

AI 倫理とガバナンス： 世界動向と産業界の取り組み

中川裕志
(理化学研究所 革新知能統合研究センター)

スライド中の図はpower point の機能でダウンロードした
creative commons のライセンスです。

AI倫理指針

国内外の組織が提案している 人工知能の倫理(古い順)

- FLI: Asilomar AI Principles (23原則) (2017)
- IEEE Ethically Aligned Design, version 2(2017/12)
 - AIおよび開発者が持つべき倫理
- Partnership on AI (2016~)
- 総務省 AIネットワーク社会推進委員会
 - AI開発ガイドライン OECDに提案(2017)
 - AI利活用ガイドライン(2019)
- 内閣府 人間中心のAI社会原則(2019/3/29)
 - AI ready な社会の在り方 G20に提案
- IEEE Ethically Aligned Design, first edition (2019/3)
 - 倫理的なAIの設計指針

国内外の組織が提案している 人工知能の倫理

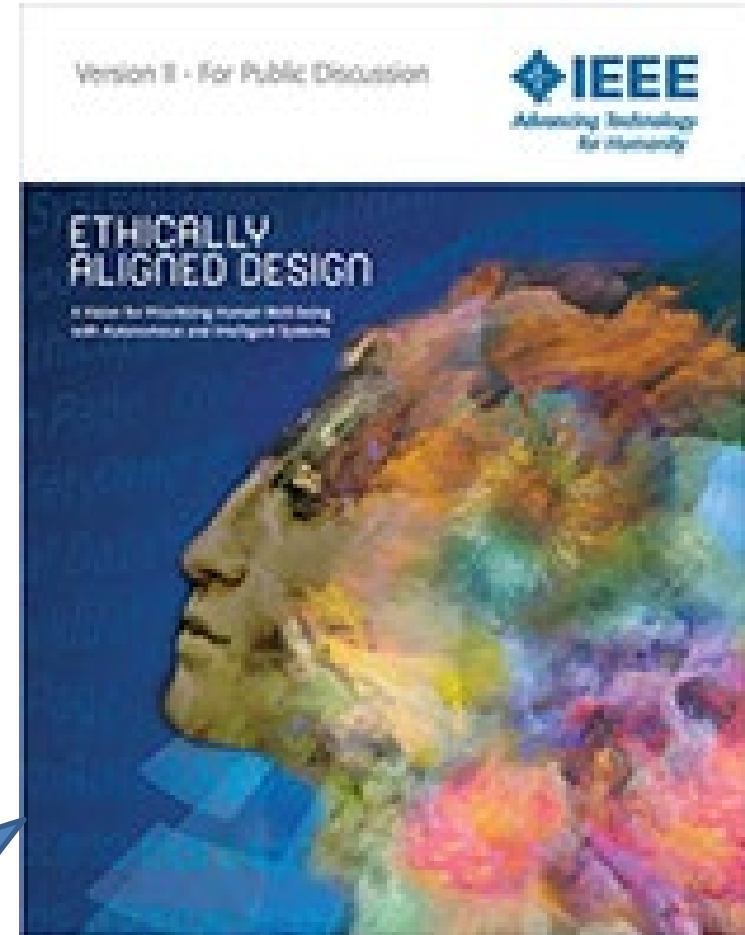
- EU: High Level Expert Group: Ethics Guidelines for Trustworthy AI (2019/4/8)
 - 倫理的なAIの設計指針
- Guidance for Regulation of Artificial Intelligence Applications:
 - USA Whitehouse. MEMORANDUM FOR THE HEADS OF EXECUTIVE DEPARTMENTS AND AGENCIES (Draft 2019/4/24)
- Recommendation of the Council on OECD Legal Instruments Artificial Intelligence
 - OECD 閣僚理事会承認 (2019/5/22)
- Beijing AI Principle (2019/5/25)
- EUROPEAN COMMISSION: White Paper on Artificial Intelligence A European approach to excellence and trust, Brussels, 19.2., 2020.

Asilomar AI Principle 見返してみる

- 一致する意見がない以上、未来のAIの可能性に上限があると決めてかかるべきではない。
- 発達したAIは地球生命の歴史に重大な変化を及ぼすかもしれない。
- 急速な進歩や増殖を行なうような自己改善、または自己複製するようにデザインされたAIは、厳格な安全、管理対策の対象するべき。
- 超知能は、広く認知されている倫理的な理想や、人類全ての利益のためにのみ開発されるべき。

IEEE Ethically Aligned Design version 2

1. Executive Summary
2. General Principles
3. Embedding Values Into Autonomous Intelligent Systems
4. Methodologies to Guide Ethical Research and Design
5. Safety and Beneficence of Artificial General Intelligence (AGI) and Artificial Superintelligence (ASI)
6. Personal Data and Individual Access Control
7. Reframing Autonomous Weapons Systems
8. Economics/Humanitarian Issues
9. Law
10. Affective Computing
11. Classical Ethics in Artificial Intelligence
12. Policy
13. Mixed Reality
14. Well-being



The final version was published

IEEE EAD (Final) on April 2019

- 1. Human Rights
 - A/IS shall be created and operated to respect, promote, and protect internationally recognized human rights.
- 2. Well-being
 - A/IS creators shall adopt increased human well-being as a primary success criterion for development.
- **3. Data Agency**
 - **A/IS creators shall empower individuals with the ability to access and securely share their data, to maintain people's capacity to have control over their identity.**
- 4. Effectiveness
 - A/IS creators and operators shall provide evidence of the effectiveness and fitness for purpose of A/IS.

IEEE EAD (Final) on April 2019

- 5. Transparency
 - The basis of a particular A/IS decision should always be discoverable.
- 6. Accountability
 - A/IS shall be created and operated to provide an unambiguous rationale for all decisions made.
- **7. Awareness of Misuse**
 - **A/IS creators shall guard against all potential misuses and risks of A/IS in operation.**
- 8. Competence
 - A/IS creators shall specify and operators shall adhere to the knowledge and skill required for safe and effective operation.

人間中心AIの社会原則

Social Principles of Human-Centric AI

2109/3/29

Council for Social Principles of Human-
centric AI

内閣府, Japan

AI社会原則

1. 人間中心の原則、
2. 教育・リテラシーの原則、
3. プライバシー確保の原則
4. セキュリティ確保の原則、
5. 公正競争確保の原則
6. 公平性、説明責任及び透明性の原則、
7. イノベーションの原則

EC High Level Expert Group

Ethical Guideline for Trustworthy AI

- 2019/4/8



INDEPENDENT
**HIGH-LEVEL EXPERT GROUP ON
ARTIFICIAL INTELLIGENCE**
SET UP BY THE EUROPEAN COMMISSION



**ETHICS GUIDELINES
FOR TRUSTWORTHY AI**

Trustworthy AI

- Lawful, Ethical, Robust
- 具体的要件
 1. Human agency and oversight
 2. Technical robustness and safety
 3. Privacy and data governance
 4. Transparency
 5. Diversity non-discrimination and fairness
 6. Societal and environmental well-being
 7. accountability

Recommendation of the Council on OECD Legal Instruments Artificial Intelligence

- 2019年5月22日のOECD 閣僚理事会で採択
- 強制力はないが、各国での立法の指針になる可能性大
 - 例： OECD Guidelines on the Protection of Privacy and Transborder Flows of Personal Data（1980）は各国のプライバシー保護法の基礎になった



OECD, Recommendation of the Council on AI, OECD/LEGAL/0449, 2019/5/23

- 技術的目標：
- inclusive growth, sustainable development and well-being
- human-centred values and fairness
- transparency and explainability
- robustness, security and safety
- accountability
-
- 政策的目標：
- investing in AI research and development
- fostering a digital ecosystem for AI
- shaping an enabling policy environment for AI
- building human capacity and preparing for labour market transformation
- international co-operation for trustworthy AI

Guidance for Regulation of Artificial Intelligence Applications: USA Whitehouse.

