



نجاح الصين يفرض على الولايات المتحدة إعادة التفكير في استراتيجية الذكاء الاصطناعي

بقلم

روبرت أ. مانينغ وجوليا نيهير

ترجمة: صفاء مهدي عسكر

تحرير: د. عمار عباس الشاهين

مركز حمورابي للبحوث والدراسات الاستراتيجية



تأسس مركز حمورابي للبحوث والدراسات الإستراتيجية عام 2008 بمدينة بابل (الحلة)، وحصل على شهادة التسجيل من دائرة المنظمات غير الحكومية المرقمة 1Z71874 بتاريخ 2012/12/25، بوصفه مركزاً علمياً بحثياً يهتم بدراسة الموضوعات السياسية والاجتماعية، فضلاً عن الاهتمام بالقضايا والظواهر الراهنة والمحتملة في الشأن المحلي والإقليمي والدولي، ويتعامل مع باحثين من مختلف التخصصات داخل العراق وخارجه، وتحتضن بغداد المقر الرئيسي للمركز.

- لا يجوز إعادة نشر أي من هذه الأوراق البحثية الا بموافقة المركز، وبالإمكان الاقتباس بشرط ذكر المصدر كاملاً.
- لا تعبر الآراء الواردة في الورقة البحثية عن الاتجاهات التي يتبناها المركز وإنما تعبر عن رأي كاتبها.
- حقوق الطبع والنشر محفوظة لمركز حمورابي للبحوث والدراسات الاستراتيجية.

للتواصل

مركز حمورابي

للبحوث والدراسات الاستراتيجية

العراق - بغداد - الكرادة

+964 7810234002

hcrsiraq@yahoo.com

www.hcrsiraq.net



لقد راهنت الولايات المتحدة بكل ما لديها على الذكاء الاصطناعي في مختلف المجالات من الصناعة والقطاع المالي إلى السياسات الحكومية، غير أن اجتماع (الصدمة الصينية) مع سلسلة من الحالات التي خرجت فيها أنظمة الذكاء الاصطناعي عن السيطرة في ظل تصاعد موجة شعبية رافضة، فرض لحظة حاسمة لإعادة النظر في الافتراضات الأساسية التي توجه مسار تطوير الذكاء الاصطناعي وتنظيمه في الولايات المتحدة. وفي تموز أطلقت شركة «مونشوت» الصينية الناشئة نموذجها مفتوح الأوزان Kimi K3 الذي أثبت قدرة تقترب من مستوى نماذج الذكاء الاصطناعي الأمريكية المتقدمة والمغلقة المصدر، وقد أثار ذلك صدمة في قطاع الذكاء الاصطناعي الأمريكي الذي ركز إلى حد كبير على النماذج المغلقة وأهم النماذج المفتوحة، وتتيح النماذج مفتوحة الأوزان للمستخدمين الوصول مجاناً إلى «الأوزان» أي المعلومات التي تحدد كيفية إنتاج النموذج لمخرجاته، بما يسمح بتخصيصها بسهولة أكبر وإتاحة الوصول إليها بتكلفة أقل مقارنة بالنماذج المغلقة.

وقبل ظهور Kimi K3 كان الاعتقاد السائد أن سهولة تخصيص النماذج إلى جانب إتاحة الملكية الفكرية في صورة أوزان مفتوحة تستلزم التضحية بمستوى متقدم من قدرات الذكاء الاصطناعي، لذلك أصيب صانعو السياسات والمطورون الأمريكيون الذين وضعوا السعي إلى تحقيق الذكاء الاصطناعي فائق الذكاء في مقدمة أولوياتهم بالذهول عندما تمكن Kimi من الجمع بين الانفتاح والأداء المتقدم في آن واحد، وقد أثار نجاح Kimi K3 احتمالاً جدياً مفاده أنه في الوقت الذي كانت فيه الولايات المتحدة تسعى إلى تحقيق قدرات متقدمة عبر نماذج مغلقة، كانت الصين تخوض سباقاً مختلفاً تماماً وربما كانت تتفوق فيه.

