

コンピューターが暴走した:AIの発展が社会に与える影響と意味 サマンサ・ベイツ(2016年5月6日)

2016年2月19日、ハーバード・ロー・スクールで第2回「コンピューターが暴走した:AIの発展が社会に与える影響と意味」ワークショップが開催された。マリン・ソーリャッチ、マックス・テグマーク、ブルース・シュナイアー、ジョナサン・ジットレインがこの非公式のワークショップを招集し、AI研究の最近の進歩について話し合った。参加者は、幅広い専門知識と視野を代表し、終日のイベントで4つの主要なトピックス——労働と経済に対するAIの影響、特に法律におけるアルゴリズムの意思決定、自律型兵器、人間レベルのAIが出現することのリスク——について議論した。

各セッションは、指名された参加者から提示されたトピックスに関連する既存の文献の簡単な概要から始まり、それに続いて2、3人の論客が発言した。その後、セッション・リーダーは、より大きなグループでの議論を切り盛りした。各セッションの終わりに、参加者はコミュニティによる更なる調査を必要とする研究質問リストに同意した。各議論の要約と、追加の研究分野に対するグループの勧告を以下に示す。

セッション#1:労働と経済

概要

労働と経済のセッションでは、2つの主要なトピックス——労働経済学に対するAIの影響と金融市場におけるAIの役割——が注目された。参加者は、資本が最終的に労働力に取って代わり、市場が金融部門と実体経済に徐々に分岐することで、資本の集中に与えるAIの影響が加速し、深刻化していることを見た。グループは、現在の失業率を見ると、AIはまだ人間の労働者に大きく取って代わるものではない、と納得した模様だが、参加者は、AIがどのようにして大きな社会的不平等を引き起こし、特にギグ経済において仕事の質を低下させるか、について話し合った。グループはまた、2010年と2014年のフラッシュ・クラッシュといった金融市場における最近の出来事は、テクノロジー、特に高頻度取引と人間の行動との間の相互作用が市場の大きな不安定性に繋がる可能性がある点を例示していると議論した。

AIが労働経済に与える影響について、グループは、特にAIが社会的流動性を妨げる場合、社会的不平等が懸念されることに同意した。例えば、テクノロジーの成長により、大学の出願や奨学金に関する情報へアクセスする学生が増加し、学生間の競争が激化している。その結果、大学や奨学金委員会は、応募者総数が増加して、成功に必要なスキル・セットを持っている候補者を見落とすかもしれない。学生がもっと一生懸命働くように動機づけるには社会的不平等が必要ではあるものの、競争の激化は、大学に自分でお金を払う経済的手段を持たない学生からチャンスを奪うかもしれない。参加者は、テクノロジーが、すべての利用者にチャンスを提供し、社会が才能を特定するのを助けるのではなく、既存の社会経済構造

を強化することがあり得るという点に懸念を表明した。

大学での応用例と同様に、テクノロジーの成長は、非常に特殊なスキル・セットを必要とする、より多くの“タレント市場“または雇用を創出することによって労働市場を再編成し始めている。これらのスキルを持った労働者が強く求められており、また彼らは高給ではあるものの、現在のシステムではこれらの候補者の多くを見落としており、労働者階級の個人が社会的地位を上げるチャンスが少なくなる可能性がある。参加者は、社会的流動性の問題を解決する最善の方法については同意しなかった。何人かの参加者は、より多くのオンライン教育、及び認定プログラムの必要性、またはテクノロジーの進歩について学ぶために労働者が定期的に学校に戻ることを要求する教育システムの再構築の必要性を強調した。別の参加者は、民間企業が情報を蓄え始める社会に移行しており、情報が広く行き渡らなくなれば、オンライン教育や教育システムの再構築は意味をなさない、と指摘した。

高頻度取引とフラッシュ・クラッシュについての議論の中で、グループは、私たちが最も懸念すべきことを決定するのに苦心した。フラッシュ・クラッシュが発生する理由は完全には理解されてないが、過去の例を見ると、フラッシュ・クラッシュは、技術的な欠陥とシステムを操作しようとする悪意のある攻撃者の組み合わせが同時に発生する“完全な経済ストーム“によって引き起こされることが示唆される。何人かの参加者は、フラッシュ・クラッシュはそれほど頻繁に発生せず、市場が非常に迅速に回復できることに驚きを表明した。それに対し、市場は急速に回復しているように見えるものの、市場におけるこれらの短期下落の影響は、個人投資家にとって壊滅的なものになる可能性がある点が指摘された。例えば、3分で退職後の貯蓄の半分を失った図書館員の話聞いた。