- AI倫理というよりはむしろ、AIシステム開発のガイダンスで、AI産業の育成が目標
 - 名宛人は産業界と読める
 - 例えば「AIアプリケーションの技術仕様を規定しようとする厳格な設計ベースの規制は、AIが進化する予想されるペースを考えると、ほとんどの場合、非実用的で非効率的」

- 指針は以下の10項目からなります。
 1. **Public Trust** in AI
 2. Public Participation
 3. Scientific Integrity and Information Quality
 4. **Risk Assessment and Management**
 5. **Benefits and Costs**
 6. Flexibility
 7. Fairness and Non-Discrimination
 8. Disclosure and Transparency
 9. Safety and Security
 10. Interagency Coordination

◆いかに規制しないかが根底

- ただし、無制限な開発への歯止めとして他の倫理指針と違うのは
- risk assessment、risk management
- リスク評価をサボると public trust を失うぞ、という言い方
 - 下記の引用を参照
- A risk-based approach should be used to determine **which risks are acceptable and which risks present the possibility of unacceptable harm, or harm that has expected costs greater than expected benefits.**
- Agencies should be transparent about their evaluations of risk and re-evaluate their assumptions



EUROPEAN
COMMISSION

Brussels, 19.2.2020
COM(2020) 65 final

WHITE PAPER

On Artificial Intelligence - A European approach to excellence and trust

- promotes the respect of

fundamental rights, including human dignity, pluralism, inclusion, non-discrimination and protection of privacy and personal data

- and it will strive to export its values **across the world**

- Effective application and enforcement of existing EU and national legislation:
- The EU should make full use of the tools to enhance its evidence base on potential risks linked to AI applications, including using the experience of the EU Cybersecurity Agency (ENISA) for assessing the AI threat landscape
 - To ensure an effective application and enforcement, it may be necessary to adjust or clarify existing legislation in certain areas, for example on liability

– *Uncertainty as regards the allocation of responsibilities between different economic operators in the **supply chain**:*

- These risks may be present at the time of placing products on the market or arise as a result of software updates or self-learning when the product is being used.

- Requirements ensuring that AI systems can adequately deal with errors or inconsistencies during **all life cycle phases**.
- the output of the AI system does **not become effective unless it has been previously reviewed and validated by a human**
 - the output of the AI system becomes immediately effective, but human intervention is ensured afterwards
- **monitoring of the AI system while in operation and the ability to intervene in real time and deactivate (by human)**
- in the design phase, by imposing operational constraints on the AI system

- the identified shortcomings will need to be remedied, for instance by **re-training the system in the EU**,
 - in case an AI system does not meet the requirements for example relating to the data used to train it,
- .
- It would allow users to easily recognise that the products and services in question are in compliance with certain objective and **standardised EU-wide benchmarks, going beyond the normally applicable legal obligations.**

USA とEU の比較: EU

- AIサービスは事前に徹底的なリスク予測を行うべき
- リスク発生の予測はAI技術に複雑さや発展の早さから技術的に抑えることは困難であることも意識
- AIシステム製造の**サプライチェーン**の各段階で倫理指針ないしは法制度に基づくリスク管理や公平性, 非差別性を徹底することを求めている

USA とEU の比較: EU

- AIシステムが実利用において再学習によって動きが変化する場合、開発者側はその都度リスク管理、公平性などを確認することを求めている.
- つまり、AIシステムの仕様を倫理指針、法制度で管理しようとする指針を打ち出している. 開発者にとっては非常な重荷になると予想される.

USA とEU の比較： EU

- 以上のEUの規則主導の指針はUSAの市場評価主導の方針と全く異なる
- EUはさらにEU域内で使われるAIシステムはEU域外の開発者が作ったものでも、EU域内においてテストすることを義務付け
 - 保護主義的な傾向が感じられる.
 - EUでは大企業中心の高リスク分野の他に中小規模開発者の保護や支援を強く打ち出していることにも特徴がある.

	≧制御	人権	公平性 非差別	透明性	アカウンタビリティ	トラスト	悪用、誤用	プライバシー	エージェント	安全性	SDGs	教育	独占禁止・協調、政策	軍事利用	法律的位置づけ	幸福
Asilomar Principles	○	○						○						○		○
人工知能学会・倫理指針	△	○					○	○		○					○	○
総務省AI開発ガイドライン	○	○	○	○	○			○		○						○
Partnership on AI		○	○	○	○			○	○	○	○	○	○			○
IEEE EAD ver2	○	○	○	○	○	○	○	○	○	○		○		○	○	○
IEEE EAD 1e		○	○	○	○	○	○	○	○	○	○	○			○	○
人間中心AI社会原則		○	○	○	○	○	○	○		○	○	○	○			○

	人権 擁護		公平性 非差別	透明性	アカウント ビリティ	悪用 スト	、誤用	プライバシー	AIエージェント	安全性	教育		独占禁止・協調、 政策	軍事利用	法的 位置づけ	
Trustworthy AI		○	○	○	○	○	○	○	○	○	○	○	○	○	△	○
OECD Recommendation		○	○	○	○	○	△	○		○	○		○			○
総務省AI活用 ガイドライン			○	○	○	○	○	○		○						○
Beijing Principle	○	○		○	○	○		○		○		○	○		△	○
Whitehouse Guidance	×		○	○	○	○	○	○		○			○		△	○
EU Whitepaper		○	○	○	○	○		○		○		○	▲		◎	

法的位置づけ: ○=AI人格権、△=AIの現行法への適法性

▲ EU域内優遇的

少し話題を変えます😊

フラッシュクラッシュ

ブラックボックス化の金融への悪影響

- ー 人工知能技術のブラックボックス化が社会にリアルな損害を与えています
- 金融取引(株の売買など)は、既にネットワークを介してエージェントベースで秒以下の売り買いされる世界です。
 - エージェントに人工知能が使われています。
- 人間(トレーダー)が介入して判断するより早く事態は進行します。
- 世界中の金融センターも似たような状況なので、なにかのトリガがかかると連鎖反応が瞬時におこり、とんでもないことになります。

