
[研究ノート]

AI 社会が個人にもたらす「価値」とそのトレードオフ：日米欧の分析をもとに
Values for Individuals in AI Society and Their Trade-off: Analysis of Guidelines in Japan, the US and Europe

国際社会経済研究所 小泉 雄介

Institute for International Socio-Economic Studies, Yusuke Koizumi

慶應義塾大学 早田 吉伸

Keio University, Yoshinobu Soda

要 旨

「個人データは 21 世紀の新たな石油である」と言われ、AI 社会では個人データの利活用による様々なサービスの提供や社会的課題の解決が期待されている。国内外のガイドライン類では、個人データの流通や利活用が社会・企業・個人にもたらす「価値」について様々な言及がなされているが、それが誰に対する価値なのかを峻別した整理が必ずしもなされていない。本稿で AI 社会において尊重すべき「価値」のうち、個人データの利活用が「個人に対して」もたらす価値について、ガイドライン類から抽出した結果、「自由」「平等」「安全」「利便」の 4 カテゴリーに分類できることが明らかになった。個人がデータ利活用と通じてこれらの価値を追求すると、個人にとってベネフィットがあるのみならず、リスクも発生するというトレードオフの関係が見られる。例えば、監視カメラを設置して不審者の画像を取得することは「安全」に役立ち、社会的にも受容されている。しかし、更なる安全を目指して監視カメラ設置台数を増やしたり公道で顔照合を行うことは、個人の「自由」に対して重大な萎縮効果をもたらす。これらベネフィットとリスクの間のバランスをどのように取るかは、国や社会の在り方によって変わらう。AI 社会において、個人にとっての諸価値間のバランスをどのように取るべきか、また個人にとっての価値と社会や企業にとっての価値との間のバランスをどのように取るべきかは、今後の重要な検討課題である。

キーワード

AI、個人データ、価値、原則、倫理

1. はじめに：研究の目的

「個人データは 21 世紀の新たな石油（価値あるリソース）である¹⁾」と言われている。AI 社会²⁾では、スマートフォン・センサー等の様々な機器を通じて個人データが取得され、それらのデータが AI（人工知能）技術等を用いた情報システム上で分析され、様々なサービスの提供や社会的課題の解決に利活用されている。また、個人データが分析され、ローン・保険・採用・教育・医療等の場面で、当該個人に関する何らかの判断や意思決定が行われることもある。

すでに国内外の政府文書やガイドライン等において、AI 社会における個人データの流通や利活用が社会・企業・個人にもたらす「価値」について、様々な言及がなされている。しかし、多くの文献においては、個人データがもたらす価値について、それが社会全体に対する価値なのか、企業に対する価値なのか、それとも個人に対する価値なのかが渾然一体となっており、それらを峻別した整理が行われていな

い。筆者は、AI 社会における何らかの個人データの利活用が仮に企業や社会全体に利するものであったとしても、それが当該個人に不利益を及ぼすものであるならば、そのような個人データ利活用が必ずしも望ましいとは考えない。社会全体や企業にとっての価値も、我々が社会生活を送る上では当然に重要であるが、個人にとっての価値が最も基本的なものと考え、本稿では AI 社会における個人データの利活用が「個人に対して」もたらす価値を取り上げることにする。

次節以降では、AI 社会において尊重すべき「価値」のうち、個人データの利活用が「個人に対して」もたらす価値について、既存文献から抽出し、整理する（第 2 節）。そして、それらの諸価値を追求した際に生じるベネフィットのみならず、リスクについても洗い出し、ベネフィットとリスクがトレードオフの関係にあることを浮き彫りにする（第 3 節）。最後に、今後の研究に向けての課題について述べる（第

4 節)。

2. AI 社会において尊重すべき「価値」の抽出

AI 社会の諸原則については、すでに国内外の政府機関等がガイドライン類を発行している。これら AI 原則・ガイドライン類の策定においては、日本政府が他国に先んじており³、総務省の AI ネットワーク社会推進会議が 2017 年に AI 開発ガイドライン案を、2018 年 7 月には AI 利活用原則案を含む報告書を公表している。また、内閣官房の人間中心の AI 社会原則会議が 2019 年 3 月に、AI 社会原則を公表している。

