

金融サービスにおける AI 利用に対する 法規制について

－ same rule 原則は新たな法規制を要求するか－

田 中 亘

要 旨

本稿は、金融サービスにおける AI 利用に対する法規制の必要性について、主に先行研究のレビューに基づく考察を行う。とりわけ、本稿では、「同様のリスクを持つ同様のサービスを提供する限り、人間が提供するときも AI が提供するときも、同様の規制を課すべきである」という same rule 原則に立つ場合、AI に対して新たな法規制を課す必要が果たしてあるのか、また、もしあるとした場合、どのような法規制をどのような理由で課す必要があるのかという問題を検討する。考察の主な結論として、現行法制は、自律的な行動をとれるのは人間だけであるという前提のもとに、人間だけを規律対象にしているため、same rule 原則を実現しようとする場合には、「AI の過失に基づいて AI 利用者が損害賠償責任を負う」というような、AI に対する新たな法規制が必要となる可能性があることを明らかにする。

キーワード：AI への法規制の必要性，same rule 原則，説明可能性，責任ルール，使用者責任

目 次

- | | |
|---------------------------------|--------------------------------|
| 1. はじめに | 可能性 |
| 2. AI の意義および金融サービスへの利用 | 3.3 AI に対する特別の法規制の正当化——別の観点から |
| 2.1 AI とは何か | 4. AI に関する法規制のあり方①－事前規制 |
| 2.2 AI の分類 | 4.1 はじめに |
| 2.3 金融サービス分野における AI の利用 | 4.2 リスクベース・アプローチ |
| 3. 金融サービスへの AI 利用に対する特別の法規制の必要性 | 4.3 説明可能性 |
| 3.1 AI について特別の法規制を必要とする根拠は何か | 4.4 事前規制の問題点 |
| 3.2 AI が人間よりも大きなリスクを惹起する | 5. AI に対する法規制のあり方②－事後規制（責任ルール） |

5.1 事後規制（責任ルール）の意義

5.2 AI に対する事後規制の課題

5.3 Ayres and Balkin (2024) の提案 —— 「意図を持たないエージェント」としての AI

5.4 Ayres and Balkin (2024) の提案の合理性

5.5 same rule 原則を実現するための法規制の改革

6. おわりに

1. はじめに

AI 技術の急速な発展に伴い、金融サービス分野を含む多様な領域で AI の活用が進んでいる（AI の定義については 2.1 で述べるが、さしあたり、一定の自律性をもって学習・推論を行うコンピュータ・システムを想定している）。AI は、社会に大きな変革をもたらす可能性を秘めているが、同時に、犯罪や情報操作に利用されたり、システミック・リスクを増大させるといった、AI の悪影響も懸念されている。AI 技術の健全な発展と利用を確保するために、AI に対する適切な規制のあり方が世界中で議論されている。

従来、日本では、AI に対する規制は、『AI 事業者ガイドライン』（総務省・経済産業省 (2024)）のような政府のガイドラインまたは業界の自主規制といったソフトローを中心に行われてきた（AI 戦略会議・AI 制度研究会 (2025) 3 頁）。もっとも、主要国では、EU における Artificial Intelligence Act（以下、「AI 法」という）の制定¹に代表されるように、法規制（ハードロー）の整備の動きも進んでいる（EU につき、Buiten et al. (2023); Kretschmer (2023); Wachter (2024), 米国につき、Brainard (2021);

Hammer (2024) 参照）。

日本でも、2024年7月に、内閣府の AI 戦略会議の下に AI 制度研究会が立ち上げられ、「リスクに応じて必要な法規制の在り方を検討」すること（内閣府科学技術・イノベーション推進事務局 (2024) 10 頁）を含め、AI に関する制度整備の検討を進めている。同研究会が 2025 年 2 月に公表した「中間とりまとめ」は、政府がイノベーションの促進とリスクへの対応の両立を図るため、AI を対象とする指針の整備や、AI に関する実態の調査・把握を行うべきであること、および、そうした対応は、実効性を確保するため法制度により実施すべきことを提唱している（AI 戦略会議・AI 制度研究会 (2025) 19 頁）。また、日本では、「現在の規則や法律で AI を安全に利用できると思う」と考える人の割合が回答者の 13% と、主要国で最も低いことを示す調査（KPMG による）も紹介している（同 4 頁）。このような政府の動きおよび一般国民の意識からすると、今後、日本においても、AI に対する法規制の動きが大きく進む可能性もあるように思われる。

本稿は、以上のような動向を踏まえて、金融サービス分野を念頭に置いて、AI について新たな法規制が必要か、また、もし必要であるとすれば、どのような法規制がどのような理由で必要

1 Artificial Intelligence Act (Regulation (EU) 2024/1689), (<https://artificialintelligenceact.eu/>). AI 法は、2024 年 8 月 1 日に発効し、その後は段階的に施行され、一部の規定を除いて 2026 年 8 月 2 日までに完全施行されることが予定されている (European Commission (2024))。

とされるのかについて、主に先行研究の調査を通じて、準備的な考察を行うことを目的とする。以下、本節では、考察に先立って、筆者の基本的な問題意識ないし研究動機を述べておきたい。

筆者は、テクノロジーと金融革新に関する研究会（第1期）の報告論文（田中（2022））。以下「旧稿」という）において、インターネット取引とロボアドバイザーに対する規制のあり方を検討したが、その中で、「同様のリスクを持つ同様のサービスを提供する限り、人間が提供しようと機械が提供しようと、同様の規制に服せしめるべきである」という、“same service, same risk, same rule”の原則（以下「same rule 原則」という）を紹介し、規制はこのような原則に立って設計することが望ましいという見解を述べた（田中（2022）58-59頁）。

Same rule 原則は、旧稿で検討対象にしたロボアドバイザー等に限らず、AI一般に対する法規制の設計方針にすることができるだけの合理性を備えていると考える。というのも、もしもAIの提供するサービスが、人間が提供する同種のサービスと同程度のリスクしかもたらないにも関わらず、人間が提供する場合よりも厳格な規制が課されるなら、AIによる技術革新を阻害することになる。また、その反対に、AIによるサービスに対し、人間の提供する同種サービスよりも穏やかな規制しか課されないなら、本来は人間が提供した方が低コストですむところにAIが利用されるといった、過剰なAI利用を生じさせてしまう。いずれも、効率性を害し、望ましくないといえよう。「中間とりまとめ」も明記するとおり、「規制はその目的を達成するために、特定の種類の技術の使用を強制したり、優遇したりすべきではない」と

