات التجارية، ما أفسح الطريق لأشياء محدودة لكنها ذات فائدة مثل خوارزميات علم البيانات التي تجعل التنبؤات أفضل مما نستطيع القيام به الآن.

فبخلاف النماذج الإحصائية المعتادة التي تركز على عامل أو اثنين معروفين بالفعل ومرتبطين بالنتائج مثل الأداء الوظيفي، تبدو خوارزميات التعلم الآلي محايدة بشأن المتغيرات التي نجحت من قبل وسبب نجاحها.

الأكثر مَرَحًا أنها تجمعهم معًا وتنتج نموذجًا واحدًا يتوقع بعض النتائج مثل من سيكون جيدًا عند توظيفه، بمنح كل متقدم للوظيفة درجة واحدة سهلة التفسير عن مدى احتمالية أدائه بشكل جيد في العمل.

كانت وعود هذه الخوارزميات عظيمة، لكن الاعتراف بمحدوديتها جذب الكثير من الانتباه خاصة فيما يتعلق بحقيقة أنه عند استخدام البيانات الأولية لبناء نموذج متحيز فإن الخوارزميات الناتجة عن تلك البيانات ستكون متحيزة على الدوام.

عندما نفوض المشرفين بمهمة التوظيف فمن المحتمل بشدة أن يكون لديهم تحيزات لصالح أو ضد بعض المرشحين للوظيفة بناءً على صفات لا علاقة لها بالأداء الجيد

أفضل مثال على ذلك في المؤسسات هو التحيز ضد المرأة في الماضي، حيث كانت بيانات الأداء في العمل متحيزة أيضًا مما يعني أن الخوارزميات القائمة على تلك البيانات ستكون متحيزة أيضًا، لذا كيف يمكن لأصحاب العمل التقدم بينما يفكرون في تبني الذكاء الاصطناعي لاتخاذ قرارات خاصة بالموظفين؟ إليكم 4 أشياء تضعونها في اعتباركم:

قد تكون الخوارزميات أقل تحيزًا من الممارسات الحالية

دعونا لا نتحدث برومانسية عن مدى سوء الحكم البشري الآن ومدى انعدام النظام في معظم الممارسات الإدارية للبشر، فعلى سبيل المثال عندما نفوض المشرفين بمهمة التوظيف فمن المحتمل بشدة أن يكون لديهم تحيزات لصالح أو ضد بعض المرشحين للوظيفة بناءً على صفات لا علاقة لها بالأداء الجيد.

فقد تفضل المشرفة (أ) المرشحين خريجي كلية معينة لأنها ذهبت إليها، بينما يفضل المشرف (ب) العكس لأن لديه تجربة سيئة مع بعض خريجها، أما الخوارزميات على الأقل ستعامل الجميع على قدم المساواة حتى لو لم يكن عادلاً تمامًا.

مقاييس جيدة لكل النتائج التي نرغب في توقعها

ربما أيضًا لا نعلم كيف نقيم العوامل المختلفة عند اتخاذ القرارات النهائية، فعلى سبيل المثال ما الذي يجعل "موظفًا جيدًا"؟ يجب أن ينجز مهامه جيدًا، يجب أن يتعاون مع زملائه بشكل جيد، أن يتناسب مع ثقافة المكان ولا يستقيل ويترك العمل وهكذا.

إن التركيز على مظهر واحد فقط حيث نملك مقاييس سينتج عنه خوارزمية توظيف تختار بناءً على مظهر واحد فقط غالبًا عندما يكون غير مرتبط تمامًا بالمظاهر الأخرى، مثل موظف المبيعات المتميز مع العملاء لكنه سيئ مع زملائه في العمل.

تقوم الخوارزميات بعمل جيد بشكل عام في التنبؤ بالنسبة للموظفين بمعدل متوسط لكنها تحقق أداءً سيئًا مع بعض المجموعات الفرعية

مرة أخرى، لا يبدو واضحًا أن ما نفعله الآن أفضل من ذلك، فالمشرف صاحب قرار الترقية قد يكون قادرًا بشكل نظري على أن يضع في اعتباره كل هذه العوامل، لكن تقييم كل منها يكون محملاً بالتحيزات وطريقة تقييمها قد تكون تعسفية، فنحن نعلم من خلال البحث الدقيق أنه كلما استخدم المديرون أحكامهم الخاصة في تلك الأمور كانت قراراتهم أسوأ.