ウォールストリート・暴走するアルゴリズム

Wired.jp 2016年9月 より 1

- だが最悪の場合、それら(AIトレーダー)は不可解でコントロール不能のフィードバックのループとなる。
- これらのアルゴリズムは、ひとつひとつは容易にコントロールできるものなのだが、ひとたび互いに作用し合うようになると、予測不能な振る舞いを-売買を誘導するためのシステムを破壊しかねないようなコンピューターの対話を引き起こしかねないのだ。

ウォールストリート・暴走するアルゴリズム

Wired.jp 2016年9月 より 3

- アルゴリズムを用いた取引は個人投資家にとっては利益となった。以前よりもはるかに速く、安く、容易に売買できるのだ。
 - だがシステムの観点からすると、市場は迷走してコントロールを失う恐れがある。
- たとえ個々のアルゴリズムは理にかなったものであっても、集まると別の論理—人工知能の論理に従うようになる。人工知能は人工の「人間の」知能ではない。
- それは、ニューロンやシナプスではなくシリコンの尺度で動く、われわれとはまったく異質なものだ。その速度を遅くすることはできるかもしれない。だが決してそれを封じ込めたり、コントロールしたり、理解したりはできない。

ウォールストリート・暴走するアルゴリズム

Wired.jp 2016年9月 より 2

- 2010年5月6日、ダウ・ジョーンズ工業株平均はのちに「フラッシュクラッシュ(瞬間暴落)」と呼ばれるようになる説明のつかない一連の下落を見た。一時は5分間で573ポイントも下げたのだ。
- ノースカロライナ州の公益事業体であるプロGRESS・エナジー社は、自社の株価が5カ月足らずで90%も下がるのをなすすべもなく見守るよりほかなかった。9月下旬にはアップル社の株価が、数分後には回復したものの30秒で4%近く下落した。

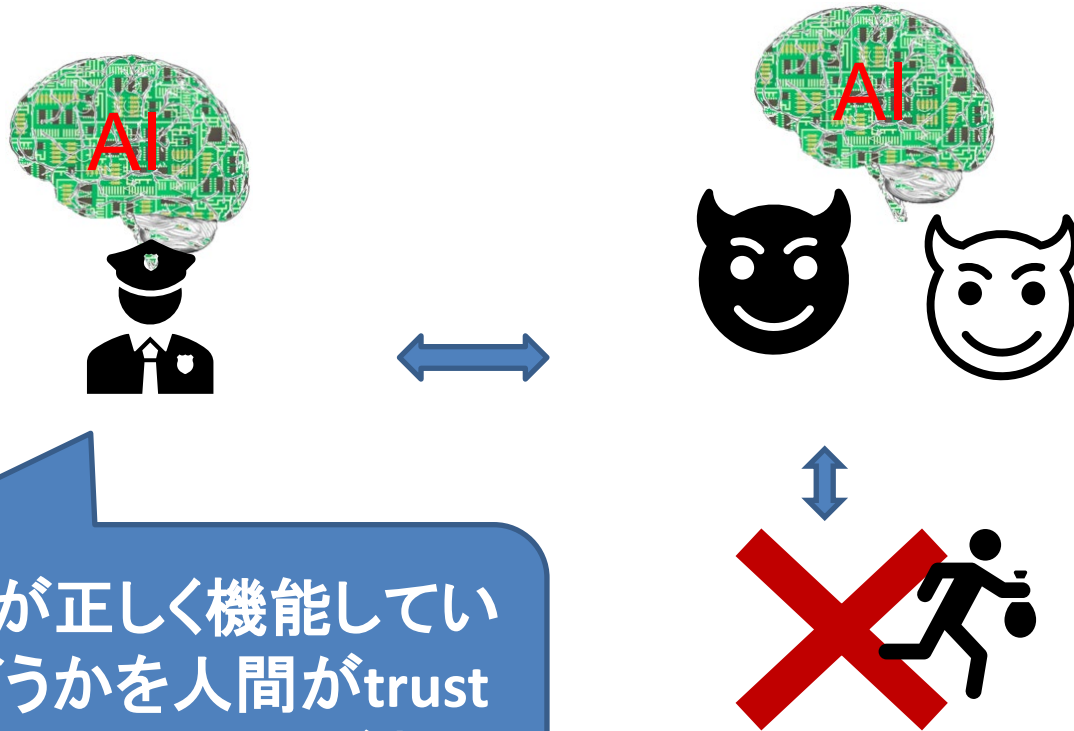
AI内部に説明能力を組み込む XAI(eXplainable AI) の研究が盛んですが

- フラッシュクラッシュのような多数のAIが相互作用する場合にはとても追いついていない。

AIトレーダーの損害検出能力

- 適格な行為者の高速トレードの誤動作（フラッシュクラッシュ）の早期発見、検証が人間にできそうもない
- 許容損害値超過を早期発見するチェック機構が必要
- ネットワーク接続された多数のAIトレーダーの行動チェック機構はAIトレーダーを外部観察、あるいはネットワークの挙動を外部観察するAIとして作る
 - チェック自体を学習機能を持つ人工知能に任せるような局面が考えられますが、果たして可能か？