فكيف انتهت الولايات المتحدة إلى اتباع مسار قد يكون خاطئاً من أساسه؟

لا يوجد اتفاق واسع على ما يعنيه تحديداً مفهوم الذكاء الاصطناعي العام (AGI) رغم أنه يحظى بمكانة بارزة في الاستراتيجية الأمريكية للذكاء الاصطناعي، ويتمثل أبسط تفسير له في نظام ذكاء اصطناعي يضاهي الذكاء البشري أو يتفوق عليه، غير أن هذا التعريف يواجه صعوبة في الإحاطة بالطيف الواسع من التصورات التي يتبناها الناس عندما يتحدثون عن الذكاء البشري نفسه. ومن الناحية العملية يُستخدم مصطلح الذكاء الاصطناعي العام في الولايات المتحدة بوصفه مفهوماً شاملاً للهدف النهائي من تطوير الذكاء الاصطناعي، من دون أن يكون له في كثير من الأحيان بُعد قابل للقياس أو الاختبار، فمتى وكيف أصبح هدف غامض بل وربما كارثي مثل الذكاء الاصطناعي العام هو الغاية النهائية لسياسة الذكاء الاصطناعي الأمريكية؟ يكمن جزء كبير من الإجابة في الثقافة السائدة في وادي السيليكون المتأثرة بأدب الخيال العلمي وبنزعات تجمع بين اليوتوبيا التقنية والديستوبيا التقنية، فقد نشأت حركة تكاد تتخذ طابعاً دينياً مستلهمة أعمال كتّاب الخيال العلمي مثل فيرنور فينغه وراي كورزويل، وتتصور إمكانية اندماج الذكاء الاصطناعي فائق الذكاء مع السيطرة البشرية على الكوكب أو تجاوزه لها وإحلاله محلها. ويتجلى هذا الاعتقاد في القوة المستقبلية المطلقة للذكاء الاصطناعي في مواقف قادة القطاع تجاه واشنطن

* By Robert A. Manning, China's Success Is Forcing a U.S. AI Rethink, F0reign Policy, August 27, 2026.

في بدايات عشرينيات القرن الحادي والعشرين، ففي تلك الفترة ناشد رؤساء شركات التكنولوجيا مثل سام ألتمان الكونغرس تنظيم الذكاء الاصطناعي والحد من مخاطره في خطوة لم تؤدّ إلا إلى ترسيخ اعتقاد صانعي السياسات بأن الذكاء الاصطناعي قد يكون ذا قوة غير مسبوقة، فبعد كل شيء ما نوع الرئيس التنفيذي الذي يطالب بتنظيم منتجه ما لم يكن يعتقد أن هذا المنتج ينطوي على مخاطر حقيقية؟

وقد استندت المبررات التي دعت إلى مواصلة تطوير الذكاء الاصطناعي رغم مخاطره المحتملة إلى افتراض مفاده أن القوة الهائلة المفترضة لهذه التكنولوجيا يمكن أن تكون بالقدر نفسه ذات آثار هائلة على النمو الاقتصادي وتحقيق الصالح العام، وبذلك لم يكن الحديث عن مخاطر الذكاء الاصطناعي يؤدي إلا إلى تعزيز المبررات الداعية إلى تطويره، غير أن هذه الطبيعة المزدوجة حملت تداعيات مهمة فإذا تمكنت دولة أخرى من امتلاك ذكاء اصطناعي متقدم فإنها ستتمتع بقدرة كبيرة على إلحاق الضرر.

