

الذكاء الاصطناعي والفلسفة المعاصرة

هيدغر، يوناس، والعفن الغروي

ماساهيرو موريوكا

[أستاذ في كلية العلوم الإنسانية في جامعة واسيدا Waseda University، اليابان]

ترجمة: د. معين رومية

1. مشكلة الإطار Frame Problem

ألقي في هذه الورقة نظرة إجمالية على المقاربات الفلسفية المعاصرة لمشكلة "الذكاء" في حقل الذكاء الاصطناعي، وذلك من خلال تفحص عدد من الأبحاث المهمة في الفينومينولوجيا وفلسفة البيولوجيا.

لا يوجد تعريف واضح للذكاء الاصطناعي. تكتب مارغريت تي. بودن Margaret T. Boden في مؤلفها "الذكاء الاصطناعي: طبيعته ومستقبله *AI: Its Nature and Future*" أن الذكاء الاصطناعي العام قد يمتلك قدرات عامة في "الاستدلال والإدراك، بالإضافة إلى اللغة والإبداع والعاطفة". ومع ذلك، لا تغفل بودن الإشارة إلى أن "القول أسهل من الفعل".¹

يشبه مفهوم بودن عن "الذكاء الاصطناعي العام" مفهوم "الذكاء الاصطناعي القوي"، الذي صاغه جون سيرل عام 1980. عرّف سيرل "الذكاء الاصطناعي الضعيف" بأنه حاسوب يتصرف كما لو أنه يفكر بحكمة، في حين أن "الذكاء الاصطناعي القوي" هو حاسوب يفكر فعليا مثل البشر. يكتب سيرل: "وفقا للذكاء الاصطناعي القوي... يُعتبر الحاسوب المُبرمج بشكل ملائم عقلا حقيقيا، مما يعني أن الحواسيب، إذا ما زُوِّدت بالبرامج الصحيحة، يُمكن القول حرفيا إنها تفهم وتمتلك حالات إدراكية أخرى."² وقد أشبع موضوع الذكاء الاصطناعي القوي بحثا في أواخر القرن العشرين؛ لكن بات من الواضح أنه لكي يُعدّ الحاسوب ذكاء اصطناعيا قويا، يتعيّن عليه حلّ كثير من المشكلات المعقدة، أصعبها المشكلة الفلسفية المسماة "مشكلة الإطار".

مشكلة الإطار هي تلك المشكلة التي تعجز فيها أنظمة الذكاء الاصطناعي عن التمييز بشكل مستقل [أي، ليس وفق تعليمات برمجية مسبقة] بين العوامل المهمة وغير المهمة عند التعامل مع شيء ما في موقف محدد. فمثلا، تبرز هذه المشكلة عندما ندع روبوتات الذكاء الاصطناعي تعمل في العالم الواقعي. طرح جون مكارثي وباتريك ج. هايز مشكلة الإطار عام 1969، وهي مشكلة فلسفية لا يُمكن اختزالها إلى مجرد مشكلة تقنية. تكتب بودن في عام 2016 أن "الادعاءات بأن مشكلة الإطار الشهيرة قد 'حُلّت' مضلّلة جدا،"³ مما يشير إلى أن كثرة من المتخصصين، حتى الآن، يعتقدون أن مشكلة الإطار لم تُحل.

لا يوجد إجماع حول تعريف مشكلة الإطار، ومع ذلك يمكننا القول إنها مشكلة تتمحور حول كيفية جعل الذكاء الاصطناعي يستظهر "المعرفة المضمرة" tacit knowledge التي يمتلكها معظم البشر الراشدين في سياق محدد. لنتخيل أن أحد الروبوتات يعمل نادلا يقدم الطعام والشراب للزبائن في مطعم. ينبغي لهذا الروبوت أن يتعلم سلسلة من المعارف والحركات اللازمة للخدمة. كم يحتاج هذا الروبوت من المعارف كي يتمكن من أداء الخدمة بكفاءة في بيئة واقعية؟ أولا، من الضروري له معرفة أن "صب كمية كبيرة من الماء في كوب يؤدي إلى انسكابه" عند التقديم. ومن الضروري أيضا معرفة أن "تحريك صفحة عليها كوب يؤدي إلى تحريك الكوب معها"، فمن دون هذه المعرفة، لن يعمل الروبوت على نقل الكوب المستخدم والصفحة في آن واحد. زد على ذلك، يجب إدخال معلومة مفادها أنه عند تحريك الكوب، يتحرك السائل بداخله معه. ولكن، لسنا بحاجة إلى إدخال معلومة مفادها أن السائل لن يتبخر أبدا بفعل حرارة الاحتكاك الناتجة عن حركة الصفحة، لأن هذه المعلومة لا صلة لها بمهمة التقديم.

يتضح إذن أنه يوجد كم هائل من المعلومات يتعين على الروبوت استظهارها عند التقديم، ويوجد أيضا كم هائل من المعلومات التي لا يحتاج إلى استظهارها. فمن يستطيع إعداد مثل هذه القائمة من المعلومات، وكيف يمكن جعل الروبوت يستظهرها؟ تقع هذه المشكلة لأن الروبوت عندما يواجه موقفا جديدا لم يسبق له تجربته، فإنه لا يستطيع أن يحكم بشكل مستقل بشأن نوع التعامل ذو الصلة أو غير ذي الصلة بهذا الموقف، ومن ثم لا يستطيع حل المشكلة التي تواجهه كما ينبغي. ومن المثير للاهتمام أن البشر يبدو أنهم قادرون على حل هذا النوع من المشكلات. فمثلا، تستطيع طالبة في المرحلة الثانوية العمل في مطعم من دون أي مشكلة إذا زودناها بمجموعة أساسية من التعليمات البسيطة. سوف تُعيد الصفحة إلى المطبخ وعليها كوب فارغ من دون أي تعليمات خاصة. وحتى لو واجهت موقفا مستجدا، فإنها ستحاول حل المشكلة باتباع نهج مرن يتناسب مع كل حالة على حدة.