AIの行動をAIが観察してテスト



このAIが正しく機能しているかどうかを人間がtrust (信用) できる必要がある

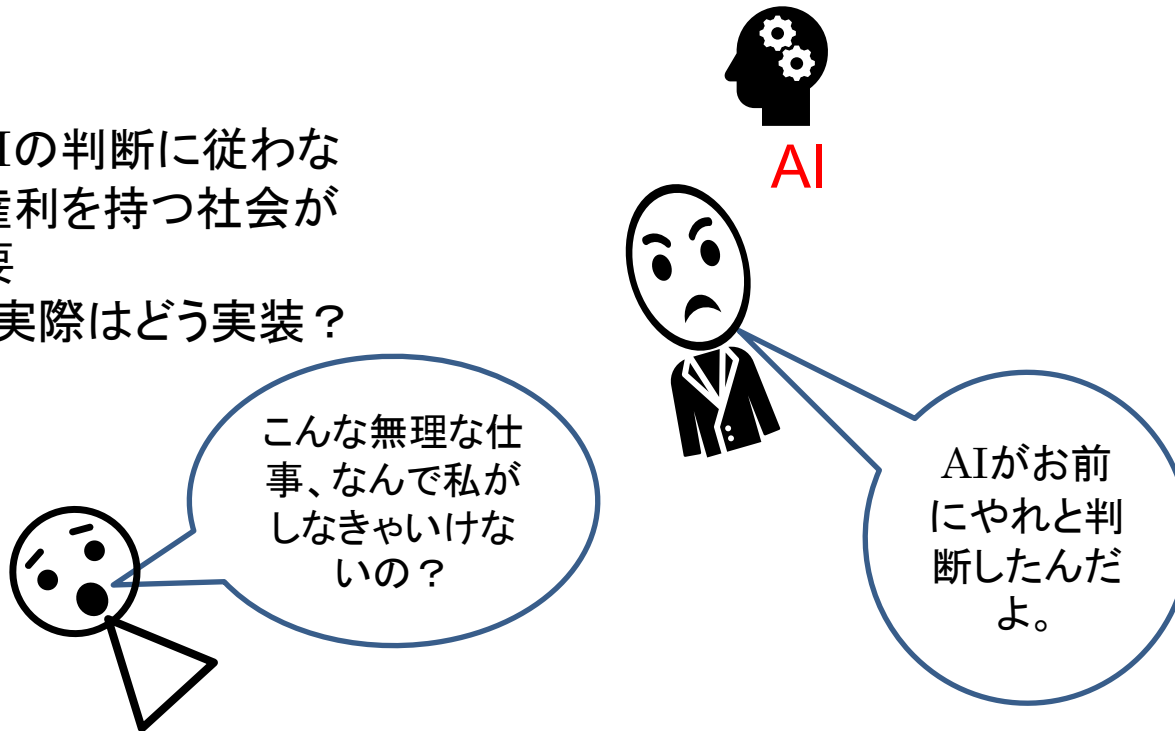
再掲: EU White paper for AI

- the output of the AI system does **not become effective unless it has been previously reviewed and validated by a human**
 - the output of the AI system becomes immediately effective, but human intervention is ensured afterwards
- **monitoring of the AI system while in operation and the ability to intervene in real time and deactivate (by human)**
- in the design phase, by imposing operational constraints on the AI system

- 早く止めすぎると、収益の機会を奪ったという文句がくる
 - 止めるのが遅すぎるとクラッシュして大きな被害を出す
- 収益ロスとクラッシュの被害の和を最小化する止め方 → 最適化問題
- ◆ 人間では無理、やはりAIに頼らざるを得ない

AIの悪用

- AIの判断に従わない権利を持つ社会が必要
- 実際はどう実装？



- 悪用された人工知能に一般人が文句を言うことができなくなってくると、実質的に言論の自由も人権もない状況になりかねません
- 人工知能に対して文句を言える社会制度を考える必要があります。
 - GDPR 22条： 計算機(人工知能)のプロファイリングから出てきた決定に服さなくてよい権利
 - 具体的には以下のようにします。
 - プロファイリングに使った入力データの開示
 - 出力された決定に対する説明責任を人間が果たす。
 - ただし、この権利の社会実装、技術による実現はかなり困難
 - 守秘義務や企業秘密の壁もあります。

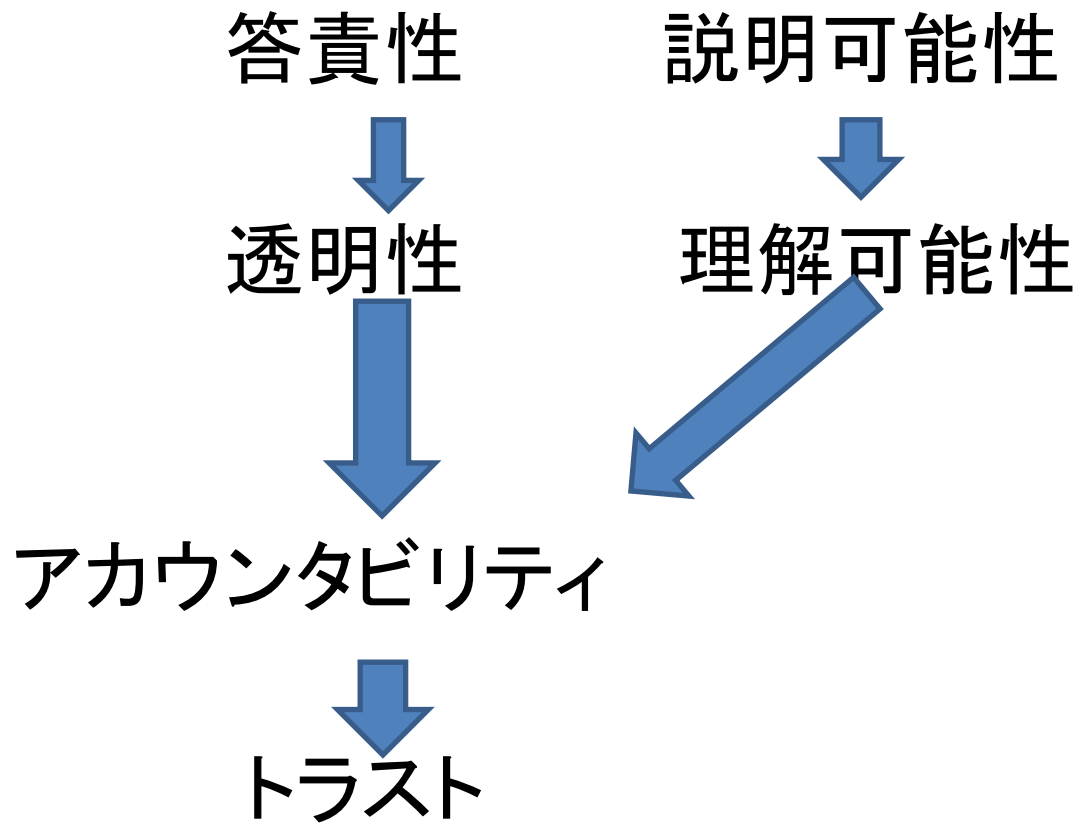
IEEE EAD ver.2 の指針

- 社内の悪用への対策として
 - 内部告発を可能にする匿名化告発(技術 Tor など)
 - 内部告発者が不利にならないように社内規則ではなく、より上位の国の法制度で内部告発者を守る
 - 告発者を経済的に支援する保険制度
- 告発者が裁判などで戦えるために
 - AIのブラックボックス化を排除→ XAI
 - AIシステムの透明性
 - + アカウンタビリティ
 - (トラストでは不十分かもしれない)

AIのブラックボックス化

- 人間の仕事を自律的AIで置き換えるにせよ、人間の能力を拡張するにせよ
 - 開発者に責任はあるはず
 - だが、既に関係者たちが把握しきれない状態に突入しているのかもしれませんが
 - 関係者： Multi-stakeholder
 - 人工知能開発者
 - 人工知能へ学習に使う素材データを提供した者
 - 人工知能製品を宣伝、販売した者
 - 人工知能製品を利用する消費者
- したがって、事故時の責任の所在を法制度として明確化しておく必要がある時期になってきています

信用できるAIへ向けての取り組み： 透明性、説明可能性、トラスト



トラストはどうやって確保するか？

- 過去の学問的成果の集積を信用してもらう
 - 数学、物理学、医学、...
- 専門家をトラストしてもらうライセンス制度
 - 専門家のスキルのトラスト：医師国家試験、司法試験など
 - ライセンスする側（国など）へのトラスト
- それでも事故は起きる
 - 補償制度の確立：保険など
- これら全部を統合したシステム体系がトラストの基本

Trust(信賴)

- (1) 能力、誠実さ、および予測可能性を扱う一連の特定の信念
- (2) 危険な状況下で、ある当事者が他の当事者に依存する意思
(証明したわけではないが...)
- 「信賴」はAIシステムのライフサイクルに関与するすべての人々とプロセスに帰することができます。

トラストの補足

- 利用者がサービス提供側をトラストするという局面ばかり考えてきたが
- サービス提供側が利用者をトラストするという問題もある
- 認証されたIDで金融、政府、通信などのサービスを受けることを保証するAI