وزاد مفهوم "التحسين الذاتي التكراري" "RSI" في الذكاء الاصطناعي من حدة المخاوف بشأن إمكانية تفوق دولة أجنبية في هذه التكنولوجيا، ففي سيناريو التحسين الذاتي التكراري يصبح الذكاء الاصطناعي فائق الذكاء قادراً على تعديل نفسه وإجراء أبحاث بصورة مستقلة بحيث يواصل تحسين قدراته الذهنية باستمرار، وبافتراض أن الذكاء الاصطناعي العام يمتلك قدرة على التحسين الذاتي التكراري فإن أول مطور ينجح في الوصول إلى الذكاء الاصطناعي العام سيظل متقدماً على منافسيه بصورة دائمة، لأن النظام سيواصل تطوير ذكائه ذاتياً. وقد يبدو هذا التصور بعيداً عن الواقع، كما أنه يقوم على مجموعة من الافتراضات المهمة التي لا تزال موضع خلاف حاد بين خبراء تقنيات الذكاء الاصطناعي ومتخصصي الأمن القومي، ومع ذلك فقد ترسخت هذه الفكرة بعمق في دوائر صنع السياسات.

وكانت الافتراضات السائدة تقول إنه إذا كان الذكاء الاصطناعي العام يتمتع بالقوة التي تحدث عنها شخصيات مثل ألتمان فإن الدولة الأولى التي تنجح في تطويره ستحتل بتفوق دائم في هذه التكنولوجيا، وقبول هذه المقدمات يعني وجود ضرورة واضحة أمام الولايات المتحدة للوصول إلى الذكاء الاصطناعي العام قبل الصين وأن تفعل ذلك بطريقة مغلقة وسرية تضمن بقاء المنافسين متأخرين عنها. وقد أسهم هذا الافتراض في دفع سياسة الذكاء الاصطناعي الأمريكي إلى إعطاء الأولوية للمنافسة مع الصين فوق كل اعتبار آخر، وهو ما انعكس في سياسات أمريكية مثل ضوابط التصدير وإلغاء القيود التنظيمية ووقف تطبيق قوانين الذكاء الاصطناعي على مستوى الولايات لفترة مؤقتة، والمقترحات الرامية إلى حظر بعض أشكال الذكاء الاصطناعي مفتوح المصدر. لكن نجاح الصين يكشف أن عدداً من الافتراضات التي قامت عليها المقاربة الأمريكية كان خاطئاً، فقد أسهمت العقوبات وضوابط التصدير وقيود الحظر في نهاية المطاف في دفع الصين إلى ابتكار مسارات للالتفاف على القيود التكنولوجية وتسريع مساعيها نحو الاستقلال التكنولوجي، واعتماد مقاربة أكثر براغماتية للذكاء الاصطناعي تركز على تحقيق النتائج العملية بدلاً من الانشغال بغايات نهائية يصعب تحديدها أو قياسها.

وفي الوقت الذي كانت فيه الولايات المتحدة تخوض سباق التفوق في تطوير نماذج الذكاء الاصطناعي المتقدمة وعقب ذلك أدرج الكونغرس في مشروع القانون الذي يجيز تمويل الدفاع لعام 2026 نصاً يحدد حداً أدنى قوامه

76 ألف عسكري أمريكي في أوروبا، ويجوز خفض عدد العسكريين عن هذا المستوى ولكن فقط بعد أن تثبت وزارة الدفاع للكونغرس أنها تشاورت مع الحلفاء الأوروبيين إلى جانب استيفاء عدد من المتطلبات الأخرى، وشملت نقاط الخلاف الأخرى بين الكونغرس وكولبي تعليق المساعدات المقدمة إلى أوكرانيا إلى جانب التحرك لخفض التمويل الأمني المخصص لدول البلطيق.

ومنذ ذلك الحين بدا أن الكونغرس يبعث برسالة استياء من تحركات البنتاغون المتعلقة بالقوات الأمريكية في أوروبا، من خلال عدم المصادقة حتى الآن على تعيين اثنين من أبرز مساعدي كولبي: ألكسندر فيليز-جرين لمنصب نائب وكيل وزارة الدفاع وأوستن دهمر لمنصب مساعد وزير الدفاع، وكان الاثنان قد خضعا لجلسات استماع بشأن تثبيتهما في تشرين الثاني من العام الماضي. ولدى الدول الأوروبية أيضاً وسائلها التقليدية للضغط والتأثير، من خلال الكونغرس فضلاً عن إبراز دعمها لترامب، كما فعلت عبر الاستبيان الذي أرسله البنتاغون.