海外に目を向けると、欧州では政府機関がガイドライン類の策定に精力的であり、2018 年 3 月には欧州委員会の諮問機関が倫理原則を公表し、それを受けて欧州委員会の AI ハイレベル専門家グループが 2019 年 4 月に AI 倫理ガイドラインを公表している。英国でも 2018 年 3 月に、議会の AI 特別委員会が AI の包括的原則を公表している。他方、米国ではこのようなガイドライン類の策定は民間主導で進められており、代表的な取組みとして、米国電気電子学会 (IEEE) が 2017 年に AI の設計・開発・実装における一般原則を、研究機関 Future of Life Institute が同年にアシロマ AI 原則を公表している⁴。

これらの文献を参照することにより、AI 社会において尊重すべき価値のうち、共通項目として、「個人にとっての価値」を抽出する。

2.1 AI ネットワーク社会推進会議の「報告書 2018」

総務省の AI ネットワーク社会推進会議では、2018 年 7 月に「報告書 2018」⁵を公表し、AI の利用者やデータ提供者が AI の利活用において留意することが期待される事項を AI 利活用原則案として取りまとめている。同報告書ではさらに、「智連社会⁶」を形成するに当たって則るべき基本理念として、以下の 8 項目を掲げている。

- ・ すべての人々による恵沢の享受
- ・ 人間の尊厳と個人の自律

- ・ イノベーティブな研究開発と公正な競争
- ・ 制御可能性と透明性
- ・ ステークホルダの参画
- ・ 物理空間とサイバー空間の調和
- ・ 空間を越えた協調による活力ある地域社会の実現
- ・ 分散協調による地球規模の課題の解決

これら 8 項目のうち、個人にとっての価値に焦点を当てているものは、「すべての人々による恵沢の享受」「人間の尊厳と個人の自律」である。前者では「平等・公平性・機会均等」に関わる内容、後者では「自由・自律・尊厳」「安全・安心」「利便・幸福」に関わる内容が挙げられている。

2.2 内閣官房の「人間中心の AI 社会原則」

内閣官房の人間中心の AI 社会原則会議が 2019 年 3 月に公表した原則⁷では、理念として尊重し、その実現を追求する社会を構築していくべき価値として、以下の 3 つを掲げている。

- (1) 人間の尊厳が尊重される社会 (Dignity)
- (2) 多様な背景を持つ人々が多様な幸せを追求できる社会 (Diversity & Inclusion)
- (3) 持続性ある社会 (Sustainability)

これら 3 つの価値のうち、個人にとっての価値に焦点を当てているものは(1)と(2)である。(1)においては「自由・自律・尊厳」に関わる内容、(2)においては「平等・公平性・多様性」に関わる内容が挙げられている。

2.3 欧州委員会 EGE の「AI に関する宣言」

欧州委員会の諮問機関である科学新技術倫理グループ (EGE) が 2018 年 3 月に公表した「AI・ロボティクス・自律システムに関する宣言」⁸では、EU 条約と EU 基本権憲章で規定された基本的価値に基づく倫理原則として、以下の 9 つを提案している。

- (a) 人間の尊厳

- (b) 自律
- (c) 責任
- (d) 正義、平等、連帯
- (e) 民主主義
- (f) 法の支配、説明責任
- (g) セキュリティ、安全性、身体・精神的完全性
- (h) データ保護、プライバシー
- (i) 持続性

これら 9 原則において、個人にとっての価値に焦点を当てているものとして、(a)では「尊厳」、(b)では「自由・自律」、(d)では「平等・公平性・機会均等」、(e)では「多様性」、(g)では「安全」、(h)では「自由・プライバシー」に関わる内容が挙げられている。

2.4 欧州委員会の「AI 倫理ガイドライン」

欧州委員会の AI ハイレベル専門家グループが 2019 年 4 月に公表したガイドライン⁹⁾は、AI が人間中心的な方式 (human-centric manner) で開発されることを保証するために遵守しなければならない 4 つの原則として、以下を掲げている¹⁰⁾。

人間の自律の尊重の原則 (The principle of respect for human autonomy)