في هذا السياق، أجرى هيتوشي ماتسوبارا Hitoshi Matsubara نقاشا شيقا عام 1990، حيث كتب: "إن الكائن الذي يمتلك قدرة محدودة على معالجة البيانات، وهذا يشمل الحواسيب والبشر، لا يمكنه أبدا التوصل إلى حل كامل لمشكلات الإطار العامة، ومع ذلك، يبدو أن البشر في بيئات الحياة اليومية ومواقفها لا يزعجون من مشكلة الإطار. وإن، ما يجب أن نفكر فيه مليا هو التالي: لماذا يبدو أن البشر هم بمنأى عن مشكلة الإطار؟^{iv} ويكون السؤال هو: لماذا يبدو الذكاء البشري في كثير من الحالات قادرا على التعامل بنجاح مع الحوادث غير المتوقعة في سياق مفتوح، على الرغم من أنهم لا يمتلكون القدرة الكافية لحل مشكلة الإطار بشكل كامل؟ وهذا يعني لماذا يبدو الذكاء البشري في كثير من الحالات قادرا على التعامل بنجاح مع الحوادث غير المتوقعة في سياق مفتوح، على الرغم من أن البشر لا يمتلكون قدرة كافية لحل مشكلة الإطار بشكل كامل.

أعتقد أن الذكاء البشري يتميز بالخصائص التالية مقارنة بالذكاء الاصطناعي: (1) [المعرفة المضمرة]، عندما يواجه الإنسان مواقف غير مألوفة، يستطيع اتخاذ قرار ملائم باستخدام "معرفة مضمرة" ومن ثم القيام بفعل معين حتى لو لم تكن لديه المعرفة الكاملة اللازمة لذلك (2) [تقدير الأهمية]، يستطيع اتخاذ قرار مستقل بشأن نوع المواجهة الأهم عند التعامل مع موقف غير مألوف والحاجة إلى النجاة منه، (3) [التجاهل]، يستطيع اختيار فعل محدد وتنفيذه فوراً، متجاهلا البدائل الأخرى الممكنة. يبدو أن الذكاء الاصطناعي لا يمتلك هذه الخصائص الثلاث، ولكي يمتلكها يجب أن يكون قادرا على حل مشكلة الإطار التي ناقشناها سابقا.

ظهرت مؤخرا سلسلة من المقاربات المشجعة بشأن حل مشكلة الإطار من حقل متداخل التخصصات ما بين أبحاث الذكاء الاصطناعي والفلسفة المعاصرة. سأتناول في الفقرات التالية نقاشين مهمين حول هذا الموضوع: "الذكاء الاصطناعي الهيدغري Heideggerian AI" وفق هوبرت درايفوس، والمقاربات البيولوجية المتأثرة بفكرة هانز يونس Hans Jonas عن "الأيض/الاستقلاب metabolism". يركز النقاش الأول بشكل خاص على "الدازين Dasein" و"الجسد"، وهما مفهومان يرتبطان ارتباطا وثيقا بالفينومينولوجيا المعاصرة. أما النقاش الثاني، فيركز بشكل خاص على شكل "الحياة" أو "الكائن الحي/المتعضي organism"، وهما مفهومان يرتبطان ارتباطا وثيقا بفلسفة الحياة حاليا وبنظرية الحياة الاصطناعية.

2. الذكاء الاصطناعي الهيدغري

هوبرت درايفوس هيدغري بامتياز، وقد انتقد فلسفيا الذكاء الاصطناعي منذ بدايات الأبحاث في هذا الحقل. أتناول هنا بحثه المنشور عام 2007 بعنوان: "لماذا فشل الذكاء الاصطناعي الهيدغري وكيف أن إصلاحه يتطلب جعله أكثر هيدغرية".

يحتاج درايفوس أيضا بأن السبب الذي يجعل الذكاء الاصطناعي يواجه مشكلة الإطار يكمن في أنه لا يميز نوع المعرفة المهمة له [ذات الصلة] في موقف مُحدد. فالحدث أو الشيء لا يكتسب معنى إلا عند وضعه في سياق مُحدد.

لقد افترضت أبحاث الذكاء الاصطناعي التقليدية النموذج الديكارتي الذي بموجبه تُضفي أذهاننا قيمة على عالم مُكوّن من أشياء وحوادث لا معنى لها. يُشدد درايفوس على أن هذا النوع من الأبحاث لن يجعل الذكاء الاصطناعي ذكاء بشريا ولن يحل مشكلة الإطار. ويولي اهتماما خاصا للتمييز الذي يقيمه مارتين هيدغر بين كينونتين هما: "الحاضرة [الوجود-أمامي] *Vorhandenheit*" و"الجاهزية [الوجود في متناولي] *Zuhandenheit*" في كتابه "الكينونة والزمان".^v فمثلا، من منظور ديكارتي، تبدو المطرقة التي تكون أمامي أداة موضوعة في حالة الحاضرة، أي إنها مجرد موجودة أمامي. في مقابل ذلك، إذا بدت المطرقة أداة حميمة مألوفة منغمسة في نسيج المعنى الذي يُحاك من العلاقة الظاهرة والمضمرة بين المطرقة والشخص (أنا) الذي يستخدمها، حينها نقول إن المطرقة تبدو لي أداة حميمة في حالة الجاهزية للاستخدام، أي إنها في متناولي. يفتقر الذكاء الاصطناعي التقليدي إلى القدرة على فهم هذا النوع من الجاهزية. فبينما يدرك كل إنسان أنه منغمس دائما في هذا النسيج من المعنى في وجوده-في-العالم، فإن الذكاء الاصطناعي التقليدي عاجز عن فهمه.^{vi}

يجادل درايفوس بأنه لكي يمتلك الذكاء الاصطناعي التقليدي القدرة على حل مشكلة الإطار ويصبح ذكاء اصطناعيا حقيقيا، يجب أن يصبح ذلك "الذكاء الاصطناعي الهيدغري" القادر على تطبيق بُعد الجاهزية. ويتفحص درايفوس دراسات عديدة أجراها باحثو الذكاء الاصطناعي في هذا الاتجاه، لكنه يخلص إلى أن أيا منها لم يُحقق الذكاء الاصطناعي الهيدغري.