البيانات التي يستخدمها الذكاء الاصطناعي تثير قضايا أخلاقية

الخوارزميات التي تتنبأ بمعدل ترك الموظفين للعمل تعتمد الآن على بيانات مواقع التواصل الاجتماعي مثل منشورات فيسبوك، ربما نرى أن جمع تلك البيانات يعد انتهاكاً لخصوصية الموظفين، لكن عدم استخدام تلك البيانات سيكون له أثر على تلك النماذج التي ستنبأ بشكل أقل جودة.

قد يكون الأمر كذلك في بعض الحالات، حيث تقوم الخوارزميات بعمل جيد بشكل عام في التنبؤ بالنسبة للموظفين بمعدل متوسط لكنها تحقق أداءً سيئاً مع بعض المجموعات الفرعية، الأمر ليس مفاجئاً فنموذج توظيف موظفي المبيعات لن يعمل جيداً لاختيار المهندسين.

الحل ببساطة أن نمتلك نموذجاً لكل مجموعة، لكن ماذا لو كانت تلك المجموعات عبارة عن رجال ونساء أو سود وبيض؟ في تلك الحالات تمنعنا القيود القانونية من استخدام ممارسات مختلفة أو نماذج توظيف مختلفة مع المجموعات السكانية المختلفة.

من الصعب أو المستحيل شرح وتبرير المعايير الكامنة خلف قرارات الخوارزميات

في معظم أماكن العمل الآن يكون هنا بعض المعايير المقبولة لاتخاذ قرارات تخص الموظفين، فهو يحصل على فرصة لأنه هنا منذ فترة طويلة، وهي إجازة في نهاية هذا الأسبوع لأنها عملت في مناوبة نهاية الأسبوع الماضي، هذه هي طريقة التعامل مع الموظفين من قبل.

في نقطة ما قد تستطيع الخوارزميات معالجة مثل هذه القضايا لكننا لسنا قريبين من ذلك في الوقت الحالي

إذا لم أحصل على الترقية أو المناوبة التي أريدها يمكنني أن اشتكي الشخص الذي اتخذ القرار، والذي يمكنه أن يشرح السبب وراء ذلك ويساعدني المرة القادمة إذا لم يكن القرار عادلاً بالنسبة لي.

لكن عندما تقوم الخوارزميات بهذا القرار فإننا نفقد القدرة على تفسير سبب اتخاذ القرار للموظفين،

فالخوارزمية تجمع كل المعلومات المتاحة ببساطة لتتنبأ بالنتائج السابقة.

سيكون من غير المحتمل بشكل كبير أن يشرح المشرف القرار إذا تطابقت تلك النتائج مع أي مبادئ، فالنموذج يقول إن هذا الأمر ناجح، لكن المشرف لا يستطيع شرح القرار أو معالجة مخاوف عدم الإنصاف.

عندما لا يختلف أداء تلك النماذج كثيرًا عما نقوم به بالفعل فإننا نتساءل هل يستحق الانزعاج الذي سيسببه للموظفين، فميزة السماح لكبير الموظفين باتخاذ القرار هو أنه يختار الجدول الذي يسهل فهم معاييرهم ويسهل تطبيقه وله فوائد كبيرة على المدى الطويل، في نقطة ما قد تستطيع الخوارزميات معالجة مثل هذه القضايا لكننا لسنا قريبين من ذلك في الوقت الحالي.

إن النماذج الخوارزمية ليست أسوأ مما نقوم به الآن، لكن مشكلات العدالة الخاصة بها من السهل تحديدها لأنها تحدث على نطاق واسع، ولحل تلك المشكلة يجب أن نحصل على معايير بيانات أكثر وأفضل وغير متحيزة، إن القيام بذلك سيساعدنا حتى لو لم نكن نستخدم خوارزميات التعلم الآلي لاتخاذ قرارات الموظفين.

المصدر: [هارفرد بيزنس ريفيو](#)

رابط المقال : <https://www.noonpost.com/38804/>