基本的人権は、人間の自由と自律の尊重を保証するためのものである。AI システムとやり取りする人間は、自分自身に対する完全かつ有効な自己決定を維持し、民主的プロセスに参加できなければならない。AI システムは人間を不当に服従させたり、支配したり、騙したり、操ったり、調整したり、管理したりすべきではない。むしろ、AI システムは人間の認知的・社会的・文化的スキルを増大化したり、補完したり、強化するために設計されるべきである。人間と AI システムの間の機能の割り当ては、人間中心の設計の原則に従うべきであり、人間の選択に興味のある機会を残すべきである。このことは、AI システムのワークプロセス全体に渡って人間の監督を保証することを意味する。また AI システムは仕事の環境を根本的に変えてしまうかもしれない。AI シス

テムは仕事の環境において人間をサポートするべきであり、意味ある仕事を創出することを目指すべきである。

害悪防止の原則 (The principle of prevention of harm)

AI システムは、人間に害を与えたり、悪い影響を与えるべきではない。AI システムは、人間の尊厳、精神的・身体的完全性 (integrity) を保護するべきである。AI システムとそれがオペレートされる環境は安全でセキュアでなければならない。AI システムは技術的に堅牢でなければならない、不正な利用がなされないことが保証されるべきである。脆弱な人々はより多くの注意を受けるべきであり、AI システムの開発・配置・利用段階に含まれるべきである。また AI システムが、雇用主と従業者、企業と消費者、政府と市民のように、権限や情報の非対称性に起因する悪影響を生み出したり増大化するような状況に対しては、特別な注意が払われなければならない。害悪を防止することには、自然環境や全ての生物を考慮することも含まれる。

公平性の原則 (The principle of fairness)

AI システムの開発・配置・利用は公平でなければならない。公平性は実質的な次元と手続き的な次元を持つ。実質的な次元は、ベネフィットとコストの平等な分配を保証すること、不公平なバイアス・差別・汚名から個人や集団が自由であることを保証することである。もし不公平なバイアスが回避されれば、AI システムは社会的な公平性を増加させることができるだろう。教育・公益・サービス・技術へのアクセスの観点からの平等な機会もまた促進されるべきである。さらに、AI システムの利用は人々を欺いたり、選択の自由を不当に損なうものであるべきではない。公平性の手続き的な次元には、AI システムやそれをオペレートする人間によってなされた意思決定に対して異議申立てしたり、有効な救済手段を求める権利が含まれる。

説明可能性の原則 (The principle of explicability)

説明可能性は、AI システムに対する利用者の信頼を構築・維持する上で非常に重要である。プロセス

は透明であり、AI システムの能力と目的は公開され、意思決定は可能な範囲で人々に対して説明可能でなければならない。このような情報が無ければ、人々は意思決定に対して十分に異議申立てを行うことができない。なぜシステムが特定のアウトプットや意思決定を行ったのか（そしてどのインプットの組み合わせがアウトプットに寄与したか）についての説明は、常に可能なわけではない。それらのケースは、「ブラックボックス」アルゴリズムとして言及され、特別な注意を必要とする。それらのケースでは、他の説明可能性の手法（トレーサビリティ、監査可能性、システム能力に関する透明性など）が必要とされるかもしれない。説明可能性が必要とされる程度は、文脈と、アウトプットが不正確だった場合の帰結の重大性に大きく依存する。

これら 4 原則において、個人にとっての価値に焦点を当てているものとして、「人間の自律の尊重の原則」では「自由・自律」「利便」、「害悪防止の原則」では「尊厳」「多様性」「安全」、「公平性の原則」では「自由」「平等・公平性・機会均等」、「説明可能性の原則」では「自由」に関わる内容が挙げられている。

2.5 英国貴族院 AI 特別委員会の「英国における AI」

同委員会の 2018 年 3 月の報告書「英国における AI: 英国は AI を活用し、そして活用できる準備ができてるか」¹¹では、AI の倫理的行動規範のための 5 つ包括的原則として、以下を掲げている¹²。

- (1) AI は人類の公益とベネフィットのために開発されるべきである。
- (2) AI は分かりやすさと公平性の原則に則って動作するべきである。
- (3) AI は個人・家族・共同体のデータの権利やプライバシーを損なうような目的で使われるべきではない。
- (4) 全ての市民は AI と共に精神的・感情的・経済的に繁栄できるように教育を受ける権利を持つ。
- (5) AI に人間を傷つけたり破壊したり欺いたりす