أولا، يتفحص درايفوس روبوتات رودني بروكس Rodney Brooks، مبتكر "هيكلية الاندراج subsumption architecture"، وهي نظام هرمي مورّع يُشبه هيكلية ممالك الحشرات، ويُستخدم حاليا في روبوت المكتسة الكهربائية "Roomba". يحرك روبوت بروكس نفسه بالإحالة المستمرة على حساساته بدلا من الاعتماد على نموذج عالم داخلي.^{vii} لكن هذه الروبوتات "تستجيب فقط للخصائص الثابتة في البيئة، وليس للسباق أو للدلالة المتغيرة"،^{viii} ولذا ينبغي القول إنها تعجز عن حل مشكلة الإطار.^{ix}

ثم يتفحص درايفوس برنامج Pengi الذي ابتكره فيل أغري Phil Agre وديفيد تشابمان David Chapman، والذي طوّر لاحقا إلى Pengo، وهي لعبة فيديو يتبادل فيها اللاعب البشري وشخصيات البطريق كرات الثلج على الشاشة. يذكر

• المصطلح يصف بنية تنظيمية (غالبا في التقنية أو الإدارة) تحاكي الطريقة التي تعيش وتعمل بها الحشرات (مثل النمل أو النحل)، ويجمع بين صفتين أساسيتين: **الهرمية**، أي وجود مستويات من القيادة أو التنظيم، لكنها ليست بالضرورة "رئيسا ومرؤوسا" بالمعنى التقليدي، بل هي تقسيم للأدوار (مثل الملكة، الجنود، والعمال)؛ و**التشتت أو التوزيع**، أي لا يوجد مركز تحكم واحد يدير كل حركة صغيرة، فكل وحدة (حشرة) تملك قدرا من الاستقلالية وتتفاعل مع محيطها، لكن النتائج النهائية تخدم النظام ككل.

♣ في سياق الذكاء الاصطناعي وعلوم الأعصاب، مصطلح "نموذج العالم الداخلي" (Internal World Model) يعني باختصار التمثيل العقلي أو الرقمي لكيفية عمل الواقع، فبدلا من أن يتفاعل الكائن (أو الخوارزمية) مع كل شيء كأنه يراه لأول مرة، فإنه يبني "خريطة" أو "محاكاة" داخلية له.

أغري أنهما استعانا في برمجة هذه اللعبة بكتاب هيدغر "الكيونة والزمان" وطرحا مفهوم "القصدية الإشارية" * deictic intentionality " فيها. لا تشير القصدية الإشارية إلى شيء محدد، بل إلى "دور قد يؤديه شيء ما في نمط تفاعل ممتد زمنيا بين الكائن وبيئته".^x وقد وصلت هذه اللعبة إلى أن تكون قادرة على التعامل مع الشيء الذي يتفاعل معه الكائن كوظيفة.^{xi}

ينقد درايفوس مقاربة أغري وتشابمان. فعلى سبيل المثال، عندما أضع يدي على مقبض الباب لأجل الخروج من الغرفة، فأنا لا أختبر الباب كمجرد باب، بل في هذا الموقف يُدفع بي نحو إمكان الخروج من الغرفة عبر هذا الباب، والباب يدعوني للخروج من خلاله. لم يبرمج الذكاء الاصطناعي الذي طوره أغري وتشابمان هذا النوع من الخبرة التي يُستدعى فيها المستخدم من خلال إمكان مُفعّل في موقف معين. هذا يبيّن أنهم انتهوا إلى إضفاء طابع موضوعي على الوظائف التي قدموها وعلى أهمية الموقف بالنسبة للمستخدم. يقول درايفوس إنه، بهذا المعنى، لم يكن لدى الذكاء الاصطناعي الذي طوره أغري وتشابمان القدرة على حل مشكلة الإطار.^{xii}

أخيرا، يتحدث درايفوس عن نظرية مايكل ويلر Michael Wheeler. يكتب ويلر في مؤلفه "إعادة بناء العالم المعرفي *Reconstructing the Cognitive World*" الصادر عام 2005 أن علم الاستعراف المُجسّد-المدمج، الذي طُبّق على أبحاث الذكاء الاصطناعي، يُشبه فلسفة هيدغر. لكن درايفوس ينتقده قائلا إنه يبحث أيضا في الموضوع الخطأ عند تناوله هذا الموضوع. وتتخلص وجهة نظر درايفوس فيما يلي: على الرغم من إصرار ويلر على أن نماذج الذكاء الاصطناعي المتجسدة -المدمجة تُعتبر هيدغرية، إلا أنها تبقى ضمن النموذج الديكارتي الذي يقوم على التمثيل الرمزي لحوادث العالم الخارجي على (مستشعرات) الذكاء الاصطناعي، وأن حل مشكلات الذكاء الاصطناعي يتم بناء على هذا التمثيل. لكنّ نموذج التمثيل هذا نفسه هو المشكلة. فلا يمكننا فهم الذكاء البشري فهما كاملا بواسطة هذا النموذج الديكارتي. ويجادل درايفوس بأن هيدغر يعتبر أن الذازين Dasein هو "وجود-في-العالم" حيث لا مجال للتمثيلات في هذه الكيفية الوجودية. ويكتب درايفوس أن "الوجود-في-العالم أساسي أكثر من التفكير وحل المشكلات، فهو ليس تمثيلا على الإطلاق".^{xiii}

عندما يحاول شخص حل المشكلات، يتلاشى الحد الفاصل بينه وبين أدواته. يكون هذا الشخص قد اندرج للتو في العالم، وبالنسبة للمتأملين الماهرين، فإنهم "ليسوا مجرد عقول، بل هم جزء لا يتجزأ من العالم".^{xiv} بمعنى أدق، نحن "متأقلمون منغمسون absorbed copers".^{xv} ومن الصعب جدا تحديد ما إذا كان حل مشكلات المنغمسين يتم داخل الذات أم في العالم، لأن التمييز بين الداخل والخارج ليس بالأمر الهين.

ينبغي أن يكون أساس الذكاء الاصطناعي الهيدغري هو الوجود-في-العالم الذي هو للذازين، أي يجب أن يكون الذكاء الاصطناعي ذا زانين، وأن تكون كيفية وجوده هي الوجود-في-العالم. لا ينبغي تسمية الذكاء الاصطناعي الذي لا يمتلك هذه الكيفية من الوجود بأنه "هيدغري"، وهذا يعني أنه لا يمكنه امتلاك القدرة على حل مشكلة الإطار.

