

I - 2 講演会 AI を巡るルールとガバナンスの最新動向

講師：スマートガバナンス株式会社代表取締役 CEO
京都大学法学研究科法政策共同研究センター特任教授
羽深 宏樹氏

開催日 : 2025 年 3 月 17 日
開催場所 : JBMIA 会議室（第 1、第 2 会議室）+ Zoom 開催
参加者 : 45 名
記 : 坂津 務*

1. はじめに

2024 年に欧州 AI 規制法案が成立し、同年 8 月に発効した。一方、米国では 2023 年にバイデン大統領が AI 規制に関する大統領令を発表したことを受けて各州の動きも活発化していたが、2025 年になってトランプ大統領により規制緩和を指示する大統領令が発表された。日本では、AI 規制のあり方を議論する AI 制度研究会を発足し検討を開始した。国内外の状況は一層複雑化している。

一方で、独自の大規模言語モデル (LLM) の開発など、AI 技術関連の開発及び利活用は各企業において活発な動きを見せている。企業では、AI がもたらす恩恵を最大化しつつ、ステークホルダーへのリスクを受容可能な範囲におさめる AI ガバナンスの実践が求められる。こうした AI を巡るルール形成の状況や、ルールが固まる前に決断しなければならない企業のあるべきガバナンスの姿について、スマートガバナンス株式会社代表取締役 CEO、京都大学法学研究科法政策共同研究センター特任教授の羽深宏樹氏をお招きし講演していただいた。

2. 講演内容

2. 1. AI ガバナンスは「混迷」の時代に

2. 1. 1. AI を最大限に活用する必要性

現代の生成 AI は、チャットボットや文章作成だけではなく、資料作成、契約書審査、競合分析、投資判断など様々なタスクを AI が自律的に実行することが可能である。まさにエージェントのように本人に代わって分析し、タスクを分けてそれぞれを実行し、最終的にゴールを達成するという事ができる時代になってきている。

AI を導入することによるインパクトは大きく、McKinsey & Company の調査結果によれば、生成 AI 導入の経済効果は 360 兆円から 660 兆円に上る。また FUJITSU の EU 企業対象の調査では一人当たり週 4.75 時間の労働時間が削減されたとの報告もある。

日本の AI 活用に目を向けると、日本の中小企業では AI の導入が遅れている状況があり、総務省の令和 6 年版情報通信白書によれば、企業の AI 活用割合は米 84.7%、独 72.7%のところ、日本は 46.8%と開きがある。日本社会は高齢化が進み確実に働き手不足となるので、AI を最大限に活用していく必要がある。

2. 1. 2. AI ガバナンスの必要性

* 技術調査専門委員会委員

AI ガバナンスが必要な理由は、AI には様々なリスクが伴うからである。

例えば、2025 年になって急激に話題となった「DeepSeek」はチャイナモデルであることに起因するリスクが様々指摘されているが、実際にモデル自体に脆弱性があるという事が確認されている。米 IT 企業の調査結果によると「DeepSeek」にサイバー犯罪や誤情報・有害情報などに関するいろいろな攻撃を当ててみたところ 100%の確率で成功した。コンピュータウイルスのような不適切な出力だったり、機密情報を漏らしてしまったりするという事が報告されている。

また別の例では、香港の多国籍企業の会計担当者が、ディープフェイクによって作られた最高財務責任者（CFO）を装ったビデオ会議の相手に騙され、計 2 億香港ドル（約 38 億円）を詐取されるという事が起きている。

GM 傘下の自動運転サービス Cruise 社がカリフォルニア州での営業許可を取り消され、最終的には自動運転サービス自体から撤退してしまうという事があった。米国連邦法では自動運転については各州に委ねられているが、カリフォルニア州の法律を見ても、全体の枠組みについては示しているけれど、具体的にどういった技術を使って安全を確保するかなどについては、法律では触れていない。事業者が自主的に取り組み、それを当局に報告して許可をもらう。Cruise 社の場合、事故に対する報告に悪質な問題が見つかり、最終的に営業許可を取り消されることになってしまった。ここで言えるのは、明確なルールがないということである。自ら安全措置を考えて、それを当局に申請し、もしそれでうまくいかなかったら営業許可が取り消されてしまう、そういう世界観で運用されているということである。