いう「規制の技術中立性」を原則として、法規制（ハードロー）を含めた規制のあり方を考えていくべきである（AI戦略会議・AI制度研究会（2025）10頁）。

以上のように、same rule 原則には合理性があると考えられるが、そうすると次のような疑問が生じてくる。すなわち、same rule 原則に従うなら、提供されるサービスの種類やリスクの程度に応じて、適切な規制のあり方を考えていけば足りるのであって、ことさら「AIに対する法規制」を新たに定める必要はそもそもない、ということにならないか。Same rule 原則の合理性を認めた場合に、なお、AIについて新たな法規制を設計する必要があるといえる場合はあるのか。もしあるとすると、どのような法規制がどういう理由で必要とされるのか。本稿では、AIに対する法規制の必要性や望ましい法規制のあり方について論じている先行研究を検討することを通じ、以上のような疑問に対する回答を探ってみたい。

以下、本稿の検討は次のように進める。2では、AIとは何かを述べた後、金融サービスにおけるAIの利用事例を簡単に紹介する。3では、AIに対して法規制を必要とする理由についての従来の見解を紹介した上で、それに対する私見を述べる。4では、法規制の基本的な分類として、事前規制と事後規制（責任ルール）が存在することを示した後、金融サービスに対する事前規制として提案されている内容と、それについての課題を述べる。5では、事後規制（責任ルール）について論じる。そこでは特に、もしも責任ルールに関して same rule 原則を実現しようとするのであれば、従来はもっぱら人間のみが行為主体になることを前提に設計されてきた責任ルールを、AIも行為主体に含

めるように（法解釈，必要があれば立法により）修正する必要があるという見解を述べる。6 は，考察の主要な結果を述べて，結びとする。

2. AI の意義および金融サービスへの利用

2.1 AI とは何か

AI（Artificial Intelligence; 人工知能）については，いまだ確立した定義は存在しないといわれるが，一般的には，人間の思考プロセスと同じような形で動作するコンピュータプログラムないしシステムを指すといえよう（総務省・経済産業省（2024）9 頁）。ただ，特に近年は，「自律性」という特徴を重視した定義づけがされるようになっているようである。すなわち，単に人間が定めたルール（プログラム，アルゴリズム）に従って予測や判定等を行うのではなく，機械が自律的に，大量のデータに基づく機械学習ないし深層学習（ディープ・ラーニング）を通じて，パターンとか法則性等についての推論を行い，それに基づいて予測や提案，決定といったタスクを行うものを，AI（または AI システム）と呼ぶようになってきている（古川（2024）43 頁）。

『AI 事業者ガイドライン』は，「AI システム」を，「活用の過程を通じて様々なレベルの自律性をもって動作し学習する機能を有するソフトウェアを要素として含むシステム」と定義し，

「AI」は，そのように定義した「AI システム自体又は機械学習をするソフトウェア若しくはプログラムを含む抽象的な概念」を意味するものとして用いている（総務省・経済産業省（2024）9 頁）。EU の AI 法における“AI system”の定義も，自律性をもって推論を行うという機能に着目したものになっている（古川（2024）43 頁）²。

2.2 AI の分類

AI は，特化型 AI と汎用型 AI に分類できる（AI 戦略会議・AI 研究会（2025）5 頁）。特化型 AI とは，音声認識，画像認識，自動運転等の特定のタスクを処理することに特化した AI のことであり，汎用型 AI とは，特化型 AI よりも大量のデータで学習され，高い汎用性を示し，様々なタスクを処理できる AI のことである（同）。近年，話題になっている生成 AI（generative AI）も，汎用型 AI に属する（同）。生成 AI とは，学習したコンテンツに基づいて，画像，文章，プログラムといった新たなコンテンツを生み出すことができる AI のことである（総務省・経済産業省（2024）9 頁。須賀（2024）8 頁，15 頁も参照）。

2.3 金融サービス分野における AI の利用

AI は，金融サービス分野において，様々な形で利用されている（以下に挙げる事例の他，Hammer（2024），pp.4-7；岡田ほか（2024）参照）。

2 EU AI Act, Article 3(1)参照（“‘AI system’ means a machine-based system that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment, and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments.”）（「『AI システム』とは，さまざまなレベルの自律性をもって動作するように設計され，展開後に適応性を示す可能性があり，明示的または暗黙的な目的のために，受け取った入力から，物理的または仮想的な環境に影響を与える予測，コンテンツ，推奨，または決定などの出力を生成する方法を推論する，機械ベースのシステムを意味する。」）。

例えば、与信審査への利用事例として、セブン銀行が、2024年10月に、自社の金融データと親会社（セブン&アイ・ホールディングス）の提供する購買データとを組み合わせ、AIの分析結果を活用した与信審査を行う取り組みについて公表している（日本経済新聞（2024/10/16））。

また、ロボアドバイザーによる投資助言や投資一任業のように（Baker and Dellaert (2018); 田中 (2022)）、投資判断や資産運用にAIを活用する動きも広がっている（Bagattini et al. (2023); Dou et al. (2024)）。2021年の調査によれば、ヘッジファンドの48%が、投資判断にAIを利用しているという（Hammer (2024), p.5）。

モニタリングやリスク管理にAIを活用する動きも進んでいる。大手カード会社がAIを用いた詐欺検出ツールを開発し、不正取引の未然阻止に活用しているとか（Business Insider Japan (2024)）、銀行がマネーロンダリング対策にAIを活用する事例（日経クロステック (2024)）等が報道されている。

近年、話題になっている生成AIについては、金融商品の広告物のコンプライアンス審査や、稟議書などの社内文書の作成、顧客からの問い合わせへの回答等への利用が既に行われている（日本経済新聞（2024/7/19））。また、金融サービスに限ったことではないが、監査業務への生成AIの利用も進んでいる（日本経済新聞（2024/6/4））。

3. 金融サービスへの AI 利用に対する特別の法規制の必要性

3.1 AI について特別の法規制を必要とする根拠は何か

次に、金融サービスへの AI 利用に対して、何らかの新たな法規制を課すとすればその根拠は何か、という問題を考えてみたい。

1で紹介した same rule 原則は、「同様のリスクを持つ同様のサービスを提供する限り、人間が提供するときも AI が提供するときも、同様の規制を課すべきである」という考え方である。逆に言えば、この原則によるときは、もしも AI による金融サービスの提供が、人間がサービス提供をするときには生じないような特別のリスクをもたらしたり、あるいは、人間がサービス提供するときにもたらすリスクよりも類型的に大きなリスクをもたらすとすれば、AI について人間よりも加重した規制を課することは正当化される、ということになるであろう。