る自律的能力を決して与えるべきではない。

これら 5 原則において、個人にとっての価値に焦点を当てているものとして、(2)では「公平性」、(3)では「プライバシー」、(4)では「幸福」、(5)では「安全」に関わる内容が挙げられている。

2.6 米国電気電子学会 (IEEE) の「倫理的に調整された設計」

IEEE の 2017 年 12 月の報告書¹³では、AI の倫理的な設計・開発・実装において参照されるべき一般原則として、以下の 5 つを掲げている。

- (1) 人権
- (2) 幸福の優先
- (3) アカウンタビリティ
- (4) 透明性
- (5) AI 技術の濫用とその認知

これらの 5 原則において、個人にとっての価値に焦点を当てているものとして、(1)では「自由・自律・尊厳・プライバシー」「多様性」「安全」、(2)では「幸福・生活の質」、(4)では「安全」「幸福」に関する内容が挙げられている。

2.7 米国 Future of Life Institute の「アシロマ原則」

2017 年 2 月公表のアシロマ原則¹⁴では、AI 開発が基づくべき原則として、以下の 23 個を挙げている。

- (1) 研究目標、(2) 研究資金、(3) 科学と政策の連携、(4) 研究文化、(5) 競争の回避、(6) 安全性、(7) 障害の透明性、(8) 司法の透明性、(9) 責任、(10) 価値観の調和、(11) 人間の価値観、(12) 個人のプライバシー、(13) 自由とプライバシー、(14) 利益の共有、(15) 繁栄の共有、(16) 人間による制御、(17) 非破壊、(18) 人工知能軍拡競争、(19) 能力に対する警戒、(20) 重要性、(21) リスク、(22) 再帰的に自己改善する人工知能、(23) 公益

表 1 国内外のガイドライン類における個人にとっての「価値」の抽出と分類

地域	日本	日本	欧州	欧州	英国	米国	米国
文献名	AI ネットワーク社会推進会議「報告書 2018」	内閣官房「人間中心の AI 社会原則」	欧州委員会 EGE「AI に関する宣言」	欧州委員会「AI 倫理ガイドライン」	貴族院 AI 特別委員会「英国における AI」	米国電気電子学会(IEEE)「倫理的に調整された設計」	Future of Life Institute「アシロマ原則」
価値							
自由	自由 自律 尊厳	自由 自律 尊厳	自由 自律 尊厳 プライバシー	自由 自律 尊厳	プライバシー	自由 自律 尊厳 プライバシー	自由 自律 尊厳 プライバシー
平等	平等 公平性 機会均等	平等 公平性 多様性	平等 公平性 機会均等 多様性	平等 公平性 機会均等 多様性	公平性	多様性	平等 多様性
安全	安全 安心		安全	安全	安全	安全	安全
利便	利便 幸福			利便	幸福	幸福 生活の質	繁栄

これらの 23 原則において、個人にとっての価値に焦点を当てているものとして、(6)では「安全」、(11)では「自由・尊厳」「多様性」、(12)では「プライバシー」、(13)では「自由・プライバシー」、(14)では「平等」、(15)では「平等」「繁栄」、(16)では「自律」に関わる内容が挙げられている。

3. 個人が各「価値」を追求することのベネフィットとリスク

前節で国内外の文献から共通項目として抽出した、AI 社会における個人データ利活用によって個人にもたらされる「価値」は、以下の 4 つのカテゴリに分類することができる（表 1 参照）。

(1)「自由」

このカテゴリには、行動の自由、思想・表現の自由、自律、尊厳、プライバシーなどが包含される。

(2)「平等」

このカテゴリには、平等、公平性、機会均等、多様性などが包含される。

(3)「安全」

このカテゴリには、身体の安全、財産の安全、安心、健康などが包含される。

(4)「利便」

このカテゴリには、利便、幸福、繁栄、生活の質（quality of life）などが包含される。

これらの価値を追求することは、個人にとって当然にメリットがあることである。ただし、個人データの特徴として、データ利活用を通じてこれらの価値を追求すればするほど、あたかも諸刃の剣のように、個人に対するデメリット／リスクも発生してしまうケースが存在する。