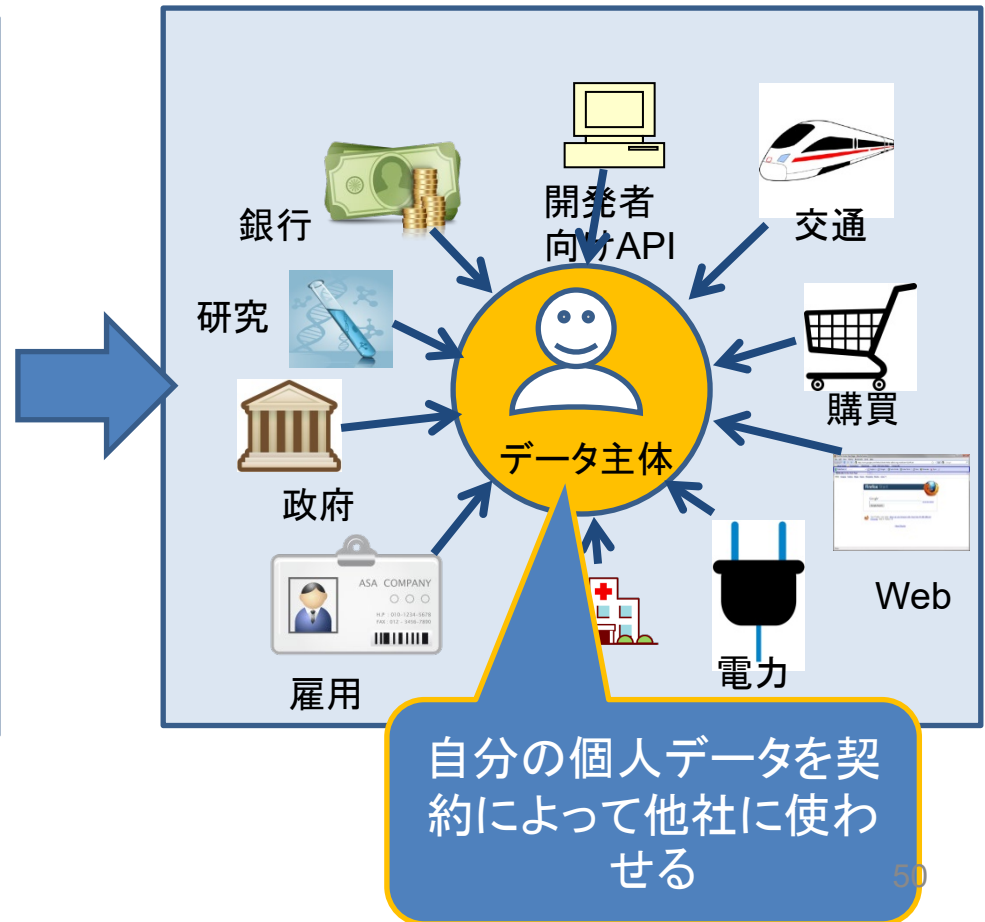
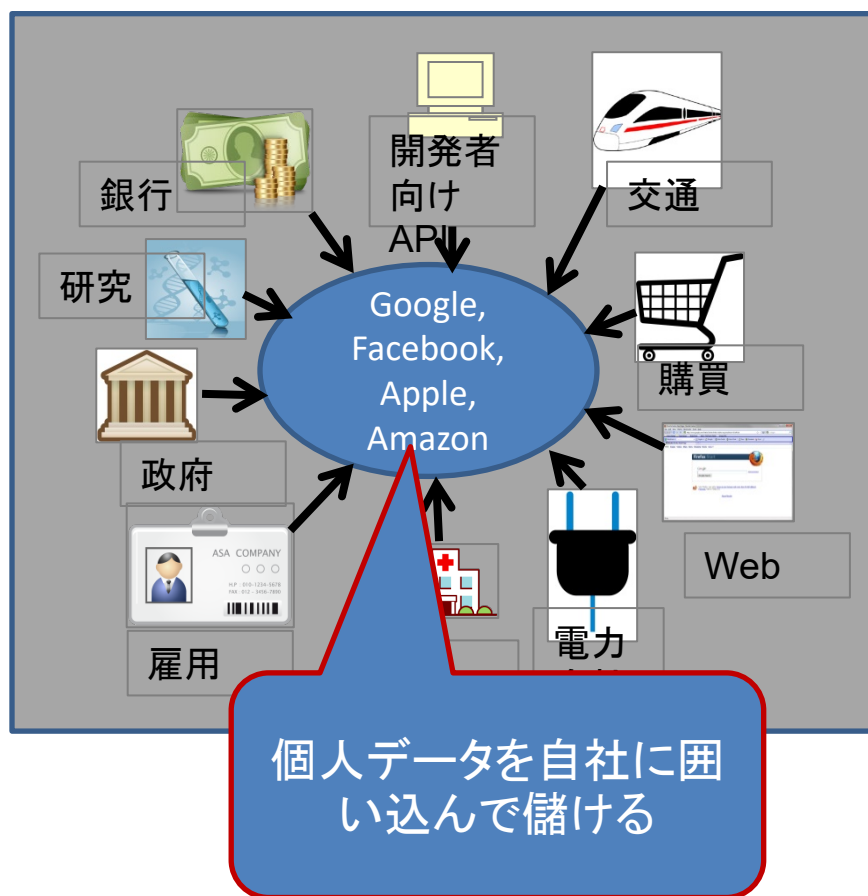
個人は信頼できるID認証にアクセス できること

- 認証はインターネットにおける個人の存在の
証拠
 - 対面認証でないので、技術的な問題が多い。
 - 特に生体認証の危険性
 - Self Sovereign Identity → 対応するサービスに
必要なことだけ identify できればよい
 - 必要最小限の個人データによる認証：比例原則