それでは、AI は、そのように特別な法規制を必要とするような特別のリスクを惹起するのだろうか。従来、AI の特別のリスクとして識者が強調するのは、次の2点である。

①不明瞭性

第1は、AI の不明瞭性（opaqueness）という問題である（Hammer (2024), p.4）。モデルの透明性の欠如（lack of model transparency）またはブラックボックス問題と呼ばれることもある（Brainard (2021), p.5）。AI は、データに基づいて自分でモデル（アルゴリズム）を構築し、それによって、投資助言、与信審査、保険

の引受けといった様々な判断をすることができる。そうすると、AI がどのようにして特定の判断に至るのが AI 事業者（AI の提供者または利用者）³にも分からないため、AI の動作を人間が適切にコントロールできない事態が生じうる（Hammer (2024), p.4; Brainard (2021), pp.5-6）。その結果、AI の不適切な動作（誤った予測や決定、他人の権利を侵害するようなコンテンツの生成など）が、人間の手で予防または是正されることなく、消費者その他の社会の人々に損害を与える恐れがある。さらに、そのような損害が、AI の不適切な動作によるものであったことが事後的にも判明しない（因果関係の立証ができない）ために、被害者が救済を受けられない恐れも指摘されている（Buiten et al. (2023), p.6）。

②データの問題

AI のリスクとしてもう一つ指摘されるのは、データに関する問題である。AI は大量のデータからパターンを抽出し、予測や決定、コンテンツの生成といった動作を行う。そうすると、使用するデータに、特定のグループ（特定の人種、性別等）に不利なバイアスが存する場合には、AI モデルがそのバイアスを増幅し、差別的な信用スコアリングやレッド・ライニング（特定グループをサービスから排除してしまうこと）といった、不適切な動作をするリスクが指摘されている（Brainard (2021), pp.4 & 7）。

AI には、不明瞭性とデータの問題という特有

のリスクがあるために、特別の規制——とりわけ、4 で後述するような事前規制——が必要になるという主張が、今日、欧米諸国における AI 規制を巡る議論において主流となっているように見受けられる（Brainard (2021); Bagattini et al. (2023), pp.17-18; Hammer (2024)）。もっとも、これには疑問がないではない。不明瞭性の問題についていえば、人間もまた、脳内でどういう思考過程を経て判断に至るのかはブラックボックスだといえる。また、判断の基礎となるデータにバイアスがあるときに誤った判断に導かれることも、人間も同じだろう。そうだとすると、なぜ、AI について、ブラックボックスやデータのバイアスをことさらに問題視するのか、という疑問が生じるように思われる。

3.2 AI が人間よりも大きなリスクを惹起する可能性

3.1 の最後に指摘した疑問に対して考えられる回答としては、不明瞭性やデータに関する問題は、確かに AI にも人間にも存するのであるが、AI は、わずかの時間に大量のデータに基づき、大量の動作を行うことができるため、不適切な動作（例えば、誤情報の拡散とか差別的な与信審査など）をしたときに社会の人々に与える損害が、人間の場合とは段違いに大きくなりうる、ということが挙げられる。

また、AI システムの開発・提供には、高度な技術と大量のデータを必要とするため、少数の提供者の寡占状態となることが見込まれる。そのため、証券市場の主要参加者が、みな同種

3 本稿では、EU の AI 法の定義（およびその訳語として日本で一般に用いられる表現）に従い、AI システムの開発者や自己の名で AI サービスを提供する者を「提供者（provider）」といい、事業のために AI を利用する者を「利用者（deployer）」という（AI 法 3 条(3)号・(4)号。私的利用のために AI を用いる自然人は、「利用者」とはならないことに注意）。そして、両者を合わせて「AI 事業者」と呼ぶことにする。金融サービス分野で AI を利用する金融機関は、主に AI の利用者であろうが、一部の AI については自社開発しているため、提供者となる場合もあるであろう。

の AI モデルを使用して取引をするようになる事態も想定される。そうすると、市場のショック時に群衆行動や連鎖反応を引き起こし、市場のボラタリティを増幅するといったシステムミック・リスクが生じる可能性が指摘されている (Bagattini et al. (2023), p.17; Baker and Dellaert (2018), p.742)⁴。さらに、AI を搭載した複数の取引システムが、自律的な学習を通じて互いの行動を予測しあうことによって、共謀的な価格操作 (相場操縦) を行い、市場価格の歪みや流動性の低下を引き起こすリスクを指摘する見解もある (Dou et al. (2024))。システムミック・リスクも共謀的な価格操作も、AI に限らず、人間が引き起こす可能性ももちろんあるが、AI がその危険を増幅する、ということはいえるだろう。

このように、AI がもたらすリスクは、必ずしも、人間がもたらすものと質的に異なるわけではないが、リスクの量において違いがあるという。その点から、AI に対して人間よりも加重された規制 (例えば 4 で述べる、事前規制) を課すことが、一定の合理性を持つといえるかもしれない。

3.3 AI に対する特別の法規制の正当化 ——別の観点から

ここでは、AI に対して特別の法規制を課すことの正当化として、3.2 で述べたこととは別の説明を考えてみたい。すなわち、不透明性 (ブラックボックス) やデータの問題によるリスクが生じることは、AI についても人間について

も同様であるが、ただ AI の場合には、適切な法規制により、人間の場合よりもそのリスクを軽減または除去できる可能性がより高いのではないか、という点である (類似の観点を述べるものとして、Kleinberg et al. (2018))。

例えば、人間が引き起こす差別は、アンコンシャス・バイアスの存在などを考えれば、法規制によってもこれを除去することは容易ではないだろう。しかし、AI の場合には、法規制により、AI の開発または提供の段階において、学習するデータに制約を課したり、または特定の属性に基づく判断 (例えば、人種や性別に基づく信用スコアリング) を禁じるプログラムを組み込むことを強制することにより、差別的な判断を完全に除去することができるかもしれない。これは、人間の思考を操作すること (例えば、差別的な傾向のある人の脳を手術してその傾向を取り除くこと) は、技術的な可能性以前に、法的・倫理的な観点からの重大な制約があるが、AI についてはその制約はなく、技術的に操作が可能である限り法的に強制することも可能である、という違いに基づく。つまり、法規制の必要性の違いというよりは、むしろ法規制の実現可能性 (enforceability) の違い (AI のほうが人間よりも、法規制によって実現できることが多い) が、人間に対しては課されないような法規制を AI に対して課すことを正当化するのではないか、ということである。今後、このような観点からの研究を進めたいと考えているが、本稿では、着想を述べるにとどめておきたい。