例えば、監視カメラを設置して不審者の画像を取得することは治安の向上（「安全」）に役立ち、社会的にも受容されている¹⁵。しかし、更なる治安向上を目指して監視カメラ設置台数を増やしたり、街角で AI を用いたリアルタイム顔照合を行ったりすることは、個人の「自由」に対して重大な萎縮効果をもたらす。

このような、「価値」毎のベネフィットとリスクの例をまとめると、表 2 のようになる¹⁶。

表 2 で挙げた諸事例においては、或る価値を追求すると当該価値に関するベネフィットが得られると同時に、当該価値または他の価値に関するリスクが発生するというトレードオフの関係が見られる。

表 2 AI 社会における個人にとって尊重すべき「価値」と、ベネフィット・リスクの例

個人データを利活用することの個人にとっての <u>価値</u>	個人にとっての <u>ベネフィット／メリット</u> の例	個人にとっての <u>リスク／デメリット</u> の例
(1) 個人の「 <u>自由</u> 」を追求すること	・ BMI ¹⁷ など機械と接続することによる人間の能力の <u>エンハンスメント</u>	・ 思考内容の <u>漏洩</u> 、人間の <u>尊厳を損ねる</u> リスク
	・ 企業に対して自分の個人データの <u>提供を拒否したり、削除</u> してもらうこと	・ 個人データを提供しない人々がアルゴリズムにおいて <u>過少代表</u> される
(2) 個人の「 <u>平等</u> 」を追求すること	・ 可能な限り多くの個人データを機械学習用データとして用いることによる <u>バイアスの排除</u>	・ 個人データ保護における <u>データ最小化の原則</u> や <u>利用制限の原則</u> への <u>抵触</u>
	・ 国民の生体情報登録による <u>公平な行政サービスの提供</u>	・ 国家による <u>国民監視</u>
(3) 個人の「 <u>安全</u> 」を追求すること	・ 監視カメラ設置や顔照合技術の利用による <u>治安の向上</u>	・ 監視・追跡による自由への <u>萎縮効果</u>
	・ ウェアラブル端末等による <u>健康見守り</u>	・ 疾病リスク等のプロファイリングによる <u>不利益・差別</u>
	・ 「ナッジ ¹⁸ 」による <u>健康的な選択肢</u> の提示	・ 極端な「ナッジ」による <u>選択の誘導</u>
(4) 個人の「 <u>利便</u> 」を追求すること	・ Web サイトにおけるパーソナライゼーション／レコメンデーションによる <u>情報の取捨選択</u>	・ フィルターバブル ¹⁹ 、選択の誘導／投票行動への介入／ステルスマーケティングなど、 <u>自律的な判断の阻害</u>
	・ 信用スコアを様々なサービスで汎用的に用いることによる <u>生活の利便性向上</u>	・ 信用スコアの数値に基づくサービス拒否や <u>社会的差別</u>
	・ 人を介さない自動処理による <u>迅速なサービス享受</u>	・ 就職・保険・ローン等における自動意思決定による <u>不利益・差別</u>

これらベネフィットとリストの間のバランスをどのように取るかは、国や社会の在り方（先進国と途上国、民主主義国家と社会主義国家など）によって変わりうる。例えば、前述の監視カメラの例で見れば、現状では先進国の方がカメラ設置台数が多い訳であるが、先進国においては（画像データの保持期間は最小限に留める、公道での顔照合は行わない等）個人の「自由」をより尊重した制度設計を行うべき

であり、他方、途上国では個人の「安全」をより重視した制度設計を行うべきと言える²⁰。

こうしたベネフィットとリスクの間のバランスとしては、以下の2つの側面が考えられる。表2で挙げたものは①の例であるが、②のバランス（トレードオフ関係）も重要な問題であろう。