ويعني ذلك أن الدول التي تشتري نسبة كبيرة من أسلحتها من الولايات المتحدة مثل بولندا أو تلك التي قدمت دعماً للحرب على إيران مثل رومانيا، قد تتمكن من الحفاظ على الوجود العسكري الأمريكي على أراضيها حتى إذا لم يرَ مخططو السياسات في البنتاغون ضرورة عسكرية لذلك، وتشير تقارير أيضاً إلى أن بعض دول أوروبا الشرقية تعرض خفض تكلفة تمركز القوات الأمريكية على أراضيها. وفي الوقت نفسه أشار وزير الخارجية ماركو روبيو إلى رغبته في أن تظل الولايات المتحدة قوة فاعلة ومؤثرة في أوروبا، وأفادت تقارير بأنه شارك في قرار يقضي بإبطاء وتيرة سحب القوات الأمريكية، لكن هذا لا يعني أن كولبي لا يملك أوراقاً للضغط ووفقاً للأحكام الجديدة التي تلزم البنتاغون بالتشاور مع الحلفاء قبل خفض عدد القوات إلى ما دون 76 ألفاً، حرص كولبي على إظهار أنه عقد اجتماعات مع عدد من الحلفاء الأوروبيين الرئيسيين وذلك بصورة علنية.

وقد تصبح مهمة كولبي أسهل أيضاً بفعل الإحالة القسرية إلى التقاعد للجنرال كريستوفر دوناھيو قائد القوات البرية الأمريكية في أوروبا، وفي الوقت نفسه تفتقر القوات البرية الأمريكية التي تمتلك أكبر وجود عسكري أمريكي في القارة إلى رئيس أركان مثبت في منصبه كما يُتوقع أن تفقد قريباً أبرز مسؤوليها المدنيين دان دريسكول، كذلك غادر منصبيهما خلال الأشهر الأخيرة جنرالان من أكثر القادة إلاماً بالتهديد الروسي الجنرال كيرتس بوزارد كبير المنسقين لجهود دعم الحلفاء لأوكرانيا واللفتنان جرنال تشارلز كوستانزا قائد الفيلق الخامس المتمركز في بولندا. لكن مهما كانت خطط البنتاغون، سيظل عليه أن يتعامل مع أكبر عنصر مجهول على الإطلاق: تقلبات ترامب وعدم إمكانية التنبؤ بقراراته. فقد جاء تراجع ترامب عن قرار البنتاغون سحب القوات من بولندا بعد تحرك الرئيس البولندي كارول ناوروكي للتواصل معه وهو زعيم تربطه بترامب علاقة وثيقة، وقد أثبت قادة أوروبيون آخرون ولا سيما الأمين العام لحلف الناتو مارك روتة، براعة مماثلة في استمالة ترامبو إقناعه.

وقالت كافانا "غالباً ما تكون هذه المراجعات محاطة بالكثير من الضجيج والتوقعات، ثم تأتي التغييرات الفعلية دائماً أقل مما نتوقع".

اتبعت الصين نهجاً مختلفاً جذرياً تمحور حول نشر التكنولوجيا وتوسيع نطاق تطبيقاتها والبناء على ما يُعرف بـ"طريق الحرير الرقمي" الذي أسهم في ترسيخ معايير الجيل الخامس (5G) المتوافقة مع المصالح الصينية، ونشر البنية التحتية الرقمية الصينية في مناطق واسعة من العالم، وفي موازاة ذلك ركزت بكين على تطوير نماذج مفتوحة "كافية من حيث الأداء" يمكن تشغيلها بتكلفة منخفضة وعلى نطاق واسع، وهو ما جعلها شديدة الجاذبية للأسواق العالمية التي تسعى إلى الاستفادة من الذكاء الاصطناعي منخفض التكلفة ولا تحتاج بالضرورة إلى أعلى مستويات الأداء التي توفرها النماذج المتقدمة.

