

操作と言論

——AIによる「言論」の憲法的価値をめぐる序論的考察——

山 本 龍 彦

- 一 はじめに
- 二 ストラウスの説得性原則と事実に関する虚偽言明
- 三 AI推薦システム——AIによる「言論」？
- 四 終わりに

一 はじめに

イギリスのジャーナリストであるケリー (Jenima Kelly) は、ソーシャルメディアを中心とする現在の情報空間を、市民が意見交換を行うアゴラではなく、観客がよく暴徒化した「ディオニュソス劇場」に擬えた⁽¹⁾。このアテナイの劇場は、葡萄酒と酩酊の神として知られるディオニソス (あるいは Bacchus) を祝す祭の会場であり、ケリーの比喩は、現在の情報空間が、AI生成物、偽・誤情報、誹謗中傷などが入り乱れ、それに酔っ払った者たちが日々乱痴気騒ぎをする混沌たる空間と化したことを悲観したものである。

確かに、アテンション・エコノミーと呼ばれるソーシャルメディアプラットフォームのビジネスモデルの下、我々のアテンションを直感的・本能的に引くような刺激的「言論」が過度に拡散・増幅するようになり、我々の理性や自律性に訴える「言論 (speech)」がその影響力を急速に失いつつあるように思える。けれども多くの憲法研究者は、この感覚を共有し、状況を憂いつつも、情報空間の交通整理に向けた積極的な議論には戸惑いを覚えるだろう。かような議論は、あらゆる情報「出力」のうち憲法上保護されるものとそうでないものを区別し、保護される「出力」のなかでもその価値に優劣をつけることに連なるからである。やがてそれは、自らの都合で情報空間に介入しようとする国家権力の免罪符にもなる。無論、この懸念は憲法字にとって正当なものである。しかし、表現の自由を殊更に重視してきたアメリカでも、かかるリスクを過度に意識した議論の萎縮や、そこからくる沈黙が、必ずしも誠実な学問的態度ではないと主張する考えが有力化している。それらは、偽・誤情報や AI による「言論」(情報出力) が遍在化した情報空間の大きな変容を前に、憲法上保護される言論、あるいは特に保護される言論を画定する必要性——「言論」の差別化の必要性——を強調するのである⁽³⁾。

本稿は、「操作 (manipulation)」や「自律性 (individual autonomy)」を鍵概念に憲法上保護される「言論」の意味を探究するストラウス (David A. Strauss)⁽⁴⁾ やピーターズ (Alec Peters)⁽⁵⁾ の議論を参照しながら、先述の「混沌」の主な原因ともなっている①事実に関する虚偽言明 (false statement of fact) と、②コンテンツの AI 推薦システム (機械学習を用いた AI が、動画等をユーザーにレコメンドするシステム。以下、「AI 推薦システム」という。) の憲法上の位置付けを試論的に考察するものである。以下、まずは、合衆国憲法修正一条の解釈としてストラウスが提示する「説得性原則 (persuasion principle)」を概観しておきたい。

二 ストラウスの説得性原則と事実に関する虚偽言明

1 説得性原則

アメリカ憲法学の泰斗であるストラウスは、表現の自由に関する法理論の中核にあり、連邦最高裁判所も基本的に支持してきた考えとして「説得性原則」を挙げる。⁽⁶⁾ 同原則は、「政府は、ある言論が、政府が有害だとみなす何かを行うよう人々を説得する可能性に基づいて、言論を抑制してはならない」、⁽⁷⁾ あるいは、「言論の説得効果 (persuasive effects) から生じる有害な結果は、言論を制限する正当化根拠の一部を構成しえない」⁽⁸⁾ (強調はいずれも筆者) という考えである。