最終的に、参加者は、AI が労働と金融市場に与える影響を管理するうえで、政策規制が主要な役割を果たすことに同意した。例えば、高頻度取引の時間制限により、迅速な取引を防ぎ、フラッシュ・クラッシュを引き起こすリスクを減らせる点が示唆された。更なる研究に向けて特定された他の分野には、フラッシュ・クラッシュではなく、それほど深刻な損失を引き起こさないために見落とされてしまうフラッシュ・トリップがあるかどうか、及びシステムがトリップする頻度を予測するモデルがあるかどうかの調査が含まれる。参加者はまた、金融システムは他の部門で発生するかもしれない問題の予測因子として機能するため、更なる研究に向けて重要な分野となる点に同意した。また、保証収入や逆所得税から、奨学金付きの教育、医療、その他の社会福祉の提供に至るまで、再分配を通じて所得の不平等を減らす様々なアプローチの賛否両論について議論した。最後に、参加者は、最適な教育システムとはどういうものか、そしてそれが社会的流動性の問題の解決にどのように役立つかを定義する必要があることに同意した。

質問と結論

- ・ 金融システムは、急速かつ頻繁に変化し、人々の生活に大きな影響を与えるため、他の部門での主要な AI の課題は何かを示すよい指標になる、と結論づけた。

- ・ グループは、フラッシュ・クラッシュの発生理由が殆ど未解明である点を懸念すべきか議論した。調査すべき研究分野の1つは、フラッシュ・クラッシュではなくフラッシュ・トリップだが、これはそれほど深刻な損害を引き起こさないため、検出がより困難である。
 - ・ 複雑なシステムのモデルとは別に、システムがトリップする頻度を予測するモデルはあるだろうか？
- ・ 2011年以降、高頻度取引のリスクを軽減するうえで進展はなかった。特にフラッシュ・クラッシュが発生する理由についての説明がないため、この問題に注意を向けた方がいいだろう。
 - ・ 迅速な取引を防ぐ時間制限を設定して、フラッシュ・クラッシュの問題を解決しなかったのはなぜか？これは、技術的な問題ではなく、政策の問題かもしれない。
 - ・ アルゴリズムが完全ではないことを考えると、金融市場はより多くの人間の監視を組み込んだ方がいいだろうか？
- ・ AIが非常に進歩し、急速かつ指数関数的に向上する社会に対して、どのような教育モデルが最適な準備となるか、それはリベラルアーツだろうか？私たちの教育モデル(約20年間の教育/専門分野と、それに続く40年間の専門分野の仕事)は時代遅れだろうか？
- ・ AIは“金融民主主義”に影響を及ぼすか(誰もが投資にアクセスして、金融資産(株式市場や債券市場など)の評価に参加できるようになるだろうか、それとも最も強力な金融AIを持っている個人/機関だけがすべての利益を獲得し、退職金は成長の見込みのないままになるのではないだろうか)？
- ・ 多くの開発途上国は、輸出ベースの成長モデルに依存して貧困から抜け出し、安価な製造業を模索してきた。先進国のロボットとAIが、発展途上国からの安価な製造業に取って代わったとしたら、貧しい国はどうやって貧困から抜け出せばよいだろうか？
- ・ いつ、どのような順序で様々な仕事が自動化されると期待すべきか、これは所得の不平等にどのように影響するか？
- ・ どのような政策が、ますます自動化された社会の繁栄を援助できるか？AIによって生み出された富をどのように扱えば、不完全雇用の人々を最もよく支援できるだろうか？教育改革、見習いプログラム、労働力を必要とするインフラ・プロジェクト、所得再分配、社会的セーフティ・ネットなどの介入の長所と短所とは何か？
- ・ 社会的流動性の問題に対する潜在的な解決策とは何か？
 - ・ 社会的流動性の問題を解決する教育システムを設計できるか？最適な教育モデルとは何か？教育はこの問題を解決するための最良の方法か？
 - ・ 私たちは最大のリソースをどこに投資すべきかを調査した方がよいだろう。私たちの最も明るい心と最大の努力をどこに向けるべきだろうか？
 - ・ 企業秘密を保持するよりも、民間部門が知識を共有する方を奨励するか？