① 個人にとっての諸価値間のバランスをどのように取るべきか。

- ② 個人にとっての価値と、社会にとっての価値もしくは企業にとっての価値との間のバランスをどのように取るべきか。

既存の文献では、後者のバランスを取ることの必要性について述べられている場合が多い。例えば、米国大統領行政府のビッグデータ報告書²¹では、「考慮すべき困難な問題」として「ビッグデータの社会的に有益な利用とプライバシーに対する被害と、より多くのデータがより多くのモノに関して必然的に収集される世界をもたらすその他の価値の間で、どのようにバランスをとるかと言う問題²²」を挙げている。個人データについては特に、(個人の) プライバシー保護と (企業による) データ利活用のバランスについて論じられることが多い。このような個人にとっての価値と、企業や社会にとっての価値のバランスをどのように取るかの問題は、功利主義に代表される倫理学の問題でもある。

4. 今後の研究に向けて

前節にまとめたように、AI 社会において個人データを利活用することで、個人にとっての価値を追求すればするほど、個人がベネフィットを享受すると同時に、まるでコインの裏表のように、個人に対するリスクが発生してしまうケースが存在する。

今後の研究課題としては、以下の3つが挙げられよう。

- (1) 諸外国の事例も参考に、AI 社会において個人にとっての価値を追求することによって生じるベネフィットとリスクの間のバランスをどのように取ればよいか。

- (2) AI 社会における個人データ利活用によってもたらされる価値のうち、社会全体にもたらされる価値や、企業にもたらされる価値についてはどのように整理すればよいか。

- (3) 社会全体や企業にもたらされる価値と、個人にもたらされる価値の間のバランスをどのように取ればよいか。

(受付日：2019年2月28日)

(受理日：2019年7月17日)

著者略歴

小泉雄介 (こいずみ・ゆうすけ)

1994年東京大学理学部地球物理学科卒、1996年東京大学教養学部科学史および科学哲学分科卒、1998年同大学院総合文化研究科中退。同年(株)NEC 総研入社。2010年より(株)国際社会経済研究所。専門領域は個人情報保護・プライバシー、電子政府 (国民ID制度)、新興国・途上国市場調査。主な著書に『国民ID導入に向けた取り組み』(共著)、『現代人のプライバシー』(共著)、『経営戦略としての個人情報保護と対策』(共著)など。日本セキュリティ・マネジメント学会会員。電子情報技術産業協会 (JEITA) 個人データ保護専門委員会客員 (2012年度～現在)。

早田吉伸 (そうだ・よしのぶ)

慶應義塾大学大学院システムデザイン・マネジメント研究科講師 (非常勤)。博士 (システムデザイン・マネジメント学)。専門は、ソーシャルイノベーション、社会情報学。

¹ World Economic Forum, *Personal Data: The Emergence of a New Asset Class*, 2011年2月, http://www3.weforum.org/docs/WEF_ITTC_PersonalDataNewAsset_Report_2011.pdf, 参照年月 2019年7月30日。

² 本稿では特にAI社会について定義しないが、例えば内閣官房「人間中心のAI社会原則」では、「AIに関わる技術自体の研究開発を進めると共

に、人、社会システム、産業構造、イノベーションシステム、ガバナンス等、あらゆる面で社会をリデザインし、AIを有効かつ安全に利用できる社会」を「AI-Readyな社会」と呼んでいる。

³ 平野晋, 「GAFA規制を考える (中) AI利活用で独走許すな」, 日本経済新聞 2019年2月20日朝刊, <https://www.nikkei.com/article/DGXXKZO4145989>

[0Z10C19A2KE8000/](https://standards.ieee.org/industry-connections/ec/autonomous-systems.html), 参照年月 2019 年 7 月 30 日。

⁴ その他、Google や Microsoft 等の民間企業も AI 原則を公表しているが、例えばアシロマ AI 原則にはこれらの企業の研究者等が署名を行っているため、本稿では一企業としての指針ではなく、複数企業が参画するような民間機関の指針を米国の代表例として取り上げた。

⁵ AI ネットワーク社会推進会議, 「報告書 2018 – AI の利活用の促進及び AI ネットワーク化の健全な進展に向けて」, 2018 年 7 月, http://www.soumu.go.jp/main_content/000564147.pdf, 参照年月 2019 年 7 月 30 日。