これは、第一義的には、言論が説得力を持ちすぎること理由に——たとえ説得の力により有害な結果を惹起しうる行為に人を走らせるとしても——政府は言論を規制できないという、政府に対する憲法上の統制規範となる。しかし、「説得」という「理性に訴えるプロセス (process of appealing to reason)」を重視するこの議論の要諦は、「説得」を人間の理性に訴えない単なる「誘導 (induce)」と区別し、後者を表現の自由の保護領域から外そうとする点、また後者に対する政府規制を積極的に許容する点にある。⁽⁹⁾ 彼の議論は、「説得」という技法——理性に訴えて人の行動や信念を変容させる力——と、非合理的な手段で行動を誘発する「誘導」という技法——理性を迂回して行動させる力——との間に線を引こうとするものである。ストラウスが「誘導」の典型として挙げたのが、①事実に関する虚偽言明と、②軽率な行動 (inconsidered action) を惹起する言明 (聞き手が、当該言明とありうる反論について考える前に行動を引き出すとする言論 (speech that seeks to elicit action before the hearer has thought about the speech and possible answering arguments)) (強調は筆者)⁽¹⁰⁾ であった。

かつて連邦最高裁も「修正一条の核心には、表現し、吟味し、受容するに値する思想および信念を、各人が自

分自身で決定すべきとする原則が存在する⁽¹¹⁾」と述べたように、アメリカでは修正一条を支える価値の一つに個人の自律性があると解するのが一般的である。⁽¹²⁾ ストラウスは、説得性原則をこの自律性によって根拠付ける。「説得」に憲法上価値があるのは、それが個人の自律性を尊重しているからであり、「誘導」に憲法上価値がないのは、それが自律性を否定しているからだ、というわけである、彼は、このことを説明するうえで、「操作的な嘘 (manipulative lie)」が他者に与える影響に着目する。

「A が B に嘘をつき、A が望むような行動を B にとらせるとする。例えば、A が実際には怠けているだけに、病気ののだと偽り、A に便宜を図るよう B を誘導する場合を考えてみよ。このような欺きは、(特段の事情がない限り) 道徳的に誤りである。それは、B をこのように取り扱うことが無礼 (disrespectful) だからである。B が、A によって自らを人間以下の存在として——A の意志の単なる道具として——扱われたと感じるのは正当である。ここにおいて、A は B を操作した。B を人間として取り扱う代わりに道具として利用したのである」⁽¹³⁾。

ストラウスのこの議論で興味深いのは、ここで A がついた「操作的な嘘」を、身体に働きかける「あからさまな強制 (outright coercion)」よりも悪質だと考えるところにある。彼は、操作的な嘘も、強制も、「被害者に対して支配力行使する手段」として被害者の自律性を蔑ろにする点で共通しているが、「被害者に、自分が支配され、操作されていることにすら気づかせない」前者は、支配への認知・認識を与える後者よりも「陰湿」だとい⁽¹⁴⁾のである。ストラウスは、あからさまな強制では未だ「被害者の精神は自由」であり、強制により作出されるこの身体的隷属状態の方が、操作的な嘘により作出される精神的隷属状態よりも(自律性を基準にすれば)まだましだともいう。繰り返すが、ストラウスにとって、被害者に、自身の目的ではなく発話者の目的を秘密裏に

追求させようとする操作的な嘘は、「人生の計画を決定し、自らの目標を定める」自律の能力⁽¹⁵⁾、すなわち、「善の観念を形成し、合理的に熟慮し、自らの目標と整合的に行動する能力」⁽¹⁶⁾に対する「特に深刻な侵害」⁽¹⁷⁾とみなされるのである。

以上のように、自律性の価値を特に重視するストラウスの説得性原則は、「説得」と「誘導」（または操作）を憲法上区別し、政府が説得の力を恐れて言論を規制しようとすることを厳しく制限する一方、「理性性に訴えるプロセス」を迂回しようとする誘導・操作の政府規制を積極的に許容するものと考えられる。

2 事実に関する虚偽言明

先述のとおり、いわゆる偽・誤情報の拡散・増幅は、現在の情報空間の深刻な課題の一つである。ストラウスの説得性原則は、事実に関する虚偽言明の憲法上の位置付けについて重要な示唆を与えうる。