أما بالنسبة للبشر، فإنهم يستطيعون تحسين مهاراتهم في التأقلم مع التغيرات المهمة في العالم بواسطة قدرتهم الفطرية المدمجة في داخلهم embodied على حل المشكلات. فمثلا، عندما نكون في غرفة، نتجاهل عادة كثيرا من التغيرات التي تطرأ عليها، ولكن إذا ارتفعت درجة الحرارة بشكل كبير، فإن نوافذ الغرفة تدعونا لفتحها، فنستجيب لهذا النداء ونفتحها. ففي هذه الحالة، تُحل المشكلة من خلال التفاعل مع الممكّنات التي تتيحها البيئة. يكتب درايفوس عن سبب قدرة البشر على حل مشكلة الإطار ما يلي: "عموما، وبالنظر إلى خبرتنا في العالم، عندما يحدث تغيير في السياق القائم، فإننا نستجيب له فقط إذا كان في الماضي ذا أهمية، وعندما نستشعر تغييرا ذا أهمية، فإننا نتعامل مع كل شيء آخر على أنه لم يتغير باستثناء

♣ تعني القصدية الإشارية توجيه الوعي أو القصد نحو شيء ما مرتبط بال لحظة والمكان اللذين يتواجد فيهما الشخص حاليا. مثال توضيحي: إذا قلت "الأسود مفترسة"، فهذه قصدية عامة (تفكير في مفهوم عام)، لكن إذا أشرت بيدك وقلت "هذا الأسد يقترب مني الآن"، فهذه قصدية إشارية. أنت لا تقصد أي أسد، بل تقصد ذلك الموجود في محيطك المكاني والزمني المباشر. هذا المصطلح مهم في الذكاء الاصطناعي والروبوتات لأنه يُستخدم لوصف قدرة الروبوت على فهم الأوامر المرتبطة ببيئته الحالية (مثل "ضع هذا الصندوق هناك"). الروبوت يحتاج لـ "قصدية إشارية" ليربط الكلمة بالشيء الفيزيائي الموجود أمامه.

ما تشير إليه معرفتنا بالعالم بأنه قد يكون قد تغير أيضا، وبالتالي يحتاج إلى التحقق. على هذا النحو، لا تنشأ مشكلة الإطار في هذا الموقف المحدد.^{xvi} ففي حالة البشر، يتم دائما تفعيل "أفتنا بالعالم" ضمنا في إدراكنا، ويتم تمييز ما هو مهم بالنسبة لنا تلقائيا عما هو غير مهم بالنسبة لنا. هذه الوظيفة تمنع ظهور مشكلة الإطار.

تُعاود مشكلة الإطار البروز عندما نضطر لتغيير السياق بأنفسنا. متى نُدرك أن المشكلات الموجودة في محيطنا تُصبح محور مهامنا في حل المشكلات؟ يقول درايفوس مشيرا إلى ميرلوبونتي، إن هذا الإدراك ناتج عن "استدعاءات" من الممكن التناح.^{xvii} باختصار، عندما يحدث أمر بالغ الأهمية بالنسبة لنا، يُمكننا إدراكه من خلال إشارات أو استدعاءات من العالم الذي نعيش فيه، ومن دون قبول هذا النموذج لن نتضمن أبدا من حل مشكلة الإطار. يخلص درايفوس إلى أنه لكي يكتسب الذكاء الاصطناعي هذه القدرة، "سيتعين علينا تضمين نموذج لجسم يشبه جسمنا إلى حد كبير في برنامجنا، شاملا احتياجاتنا ورغباتنا وملذاتنا وأمانا وطرق حركتنا وخلفيتنا الثقافية، وما إلى ذلك."^{xviii}

يبدو أن الذكاء الاصطناعي الهيدغري الذي طرحه درايفوس يجب أن يمتلك "جسدا" شبيهها بالجسد البشري وأن يعيش في داخل ذلك الجسد. ولكن، هل من الممكن لروبوتات الذكاء الاصطناعي الحالية، المصنوعة من رقائق السيليكون والمعادن والبلاستيك، أن تلبّي هذه المتطلبات العالية؟ أنتقل في القسم التالي إلى تفحص المقاربات البيولوجية التي تختلف تماما عن الذكاء الاصطناعي الهيدغري.

3. الذكاء الاصطناعي والأبيض^v

يعتقد بعض الباحثين أنه لكي يمتلك الذكاء الاصطناعي القدرة على حل مشكلة الإطار، يجب أن يكون ضريبا من ضروب "الكائن الحي/المتعضي"، أو شكلا من أشكال الحياة، وذلك قبل أن يستطيع اكتساب جسم شبيهه بالجسم البشري. عندما تواجه الكائنات الحية صعوبات مميتة، فإنها تحاول البقاء على قيد الحياة باستخدام كل وسيلة ممكنة. تمتلك الكائنات الحية مثل هذه القدرات فطريا. ويعتقد هؤلاء الباحثون أن هذه القدرات الفطرية التي تمتلكها الكائنات الحية للبقاء على قيد الحياة هي ما يجب أن تكون الأساس اللازم لحل مشكلة الإطار.

أكد هانز يونس، الذي كان تلميذا لهيدغر، على أهمية مفهوم "الأبيض" في مجال فلسفة البيولوجيا، وقد كان لهذا المفهوم تأثير كبير على المقاربات المذكورة أعلاه. نشر يونس كتابه "ظاهرة الحياة" *The Phenomenon of Life* عام 1966 (ونسخته الألمانية عنوانها "الكائن الحي والحرية" *Organismus und Freiheit* عام 1973) وأسس فلسفةً بيولوجيةً أصلية. يعتقد يونس أن "الحرية" نشأت عندما ظهرت قديما على الأرض ميكروبات ذات أغشية خلوية. تتغذى هذه الميكروبات عبر الغشاء وتطرح الفضلات من خلاله أيضا. تستطيع الميكروبات الحفاظ على حياتها بهذا النوع من الامتصاص المستمر وإخراج الجزيئات الدقيقة عبر الغشاء. تُستبدل جميع المواد المكونة للخلية بمرور الوقت. ومع ذلك، تحتفظ الخلية الحية بهويتها في بُعد مختلف. ويرى يونس في ذلك تحرر الحياة من بُعد المادة. ويعتقد أن هذا التحرر هو "الحرية" التي يكتسبها الشكل الحي.

من جهة أخرى، تكون الحياة مقيّدة بهذه العملية التي قوامها استبدال الجسيمات الدقيقة عبر الغشاء. فإذا توقفت هذه العملية، فإن الحياة محكوم عليها بالزوال، لأنها هي التي تضمن بقاءها. وبهذا المعنى، تعتمد الحياة على المادة. ويُطلق يونس على هذا النوع من الحرية اسم "الحرية التابعة" *dependent freedom* أو "الحرية المُفتقرة" *impoverished freedom*.^{xix} وباستمرار، يُهدد هذا الخطر المُحتمل بقاء الحياة. فالحياة مُقدر لها البقاء من خلال الاستبدال المُستمر لمحتوياتها. وإذا أهملت جهود استبدال المواد، فإنها ستواجه الموت. فالحياة هشة بطبيعتها، لأنها ستفنى سريعا إذا لم تُبذل جهود مُستمرة للبقاء.