4 「中間とりまとめ」でも、AI がシステムミック・リスクを生じさせる可能性に注意が払われている。「システムミック・リスクについては、将来的に、複数の AI システムが連動する大規模な AI システム群が社会システムを支える状況となる可能性があり、その際、当該 AI システム群が予期せぬ挙動をした場合、社会全体に大きな混乱をもたらす可能性があるため、適切に対処することが重要である」(AI 戦略会議・AI 制度研究会 (2025) 19頁)。

4. AI に関する法規制のあり方① ー事前規制

4.1 はじめに

3.2 で述べたような疑問はあるものの、一般的には、AI は、不明瞭性およびデータの問題という特有のリスクを持つことから、AI に対して特別の法規制を課すことには合理性があるとする見方が多いと思われる (3.1 参照)。特に、AI を開発もしくは提供し、または AI を利用して事業を行う者に対し、事前に AI のリスクに対処するため一定の行為を要求する規制 (事前規制) を課す必要があるとする見解が有力であり (Hammer (2024)), EU のように、既に体系的な事前規制 (AI 法) を制定した法域もある。

本節 4 では、AI に対する事前規制についての議論を概観する。なお、法規制の選択肢としては、事前規制の他に、AI が実際に損害をもたらしたときに AI 事業者には損害賠償責任を課すという規制 (事後規制または責任ルール) も重要であるが、これについては、次節 5 で検討する。

4.2 リスクベース・アプローチ

事前規制の方法として一般にとられているのは、AI の惹起するリスクをその程度 (強さ・範囲) に応じて分類し、よりリスクの程度の大きな AI 利用に対してより厳格な事前規制を課すという、リスクベース・アプローチである。

例えば、Hammer (2024) は、金融サービス

に関連して AI が惹起するリスクを、高度のもの、中程度のものおよび低度のものの 3 種類に分類する。高度のリスクの例としては、偏見や差別、金融市場に対する重大なインパクト、顧客財産に重大な損失を与えること等が挙げられる。中程度のリスクには、顧客の個人情報やプライバシーの侵害、顧客行動の不当な操作 (manipulation) が含まれる。プロセスの自動化・最適化に伴うリスクは、低度のリスクに分類される。そして、高いリスクを惹起する AI ほど、それを利用する金融機関に対して高度な説明可能性 (4.3 で後述) を要求することを提唱する (Hammer (2024), pp.11-12)。

EU の AI 法 (前掲注1) は、リスクベース・アプローチに基づく包括的な事前規制の代表例である。AI 法は、AI システムのリスクを、「許容できないリスク (unacceptable risk)」、「高リスク (high risk)」、「限定リスク (limited risk)」 (または「透明性リスク (transparency risk)」、「最小リスク (minimal risk)」) に分類し、リスクの程度に応じて段階的な規制を課している (European Commission (2024); 鈴木・森田 (2024); 古川 (2024) 41 頁。図表 1 参照)⁵。

「許容できないリスク」を持つ AI 利用は、禁止される (AI 法 5 条)。例えば、サブリミナ

図表 1 EU AI 法のリスク分類

リスクの程度	規制方針
「許容できないリスク」	禁止
「高リスク」	厳格な規制
「限定リスク」 (「透明性リスク」)	透明性の義務
「最小リスク」	特段の規制なし

出典: European Commission (2024)

5 本文で以下に説明する規制に加えて、AI 法は、汎用型 AI (同法にいう general-purpose AI models) に対する特別の規制 (AI 法 51~56 条) などを定めているが、説明は省略する。

ル技術を用いるもの（同条1項(a)）や、人の社会的行動を評価するソーシャルスコアリングのうち一定のもの（同項(c)）が、それに当たる（古川（2024）44-46頁）。

次に、「高リスク」のAIシステムに対しては、リスクマネジメントシステムの導入・実施、データガバナンス、記録の保存、人間による監視などの「要求事項(requirements)」の遵守義務（同法8～15条、16条(a)）や、サービス開始前にEUのデータベースに登録する義務（同法16条(i)・49条）といった、詳細な事前規制が課される⁶。どのようなAIシステムが「高リスク」に分類されるかは、AI法6条および付属書(Annex)が定める（羽深・古川（2024）96-98頁）。金融サービスに関係するものとしては、「信用性の評価またはクレジットスコアを設けるために使われるAIシステム（財産的詐欺の検出目的のために使われるものを除く）」（AI6条(2)・付属書Ⅲ5項(b)）や、「生命保険、健康保険において自然人に関するリスクアセスメントおよび価格決定に使われるAIシステム」（同項(c)）が挙げられる。

AI法はまた、チャットボットのようにユーザーがAIシステムと直接やりとりする場合（「透明性リスク」をもたらす場合）には、ユーザーが接しているものがAIであることを通知することなどの透明性義務を課している（AI法50条。European Commission（2024））。

以上に規律するもの以外のAIシステムは、「最小リスク」しかもたらさないものとして、AI法では特に規制していない（European Commission（2024））。

4.3 説明可能性

AIに対する事前規制において、特に重視されているのが「説明可能性」(explainability)の要求である（Brainard（2021）, pp.6-7; Hammer（2024）, pp.9-12）。これは、AIモデルがどのようにして決定や予測に至るのかを理解し解釈できるようにするということである。金融機関に対し、AIモデルの明確な説明を要求することによって、規制当局が、AIモデルの潜在的リスクをよりよく評価することを可能にするとともに、AIがバイアスや誤った判断を生じたり、システミックリスクを引き起こしたりすることを防止できるようになることが期待されている。説明のための具体的な方法としては、重要度分析（モデルの出力に最も大きな影響を与える入力変数を特定すること）、感度分析（入力の変更が予測にどのように影響するかを調べること）、ルール抽出（モデルの意思決定プロセスを解釈可能なルールのセットに抽出する試み）といったものが考案されている（Hammer（2024）, p.10）。

Hammerは、説明可能性の要求は、既に多くの法域の金融規制に取り入れられているとして、バーゼル規制、EUならびに米国・英国の関連する規定を挙げている（Hammer（2024）, Appendix, pp. 15-21）。

4.4 事前規制の問題点

もっとも、事前規制については問題点ないし課題も指摘されている。規制担当者が、AIシステムのリスクに対処するためにどのような行為を事業者に要求することが適切であるかを事

6 高リスクのAIシステムについての条文は、AI法の条文全体の85%以上を占める（羽深・古川（2024）95頁）。規制の詳細については、羽深・古川（2024）、古川・羽深（2024）、有坂（2024a）（2024b）、吉永（2024）参照。