ストラウスは、端的に「事実に関する虚偽言明は理性性に訴えるものではないため、その使用は説得を構成せず、したがって説得性原則により保護されない」⁽¹⁸⁾と述べる。名誉毀損的表現に関する一九七四年の *Gertz v. Robert Welch, Inc.* 事件判決⁽¹⁹⁾で、連邦最高裁が、「事実に関する虚偽言明には憲法上の価値はない。意図的な嘘も、不注意な誤りも、公共問題に関する『抑制されない、活発で、広く開かれた』議論⁽²⁰⁾という社会的利益を促進しない」と述べたのも、ストラウスによれば同趣旨とされる。彼の議論を継承・発展させるピーターズも、*Gertz* 判決は「『事実に関する』虚偽言明は、現実には相反する行動様式をとるよう聞き手を操作すること (manipulating) で聞き手の自律性に干渉し、以て合理的な自己決定を阻害する」ことを問題視したものであって、「『抑制されない』議論の価値を、實際上その抑制されない性質そのものに見ていたわけではない」と指摘している⁽²¹⁾。

周知のとおり、日本の最高裁は、名誉毀損罪に関する事案で、「刑法二三〇条ノ二第一項にいう事実が真実で

あることの証明がない場合でも、行為者がその事実を真実であると誤信し、その誤信したことについて、確実な資料、根拠に照らし相当の理由があるときは、犯罪の故意がなく、名誉毀損の罪は成立しない」(強調は筆者)⁽²²⁾と述べてきた(民事上の不法行為についても同旨)。この一節からは、誤信することに相当な理由さえあれば、事実に関する虚偽言明も憲法上の保護を受けるかのように見える。しかし、この「相当性の法理」も、ストラウスの指摘のように、事実に関する虚偽言明それ自体に憲法上の価値を認めているわけではないと考えるのが自然である。ここでは、理性に訴える真実の言明を萎縮させないために、事実に関する一定の虚偽言明を政策的・予防的に保護していると解されるからである⁽²³⁾。加えてストラウスの議論からは、「確実な資料、根拠に照らし」て真実であると誤信した者——誠実に間違えた者——には、受け手を操作する意図はなく、道徳的にも厳しく非難される謂れないことになる。むしろ彼(女)らは、受け手の理性に訴えようと、ぎりぎりまで真実を探索しようと試みた可能性がある。こうみると、日本の判例における相当性の法理も、ストラウスの説得性原則から説明可能であるように思われる。

以上、ごく簡単にストラウスの説得性原則をみたが、偽・誤情報が蔓延る現在の情報空間を前に、事実に関する虚偽言明が憲法上の価値を本来的に持つわけではないというその論旨を改めて確認しておくことの意義は小さくないように思われる。

三 AI 推薦システム——AI による「言論」?

1 問題の所在——Moody 判決を踏まえて

冒頭でも触れたように、近年は、アテンションやエンゲージメントの獲得を最大の使命とした AI 推薦システ

ムが、刺激的な偽・誤情報や誹謗中傷を拡散・増幅させたり、フィルターバブルやエコーチェンバーを生み出してユーザーの政治的傾向を極端化させ、分断を増長させたりしているなどと指摘される。加えて最近では、エンゲージメントの最大化を目的とした同システムによる動画等の（ユーザーの嗜好等に合わせた）フィードにより、SNSへの依存・中毒や、メンタルヘルスの悪化などを引き起こす可能性も報告されている⁽²⁴⁾。

周知のように、欧米では、このような課題を踏まえて、AI推薦システムに関する透明性確保や情報摂取に関するユーザーの主体性確保のための種々の規制が行われつつあるが、同時に、AI推薦システムによる情報の選別・提供行為は、プラットフォーム事業者の編集権行使に当たり、こうした規制が事業者の表現の自由を侵害するのではないかとの議論が活発化している。特にアメリカでは、二〇二四年七月、ソーシャルメディアプラットフォームのコンテンツモデレーションが恣意的である（実際には保守派の見解を不当に扱っている）として、その方法に介入する（特定の理由・見解に基づくコンテンツ表示順位の変更や削除等を禁止する）フロリダ州法およびテキサス州法⁽²⁶⁾の合憲性を争った *Moody v. NetChoice, LLC* 事件判決⁽²⁷⁾が下されたこともあって、憲法上大きな論争となっている。ここで本質的に問われているのは、AI推薦システムによる情報の選別・提供行為が修正一条によって保護されるのか、保護されるとしても、新聞社等の人間による編集と同等の保護を受けるか、という点である。