^v Metabolism الأبيض أو الاستقلاب هو مجموعة من التفاعلات الكيميائية التي تحدث داخل خلايا الكائن الحي لتحويل الغذاء إلى طاقة ولتكوين وحدات بناء ضرورية للنمو والترميم.

عندما كان يوناس يدرس هذه الفكرة، لم يكن يتخيل وجود الذكاء الاصطناعي. وكانت أفكاره بشأن الحياة والحرية تُناقش فحسب ضمن دائرة ضيقة من فلاسفة البيولوجيا آنذاك. ولكن، بعد وفاته عام 1993، سرعان ما بدأت فلسفته تُناقش خارج حقل البيولوجيا.

كان فرانسيسكو ج. فاريللا Francisco J. Varela أحد الفلاسفة الذين سلطوا الضوء على فكر يوناس، إذ دافع عن مفهوم "الصنع الذاتي" autopoiesis في البيولوجيا، و"التفاعلية" enactivism في الفينومينولوجيا والذكاء الاصطناعي. في الورقة البحثية الرائدة "الحياة بعد كانط" Life after Kant التي كتبها أندرياس وبيير وفاريللا ونُشرت عام 2002 (بعد وفاة فاريللا)، سعى الاثنان إلى الربط بين مفهوم الصنع الذاتي عند فاريللا ومفهوم الأيض عند يوناس. كتب أن "الصنع الذاتي هو الأساس التجريبي الضروري لنظرية يوناس في القيمة"^{xx} وأن هاتين الفكرتين (الصنع الذاتي والأيض) "ليستا متزامنتين فحسب، بل متكاملتين تماما. سعى الاثنان في سبيل هيرمنيوطيقا للكانت الحي، أي فهم غاية الكائن الحي ومعناه من داخله".^{xxi} وفي كلتا النظريتين، كان غشاء الخلية وأيضها هما المصطلحان الرئيسيان. تحرّى يوناس وفاريللا فكرة أن الخلية الواحدة ذات الغشاء تحتوي على "غاية ذاتية/متأصلة"، وأن هذا الكائن الخلوي له "غرض أساسي في الحفاظ على هويته، وتأكيد الحياة".^{xxii} كان لانتباه فاريللا إلى فلسفة البيولوجيا عند يوناس، ولاسيما تركيزه على الأيض، أثر بالغ على بعض باحثي الذكاء الاصطناعي.

تُعدّ ورقة توم فرويز Tom Froese وتوم زيمكي Tom Ziemke بعنوان "الذكاء الاصطناعي التفاعلي: دراسة التعصّي المنظومي للحياة والعقل of Life and Mind"، المنشورة عام 2008، محاولة في حقل الذكاء الاصطناعي لتطوير فكرة يوناس عن الأيض.

يفسر فرويز وزيمكي مشكلة الإطار على النحو التالي: تكمن المشكلة في كيفية تصميم نظام اصطناعي بحيث تظهر سمات العالم ذات الصلة كعناصر مهمة من منظور النظام نفسه، وليس فقط من منظور المصمم أو المراقب البشري.^{xxiii} ويشيرون إلى ورقة درايفوس البحثية شددت على أن مشكلة الإطار لم تُحل بعد، ويضيفون أن إسهامات الفينومينولوجيا والبيولوجيا النظرية ضرورية لحل هذه المشكلة.

ويقول فرويز وزيمكي إن العلوم المعرفية شهدت في السنوات الأخيرة "انعطافة نحو فكرة الدمج embodied turn". ومع ذلك بقينا نجهل كيفية تعليم الذكاء الاصطناعي أن يفهم "بشكل مستقل" المشكلات المهمة بنفسه. ويركزون على فلسفة يوناس في البيولوجيا. يكتبون أن "وجود ما يمكن وصفه من قبل مراقب خارجي بأنه سلوك "موجه نحو هدف" لا يستلزم بالضرورة أن النظام قيد الدراسة نفسه يمتلك تلك الأهداف - فقد تكون عَرَضِيَّة (أي مفروضة من الخارج) وليست داخلية (أي مُؤَلَّدة في الداخل)..."^{xxiv} إذا كان لروبوت الذكاء الاصطناعي "أهدافه" الخاصة، فينبغي أن تُرَوِّد هذه الأهداف تلقائيا من داخل الروبوت. ويحتاجون بأن السؤال الذي ينبغي طرحه هو: ما نوع الجسد التي يجب أن يمتلكه الروبوت لإنجاز مثل هذه المهمة؟

يشير فرويز وزيمكي إلى فكرة يوناس عن الأيض، ويناقشان الفرق بين النظام الاصطناعي والنظام الحي. يتمثل نمط وجود النظام الاصطناعي في "الوجود- بـ الوجود being by being". يستطيع النظام الاصطناعي أن يتصرف، لكن هذا التصرف ليس بالضرورة من أجل بقائه. على النقيض من ذلك، يتمثل نمط وجود النظام الحي في "الوجود- بـ الفعل being by doing". يجب على النظام الحي أن يخرط في عمليات معينة تُعرف بـ "عمليات التكوين الذاتي self-constituting operations"، أي الاستبدال المستمر للمواد الدقيقة عبر غشاء الخلية. إذا توقف النظام الحي عن عمليات الاستبدال، فإنه يموت، ويختفي من الوجود. الفعل أو التصرف ضروري للنظام الحي، ولكنه ليس كذلك للنظام الاصطناعي. هذا هو الفرق الحاسم بين النظام الاصطناعي والنظام الحي، وهذا تحديدا ما أراد يوناس التأكيد عليه بعبارة "الحرية التابعة". ناقش يوناس هذا الموضوع في سياق علم التحكم الآلي ونظرية المنظومات العامة في ستينيات القرن العشرين، وأعاد فرويز وزيمكي إحياء فكرته في عصر الذكاء الاصطناعي في القرن الحادي والعشرين.

يصعب جدا ابتكار ذكاء اصطناعي قادر على الأيض. لكنهم يجادلون بأن السبب الجوهري لعجز الذكاء الاصطناعي عن حل معضلة الإطار يكمن في افتقاره إلى نمط الوجود البيولوجي، أي "الوجود- بـ الفعل". فعلى سبيل المثال، حتى لو أوقفنا تشغيل الذكاء الاصطناعي، ثم أعدنا تشغيله، فسيستمر في العمل من دون أي مشكلات تُذكر. أما الكائن الحي، فبمجرد

موته لن يعود للحياة أبدا. ^{xxv} هذا الشعور المُلح بأن كل شيء سينتهي بموته يُميّز نمط وجود الكائن الحي. ويرى هؤلاء أن هذا هو مفتاح حل معضلة الإطار.