前に知することは難しいために、事前規制が、費用便益の観点から過剰または過小となったり、不明確な内容になったりする恐れがある。4.3 で述べた AI モデルの説明可能性についても、これを過度に要求した場合には、事業者をして、複雑だが効果的な AI システムの導入を避け、単純な AI システムに逃避してしまう恐れがあることが指摘されている (Hammer (2024), p.12)。

Kretschmer らは、EU の AI 法の事前規制 (特に、高リスクの AI システムに対するもの) について、主に規制が複雑かつ過剰であるという観点から批判を加える。例えば、同法14条は、高リスクの AI に対するリスク管理の一貫として「人間による監督 (human oversight)」を要求している。これに対して、Kretschmer らは、AI システムの作業は膨大であるため、人間が個々の結果をすべて評価することはできないことや、AI の反応は速すぎて人間が監督していても誤作動を防げないことが多いことなど、人間が AI を監督することによる負担 (心理的な苦痛を含む) に見合う効果が見込めるかについて疑問を提起している (Kretschmer et al. (2023), p.9)⁷。

このような批判的検討を踏まえて、Kretschmer らは、「特定の行為に対して事前の制約を課すことよりも、有害な帰結に対する責任という形の制裁を設計の方が望ましく思われる」と述べて、AI のリスクに対しては、事前規制よりはむしろ事後規制 (責任ルール) を適切に設計することによって対処することを

提唱する (Kretschmer et al. (2023), p.10)。そこで、次に、AI に対する事後規制 (責任ルール) についての議論を見ていこう。

5. AI に対する法規制のあり方② －事後規制 (責任ルール)

5.1 事後規制 (責任ルール) の意義

本節において検討対象とする「AI に対する事後規制」とは、AI が顧客等の人々に損害をもたらしたときに、AI 事業者 (提供者または利用者)⁸にその損害を賠償する責任を負わせる法制度、つまり、法と経済学において責任ルール (liability rules) の名で分析されている法制度を指す (Calabresi and Melamed (1972))。適切に設計された責任ルールは、人々に損害をもたらす危険のある活動をする者に対し、損害防止のための注意を払うインセンティブを与えること等を通じ、効率性ないし社会厚生を改善する機能を有する (Calabresi (1970) ; Shavell (2004) [邦訳 2010], chaps.8-12)。法と経済学の有力な研究者の中には、国 (政府) は適切な事前規制を行うために十分な情報を持たないことが多いとして、事前規制を減らし、責任ルールをより多く利用することが望ましいと主張する者もいる (Shavell (2004), p.591 [邦訳 2010, p.686])。

5.2 AI に対する事後規制の課題

もっとも、AI のもたらすリスクに対し、AI

7 本文で紹介した Kretschmer et al. (2023) とは逆に、AI 法には多くの抜け穴 (loopholes) があり、AI のリスクに適切に対処できていないとする批判も存在する (Wachter (2024))。

8 提供者・利用者の定義は、注3参照。なお、AI の起こした損害について、提供者と利用者にとどのように責任を分担させるかという点は、責任ルールの設計において重要な問題であるが、紙幅の関係から本稿では扱わない。Kretschmer et al. (2023), pp.10-13参照。

事業者に対する事後規制（責任ルール）で対処することについては、課題も指摘されている。

すなわち、現行法制は、特別法によって厳格責任（無過失責任）が定められている場合（例えば、製造物責任法3条参照⁹）を除き、一般に損害賠償責任は、加害者が故意又は過失（以下、単に「過失」という）によって起こした損害についてのみ生じるものとする、過失責任（negligence/fault-based liability）の原則を採用している（一般不法行為責任を定める民法709条参照¹⁰。欧米主要国も同様である〔Buiten et al. (2023), p.3〕）。もっとも、事業のために他人（被用者）を使用する者（使用者）については、被用者がその事業の執行について第三者に加えた損害について責任を負うものとする、使用者責任のルール（民法715条1項）が存在する。同項ただし書は、使用者が被用者の選任・監督につき相当の注意を払ったことを証明すれば免責される旨を規定するが、実際は、裁判で免責が認められることはほぼないといわれており、その限りでは、使用者に無過失責任を課すように運用されている（窪田（2018）205頁）。けれども、使用者責任の成立のためには、被用者の第三者への加害が不法行為と評価されるものであることが前提であると解されている（同206頁）。そのため、被用者の過失によって第三

者に損害が生じたと認められること（民法709条参照）が、使用者責任の成立のためにはやはり必要であると一般には解されている。

このように、加害者（加害者が使用者である場合、その被用者を含む）という「人間」の過失によって損害が起きたことを責任の成立要件にしている現行法制のもとでは、AIが自律性を獲得するほど——データからAI自身がアルゴリズムを構築し、それによって予測、決定、さらには新たなコンテンツ生成等を行えるようになるほど——、AIが起こした損害について、それが人間の過失によって起きた損害であると認めることがより困難となり、責任ルールがうまく機能しなくなる恐れが指摘されている（Buiten et al. (2023), p.7）。

EUは、上記の問題に対処するため、2024年に、製造物責任指令を改正し¹¹、規制対象の「製品」の範囲にソフトウェア等の無体物を含めることにより（同指令4条(1)）、同法の無過失責任のルール（製品の欠陥から生じた損害については、製造者は過失の有無を問わず、責任を負う。同指令5条）をAIにも適用できるようにした（Buiten et al. (2023), pp.4-5）。また、2022年に公表されたAI責任指令案¹²では、高リスクAIシステムが損害をもたらしたことが疑われる場合に、裁判において被告たるAI事

9 平成6年（1998年）7月成立の製造物責任法は、製造物の「欠陥」から生じた損害については、製造業者は、原則として（例外は同法4条）、欠陥についての過失の有無を問わず賠償責任を負うものとする（同法3条）。これは、特別法による無過失責任の代表例である。しかし、現行法上、「製造物」は「動産」に限定され（同法2条1項）、ソフトウェアのような無体物は含まないと解されているため（窪田（2018）269頁）、AIの誤作動が起こした損害についてAI提供者に同法の責任を問うことは難しい。なお、この点に関し、EUにおける近時の製造物責任指令の改正について、後掲注11と対応する本文参照。

10 民法709条参照（「故意又は過失によって他人の権利又は法律上保護される利益を侵害した者は、これによって生じた損害を賠償する責任を負う」）。

11 DIRECTIVE (EU) 2024/2853 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL of 23 October 2024 on liability for defective products and repealing Council Directive 85/374/EEC, (https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=OJ:L_202402853).