セッション#2:アルゴリズムと法律

概要

議論はおもに刑事司法制度におけるアルゴリズムの使用が注目された。より具体的に、参加者は、各囚人の仮釈放を決定するアルゴリズムの利用、最も多く犯罪が行われる場所、誰が犯罪に手を染めそうか、誰が犯罪の犠牲者になりそうかを予測するアルゴリズムの利用、囚人に対する保釈金を決定するアルゴリズムの利用を検討した。

アルゴリズムによる偏見は、参加者によって表明された主要な懸念事項の 1 つとなった。偏見は、人種、犯罪歴、性別、社会経済的状況といった特定のデータ・タイプを使用して、アルゴリズムの出力を通知することによって発生する。しかし、これらのデータ・タイプを除外すると、出力がより少ないデータによって決定されることを意味するため、偏見の問題が悪化する可能性がある。更に、特定のデータ・タイプが、個人に関する他の特性を推測するためにリンクされ、使用できるため、データセットから特定のデータ・タイプを削除するだけでは、アルゴリズムの偏見を解決するのに充分でないかもしれない。データ・タイプ、特に犯罪歴を排除することは、犯罪者がいずれ再犯に手を染めることを思いとどまらせるなど、刑事司法制度の基本的な目的のいくつかを弱めてしまうかもしれない。裁判官は、判決を下す際に、囚人が再犯者であるかどうかを検討する。犯罪歴がデータセットから削除された場合、アルゴリズムは司法上の意思決定のこの側面をどのように計算すべきだろうか？グループの何人かのメンバーは、偏見のないアルゴリズムの作成に成功したとしても、システムを公平にするために、司法システム全体がそのアルゴリズムの使用を要求する必要があると指摘した。普遍的な採用がなければ、特定のコミュニティが他のコミュニティよりも公平に扱われないかもしれない。例えば、裕福なコミュニティがアルゴリズムの使用を拒否し、不利な立場にあるグループのみがアルゴリズムによって評価された場合はどうなるのか？他の参加者は、人間の裁判官は不公平で偏見がある可能性があり、少なくとも司法の意見が一貫していることを保証するアルゴリズムは、すべての人に利益をもたらすかもしれない、と反論した。提示された代替解決案は、裁判結果のチェックや赤旗として機能する、パイロット支援システムに似た、裁判官用の支援ツールとしてアルゴリズムを導入することであった。

参加者によって表明されたもう 1 つの懸念事項には、正義、説明責任、正当性に対する人間の認識に関連する問題が伴った。議論は、人間が必要とする報いや、正義が行使されたという感覚に触れた。同様に、グループは、司法上の決定を下す、または支援するためのアルゴリズムの使用が、説明責任を裁判官から遠ざけるかもしれない点について議論した。裁判官がアルゴリズムの出力に異議を唱えるための他の情報源を持っていなければ、裁判官はアルゴリズムに頼って決定を下すかもしれない。更に、アルゴリズムの出力が公式の承認印を意味するのであれば、裁判官は判決の責任を容易に問われることはないだろう。もう 1 つの懸念は、所与の事件の判決が正しいかどうかを判断するテクノロジーを信頼することにより、裁判官が法廷制度に個人的に没頭する機会が減り、最終的に法廷の正当性を損なう可

能性であった。

議論された他の主要なトピックスは、犯罪防止プログラムに情報を提供するためのアルゴリズムの使用であった。シカゴ警察は、アルゴリズムを使用して、犯罪に手を染める可能性のある個人、犯罪の被害者になる可能性のある個人を特定し、これらの個人に社会福祉からの支援を積極的に提供している。カリフォルニア州リッチモンドでの同様のプログラムは、地域社会の問題を抱えた若者の間で、よい行動を促すための金銭的インセンティブを提供している。シカゴ警察とは異なり、このプログラムは、最大リスクにさらされている個人を特定するためのアルゴリズムではなく、人間の調査に基づいている。グループは、この場合の予測アルゴリズムがどのように識別プロセスを自動化できるかについて議論した。しかし、どちらの例でも、州は特定の集団の人々にのみ支援を提供しているため、グループはこれらのプログラムの前提そのものが公正であるかどうかにも検討した。繰り返しになるが、偏見は、これらのプログラムの対象となる個人に影響を与えるかもしれない。