ويقولون إنه ينبغي علينا الانتباه إلى حقيقة أن الكائن الحي يُولد ويحافظ بفاعلية على "الهوية المنظومية في ظل شروط هشة". ^{xxvi} ويُطلقون على هذا النمط من الوجود اسم "الاستقلالية التكوينية constitutive autonomy"، وفقا لتسمية فاريل. ويقولون إن "الأنظمة المستقلة تكوينيا تجلب معها هويتها الخاصة ونطاق تفاعلاتها، وبالتالي تعين مشكلاتها التي يجب حلها" وفقا للمُتاح لها في العمل. ^{xxvii} ويُجرون تحليلا نظريا للذكاء الاصطناعي ذي الاستقلالية التكوينية، ويحاولون إيجاد توليفة ممكنة بين الذكاء الاصطناعي والحياة الاصطناعية.

أولا، يُشيرون إلى إمكان "تضاييف الميكروب-الروبوت". ^{xxviii} فمثلا، إذا استطعنا عكس حالة الميكروبات المُدمجة في الروبوت على وحدة تحكم الروبوت، يُمكن تسجيل الحركة المستقلة للميكروبات في ذكاء الروبوت في الوقت الفعلي. ويجادلون أيضا بأنه من خلال دمج مبدأ آلة الرّصف * tessellation automaton في الروبوت، يمكننا الاستفادة من خاصيتها المتمثلة في أنه على الرغم كون مبدأ الإنتاج ليس ذكيا، لكن النتيجة تبدو ذكية عند النظر إليها من الخارج. ^{xxix} ويؤكدون أن مثل هذا النظام لم يُبتكر من قبل، وأنه "من أجل تطوير نظرية أفضل للجذور البيولوجية للفاعلية القصدية، نحتاج أولا إلى فهم أفضل للذكاء على مستوى البكتيريا. وفقط بالعودة إلى بدايات الحياة نفسها، تتوافر لدينا فرصة لوضع نظرية راسخة للفاعلية القصدية وعملية التعرّف". ^{xxx}

يبدو لي أن حجتهم القائلة بأنه لتطوير ذكاء اصطناعي تلقائي، يتعين علينا العودة إلى مستوى البكتيريا، حجة محفزة ومعقولة. تؤكد مارغريت أ. بون على أهمية الأيض من خلال الاستشهاد بهانز يونس، وتخلص إلى أنه إذا كان الأيض هو الشرط الضروري للعقل، فلا بد أن الذكاء الاصطناعي القوي مستحيل لأن عملية الأيض "يمكن محاكاتها ونمذجتها بواسطة الحواسيب، ولكن لا يمكن إنشاء مثل لها". ^{xxxi} ربما يكون نموذج الأيض لدى يونس المفتاح الأعمق لفهم الذكاء الاصطناعي.

4. العفن الغروي والحاسوب-الحيوي

بدأت بالفعل مساعي استكشاف الذكاء بالرجوع إلى مستوى البكتيريا. ومن البحوث في هذا الصدد تجدر الإشارة بشكل خاص إلى حاسوب العفن الغروي slime mold computer، الذي درسه توشيوكي ناكاجاكي Toshiyuki Nakagaki وريو كوباياشي Ryo Kobayashi. سبق للاثنتين أن اكتشفا في عام 2000 أنه عند وضع الطعام في موضعين منفصلين على مائدة زجاجية صغيرة ونشر العفن الغروي الجائع على سطحها، فإن العفن يعمل على تحويل خليته ببطء بغية رسم أقصر مسار بين الموضعين. ^{xxxii} تبدأ أطراف العفن الموجودة على مسار مسدود بالانسحاب منه، والانضمام إلى المسار الرئيسي المتصل بالطعام، مما يساعد على زيادة سماكة المقطع العرضي للمسار الرئيسي الذي رسمه العفن. هذا يعني أن العفن الغروي يُجري بنفسه نوعا من الحساب، ويكتشف أقصر مسار بين الموضعين، ويُعدّل جسمه ليتخذ الشكل الأكفأ لهذا المسار. يُظهر هذا الفعل الذي يقوم به العفن الجائع بوضوح الطريقة الأساسية التي تسعى بها الكائنات الحية للحفاظ على وجودها في وضع "هش".

* مبدأ الرصف أو التبليط، مبدأ في الرياضيات يصف عملية تغطية سطح ما باستخدام أشكال هندسية متكررة بالطريقة الأمثل، أي من دون ترك أي فراغات أو حدوث أي تداخل بين الأشكال.

في ورقتهم البحثية المنشورة عام 2011 بعنوان "أداء معالجة المعلومات في كائن حي بدائي من العفن الغروي الحقيقي" (باللغة اليابانية)، يجادل ناكاجاكي وكوباياشي بأن هذا النوع من أفعال البقاء لدى العفن ناتج عن "حسابات" يجريها العفن بنفسه.^{xxxiii} أي أن فعل البقاء ينشأ فطرياً وتلقائياً داخل العفن الغروي، حيث تُجرى حسابات للبحث عن الحل الأمثل، ويُحوّل العفن نفسه وفقاً لهذا الحل. يُمكن تسمية هذا بآلة حاسبة حيوية، وأفترض أن مشكلة الإطار قد تُحل في حالة العفن الغروي هذه. إذا أعطينا العفن الغروي مهمة صعبة جديدة وتتبعناها، فسيعيد بالتأكيد النظر في استراتيجيته للوضع الجديد وسيحاول تحويل جسمه نحو حل جديد ملائم. يبدو أن العفن الغروي يمتلك القدرة على التكيف المستمر مع التغيرات البيئية غير المعروفة له مسبقاً، وذلك بتحويل جسمه بطريقة مستجدة عفوية.

وضع ناكاجاكي وكوباياشي نموذج محاكاة رياضي (حلّال فيزاروم[®] physarum solver) لتتبع حركة العفن الغروي، ودرس سلوكه. ونتيجة لذلك، اكتشفاً أن الحسابات التي يجريها العفن الغروي ليست دقيقة أو مثالية، بل هي "تقريبية وسريعة". ويجادلان بأن هذا الحساب "التقريبي والسريع" يُعد "سمة بارزة للحوسبة الحيوية".^{xxxiv} أعتقد، مع أنهما لم يذكرنا ذلك صراحة، أن هذه السمة قد تكون مفتاحاً لا ابتكار ذكاء شبيه بالذكاء البشري، ذكاء قادر، عند مواجهة بيئة جديدة، على تحديد العامل الأهم بسرعة والتصرف بحدة، متجاهلاً العوامل الأخرى غير المهمة. هذا النوع من الذكاء هو الضروري لحل مشكلة الإطار.