12 Proposal for a DIRECTIVE OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL on adapting non-contractual civil liability rules to artificial intelligence (AI Liability Directive)." COM (2022) 496 final 2022/0303 (COD) (28 Sep 2022), (<https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=celex%3A52022PC0496>)

業者に証拠提出義務を課すとともに（同指令案 3 条）、被告が証拠提出を怠った場合には、被告の過失と損害との間の因果関係を推定する旨の規定（同 4 条(1)）を置くことが提案されている（Buiten et al. (2023), p.5）。

このような EU の近時の立法（草案段階も含む）は、日本における事後規制（責任ルール）の将来設計を考える上でも、参考になるであろう。ただ、本稿では、以下、AI 時代に即応する責任ルールの構築を目指すものとして理論的により興味深いと思われる Ayres and Balkin (2024) の提案を紹介し、検討することにした。

5.3 Ayres and Balkin (2024) の提案 ——「意図を持たないエージェント」 としての AI

Ayres と Balkin は、AI の利用により生じた損害について誰がどのような要件のもとに責任を負うかという問題の解決策として、AI を「意図を持たないエージェント（agents without intentions）」として法的に取り扱うことを提唱する（Ayres and Balkin (2024), *1）。

Ayres と Balkin が検討対象とする米国法においては、一般に、過失責任のような法的責任の判断基準は客観化されている。そのため、AI のように意図を持たない存在であっても、客観的に見て合理的な損害回避措置（reasonable care）をとっていたか否かを問うことにより、その過失の有無を判定することが可能である（Ayres and Balkin (2024), *3）。そうであれば、

AI をエージェント（代理人）とみなし、金融機関のように AI を利用して事業をする主体（AI 利用者）をプリンシパル（本人）とみなして、エージェントの行為についてプリンシパルに責任を負わせるエージェンシー法理を適用することができる¹³。とりわけ、エージェンシー法理の下では、エージェントがプリンシパルの支配（control）下で職務を行う場合は、エージェントは被用者（employee）、プリンシパルは使用者（employer）とされ¹⁴、被用者が雇用の範囲内（scope of employment、日本の民法715条 1 項の「その事業の執行について」にほぼ相当する）で犯した不法行為については、使用者は、無過失の代位責任（vicarious liability、Respondeat superior ともいう。日本法の使用責任にほぼ相当）を負うものとされている¹⁵。そこで、AI を AI 利用者のエージェント（さらには被用者）とみなせば、AI が客観的に見て過失ありと判断されるような動作によって第三者に損害を与えた場合は、AI 利用者の賠償責任が認められることになる（Ayres and Balkin (2024), p.*3）。

Ayres と Balkin がこのような責任ルールを提唱するのは、「人間のエージェントを AI のエージェントに置き換えることによって、注意義務の軽減が得られるものとするべきではない」という考慮に基づくものである（Ayres and Balkin (2024), p.*2）。これは、same rule 原則の考え方に他ならない。

13 米国のエージェンシー法理については、Restatement (Third) of Agency (2006) 参照。同リステイメントの邦訳および解説として、樋口・佐久間 (2014) 参照。また、Cox and Eisenberg (2019), Chap.1 も参照。

14 Restatement (Third) of Agency, section 7.07(3)。

15 Restatement (Third) of Agency, section 7.07(1)(2); Cox and Eisenberg (2019), pp.30-33。米国法（各州の判例法理）の下で、使用者の代位責任が成立する範囲は、日本法の下で使用者責任が成立する範囲とは必ずしも一致しない（前者のほうが後者より狭いように見える部分がある）が、詳しい対比は省略する。樋口・佐久間 (2014) 第 9 章【樋口範雄】参照。

5.4 Ayres and Balkin (2024) の提案の合理性

Ayres と Balkin の提案は、日本法の下でも採用を検討するに十分値するほどの合理性を持っていると考える。

5.1 で述べたとおり、日本法の下でも、被用者が使用者の事業の執行について第三者に不法行為による損害を与えたときは、使用者は、事実上無過失の賠償責任を負うものとされている(民法715条1項。同項ただし書による免責がほぼ認められていないことは、既に述べた)。このような使用者責任は、使用者に対し、被用者の選任・監督を通じて損害回避のための行動(注意)をとらせたり、損害をもたらす危険のある事業活動の水準を抑えるインセンティブを与えることにより、効率性を改善する機能を持つ(Shavell (2004), p.233 [邦訳 (2010) 266-267頁])。特に、使用者責任を過失責任でなく無過失責任とすることは、損害回避のために被用者にどのような行動をとらせることが適切であるかについては国(裁判所)よりも使用者の方が情報を持つことが多いと考えられる点(Shavell (2004), p.234 [邦訳 (2010) 268-269頁])や、注意水準だけでなく活動水準をも効率的にするインセンティブを与えることができる点(Shavell (2004), chap.8, sec.4 [邦訳 (2010) 221-228頁] 参照)で、合理性を持っている¹⁶。そして、same rule 原則に則り、このような使用者責任は、事業者が用いる者(もの)が人間であろうと AI であろうと、変わらず適用され

るべきである。

そして、日本法においても、過失とは、精神の緊張を欠いたというような主観的なものではなく、損害の発生という結果を回避するためにとるべき行動をとらなかったこと(結果回避義務の違反)というように、客観的な行為態様であるとする理解が一般的となっている(平井(1994) 28頁, 窪田(2018) 46頁)。このような過失の理解によれば、AIの動作(予測, 決定, コンテンツ作成等)についても、同じ状況下で人間が同様の動作をしたとすれば過失ありと判断される場合には、AIがした場合でも過失ありと解することにより、「AIの過失」を観念することが可能となる。

例えば、チャットボットによる顧客の質問に生成 AI が答えるサービスは、多くの金融機関で既に実用化されているが(岡田ほか(2024))、もしも AI が虚偽の回答をし、それによって顧客に損害が生じた場合には、人間が虚偽の回答をした場合と同様に不法行為と評価され、AI 利用者である金融機関の顧客に対する損害賠償責任が生じるものとすべきである。また、将来的には、AI による顧客への金融商品の説明や、さらには金融商品の勧誘が一般的に行われるようになることも想定されるが、その際に、AI が説明義務(金融商品取引法[金商法] 37条の3・金融商品取引業等に関する内閣府令117条1号, 金融サービス提供法 4条1項, 飯田(2023) 398-399頁参照)に違反する態様の説明をしたり、あるいは適合性の原則(金商法40条1項)に著しく反する態様の勧誘¹⁷をした場合には、

16 本文で述べたような使用者責任(英米法の vicarious liability)の存在意義(機能)についての詳しい分析として、Sykes (1984) 参照。

17 適合性の原則を定める金商法40条1項自体は業法上の規制であるが、金融商品取引業者等が、適合性の原則を著しく逸脱した金融商品取引の勧誘をする場合には、不法行為法上も違法となり、損害賠償責任が生じることが判例により認められている(最判平成17・7・14民集59巻6号1323頁)。飯田(2023) 407-410頁参照。