グループは、アルゴリズムが、人間の能力を活用し、“友達“ではなく“ツール“として機能するように構築された場合に最も役立つ可能性がある」と結論づけた。より具体的には、参加者は、所与のアルゴリズムがどのように機能するか、及びアルゴリズムが決定を行うためにどのタイプのデータを使用するかについて、より透明性を高める必要があることに同意した。正当性の問題と、支援プログラムが人間の意思決定の代わりになる可能性を除けば、プロセスがどのように機能するかについて透明性があり、アルゴリズムにその理由を説明するよう求めることができるならば、裁判官が支援ツールとしてアルゴリズムを使う点において、参加者はより安心感を覚えるようであった。更に、何人かの参加者は、意思決定プロセスに関する透明性が将来の犯罪を思いとどまらせるかもしれないと指摘した。

質問と結論

- ・ グループは、アルゴリズムが、人間の能力を活用し、“友達“ではなく“ツール“として機能するように構築された場合に最も役立つ可能性がある」と結論づけた。
- ・ プロセスがどのように機能するかについて透明性があり、アルゴリズムにその理由を説明するよう求めることができるならば、裁判官が支援ツールとしてアルゴリズムを使う点において、参加者はより安心感を覚えるようであった。
- ・ AI システムの透過性を高めるにはどうすればよいか？
- ・ アルゴリズムを設計するために、どのデータを使用すべきか？テクノロジーの有効性を損なうことなく、アルゴリズムの偏見をどのように排除できるか？
- ・ 説明責任や正義感といった意思決定の人間的要素を組み込んだアルゴリズムをどのように設計できるか？
- ・ 証拠を評価するアルゴリズムの設計に投資すべきか？特定のケースに利用可能な証拠データをどのように適用するか？

セッション#3:自律型兵器

概要

グループは、AWS を”ターゲットを独立的に選択して交戦できる武器システム”と定義した。しかし、どの武器を自律的と見なすべきかを決定する段になると、特に、最近まで自律的と考えられていたが、最早そう思われていない軍事利用兵器(ドローンはターゲットを選択して交戦しないものの、ドローンの特徴は、ある種のミサイル(例:熱感知ミサイル)同様、自律的である)があるため、グループ内で意見の不一致が生じた。

AWS を規制する最善の方法や、このテクノロジーの開発と利用を米国内、或いは国際的に禁止するかについての議論があった。見えてきた包括的な疑問は、“賢すぎる“という理由から AWS を懸念しているのか、将来の AWS の性能を心配しているのか、或いはテクノロジーが戦争法に準拠するには“愚かすぎる“という理由から AWS を懸念しているのか、ということであった。例えば、この分野の専門家の何人かは、攻撃の比例性と必要性を判断するプロセスの概要を示す既存の法的要件を自動化できるかどうか、疑問視している。会場で投票したところ、殆どすべての参加者が、多大な損失を招く自律型兵器の使用を禁止する国際条約を交渉するよう政府に働きかけることを望む一方で、AI 開発に制限を設けるべきではないと考える参加者が 1 人、また他国の政府が国際規制の交渉に同意したかどうかに関わらず、米国は AWS の研究を進めるべきではないと考える参加者が 1 人いた。米国は AWS 研究を全面的に禁止すべきか、一連の国際規制を交渉すべきかで葛藤する参加者もいたので、彼らの票は両方のカテゴリーでカウントされた。何人かの参加者は、AWS は非国家主体に利益をもたらす可能性が高いため、禁止は効果的ではないと論じた。大量生産された AWS が安価になり、闇市場ですぐに利用できるようになるので、国際的な禁止がないと、AWS の軍拡競争を引き起こし、最終的には非国家主体を強化し、国家主体を弱体化させるだろうと主張する参加者もいた。禁止に賛成する参加者はまた、攻撃目標の決定を機械に任せることによって発生する道徳的懸念、法的懸念、説明責任の懸念を強調した。

防衛の観点から、他の主体(国家、及び非国家)がこのテクノロジーを開発し続けるので、米国は自律型兵器を開発するしかない、と主張する人もいた。そういうわけで、米国は研究を防衛的な AWS のみに限定することが提案された。しかし、米国が防衛的な AWS のみを開発したとしても、予測不可能の望ましくない 2 次効果が存在するかもしれない。更に、AWS の開発の継続を拒否することで、米国は市場のリーダーとして、他の国々にもこの分野の AI 研究を放棄するように促すことができるかもしれない。