يكتب كوباياشي أيضاً في بحثه المنشور عام 2015 بعنوان "التحكم اللامركزي المستقل في الكائنات الحية: من العفن الغروي إلى الروبوت" (باللغة اليابانية) أنه بينما تستطيع معظم الروبوتات التحرك بشكل صحيح في البيئات المتوقعة، تستطيع الحيوانات التحرك بمواظبة حتى عند مواجهة بيئات غير معروفة. ويجادل بأن الحيوانات لديها القدرة على حل مشكلة الإطار،^{xxxv} وهذا لا يشمل الثدييات والحشرات فحسب، بل أيضاً العفن الغروي. ويشير كوباياشي إلى أن الحشرات والعفن الغروي قادرة على التحكم اللامركزي التلقائي في أجسامها. يشير هذا إلى أنه لحل مشكلة الإطار، سيكون تطوير نظام جسدي لا مركزي تلقائي أفضل من نظام تحكم مركزي كالجهاز العصبي المركزي. يقول كوباياشي إن روبوته ذاتي الحركة الشبيه بالثعبان قد يمتلك مثل هذا النظام اللامركزي، ويقترح تطوير نظام تحكم يجعل البيئة "صديقة" له.

تشير هذه الدراسات إلى أن الحل التلقائي والفطري لمشكلة الإطار المتوافر لدى البشر لا يتم بواسطة الجهاز العصبي المركزي، بل بواسطة نظام تحكم لا مركزي موجود في كل جزء من الجسم متحرر من سيطرة الجهاز العصبي المركزي. وتجدر الإشارة إلى أن هيكلية الاندراج التي صممها برونكس لم تنجح في حل هذه المشكلة.

قد يكون مفهوم "تضاييف الميكروب-الروبوت" الذي طرحه فرويز وزيمكي حلاً محتملاً. يقترح الاثنان إدخال مستعمرة من الميكروبات في جسم روبوت. ولكن، ألا يوجد احتمال آخر في الاتجاه المعاكس- أي إمكان إدخال أجسام اصطناعية مثل الذكاء الاصطناعي فائق الصغر، أو الروبوتات فائقة الصغر، أو أجزاء مصنعة من الحمض النووي DNA أو الحمض النووي الريبي RNA، في خلايا الميكروبات؟ قد يكون من الممكن إنشاء أنظمة تضاييفية لمجموعة من هذه الأجسام الاصطناعية فائقة الصغر مع الميكروبات.

لنأخذ مثال العفن الغروي. يمكننا تزويد هذا العفن بقدرة حسابية مُحسّنة بواسطة روبوتات دقيقة جداً، أو معالجات دقيقة جداً، أو أحماض نووية مُصنّعة. ففي هذه الحال، لن يستطيع هذا العفن المُحسّن صناعياً حلّ مشكلة أقصر مسار تلقائياً فحسب، بل يستطيع أيضاً حلّ مهام أكثر تعقيداً عن طريق الحساب. ومن ثمّ يمكننا افتراض أن هذا العفن يمتلك القدرة على اكتشاف المشكلة المهمة لبقائه وحلّها تلقائياً. ولما كان العفن الغروي، ككائن حي، يمتلك القدرة على حلّ مشكلة الإطار، فإن العفن المُحسّن صناعياً بقدرة حسابية لا بد أن يمتلك أيضاً القدرة على حلّ هذه المشكلة. يُمكن اعتبار العفن المُحسّن

® نموذج رياضي وخوارزمية مستوحاة من الطبيعة تُحاكي السلوك الذكي للعفن الغروي المعروف باسم Physarum polycephalum في بحثه عن الطعام.

صناعيا نوعا من الحواسيب الحيوية. وفي سياق الحواسيب، تُعدّ الحواسيب الحيوية مفتاح حلّ مشكلة الإطار. هذه هي الخلاصة الأولى لهذه الورقة.

انبثقت من نقاشنا إحدى المشكلات الفلسفية وهي: إذا أمكن حل مشكلة الإطار بواسطة التحكم اللامركزي التلقائي لدى لكائن الحي، فهذا يقتضي أنه يمكن حلها من دون تحقيق الذكاء الاصطناعي الهيدغري الذي اقترحه درايفوس، والذي يوجد في نمط الوجود- في-العالم. قد لا تشكل مسألة الإطار مشكلة على مستوى الجهاز العصبي المركزي الذي يُجري عمليات معالجة الرموز، بل مشكلة على مستوى أنظمة التحكم اللامركزية التلقائية المستندة إلى عمليات الأيض. وباعتبار أن درايفوس كان سيفترض مسبقا تحكم الجهاز العصبي المركزي، فقد تكون فكرته خاطئة تماما. يرى بعض الباحثين أن التطور الحاصل في حقل التعلم العميق Deep Learning قد ينجح في حل مشكلة الإطار، لكن قدرة التعلم العميق لا تزال غير واضحة. تعتمد المناقشات السابقة على فهمنا لجوهر مشكلة الإطار. هذا هو السؤال الذي ينبغي على الفلاسفة معالجته بشكل مباشر.

حاولنا أعلاه تقديم لمحة عامة عن العلاقة بين مشكلة الإطار، والذكاء الاصطناعي الهيدغري، والذكاء الاصطناعي القائم على الأيض ونظام التحكم اللامركزي التلقائي الذي يُنتج العفن الغروي، واقترحنا حلا مُحتملا لمشكلة الإطار في المستقبل باستخدام الحواسيب الحيوية. نجد في ذلك عدة مواضيع مُحفزة للفلسفة المعاصرة. سيهتم باحثو الفلسفة بظهور اسمي هيدغر وأحد تلاميذه البارزين، يوناس، في مناقشتنا لمشكلة الإطار. لسْتُ متخصصا في أبحاث الذكاء الاصطناعي، لذا يُرجى إبلاغي في حال وجود أي تعابير مُضللة أو استخدامات غير صحيحة للمصطلحات العلمية التقنية في هذه الورقة.