人間の被用者がそうした説明や勧誘をした場合と同様に、AI 利用者である金融機関の顧客に対する損害賠償責任が生じるものとすべきである。

5.5 same rule 原則を実現するための法規制の改革

問題は、「AI の過失（説明義務、適合性の原則等、各種法令上の義務の違反も含む）に基づいて AI 利用者が損害賠償責任を負う」という責任ルールが、日本の現行法の解釈として実現可能かどうかである。

この点については、民法715条1項の使用者責任にいう「被用者」は、従来は人間であることが当然の前提とされてきたと思われる。しかし、同項は、「被用者」について特に定義を置いていないし、また、雇用契約がなくても民法715条の使用者・被用者関係は成立しうるといように、被用者概念は以前から柔軟に解されてきた（窪田（2018）208頁）。そのことを考えれば、技術の進展により AI がますます自律性を獲得しつつある今日では、AI も被用者に当たりうるという解釈も可能であるように思われる。

もう一つ、現行法の解釈として実現可能と思われる選択肢は、AI に過失（各種義務の違反を含む）と評価できる動作があった場合、それを直接、（通常は法人である）AI 利用者自身の過失であると捉えて、AI 利用者自体に民法709条の不法行為責任の成立を認めることである。「法人自体の不法行為」を認めうるかについては、従来、裁判例が対立し、学説上も議論のあるところであるが（平井（1994）226-227頁）、法人の個々の被用者の過失を立証できない（そのため、民法715条1項の責任追及が難しい）

場合に、なお被害者の救済を可能にするための法解釈として、一定の範囲でこれを認める見解も有力である（橋本（2016））。過失概念が客観化している今日にあっては、法人において損害回避のためとるべき措置がとられていなかった（結果回避義務の違反があった）と認められる場合は、それを誰か特定の人間の責めに帰せしめるまでもなく、法人自身に過失ありと認めてもよいように思われる。今後、自律性のある AI 利用が拡大するにつれて、このような法解釈が受け入れられていく可能性はあると考える。

もっとも、民法715条1項の「被用者」には人間以外の存在が含まれるという解釈は、従来、誰も主張していないものであり、果たしてこのような解釈が裁判実務に採り入れられる可能性があるのか、また、仮に採り入れられるとしてもそれがいつのことになるかは、まったくの未知数であると言わざるを得ない。また、法人自体に不法行為責任が成立するかについても、過失の前提である行為義務違反の有無を判断するには特定人（自然人）の行動を問題にしなければならないなどとして、否定する見解も有力である（平井（1994）227頁）。

もしも、「AI の過失に基づいて AI 利用者の損害賠償責任が生じる」というルールを現行法の解釈として実現することが難しいとすれば、same rule 原則の実現のためには、このようなルールを明文で認めるための立法が必要になるだろう。これは、same rule 原則が、AI に対する新たな法規制を要求する場合の一つといえるかもしれない。つまり、本来は、同様のリスクを持つ同様のサービスを提供する限り、人間であっても AI であっても同様の法規制の適用を受けるべきなのであるが、現行法は、自律的な行動をとることができるのは人間だけである

という（これまでは疑われることのなかった）前提のもとに、人間が行為主体となる場合だけを規律対象としたルールを置いている。そのため、same rule 原則を実現するためには、特に AI について新たな法規制が必要になる場合がある、ということである。

6. おわりに

本稿は、「同様のリスクを持つ同様のサービスを提供する限り、人間が提供するときも AI が提供するときも、同様の規制を課すべきである」という same rule 原則に立つ場合、AI に対して新たな法規制を課す必要性が果たしてあるのかという問題意識のもとに、金融サービスにおける AI 利用に対する法規制をめぐる議論を概観した。上記の問題意識に対し、本稿の検討から得られた一応の回答をまとめると、次のようになる。

- ① AI については、不透明性とデータの問題という特有のリスクがあるために、特別の法規制が必要になるという見解が有力であるが、これらのリスクは人間についても存するものである。ただ、大量のデータを瞬時に処理できると言う AI の性質から、それらのリスクが人間よりも増幅される可能性がある。この点から、AI について特別な法規制を課することが正当化されるかもしれない (3.2)。
- ②①とは別の観点として、AI は人間の場合よりも、法的規制によって上記のリスクを除去または軽減することがより容易にできるかもしれない。このような、法規制の必要性というよりは実現可能性の違いを反映して、AI について人間にない法規制を課することが正当

化される可能性がある (3.3)。

- ③現行法制は、自律的な行動をとれるのは人間だけであるという前提のもとに、人間だけを規律対象にしているルールを置いている場合がある。そのため、same rule 原則を実現しようとする、人間だけでなく AI も規律対象とするために新たな法規制が必要とされることがありうる。とりわけ、「AI の過失に基づいて AI 利用者に損害賠償責任を課す」というルールを実現するためには、新たな立法が必要とされるかもしれない (5.5)。

本稿の考察は、限られた文献資料の調査に基づく準備的なものにすぎない。今後、金融サービスにおける AI 利用についてさらに研究を重ね、適切な規制のあり方について考察を深めていきたい。

引用文献

- AI 戦略会議・AI 制度研究会 (2025)「中間とりまとめ」(2025年2月4日) (https://www8.cao.go.jp/cstp/ai/interim_report.pdf).
- Business Insider Japan (2024)「生成 AI を業務現場で活用「すでに収益改善が始まっている」注目 の 9 銘柄。バンク・オブ・アメリカ分析」Business Insider Japan PREMIUM 2024年08月21日 [日経テレコン21で閲覧]
- 有坂陽子 (2024a)「EU AI 法概説 (第 4 回) プロバイダー等の義務 (16条~22条)」NBL1273号 89-97 頁.
- 有坂陽子 (2024b)「EU AI 法概説 (第 5 回) 輸入事業者・販売事業者等の義務 (23条~39条)」NBL 1274号76-83頁.
- 飯田秀総 (2023)『金融商品取引法』新世社
- 岡田拓郎・佐藤市雄・山本俊之 (2024)「金融 業界・