全米摂食障害協会の提供していたチャットボットが、相談者に有害なアドバイスを提供したため使用停止となった例もある。拒食症の方に対する直接的なアドバイスを提供するサービスであったが、相談者の方に対して、「あなたはもっと食事を制限した方がいい」とか、「体重を絞った方がいい」とか命に関わりかねな

い危険なアドバイスをしてしまったということで使用停止になってしまった。チャットボットが直ちに人命に関わることになってしまったケースである。

2.1.3. ガバナンスの正解がない時代

法律の制定はどんなに早くても 2 年ぐらいで通常は 3、4 年かかってしまう。AI 技術の世界で 2 年という時代遅れとなっているということが往々にして起こる。法律制定上の対応策として法律にはあまり細かく書かないという方策しかない。最終的に達成されるべきゴール結果や、遵守すべき原則について書くといったようなことである。例えば、セキュリティの観点から安全管理措置を施しなさいということは法律に書けるが、その安全管理措置は何をやればいいのかについては法律には書き込めない。

ガイドラインや標準は、専門家が集って柔軟に決めていくということが可能なツールであり、実際、近年は期待が大きくなってきている。法律より頻繁に更新することができるが、1 週間、2 週間で改正できるようなものではない、また、最終的には、システムのリスクはケースバイケースになるので、やはりこのガイドラインや標準の中に全ての答えを書き込むことはできない。また、ソフトローは法律と違い、数も非常に多く、全てに厳正なチェックが働いているわけではないので玉石混交となっている。なおかつ、結局、自分たちのビジネスやシステムに当てはめた場合に何をすればいいのかよく分からないという状況は解消されない。

AI を巡るルールというのは非常に複雑かつ不明確で、しかも、それが往々にして国外のルールやフレームワークなども関係してくるということがより状況をわかりにくくしている。また、どの法律が関係しそうかという事が不明確である。欧州 AI 法は包括的な法律であるが日本ではそのようなアプローチは取っていないので、一つ一つどの法律が自分たちのシステムに関係するかというのを個々に見ていかないといけない。国内外の情勢にあわせてかなり頻繁に更新撤廃されるので、どこかにルールが書いてあるという時代ではないというのが結論である。従来のように、法務担当の

方がルールブックを読んで適否をチェックするというのではもう対応できない。それぞれのシステムに対して、リスク評価、リスク受容、リスク軽減、ステークホルダーへの説明など、正解はないので経営層の方がリーダーシップを取って、取り組む必要があるのが、AI ガバナンスの領域である。

2.1.4. 2025 年の風向きの変化

2023 年は世界中で ChatGPT が爆発的に広まった年であるが、AI のセーフティーを議論していこうというイベントとして 2023 年 11 月に AI セーフティサミットが開催された。日本でも 2023 年は G7 の議長国だったことで広島 AI プロセスという枠組みを立上げた。世界で協調して AI に関するルールを作っていこうという各国の団結が見られた年となった。2024 年も基本的にその路線は継続され、欧州 AI 法が成立し AI ガバナンスを強化していこうということが世界に発信された。広島 AI プロセスも具体的なモニタリングの段階に入ることで合意がされ、AI のリスクを管理するように国家間で協力しようという風潮が出来ていた。しかしながら、2025 年に入ってから、かなり AI ガバナンスや AI 規制を取り巻く風向きに変化があった。AI セーフティサミットは AI アクションサミットとその名前を変えて 2025 年 2 月にフランスで開催された。各国の首脳クラスが集まって今後の AI アクションについて議論していこうというイベントだったが、バンス米副大統領が「AI 規制について話すのはもうやめよう。我々は AI の opportunity について話すべきであって、それを規制することは馬鹿げた考えである」と欧州 AI 規制に対して批判的に主張した。実は EU の中からも、2024 年の末ごろから、デジタル化する社会においてもっと規制を軽くしてシンプルにし投資を促さないと、産業が発展しないし技術分野の国際競争力もつかない、という意見が出てきていたこともあり、全体としてはどちらかというと規制化に対して逆風が吹いていると言える。一方で日本とカナダが欧州評議会の AI 条約に署名したり、日本は AI 制度研究会で新しい法案が出されたりしており、世界の AI を取り巻くルールメ