AWS を規制する最善の方法については様々な意見があるが、グループは、これらの規制の施行は難しいかもしれない点に同意した。例えば、武器が自律的であるかどうかを配備者だけが知っている場合、制限を実施することは不可能かもしれない。更に、自律型兵器の使用に関する責任を個々の配備者または国民国家に負わせるのだろうか?もし私たちが国民国家に責任を負わせることを選択したら、他国が規制を遵守することを私たちが信用できないかもしれないので、AWS の開発と使用を規制することは困難だろう。その規制には、民

間用に開発されたテクノロジーが軍事用に転用される可能性があることも考慮に入れる必要があることが指摘された。最後に、かつては自律的だと定義されていたかもしれない高度な兵器の過去の例を見てみると、米国が歴史的にケースバイケースで規制上の懸念に対処してきたことは明白なので、規制に関する疑問は大きな懸念事項ではないと主張する参加者もいた。

軍隊で AWS を使用することのコストと利点も、複数の視点を取り入れた議論を生み出した。何人かの参加者は、殺戮の点でより効率的であるかもしれない AWS ならば、戦争での死者数を減らすことができるだろうと主張した。非常に多くの命を守ることができるなら、AWS の研究を禁止することを、どのように正当化するのか? 反対に、人命の喪失により、各国は将来的に戦争へ参入することを思いとどまると推論する参加者もいた。更に、参加者は戦争の民主的な側面を検討した。AWS は、戦争を行う権力を一般市民から政治家や社会的エリートに移すことにより、民主主義の中に不均一な権力構造を生み出すかもしれない。また、AWS を使用すると、殺戮行為から厄介で個人的な側面が排除されるため、殺戮が増えるのではないか、という点について議論した。ルワンダでの大量虐殺といった歴史的事件は、殺戮の個人的な側面が暴力を思いとどまらせることはなく、従って AWS を使っても必ずしも死者数は増減しないことを示唆している、と感じる参加者もいた。

長期研究に関する疑問には、AWS の研究と使用を米国内で禁止すべきか、国際的に禁止すべきか、また禁止しない場合、研究を防衛的な AWS に限定すべきか、という内容もあった。更に、AWS の開発と使用に関するグローバルな規制の決定を担当する国際条約または国際団体があるならば、それらの規制はどのようなものになり、どのように施行されるだろうか? 更に、その規則には、説明責任と施行に関する懸念への対処に向けて、殺戮スイッチや透かしといった特定の AWS 機能を組み込むことに関する条項を含めた方がいいだろうか? また、AWS が禁止されているかどうかに関係なく、民間用の AI システムが軍事用に転用されるかもしれない可能性について、どのように説明すべきだろうか?

質問と結論

- ・ 殆どの参加者(すべてではない)は、世界の大国が AWS の研究と使用を導く何らかの形式の国際規制を起草した方がよい点について同意した。
 - ・ 契約条項はどうあるべきか? 禁止または規制はどのようにすべきか? 意思決定者の合意形成方法は?
 - ・ 特に非国家主体が AWS を手に入れて使用する可能性が高いことを考慮すると、規制はどのように施行されるべきか?
 - ・ AWS の研究開発の説明責任を、どのように国民国家に負わせるか? 例えば、国家は厳格な責任を問われるべきだろうか?
- ・ または、AWS の研究に関する懸念事項の一部を説明する現行の規制はすでにあるだろうか? 米国は、ケースバイケースで兵器開発を評価し続けるべきだろうか?

- ・ 特定の種類の AWS の開発と使用を禁止すべきだろうか?攻撃的な AWS だけか?致命的な武器だけか?禁止は効果的だろうか?
- ・ AWS に、殺戮スイッチといった特定の安全/防御機能の実装を求めるべきだろうか?
- ・ たとえ米国が研究を防衛的な AWS に限定したとしても、望ましくないかもしれない予想外の結果をどのように説明するか?例えば、武器が防衛的であるかどうか、どのように定義するか?もし“防衛”兵器が他国の領土に送られ、攻撃されたので反撃に出たとしたら、どうか?それは攻撃的兵器、または防衛的武器と見なせるか?
- ・ 民間 AI の軍事的転用をどう防止するか?
- ・ この議論には、より多様な視点を盛り込む必要がある。これらの問題を解決するには、弁護士、哲学者、経済学者からのより多くの意見が必要となる。