ومن اللازم التنبيه إلى البحث في إنتاج العفن الغروي المُحسن صناعيا ينطوي على مخاطر جسيمة. يجب علينا منع الانتشار غير المنضبط للعفن الغروي المُحسن صناعيا، لأن هذا البحث يهدف إلى منح قدرات حسابية عالية المستوى. إذا ما انبثقت هذه المواد في البيئة، فقد تُلحق أضرارا جسيمة بالبشر والنظم البيئية، لذا يجب إجراء البحث وفقا لأعلى معايير السلامة البيولوجية في مرافق مُجهزة بالقدرة على الاحتواء المادي المنصوص عليه في بروتوكول قرطاجنة. في المقام الأول، لا يُمكننا تخيل سلوك العفن الغروي عند تعزيز قدراته الحسابية. قد يكون هناك خطر من أن ينتشر العفن الغروي المُعزز صناعيا، ذو الذكاء العالي، على نطاق واسع ويغطي الأرض بأكملها بحثا عن الغذاء. في حالة الميكروبات السامة، لا ينبغي السماح بإجراء أبحاث تهدف إلى منحها قدرات حسابية عالية.

يُمكن أيضا اعتبار هذا النوع من الدراسات بحثا في مجال التحسين باستخدام أجسام اصطناعية تستهدف الميكروبات. لذلك، يرتبط هذا الموضوع بالمناقشات الأخلاقية الحيوية حول التحسين.

لقد دعم الذكاء الاصطناعي أبحاث التكنولوجيا الحيوية بطرق عديدة، وسيظهر في المستقبل وضع مختلف تماما، حيث يندمج بحث الذكاء الاصطناعي مباشرة مع التلاعب بالكائنات الحية في مجال التكنولوجيا الحيوية. ينبغي أن نجري نقاشا معمقا ومتعدد التخصصات قبل أن يصبح هذا واقعا. ولعلنا نستنتج أن الفجوات بين أبحاث الذكاء الاصطناعي وعلم الأحياء والفلسفة قد تقلصت بشكل ملحوظ.

مصدر المقالة:

Journal of Philosophy of Life Vol.13, No.1 (January 2023):29-43.

Artificial Intelligence and Contemporary Philosophy

Heidegger, Jonas, and Slime Mold

Masahiro Morioka*

والمقالة ترجمة لورقة بحثية باللغة اليابانية منشورة بالعنوان نفسه في مجلة - 51: Vol.70, (2019) 『哲学』
68.

الحواشي

-
- ⁱ Boden (2006), p.22.
ⁱⁱ Searl (1980), p.417.
ⁱⁱⁱ Boden (2016), p.55.
^{iv} Matsubara (1990), p.179.
^v Section 18 and others.
^{vi} Dreyfus (2007), pp.248, 251.
^{vii} Dreyfus (2007), p.249.
^{viii} P.250.
^{ix} PP.249-250.
^x P.252.
^{xi} PP.251-252.
^{xii} P.253.
^{xiii} P.254.
^{xiv} P.255.
^{xv} P.255.
^{xvi} P.263.
^{xvii} P.264.
^{xviii} P.265.
^{xix} Jonas (1973), S.150.
^{xx} Weber and Varela (2002), p.120.
^{xxi} P.116.
^{xxii} P.117.
^{xxiii} Froese and Ziemke (2008), p.467.
^{xxiv} P.472.
^{xxv} P.485.
^{xxvi} P.480.
^{xxvii} P.481.
^{xxviii} P.492.
^{xxix} P.494.
^{xxx} P.495.
^{xxxi} P.144.
^{xxxii} Nakagaki, T. et al. (2000).
^{xxxiii} Nakagaki and Kobayashi (2011), p.483.
^{xxxiv} P.491.
^{xxxv} Kobayashi (2015), p.236.
-

المراجع

- Boden, Margaret A. (2016). *AI: Its Nature and Future*. Oxford University Press.
- Dreyfus, Hubert L. (2007). "Why Heideggerian AI Failed and How Fixing it Would Require Making it More Heideggerian." *Philosophical Psychology* 20(2):247-268.
- De Jesus, Paulo (2016). "Autopoietic Enactivism, Phenomenology and the Deep Continuity Between Life and Mind." *Phenomenology and the Cognitive Sciences* 15:265-289.
- Froese, Tom, and Gallagher, Shaun (2010). "Phenomenology and Artificial Life: Toward a Technological Supplementation of Phenomenological Methodology." *Husserl Studies* 26:83-106.
- Froese, Tom, and Ziemke, Tom (2009). "Enactive artificial intelligence: Investigating the systemic organization of life and mind." *Artificial Intelligence* 173:466-500.
- Heidegger, Martin (2006). *Sein und Zeit*. Max Niemeyer Verlag.
- Jonas, Hans (1973, 1977). *Das Prinzip Leben*. Suhrkamp. (*Organismus und Freiheit: Ansätze zu einer philosophischen Biologie*).
- Kiverstein, J. and Wheeler, M. (eds.) (2012). *Heidegger and Cognitive Science*. Palgrave Macmillan.
- Kobayashi, Ryo (2015). 小林亮「生物に学ぶ自律分散制御：粘菌からロボットへ」『計測と制御』54(4):236-241.
- Matsubara, Jin (1990). 松原仁「一般化フレーム問題の提唱」J・マッカーシー、P・J・ヘイズ『人工知能になぜ哲学が必要か』哲学書房, pp.175-245.
- Nakagaki, T., Yamada, H., and Tóth, A. (2000). "Path finding by tube morphogenesis in an amoeboid organism, *Nature* 407:470.
- Nakagaki, Toshiyuki, and Kobayashi, Ryo (2011). 中垣俊之・小林亮「原生生物粘菌による組合せ最適化法：物理現象として見た行動知」『人工知能学会誌』26(5):482-493.
- Searle, John R. (1980). "Minds, brains, and programs." *The Behavioral and Brain Sciences* 3:417-457.
- Shimonishi, Kazeto (2015). 下西風澄「生命と意識の行為論：フランシスコ・ヴァレラのエナクティブ主義と現象学」『情報学研究』89:83-98.
- Weber, A. and Varela, F. J. (2002). "Life After Kant: Natural Purposes and the Autopoietic Foundations of Biological Individuality." *Phenomenology and the Cognitive Sciences* 1:97-125.

Wheeler, Michael (2005). *Reconstructing the Cognitive World: The Next Step*. The MIT Press.

Wheeler, Michael (2008). "Cognition in Context: Phenomenology, Situated Robotics and the Frame Problem." *International Journal of Philosophical Studies* 16(3):323-349.