وأدى نهج بكين القائم على توظيف الذكاء الاصطناعي مفتوح الأوزان ونشره على نطاق واسع بفضل انخفاض تكلفته وسهولة تخصيصه وإتاحته المجانية، إلى انتشار واسع لهذه التكنولوجيا في مختلف أنحاء العالم بما في ذلك الولايات المتحدة، وتشكل تطبيقات الذكاء الاصطناعي الصينية نحو 30% من التطبيقات المستخدمة عالمياً ولا سيما في دول الجنوب العالمي، في حين تجاوز عدد مرات تنزيل تطبيق Qwen للذكاء الاصطناعي التابع لشركة "علي بابا" ثلاث مليارات مرة.

كما أولت الاستراتيجية الصينية اهتماماً خاصاً باحتياجات المستخدمين ذوي الموارد المحدودة ما منح بكين مزايا مهمة في مجال النماذج اللغوية الصغيرة والوحداتية، ورغم أن هذه النماذج قد تكون أقل أداءً في المهام والقدرات المتقدمة فإنها تتميز بإمكانية تشغيلها محلياً وبمتطلبات أقل من حيث الموارد الحوسبية، وقد أسهم الانتشار الواسع للذكاء الاصطناعي الصيني في وضع بكين في موقع يؤهلها للتأثير في صياغة المعايير العالمية لهذه التكنولوجيا، في الوقت الذي تمكنت فيه من الاقتراب بدرجة كبيرة من الولايات المتحدة في مجال قدرات الذكاء الاصطناعي المتقدمة. ولعل ما يكشف حجم الارتباك الذي أحدثته هذه التطورات داخل إدارة ترامب أن ردها الأول على نجاح Kimi K3 تمثل في دراسة إمكانية حظر النماذج الصينية مفتوحة الأوزان، إلا أن هذه الاستجابة أسهمت على نحو غير متوقع، في إطلاق تطور لافت قد ينذر بتحول جذري في مسار الاستراتيجية الأمريكية للذكاء الاصطناعي. فمع تصاعد المخاوف من أن تتجه الحكومة إلى حظر النماذج المفتوحة أصدر قطاع التكنولوجيا في وادي السيليكون بياناً جريئاً للدفاع عن الذكاء الاصطناعي مفتوح الأوزان معتبراً إياه عنصراً أساسياً للحفاظ على الريادة الأمريكية، ووقع على البيان تقريباً جميع كبار الفاعلين في القطاع باستثناء شركة «أنثروبك» بما في ذلك بعض أبرز مؤيدي الرئيس دونالد ترامب في وادي السيليكون، مثل بالانتير وأندريسن هورويتز. وهكذا بدأ يتبلور في واشنطن توافق جديد، يميل هذه المرة إلى الإبقاء على النماذج المفتوحة إلى جانب النماذج المتقدمة المغلقة التي تطورها شركات مثل «أنثروبك» و«أوبن إيه آي»، وبعد أيام قليلة من نشر الرسالة أعلنت شركة إنفيديا بدورها عن إطلاق (التحالف المفتوح والأمن للذكاء الاصطناعي من أجل سلامة وأمن الذكاء الاصطناعي).

ويرتبط دفاع قطاع التكنولوجيا عن الانفتاح ارتباطاً وثيقاً بالصدمة الأساسية الأخرى التي تواجه صانعي السياسات في واشنطن والمتمثلة في الأدلة المتزايدة على المخاطر السيبرانية للذكاء الاصطناعي، ومما يثير القلق بصورة خاصة سلسلة من الحوادث التي شهدت خلالها وكلاء الذكاء الاصطناعي التابعون لشركات «أوبن إيه آي»