トラストする対象

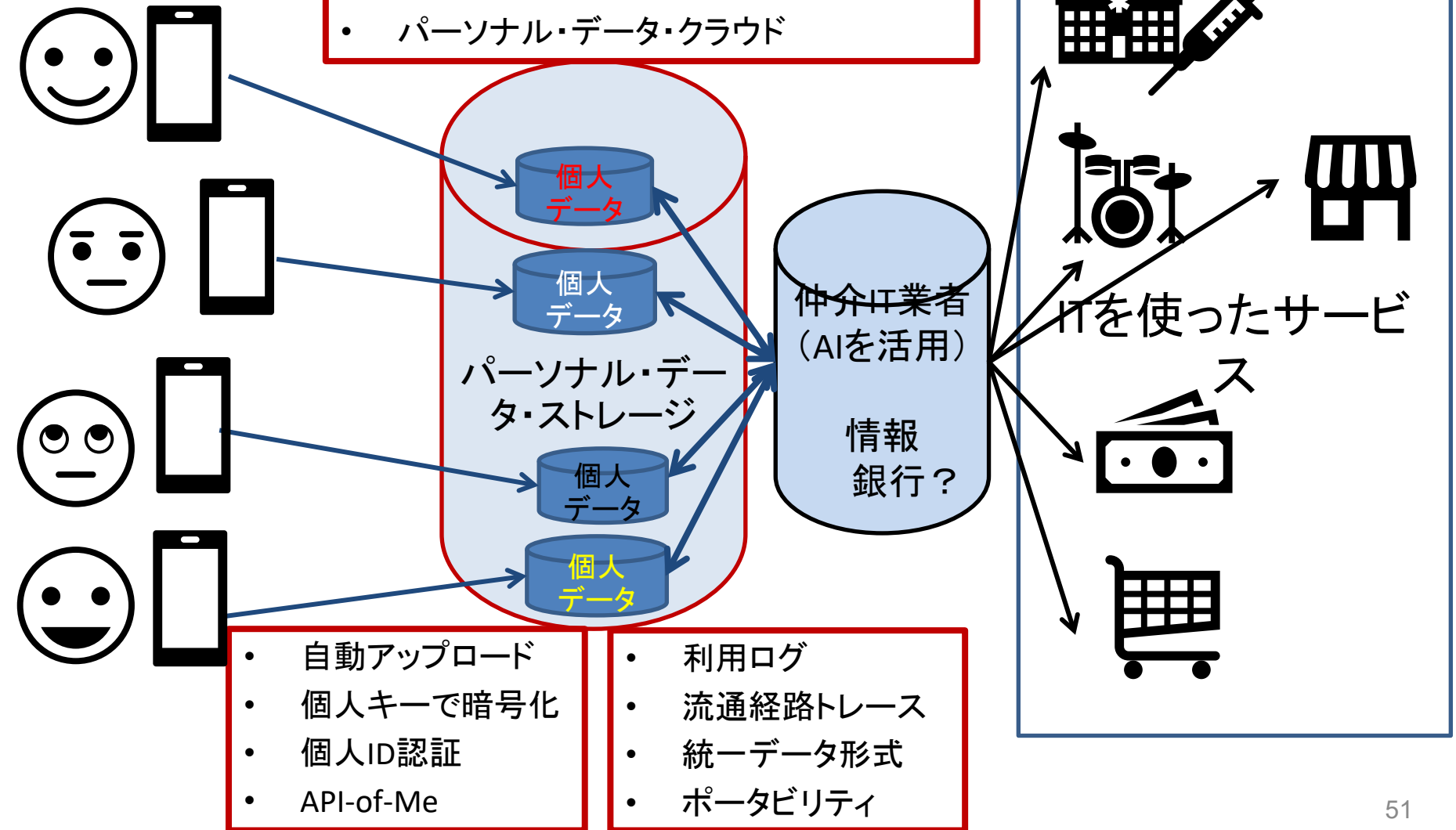
- 組織 （利用者からみたら）
 - 情報銀行、ITサービス企業
- 個人 （サービス業者からみたら顧客）
 - 利用者の個人をトラストできるか
 - Self Sovereign Identity
- データ
 - 流通する個人データ、サービスの結果

個人データは個人が管理したほうが 質が良く、データ自体をトラストできる のではないかな？



パーソナル・データ・ストレージ(PDS)

- ・ パーソナル・データ・ストア／ボールド
- ・ あるいは
- ・ パーソナル・データ・クラウド



主要な技術的ポイント

- パーソナルクラウド
- インターネットにおける Identity 認証
- 個人データのポータビリティ
- Block Chain による個人データの真正性認証
- プライバシー保護(暗号化,複数当事者による計算:
MPC , etc.)
- 公平性、透明性の確保手段

データポータビリティ

- GDPR20条
- データ主体は、個人データを機械可読性のある形式で受け取る権利があり、
- 当該データを、個人データが提供された管理者の妨害なしに、他の管理者に移行する権利がある。
 - ただし、20条3項「職責によって収集した個人データ(例えば医療データ、医療カルテ)にはデータポータビリティは適用しない」
- Googleの個人データAPI
- 日本
 - 銀行API、個人医療履歴、

個人データ個人管理の問題点

- 個人データがどう使われているかにsensitiveな人は多いのか？
 - 痛い目を見るまで分からない
 - ポイントの餌に釣られる？ 目先の利益を優先する人々が大多数
 - だからこそ、きちんと規制すべきという意見もあるが
.....
- 個人データを自分で管理するスキルがない人が大部分
 - 近代的個人の消失につながるのか？