セッション#4:人間レベルの AI が出現することのリスク

概要

最後のセッションでは、生命の未来研究所が資金提供し、超人的知能を備えた将来の AI を有益なものにすることに焦点を当てた、進行中の研究についての発表が行われた。グループのメンバーは、もし人間レベルの AI の開発に成功するとしたら、それはいつなのかについて、様々な見解を表明した。狭義の AI に限定されている現在の AI 技術を集約させて、高度な汎用 AI システムを構築するためには、解決すべき課題がたくさんあるとの議論があった。逆に、スケジュールは不確かなままだが、人間レベルの AI は達成可能であると主張する参加者もいた。彼らは、この技術の社会的、経済的、及び法的な結果に備えるため、更なる研究の必要性を強調した。

議論された 2 番目の主要なポイントは、人間レベルの AI の潜在的な危険性であった。高度な AI システムが人間の作成者の制御を超えて行動する能力を持つ可能性について、参加者の意見は一致しなかった。それに応えて、他のグループメンバーは、高度な汎用 AI システムは、狭義の AI システムのように特定の一連の指示に従うのではなく、新しいタスクを学習するように設計されているため、タスクを完了するために自衛本能やリソースの必要性から行動するかもしれず、制御が難しいかもしれないし、汎用 AI と狭義の AI の違いを考慮すると、テクノロジーを設計する際、汎用 AI システムの社会的、感情的な側面を考慮することが重要になるだろう。例えば、テクノロジーが人間の価値観を学び、尊重することを可能にするシステムを構築する最良の方法についての議論があった。同時に、人間レベルの知能を持たない現在の AI システムも、“野生に“放たれた場合、人間に脅威を与える可能性があることが指摘された。動物と同様に、悪意のある主体によって操作されたり、予測不可能で意図しない結果をもたらす可能性のあるコードを備えた AI システムは、人間にとっても同様に危険であることが判明するかもしれない。人間レベルの AI の制御を監視、維持するための推奨事項のいくつかを、既存の AI システムに適用することが提案された。

結論として、グループは、安全設計の倫理に関連する語彙を含む、ロボットと AI の用語を

定義する必要性など、更なる研究を必要とするいくつかの領域を特定した。例えば、ロボットを人間のような特性を持つ生物として定義した方がよいだろうか、それとも物体または人間の支援ツールとして定義した方がよいだろうか?法律、哲学、コンピューター・サイエンスなど様々な分野の専門家は、高度な AI によって生み出されるかもしれない様々な課題の解決策を協力して考案するに当たって、共有語彙を必要とするだろう。更に、人間レベルの AI はすぐに超人的 AI になる可能性があるため、一部の参加者は、この分野において、安全性、制御、及び AI の価値観と人間の価値観の調整に関連する更なる研究の必要性を強調した。また、AI の社会的、感情的要素に関する考え方に影響を与えるかもしれないので、社会性に対する私たちの理解が変化しているかどうかを調査することが提案された。最後に、弁護士は、将来の AI システムに組み込むべき指導的な倫理原則があるかどうかを判断するために、法的責任と不法行為法を調査すべきである。

質問と結論

- ・ グループは、人間レベルの AI のスケジュールと潜在的な結果に関するあらゆる視点からの意見を提示した。まだ検討する必要がある多くの質問が残されている。
- ・ グループは、安全設計の倫理に関連する語彙を含む、ロボットと AI の用語を定義する必要性を特定した。例えば、ロボットを生物として定義した方がよいだろうか、それとも物体として定義した方がよいだろうか?法律、哲学、コンピューター・サイエンスなど様々な分野の専門家は、高度な AI によって生み出されるかもしれない様々な課題の解決策を協力して考案するに当たって、共有語彙を必要とするだろう。
- ・ グループは、AI システムの社会的、感情的な側面を検討した。人間と機械の協力関係の本質とは何か?AI の社会的、感情的要素に関する考え方に影響を与えるかもしれないので、社会性に対する私たちの理解が変化しているかどうかを調査した方がよい。
- ・ 人間レベルの知能を持つ AI の研究と使用を監視する法的枠組みをどのように設計すべきだろうか?
- ・ 弁護士は、将来の AI システムに組み込むべき指導的な倫理原則があるかどうかを判断するために、法的責任と不法行為法を調査すべきである。
 - ・ 人間の価値観、法律、社会規範が時間とともに変化していく中で、これに適応できるテクノロジーをどのように構築することができるだろうか?