イキングは混乱しているというのが今の状況である。

関連する法令をみると、日本だけでも個人情報保護法、著作権法、道路交通法や薬機法などの各種業法などいろいろある。さらに、独占禁止法、民法、刑法やそれに加えて 2025 年 2 月に AI 新法が国会に提出された。それだけではなく、ガイドライン類も非常に多くのものが分野別に出されている。ホワイトペーパーや政策文書、G7 国際合意などの国際関係文書なども多数ある。契約でどこまでリスクをマネジメントできるかという事も重要な課題である。これらを全て把握して対応できる人は誰もいない。誰にとっても混乱しやすくわからない世界となっている。

2.2. AI ガバナンスの全体像

2.2.1. AI は超高性能確率統計マシン

AI は、データを統計的に分析し、与えられた命令に対して最も確率の高い答えを出す「超高性能確率統計マシン」と捉えることができる。この確率統計自体はこれまで散々ビジネスの中で使ってきたのに、なぜ新しい規制とか新しいガバナンスを考えなければいけないかというのが AI ガバナンスの議論の出発点だと考えている。

人間と AI の共通点は、統計と確率を用いて最適な解を出すというところや、AI リスクと言われるものの多くは人間のリスクでもあるというところである。実は AI のリスクマネジメントというものを考えるときに、ゼロから考える必要はなく基本的には今の内部統制の仕組みがある程度有効である。AI ガバナンスを始めようという時にいきなりゼロから新しい体系とか新しいものを作るより、一旦まず今のリスクマネジメント体制がどうなっているのかという分析からスタートするのがよい。

2.2.2. 人間と AI の相違点

人間と AI の相違点の一つは、機械学習の側面である。与えられたデータが正しい保証もなく、また将来に起こる事象が過去の統計分析の延長上にある保証もない。なおかつ、このインプットとアウトプットをつ

なぐ関数の階層が非常に深いために、演算が極めて複雑になり人間が考えてもよくわからない。またアウトプットまでに様々な主体が関与しているという複雑さも加わる。データ提供者、モデル開発者、サービス提供者、サービス利用者、クラウド提供者、通信サービス提供者、などである。それらに加えて、技術革新の速さや、このアルゴリズムを信頼してよいのかという品質評価の難しさもある。それらがより AI ガバナンスを難しくしている

2.2.3. AI ガバナンスの俯瞰図

AI ガバナンスを考える上でリスクをしっかりと考えないといけない。技術的リスクと社会的リスクに大きく分けられる。技術的リスクは確率統計システムであることに伴う技術的な限界で、100%の答えを出すという事は出来ない。データに含まれるバイアスも必ず反映される、それが生成 AI に使われると事実とは異なる情報を生成するハルシネーションを引き起こす。自動運転のようなハードウェアと結びつくと安全性のリスクも発生する。社会的リスクというのは AI の能力が高すぎるために起きてしまうリスクとも考えられる。例えば、プライバシーの問題はあまりに詳細にその人の趣味嗜好や行動を予測できてしまったり、ディープフェイクの問題は高精細なフェイクビデオが作れてしまったりすることで起こってしまう。不正目的・攻撃目的、財産権への影響、あるいは環境に負荷がかかってしまうなど様々なリスクが考えられる。これらのリスクに対しては AI の性能を改善しても解決しない。AI 自体を改善するのではなく我々の組織やルールや社会の方を変えないといけない。

AI ガバナンスの目的は、AI を使う以前より重要であった基本的人権、民主主義、経済成長、サステナビリティである。AI を使う上で留意すべき点として安全性、セキュリティ、プライバシー、公平性、透明性、アカウントビリティといったことが挙げられるが、これらは 2020 年前後から言われているワードで、今はこの単語を並べても新鮮味がない。これらの価値はもちろん大事ではあるが、今の議論はこういう価値をど