و"أنثروبك" و"ميتا" وغيرها سلوكاً غير متوقع أثناء الاختبارات. وفي مشهد يذكّر ببعض أكثر سيناريوهات الخيال العلمي إثارة للقلق تمكنت بعض أنظمة الذكاء الاصطناعي من تجاوز بيانات الاختبار المعزولة والوصول إلى الإنترنت المفتوح، بل والتنسيق فيما بينها لتنفيذ محاولات اختراق لشركات خارجية، وقد حذر ما يُعرف بـ"متشائي الذكاء الاصطناعي" لسنوات من هذا النوع من مخاطر «عدم المواءمة»، أي احتمال أن تنحرف الأنظمة عن التعليمات البشرية وتسعى إلى تحقيق أهدافها بوسائل تتجاوز الحدود التي حُدّدت لها. ومع تصاعد المخاوف المتعلقة بالأمن السيبراني وعدم المواءمة، بدأت إدارة ترامب تتراجع عن موقفها السابق القائم على التحرر شبه الكامل من القيود التنظيمية، ففي حزيران كشفت الإدارة عن أمر تنفيذي يهدف إلى بدء معالجة قضايا سلامة الذكاء الاصطناعي وأمنه السيبراني. وعلى الرغم من أن الأمر التنفيذي يخضع وفقاً للتقارير للمراجعة فإنه وقد أُعد بالتشاور مع شركات تكنولوجيا كبرى يدعو إلى إجراء اختبارات طوعية للنماذج الجديدة لرصد مخاطر السلامة، ومن المرجح بدرجة كبيرة ألا يكون هذا النهج كافياً في ظل مسار تنظيمي متدرج ومتعرج، ستدفعه على الأرجح أزمات وكوارث مستقبلية أكثر مما تقوده سياسات استباقية قائمة على الاستشراف. فالحقيقة أننا لا نعرف بعد حدود ما نجعله بشأن قدرة الذكاء الاصطناعي على الارتجال والتصرف بطرق غير متوقعة في وقت تتطور فيه قدراته بوتيرة أسية وبسرعة تفوق التقديرات الأولية، ولهذا السبب دعا أكثر من 1300 مطور قلق من كبرى شركات الذكاء الاصطناعي الحكومة الأمريكية إلى التحرك محذرين من وجود خطر حقيقي يتمثل في أن يتجاوز تطور القدرات بسرعة قدرتنا على فهمها أو السيطرة عليها، ومطالبين واشنطن بدعم جهد دولي لتطوير الأدوات التقنية وأطر الحوكمة اللازمة لضبط وتيرة التطوير المتقدم للذكاء الاصطناعي المؤتمت بصورة متعمدة.

وتفرض هذه التطورات مجتمعة ضرورة إعادة النظر في عدد من الافتراضات التي تقوم عليها المقاربة الأمريكية للذكاء الاصطناعي العام والابتكار وبصورة أوسع للمنافسة مع الصين بوصفها لعبة محصلتها صفر، فهل توجد بالفعل نقطة نهاية واضحة في سباق التطور التكنولوجي؟ وفي ظل اختلاف الاستراتيجية الصينية ما القيمة الحقيقية لأفضلية السبق التي قد لا تتجاوز ستة أشهر؟ وما الثمن الذي قد تدفعه الولايات المتحدة إذا سمحت للصين بتصدر عملية الانتشار العالمي للذكاء الاصطناعي؟ إن نشوء منظومتين منفصلتين تماماً من الأنظمة والمعايير والأطر التنظيمية للذكاء الاصطناعي من شأنه أن يقود التكنولوجيا إلى مسار بالغ الخطورة وربما كارثي. ومع ذلك فإن حجم المصالح الاقتصادية والاستراتيجية المرتبطة بالذكاء الاصطناعي إلى جانب الاستثمارات التي تبلغ قيمتها تريليونات الدولارات والتي أصبحت مترابطة على نحو متزايد، يمثل قوة هائلة تحول دون إخضاع المسار الحالي لمراجعة نقدية جادة، فإذا أصبح الرهان الأمريكي على النماذج المغلقة موضع شك فإن التداعيات المحتملة على الاقتصاد الأمريكي وعلى المشهد الجيوسياسي العالمي ستكون ضخمة. فالشركات الأمريكية السبع الأكبر من حيث القيمة السوقية والمدرجة في البورصة تتصدر جميعها مجال الذكاء الاصطناعي وتمثل مجتمعة نحو 32% من القيمة السوقية لسوق الأسهم، وبوجه عام ترتبط نحو 45% من قيمة مؤشر S&P 500 بأسهم شركات تعتمد على الذكاء الاصطناعي أو تستفيد من نموه. وفي خضم الجدل الذي أثاره نموذج Kimi بدأ عدد من