⁶ 同報告書では、「智連社会とは、AI ネットワーク化の進展の結果として、人間が主体的に技術を使いこなすことによって AI ネットワークと共生し、データ・情報・知識を自由かつ安全に創造・流通・連結して「智のネットワーク」を形成することにより、あらゆる分野におけるヒト・モノ・コト相互間の空間を越えた協調が進展し、もって創造的かつ活力ある発展が可能となる社会」と定義している。

⁷ 内閣官房人間中心の AI 社会原則会議, 「人間中心の AI 社会原則」, 2019 年 3 月, <https://www.cas.go.jp/jp/seisaku/jinkouchinou/pdf/aigensoku.pdf>, 参照年月 2019 年 7 月 30 日。

⁸ European Group on Ethics in Science and New Technologies, *Statement on Artificial Intelligence, Robotics and 'Autonomous' Systems*, 2018 年 3 月, https://ec.europa.eu/research/ege/pdf/ege_ai_statement_2018.pdf, 参照年月 2019 年 7 月 30 日。

⁹ The European Commission's High-Level Expert Group on Artificial Intelligence, *Ethics Guidelines for Trustworthy AI*, 2019 年 4 月, <https://ec.europa.eu/digital-single-market/en/news/ethics-guidelines-trustworthy-ai>, 参照年月 2019 年 7 月 30 日。

¹⁰ 抄訳は筆者による。

¹¹ House of Lords Select Committee on Artificial Intelligence, *AI in the UK: ready, willing and able?*, 2018 年 4 月, <https://publications.parliament.uk/pa/ld201719/ldselect/ldai/100/100.pdf>, 参照年月 2019 年 7 月 30 日。

¹² 翻訳は筆者による。

¹³ The IEEE Global Initiative on Ethics of Autonomous and Intelligent Systems, *Ethically Aligned Design (Version 2)*, 2017 年 12 月,

<https://standards.ieee.org/industry-connections/ec/autonomous-systems.html>, 参照年月 2019 年 7 月 30 日。

¹⁴ Future of Life Institute, *Asilomar AI Principles*, 2017 年 2 月, <https://futureoflife.org/ai-principles-japanese/>, 参照年月 2019 年 7 月 30 日。

¹⁵ 小泉雄介, 「監視社会とプライバシー: リトルブラザーの共存する世界へ」, 『日本セキュリティ・マネジメント学会誌』, Vol.32, No.2, 2018 年 9 月, 18-19 ページ, https://www.ise.com/jp/information/report/2019/20190116_JSSM.pdf, 参照年月 2019 年 7 月 30 日。

¹⁶ これらの例は各種文献を参考にしているが、基本的に著者が考案・整理したものである。

¹⁷ BMI (Brain Machine Interface) は、脳と機械をつなげる技術の総称。頭皮や脳内にセンサーや電極を取り付けて脳波等の活動を読み取ることで人間の思考を機械への命令に変換する技術などが研究されている。

¹⁸ 人々が自らの最善の利益を追求できるように配慮するが、あくまで強制はせず、各人が異なる選択肢を選ぶ自由を保障するもの。セーラーとサンステーションはこの立場を「ナッジ」と呼んでいる。児玉聡, 『功利主義入門』, ちくま新書, 2012 年, 121 ページ。

¹⁹ ニュース等が配信されるポータルサイトの構成や検索サイトの検索結果等は、プロファイリングを通じて推測された個人の趣味嗜好や政治的傾向等に応じて変わりえる。ネット空間では、その人の好みに合わない情報は予測エンジンのフィルターによって排除されるため、個人は自らの好みに合った情報のみに取り囲まれるようになる。イーライ・パリサーはこのような現象を「フィルターバブル」と呼んでいる。イーライ・パリサー (井口耕二訳), 『フィルターバブル』, 早川書房, 2016 年, 23 ページ。

²⁰ 小泉前掲書, 21-22 ページ。

²¹ Executive Office of the President, *Big Data: Seizing Opportunities, Preserving Values*, 2014 年 5 月, https://obamawhitehouse.archives.gov/sites/default/files/docs/big_data_privacy_report_may_1_2014.pdf, 参照年月 2019 年 7 月 30 日。

²² 翻訳は次世代パーソナルサービス推進コンソーシアムによる。
https://www.coneps.jp/contents/product_002.pdf, 参照年月 2019 年 7 月 30 日。