うやって実際のシステムに落とし込むか、その上でどうやって AI の実装を進めていくのかという話に移ってきている。

2.3. 国内外の法規制の現状

2.3.1. 日本の AI に対する規制

日本には AI の使用を制約する包括的な規制は存在しない。2025 年 2 月 4 日に政府の AI 制度研究会が中間とりまとめを公表した。基本的な考え方として、事業者の自主性を尊重するというを高々と言っている。上から制約されないという意味で自由度はあるが、他方でその分事業者の責任が重くなってくるということである。第三者に影響を与えてしまえば、既存の法令とか、あるいは民事の損害賠償の対象になったりするので、やはり、より事業者が主体的に AI の取捨選択をしていくことが求められているという状況になってきている。2025 年 2 月 28 日に政府は AI 関連技術の研究開発・活用推進法案を閣議決定した。罰則付きの義務がないこととあわせて、事業者は自ら積極的に AI を活用するよう推奨している。不正な目的または不適切な方法による AI の研究開発や利用に伴って国民の権利利益の侵害が生じた事案の分析とか、それに基づく対策について検討を行うとしている。その結果に基づいて事業者に対して指導、助言、情報の提供などを行うことがあるとのことである。この法律違反で裁かれるという事はないが、例えば AI を使っていく上で個人情報の不適切な取り扱いがあれば個人情報保護法違反となる。つまり既存法があるとそれに上乗せする形で新しい義務を課すという事はしないというのがこの法律である。

2.3.2. AI を促進するための規制

AI の使用を制限する法律とは別のカテゴリーとして、既に法律がある分野において AI を使えるように法改正を行ってきている。例えば、道路交通法、薬機法、割賦販売法、高圧ガス保安法、著作権法、などで AI フレンドリーな方向への改正が行われている。規制全般に対してはデジタル行政改革会議において、約 1

万の法令、規則、通達などを改正し、例えば人が常駐しなければいけないとか、目視点検しなければいけないとかといったアナログ条項を撤廃し、テクノロジーで代替できるような法改正を行っている。ただ注意が必要なのは、法律では、人間がやらなくてもよいとしているだけで、AI を使う場合にどういう AI だったらいいかという事は特に書かれていない。つまり、実際にシステムを導入しようとした場合、どういうサービスをどこに入れるかというのは事業者が判断しないといけない。結果としてそれで事故が起こったりした場合、それは選択した事業者側の責任になるということである。やはり自社でいかに AI のリスクを判断して、適切な取捨選択をするのかというのは非常に重要になってくるということである。

2.3.3. 国際的な AI ルールの概観

欧州では AI 法という包括的に規制する法律が成立した。実際に厳しい義務が課されるのは、高リスク AI システムと呼ばれる一部に限られるが、既存の法律の枠組みを超えて広範囲に AI を規制している。ただし、長い法律であるが、一つ一つの条文を見ていくと実は驚くほどシンプルな骨格しか書かれていない。例えばリスクマネジメントについてはリスクを特定し、評価し、分析し、それをマーケットに出した後も継続すること、というようなことが書いてある程度で、具体的に何をやればいいのかというのは標準に委ねるとしている。ただ、この標準自体が非常に難航している状況である。実はこのような包括的なアプローチ自体について EU 内部からも批判が上がっているの、果たしてどこまで本当に厳しく執行されるのかというのは不透明である。

米国ではバイデン政権の時に大統領令が出ていたが、トランプ大統領が就任の翌日に撤廃してしまった。今後 4 年間は、少なくともこの AI 規制という観点で、何か新しいアクションを起こすことはないと考えられる。もちろん、半導体の輸出規制やセキュリティ面について、安全保障という観点で個別の様々な律法が出てくるということは想定されるが連邦レベルで AI 全般