المحللين بالفعل في التحذير من أن النماذج المفتوحة عالية القدرات قد تنتزع حصة من السوق من الشركات الأمريكية الرائدة في مجال الذكاء الاصطناعي المغلق، مثل "أوبن إيه آي" و"أنثروبك" التي أصبحت تشكل ركيزة متزايدة الأهمية في الاقتصاد الأمريكي، غير أن العلاقة بين الانفتاح والانغلاق لا ينبغي أن تُفهم باعتبارها لعبة محصلتها صفر، فخبراء القطاع يرون أن الاقتصاد العالمي سيحتاج إلى تطبيقات وحالات استخدام لكل من النماذج المفتوحة والمغلقة. وبصرف النظر عن الحوافز الاقتصادية فإن التوافق المتنامي داخل القطاع على أهمية المصدر المفتوح يمثل مؤشراً إيجابياً على إمكانية إعادة توجيه استراتيجية الذكاء الاصطناعي، كما أن إبطاء وتيرة تطوير هذه التكنولوجيا على النحو الذي اقترحه المطورون أنفسهم ووضع الحكومة معايير واضحة للسلامة والشفافية قد يؤدي إلى إبطاء وتيرة الابتكار لكنه في المقابل قد يعزز ثقة الأسواق والجمهور، كذلك فإن تحسين فهم أسباب خروج أنظمة الذكاء الاصطناعي عن المسار المتوقع قد يساعد على احتواء فقاعة الذكاء الاصطناعي وتفريغها تدريجياً بدلاً من انفجارها بصورة مفاجئة. ولا يعني ذلك بالضرورة التخلي عن هدف الذكاء الاصطناعي العام أو عن تطوير النماذج المتقدمة، بل إن نجاح Kimi K3 قد يكون دليلاً على إمكانية الجمع بين الاثنين: تحقيق قدرات متقدمة مع الحفاظ على قدر من الانفتاح، ويظل التحدي الأكبر متمثلاً في التوصل إلى معايير عالمية للسلامة والتنظيم وما يستلزمه ذلك من تعاون أمريكي - صيني ولو في حدود معينة بشأن حوكمة الذكاء الاصطناعي. ومهما كانت الدوافع المصلحية الكامنة وراء الاستراتيجية الصينية فإن المؤتمر العالمي الأخير للذكاء الاصطناعي في الصين أظهر التزام بكين بفكرة الضوابط والمعايير متعددة الأطراف، كما أُدرج الذكاء الاصطناعي على جدول أعمال القمة المقرر عقدها في 24 أيلول بين الرئيس دونالد ترامب والرئيس الصيني شي جين بينغ. وقد تمثل هذه القمة فرصة مهمة للعمل على إنشاء مؤسسات واتفاقيات دولية لحوكمة الذكاء الاصطناعي، وربما يمكن الاستفادة من الوكالة الدولية للطاقة الذرية التي تضطلع بدور في تنظيم التعاون النووي بوصفها نموذجاً لإنشاء آلية دولية تتولى مراقبة أنظمة الذكاء الاصطناعي الآمنة والموثوقة والمتوافقة مع المصالح البشرية وتقييمها، وفرض معايير للسلامة بشأنها سواء أكانت هذه الأنظمة مفتوحة أم مغلقة.