- 金融機関における生成 AI の活用：パネルディスカッション」金融法務事情2231号23-34頁。
- 窪田充見（2018）『不法行為法（第2版）』有斐閣。
- 須賀千鶴（2024）「生成 AI で変わる日本社会：基調講演」金融法務事情 2231号 8-22頁。
- 鈴木翔平・森田祐行（2024）「欧州 AI Act の概要」（2024.05.27）（TMI 総合法律事務所ブログ）（<https://www.tmi.gr.jp/eyes/blog/2024/15787.html>）。
- 総務省・経済産業省（2024）『AI 事業者ガイドライン（1.01版）』（令和 6 年11月22日）（https://www.meti.go.jp/shingikai/mono_info_service/ai-shakai_jisso/pdf/20241122_1.pdf）。
- 田中亘（2022）「リテール分野における技術革新の社会的意義および規制のあり方：インターネット取引とロボアドバイザーを中心に」証券経済研究119号45-64頁。
- 羽深広樹・古川直祐（2024）「EU AI 法概説（第2回）ハイリスク AI 分類とハイリスク AI に対する要求事項（6 条～10条）」NBL1270号95-102頁。
- 樋口範雄・佐久間毅編（2014）『現代の代理法：アメリカと日本』弘文堂。
- 平井宜雄（1992）『債権各論Ⅱ不法行為』弘文堂。
- 古川直裕（2024）「EU AI 法概説（第1回）総則と禁止 AI（1 条～5 条）」NBL1269号40-46頁。
- 古川直裕・羽深広樹（2024）「EU AI 法概説（第3回）ハイリスク AI に対する要求事項（11条～15条）」NBL1271号101-108頁。
- 内閣府科学技術・イノベーション推進事務局（2024）「AI 政策の現状と制度課題」AI 戦略会議（第11回）・AI 制度研究会（第1回）（令和 6 年 8 月 2 日開催）配付資料 1（https://www8.cao.go.jp/cstp/ai/ai_senryaku/11kai/shiryol.pdf）。
- 日経クロステック（2024）「りそな HD がマネロン対策に AI 活用、『FATFS 次審査』もにらむ」2024年 6 月27日 [日経テレコン21 で閲覧]
- 日本経済新聞（2024/6/4）「監査法人、生成 AI で業務効率化あずさは22万時間削減」日本経済新聞電子版2024年 6 月 4 日
- 日本経済新聞（2024/7/19）「野村、広告審査に生成 AI みずほは月2000時間業務削減」日本経済新聞電子版2024年 7 月19日
- 日本経済新聞（2024/10/16）「セブン銀行、カードローンと信審査に『7iD』の購買データを活用する取り組みを期間限定で実施」日本経済新聞電子版2024年10月16日
- 橋本佳幸（2016）「『法人自体の不法行為』の再検討－総体としての事業組織に関する責任規律をめぐって」論究ジュリスト16号50-60頁。
- 吉永京子（2024）「EU AI 法概説（第6回）ハイリスク AI の標準、適合性評価、証明書、登録（40 条～49条）、特定 AI システムのプロバイダーおよびデプロイヤーに対する透明性義務（50条）」NBL（1275）：67-78。
- Ayres, I. and Balkin, J. M. (2024), "The Law of AI is the Law of Risky Agents Without Intentions," *University of Chicago Law Review Online*, *1-11 (11/27/2024), https://lawreview.uchicago.edu/sites/default/files/2024-11/Ayres_Balkin_Law%20of%20Risky%20Agents.pdf.
- Bagattini, G., Benetti, Z., and Guagliano, C. (2023), "Artificial Intelligence in EU Securities Markets," *ESMA Report on Trends, Risks and Vulnerabilities Risk Analysis* (1 February 2023), available at https://www.esma.europa.eu/sites/default/files/library/ESMA50-164-6247-AI_in_securities_markets.pdf.
- Baker, T., and Dellaert, B. (2018), "Regulating Robo Advice across the Financial Services Industry," *Iowa Law Review* 103(2) : 713-750.
- Brainard, L. (2021), "Supporting Responsible Use of AI and Equitable Outcomes in Financial Services," at the AI Academic Symposium hosted by the Board of Governors of the Federal Reserve System, Washington, D.C., Virtual Event (January 12, 2021), available at <https://www.federalreserve.gov/newsevents/speech/brainard20210112a.htm>.

- Buiten, de Streel, Alexandre, d.S., and Peitz, M. (2023), "The Law and Economics of AI Liability," *Computer Law & Security Review* 48: 105794.
- Calabresi, G. (1970), *The Costs of Accidents: A Legal and Economic Analysis*, Yale University Press.
- Calabresi, G., and D. Melamed, D. (1972), "Property Rules, Liability Rules and Inalienability: One View of the Cathedral," *Harvard Law Review* 85: 1089-1128.
- Cox, J.D., and Eisenberg, M.A. (2019), *Business Organizations : Cases and Materials* (12th ed.), Foundation Press.
- Dou, W. W., Itay, G., and Ji, Y. (2024), "AI-Powered Trading, Algorithmic Collusion, and Price Efficiency," Jacobs Levy Equity Management Center for Quantitative Financial Research Paper , The Wharton School Research Paper (May 30, 2024), Available at SSRN: <https://ssrn.com/abstract=4452704> or <http://dx.doi.org/10.2139/ssrn.4452704>.
- European Commission (2024), "AI Act," <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>.
- Hammer, S. (2024), "Navigating the Neural Network: Artificial Intelligence in Finance and Recalibration of the Regulatory Framework," Penn Program on Regulation Research Paper Series No.24-01 (October 2024), <https://pennreg.org/wp-content/uploads/2024/10/Hammer-Navigating-the-Neural-Network.pdf> [forthcoming: in P Hacker, P., Engel, A., Hammer, S., and Mittelstadt, B. (eds), *The Oxford Handbook on the Foundations and Regulation of Generative AI* (Oxford University Press)]
- Kleinberg, J., Ludwig, J., Mullainathan, S., and Stein, C.R. (2018), "Discrimination in the Age of Algorithms," *Journal of Legal Analysis* 10(2) : 1-62.
- Kretschmer, M., Kretschmer, T., Peukert, A., and Peukert, C. (2023), "The Risks of Risk-based AI Regulation: Taking Liability Seriously" (September 27, 2023). Available at SSRN: <https://ssrn.com/abstract=4622405> or <http://dx.doi.org/10.2139/ssrn.4622405>
- Shavell, S. (2004), *Foundations of Economic Analysis of Law*, Belknap Press of Harvard University Press [田中亘・飯田高訳 (2010) 『法と経済学』日本経済新聞出版社] .
- Sykes, A. O. (1984). "The Economics of Vicarious Liability," *Yale Law Journal* 93(7) : 1231-1282.
- Wachter, S. (2024), "Limitations and Loopholes in the EU AI Act and AI Liability Directives: What This Means for the European Union, the United States, and Beyond," *Yale Journal of Law & Technology* 26(3) : 671-718.

(東京大学社会科学研究所教授)