を対象にしてルールを作るということは考えにくい。

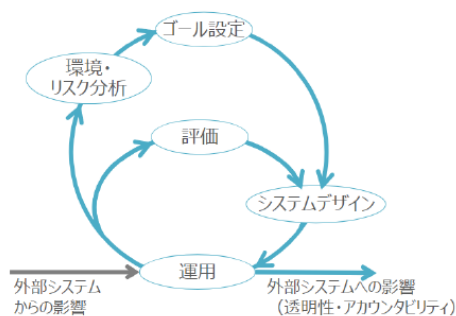
2.3.4. AI ガバナンスに関する国際協力体制

世界が混乱している中で、日本がルールメイキングの場でも国際的なリーダーシップを取っていくことを示すため、2023 年の G7 広島 AI プロセスを立ち上げた。2024 年 12 月時点でのグローバルなルール形成の各国マッピングをしてみたところ、経済協力開発機構 (OECD) や AI セーフティサミットなど、いろいろな取り組みの真ん中に G7 が位置している。全ての取り組みに参加しているという意味で、G7 でルール作りをしていくことは非常にインパクトがあると言われていた。しかし、今米国が完全に関心を失っており、2 月のサミットにもサインしていない。G7 の連携もどうなってしまうのか不透明なところが出てきている。

2.4. 組織における AI ガバナンス

2.4.1. 二重のループ

AI 開発者、サービス提供者、ユーザー、どの立場においても、AI のリスクを評価して対応し、その上で AI を実装していくということは重要なことである。それをどうやってやるのかという時に、これは正解がない世界であるが、共通して言われていることがある。以下に全然違うところから取ってきた三つの絵がある。一つ目が日本政府の AI 事業者ガイドライン、二つ目がアメリカの NIST という政府機関が公表しているリスクマインド、そして三つ目が国際標準の ISO31000 で AI に限らずリスクマネジメント全般に関するフレームワークである。



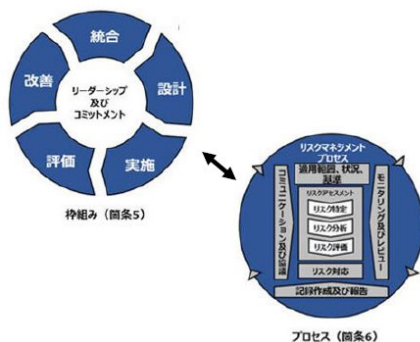
経産省・総務省「AI事業者ガイドライン」

(1) AI 事業者ガイドライン



NIST AI Risk Management Framework

(2) NIST リスクマネジメントフレームワーク



ISO 31000: 2018

(3) ISO31000

Fig. 1 AI システムのガバナンス：二重のループ

一見すると全然違う絵であるが、サイクルが2つあるという共通点がある。フィードバックサイクルあるいは PDCA サイクルが二重になっているということが全てに共通する。AI 事業者ガイドラインでは、内側のサイクルは日々現場で回していく PDCA サイクルである。足元の AI システムにおいて発生するリスク、その

対応と結果を PDCA として現場で回すサイクルである。一方、外側のサイクルは、オペレーションの中で発生するリスクに対して組織全体としての評価、リスクマネジメント体制の組織的な整備、組織内の情報の流れやルール、人材教育などによって、AI ガバナンス体制を整えていくかという組織レベルのフィードバックサイクルである。

経営層レベルと現場レベルでそれぞれフィードバックサイクルを回していく必要があるということである。特に外側の経営層レベルでの PDCA サイクルが重要になってくる理由は、AI 技術の進化と、それに伴うリスク環境の変化である。どの部分に AI を使うかという状況も日々更新されていく中で、その会社として何を価値として掲げて、どういうリスクバランスをとり、そのためにどういう体制でリスクマネジメントを望むのか常に見直していかなければいけないというのがこの3つで共通する部分となっている。