パーソナルAIエージェントとガバナンス

- **背景**: 既存のガバナンスの枠組みがBrexit、トランプ現象、中国の台頭などで揺らいでいる現状への危機感

➤ 新しい方向性

(1) デジタル・レーニズム

(2) GAFAのような国境を超えるITプラットフォームによる情報支配

(3) 既存の民主主義を基礎にするガバナンスの拡充

(3)は望ましいが、多くの人間は近代的法制度、政治制度が前提にした完全な自我と自由意志に基づいて行動する主体には程遠い

パーソナルAIエージェントとガバナンス

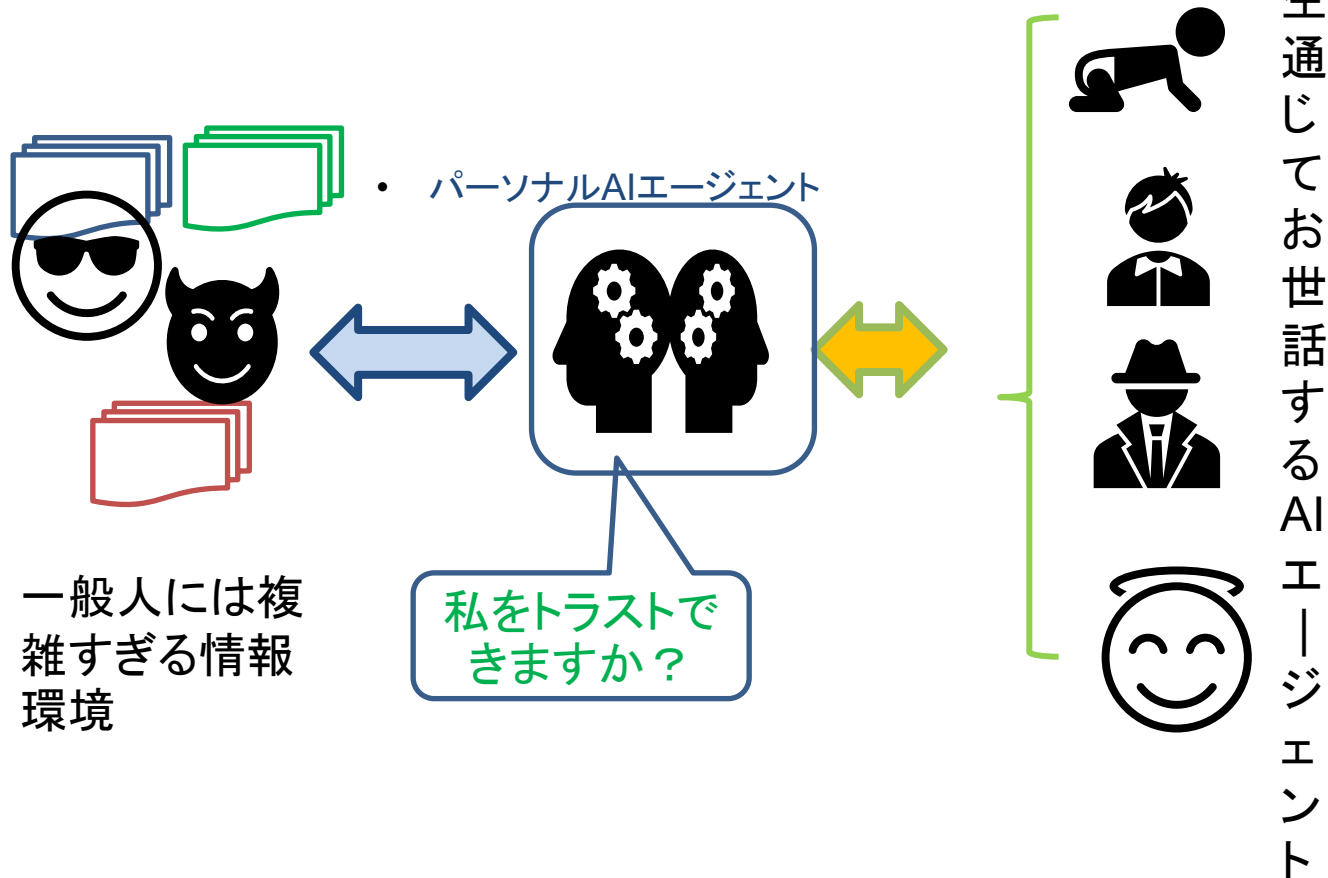
”(3)既存の民主主義を基礎にするガバナンスの拡充“

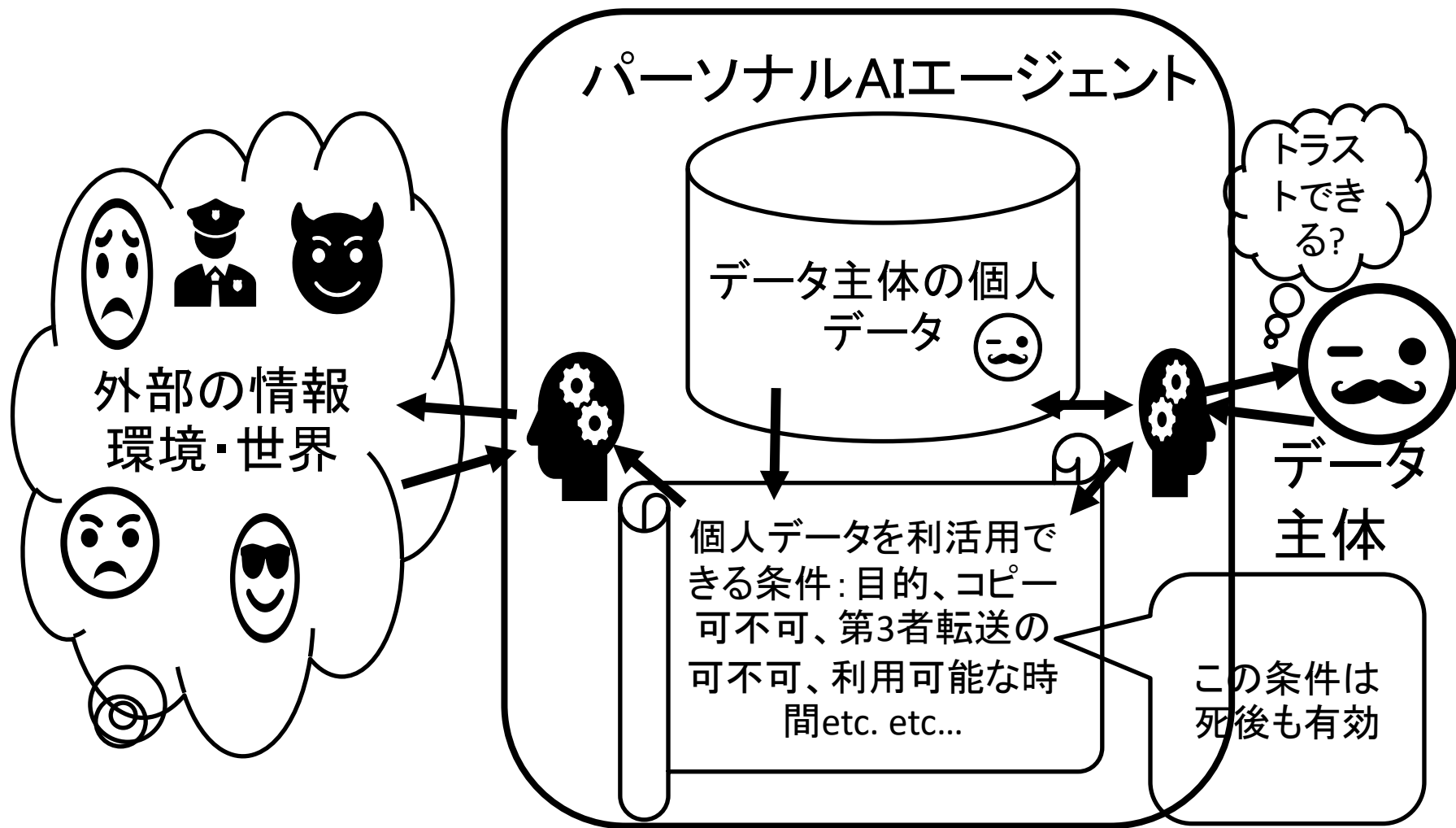
- **提案:** 生身の人間には対処しきれない複雑な情報環境をパーソナルAIエージェントが支援し、人間の情報能力を増強することで対処する
- こうして(3)に近づこうとする枠組みが民主主義国家に住む人々にとっては最も受け入れやすくかつWell beingに資する。