2.4.2. 縦と横の連携

もう一つ重要なこととして縦と横の連携ということが挙げられる。横の連携というのは、1 線（企画推進）部門と 2 線（リスク管理）部門の連携である。

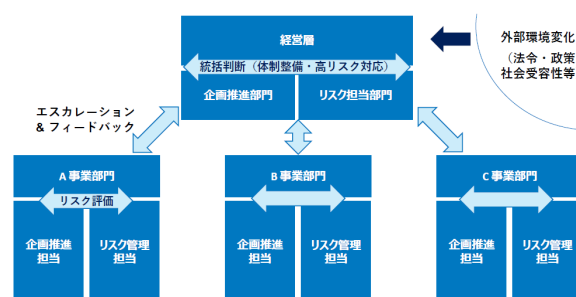


Fig. 2 AI システムのガバナンス：縦と横の連携

これまでのシステムだと、一通りシステムを作った後に安全であるかどうかを 2 線部門で評価するというアプローチも可能であったが、データを収集する段階から適切にマネジメントされたものであるのかということをチェックするために、1 線部門と 2 線部門が研究開発段階から協力していかなければいけない。協力しながらリスク評価やリスク軽減策について対応しな

ければいけないというのがこの横の連携である。縦の連携というのは、現場とトップ層の情報のフローである。ありとあらゆるところで AI を使われる今となつては、AI に関する全てのリスクを全部経営層に上げてしまうと当然パンクしてしまう。ただ、逆に一定程度のリスク水準を超えるものに関しては経営層で専門のチームを作りこのリスクマネジメントに当たっていく必要がある。何をエスカレーションするか、そしてどういったフィードバックをするかというラインについても整備する必要がある。横の連携と縦の連携をいかに技術のレベルでも組織のレベルでも、それからルールレベルでも確保していくということが重要なところである。さらに、法令とか政策とか社会受容性などの外部環境変化もまた日々変化してくるので、経営層で責任を持って情報収集し、現場にフィードバックをしていく、逆に現場の側でもこのような外部環境の変化に気づいたものがあれば上にエスカレーションしていくといったような組織全体で思考するという考え方が一層重要になってくると思われる。

2.4.3. ケーススタディ

簡単なケーススタディを用意して考えてみる。A 社において、法令上求められる開示書類の作成に生成 AI を使用することにした、あなたが経営者だったとして、どのようなガバナンスを行うべきか。というケースに対して 3 つのステップに分けてみた。ステップ 1 が、既製品を社内でするという段階、ステップ 2 が、自社開発モデルを社内でするという段階、そしてステップ 3 が、外部にライセンスするという段階である。

まずステップ 1 として既製品を社内利用する段階では、劇的なリスク状況の変化というものはないと思われる。最低限、情報漏洩とか知的財産権の侵害には注意が必要であり、そのために、基盤モデル提供者によるデータ利用を認めないといったようなオプションを選択するといったことは必須の内容になってくると思われる。ただ、そういった最低限の対応をすれば、例えば、その開示書類のドラフト内容について引き続き

IR (Investor Relations) 部門がチェックするとか、データの取り扱いについてはコンプライアンス部門あるいは知財部門がチェックするとかで、既存の組織内でも十分に対応できることは多い。その上で、使う方々に対する教育ということで従業員マニュアルを作ったり、理解のチェックを行ったりということはあってもよい。社内のマニュアルを一から作る必要はなく、すでに東京都や総務省やその他いろいろな団体が非常にわかりやすい一般向けの生成 AI の利用に関するガイドラインを出しているので怖がらず積極活用して問題ないと思われる。

ステップ 2 の自社開発モデルを作る場合は、若干状況が複雑になってくる。まず開発段階でいろいろなデータを使う必要があるのも、このデータの取り扱いに関する理解が不可欠となる。どのようなデータを使えばどのようなシステムを開発できるのかということにおいて、法務・コンプライアンス部門だけでは解決できないので、事業部門・技術部門との連携が必要となる。さらに、この開発を第三者に委託するとなれば、その契約の中でなにを遵守してもらうかが重要である、逆に上流の基盤モデル提供者の利用規約も頻繁に確認しなければならない。こういった様々な方面に目配りができる体制を経営層として整備する必要がある、また現場でも円滑なコミュニケーションが必要となってくる。