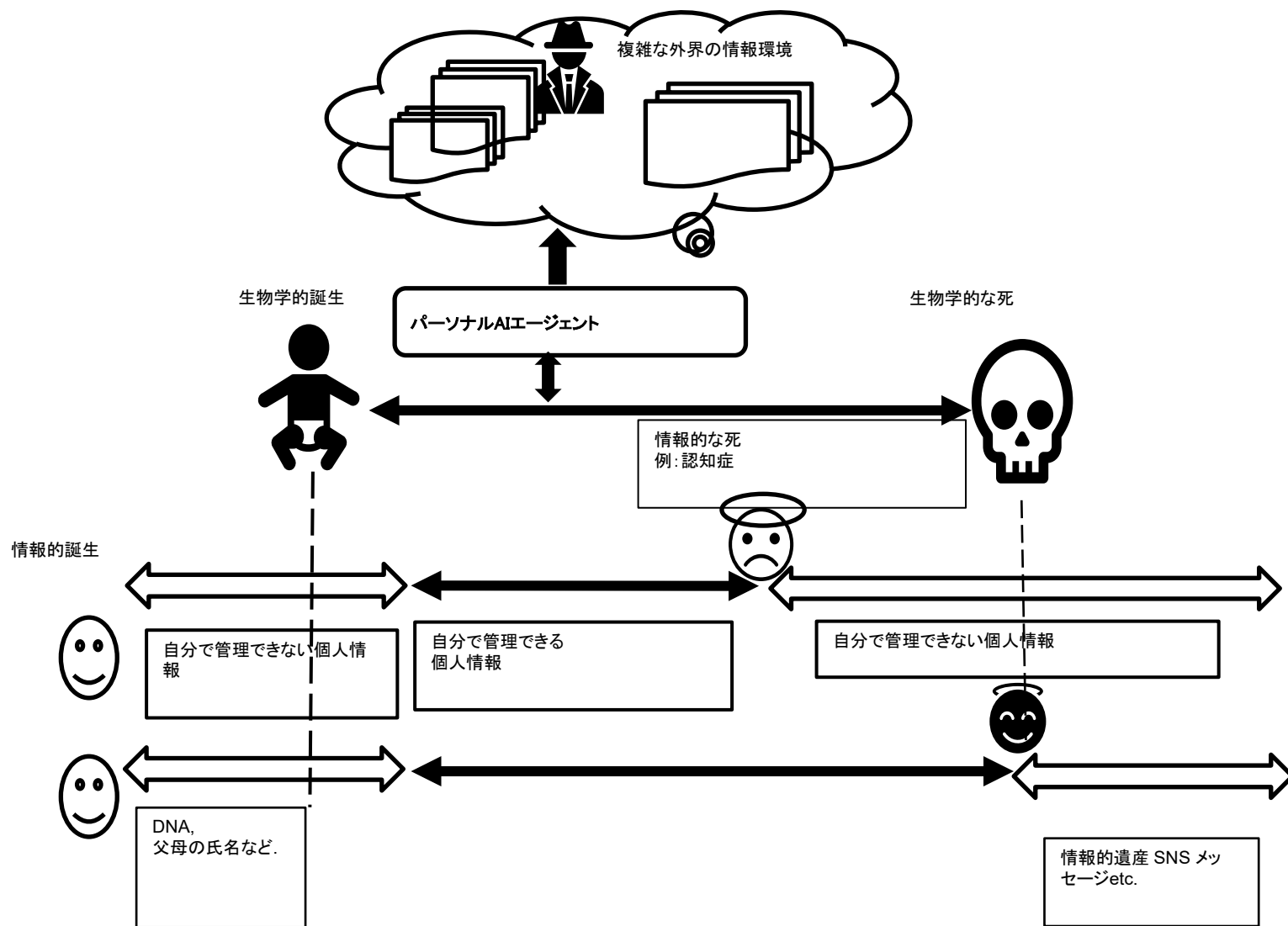
* IEEE EAD : Personal Data Agent

パーソナルAIエージェント*

- 誕生から死まで継続的にサポート -







個人データは自分で管理できない期間のほうが長い

死後の個人データの諸問題

- 著作人格権
 - 一身専属性
 - パブリシティ権
 - 著作隣接権(儲けるネタ)
- 倫理的感覚
 - 死者を冒涇 (AI美空ひばり問題)
 - 死者の知的能力を備えたアバター
 - 不変
 - 可変 → 生存者とやり取りすることによって変化、成長
 - ゲーム能力など

ご清聴ありがとうございました

付録

個人データ保護の流れ

- インターネットの普及によりデータ主体である個人の個人情報、個人データはネットに氾濫している。
- 従来のプライバシーとは違う概念のプライバシー
- EUの動き
 - OECD8原則(EUだけではない) → 種々のプライバシー保護法制の基礎
 - 以前のEUデータ保護指令 → 法律は国毎に違う
- EU全域で統一したプライバシー保護の法律として
- **GDPR(General Data Protection Regulation)**が2016/4/14(欧州議会採択)、2016/5/24施行、2018/5/25適用

個人データ保護の流れ

- EU:GDPR(General Data Protection Rule)
 - 匿名化処理によって個人情報ではなくなり自由流通できると思う人もいるかもしれないが、
 - EUではこのような匿名化手法は基本的に存在しないとしている。
 - つまり、個人データ流通は以下が想定されている
 1. 事業者の説明責任と
 2. データ主体の同意による
 - 個人データを一業者が囲い込むことを許さず、データ主体個人の意志で個人に**統一的な様式**で還元できる**データポータビリティ**を保障 (GDPR 第20条)
 - 個人は還元された**統一的な様式**のデータを類似のより良いサービス提供者に**送って乗り換えられる**

気を許すか否か 人工知能とプライバシー

- 対話ができる人工知能は古くから研究されてきました
- MITのワイゼンバウムが1966年に開発したELIZA
 - ほとんどオウム返し+アルファくらいだが、個々の文はいかにも人間っぽい。
- ELIZAにおいて見られた現象
- 相手が人工知能であると気を許して個人情報をお話してしまいがち
 - 一方、ロボットも見た目が人間っぽい
- 相手が人間であると気を許して個人情報をお話してしまいそうです
 - 人間らしい外観が有名な阪大の石黒先生のアンドロイドにはけっこう感情移入する人が多いらしいです
 - チューリングテストをパスしたAIにもプライバシーを語ってしまいそう

- つまり、人工知能ロボットが会話内容や外観で人間らしくなると、
 - 相手が計算機だという、嘘はつかないだろうという安心感
 - 人間っぽい外観での同類意識
- などから、プライベートな事柄をしゃべりがち
- その気になれば、人工知能ロボットはプライバシー情報を収集しやすい状態なのです。

- もし、コンパニオンアニマルのロボットがインターネットに接続され
- その結果、マルウェアに汚染され
- なつかれている高齢者の主人が、自分の金融資産管理などについての情報を漏らしてしまい、盗まれてしまうかも
 - デジタル狂犬病

プライバシー保護も人工知能の能力として持つべき

- 何がプライバシーか、という人間にとっての常識を人工知能が持てるか？
- そうして人工知能ロボットが収集したプライバシー情報を保護するには、やはり人工知能の能力が必要

- なにがプライバシーかは個人毎、状況毎に異なる
- 浮気中の人にとっては宿泊したホテルや夕食を採ったレストランもプライバシー

－果たして人工知能に判断できるでしょうか？

－こういったことができる人工知能になると、それこそシンギュラリティにかなり近づいているといえそうです。

－つまり、人間に好意や悪意を持った行動をできる一歩手前まで来ているのです。

同意の取り方

- AIを理解しないデータ主体からデータ収集するにあたって、有効な同意をとるAIの開発
- 同意は1回取ればおしまいではなく、状況変化において随時とれるようなAIシステム
- 被雇用者<<雇用者
という力関係をバランスさせるような同意のためのAIシステムが必要
 - 労使関係では、明確な同意のない場合を多数見られるので、それを改善すべき
 - 力関係が大きくバランスを欠く非対称関係における同意の有効性への疑義あり

同意の取り方：情報弱者対応

- 高齢者や非健常者などの情報弱者がプライバシー保護で置いてけぼりを食わないような同意ないし利用監視システムをAIを利用して構築
- 子供は親が代理するにしても、それ以外の情報弱者対策は喫緊の課題 → AIの出番