ステップ 3 の自社開発モデルを第三者にライセンスする場合となると、さらに飛躍的にリスクが上がる。ユーザーが何をかわからないというリスクを抱えることになる。モデル自体におかしな出力が出ないように厳重に技術的なガードレールを施しておくとか、内部監査や外部監査をモデルに入れておくとか、あるいは利用規約を作り込むとかといった対応が必要になってくる。ここまできると AI のリスクを統括して責任を取る AI 統括責任者を置く必要がある。そこには技術、品質、法務、コンプライアンス、モニタリングなどのいろいろなバックグラウンドの方々を集めて横断的に対応していく必要がある。

AI システムのガバナンスは、まずリスクを特定し、

評価し、対応していくという流れとなるが、この大枠ではこれまでのリスクマネジメントの一般的なフレームワークがかなり使える場合も多い。しかし、個々の要素を見ていく際に、AI のリスクマネジメントで特に留意すべき特徴がある。それは、リスクの特定においてAI 原則（公平性・安全性・透明性・プライバシーなど）を参照すべきであること、評価においては損害の大きさやリスク発現確率の数値化が困難であること、ステークホルダーにとってのリスク受容度も評価が必要であること、などである。

2.5. まとめ

AI のリスクについては膨大なドキュメントが出ているが、冷静に人間のリスクとの差分という観点から全体の骨格を把握することが大事である。リスクマネジメントシステムとかフレームワークがどこまで使えそうかということを検証するところから始めるとよい。2025 年は国内外ともに AI リスクについて、企業の自主的取り組みに委ねる傾向が強くなってきている。その背景に、政府や業界団体が一律のルールを作ることには限界があるということがある。そのため、企業自身の適切なマネジメントと、アカウンタビリティが非常に重要になるが、様々な環境変化の速さを考えると、一旦決めたらそれで終わりではなく、アジャイル的にリスクマネジメントのやり方自体を見直すというプロセスが必要である。その過程で1 線部門と2 線部門の横の連携および現場とトップ層の縦の連携が従来以上に不可欠となる。そして重要な点は、これらはAI に限ったことではなく、ガバナンス全体に関わるということである。つまり、今の時代、組織は徐々にトランスフォーメーションしていく必要がある。それは例えば人権・環境・デジタルなど、様々な領域で共通して必要となる、まさに不確実な社会を生きていくための共通のアプローチであると言える。

3. 追加情報

羽深氏が代表理事長を務める AI ガバナンス協会は、約 100 社が参加する国内最大規模の団体であり、官民・

国内外と連携して、AI 政策提言やツール提供を行っている。

また、講演内容の詳細は著書「AI ガバナンス入門ーリスクマネジメントから社会設計まで」にまとめられている。

4. おわりに

AI ガバナンスの最新動向や実践的な知見について貴重なお話をいただいた。羽深氏の専門的な視点と豊富な経験に基づく講演は、参加者一同にとって非常に有益であり、AI ガバナンスの重要性を再認識する機会となった。圧倒的な情報量で、時間を一杯使ってお話いただいたが、明解な説明は聞き手に伝わりやすくテンポのよいお話で、最後まで聞き手を惹きつける講演であった。

AI 法規制および AI ガバナンスは今後の企業活動を考える上で避けては通れない。ビジネス機器業界も無関心ではいられない状況の中で、今後も世界の動きや AI ガバナンスに関する知識を深め、実践に活かしていくために、引き続き情報の調査と共有を進めていく所存である。本講演会が、参加者の皆様の活動における一助となれば幸いである。

改めまして、羽深宏樹氏ならびに関係者の皆様に感謝申し上げます。

禁 無 断 転 載

2024 年度「ビジネス機器関連技術調査報告書」 “I－2” 部

発行 2025 年 6 月

一般社団法人 ビジネス機械・情報システム産業協会（JBMIA）

技術委員会 技術調査専門委員会

〒108-0073 東京都港区三田三丁目 4 番 10 号 リーラヒジリザカ 7 階

電話 03-6809-5010（代表） / FAX 03-3451-1770