日本でも既に紹介されているように⁽²⁸⁾、ケーガン判事 (Kagan, J.) の手になる *Moody* 判決の多数意見は、「修正一条により保護される」表現活動には、他者が元々創造した言論を整理した編集物 (curated compilation) を提供することも含まれる⁽²⁹⁾としたうえで、（人間による）「編集者 (editors) ……と同様に、主要ソーシャルメディアプラットフォームは、フィードを整理 (curating) する際に、『多様な声』を組み合わせ、独自の表現的提供者物 (expressive offering) を創造する事業を行っている」⁽³⁰⁾——「編集のコントロールと判断を行う」のは、「新聞社も

プラットフォームも変わらない⁽³¹⁾——と述べており、ここだけを見れば、既に連邦最高裁は、先述の問いについて最終的な答えを与えたようにも思える。⁽³²⁾しかし、この多数意見は、実際には原審が文面上違憲の主張を審査するために必要な検討を行っていないとのきわめて形式的な理由で差し戻しており、当該差し戻しの理由となった形式的問題に関する判示以外の部分は「全て拘束力のない傍論 (nonbinding dicta) である」⁽³³⁾とも評されている。そして、「本質的に表現的である場合にのみ修正一条の保護を受ける」としたうえ、「典型的なソーシャルメディアのフィードでさえ、この判断『本質的に表現的であるか』は見た目以上に複雑である」⁽³⁴⁾とし、①「アルゴリズムが単に各ユーザーが好むと判断した内容……を自動的に提供する場合」や、②プラットフォーム所有者が「憎悪的」コンテンツの削除を AI に指示はするが、何が「憎悪的」で何を削除すべきかの具体的判断を AI の大規模言語モデルに依拠するような場合には、「修正一条の権利を有する人間が……『本質的に表現的な選択』を行った」とは言えない、可能性があるとするバレット判事 (Barrett, J.) の同意意見、さらには、人間による編集と AI 推薦システムによる情報選別・提供との同等性を根本的に疑問視するアリート判事 (Alito, J.) の結果同意意見があるように、AI 推薦システム等の技術的特徴を軽視し、安易に人間による編集との同等性を認めたケーガン判事の多数意見には連邦最高裁のなかでも厳しい批判が寄せられている。⁽³⁶⁾したがって、Moody 判決が出された以降も、先述の問いはなお開かれていてと解するのが相当である。

そこで以下、ストラスやピーターズの議論などを踏まえて、かかる論点について若干の検討を加えてみたい。

2 AI 推薦システムの実態

まずは、ソーシャルメディアプラットフォームによる AI 推薦システムの実態について概観しておこう。ソーシャルメディアプラットフォームが二〇一〇年代半ばから本格的に導入したのが、機械学習モデルの AI である。

例えば「Twitter (現 X) は、二〇一六年にタイムラインの表示順を時系列順から AI によるランキング順に変更しているが、⁽³⁷⁾「今では全ての主要プラットフォームが、表示するコンテンツを選択するために何らかの深層学習モデルに依存している」⁽³⁸⁾という。AI 推薦システムとは、単純に言えば、ウェブ上の行動履歴等を含む膨大な個人データ (シグナル) に基づき、どのコンテンツが当該ユーザーにとってどの程度関連性があるかを判断し、当該スコアに応じて自動的にコンテンツを当該ユーザーにフィードする AI システムであるが、重要なのは、それが機械学習モデルに依拠している点である。ここでは、その特徴に関するオースティンとレビイ (Austin & Max Levy) の整理を紹介しておきたい。

「1. 従来のコードは、プログラマーが記述したルールのみを実行する。従来のプログラミングでは、プログラマーがアルゴリズムの従うべきルールを決定し、当該ルールをコードに翻訳する。プログラマーは、コードが実行される際、記述されたルールが実行されているかを確実に理解できる。

2. 機械学習アルゴリズムは、予測を行うための独自のルールを記述する。機械学習アルゴリズムの唯一の機能は予測を行うことである。機械学習アルゴリズムは、複雑な数学を用いて自動的に実行される独自のルールを記述することで予測を行っている。このプロセスにおけるプログラマーの主たる役割は、アルゴリズムに何を予測すべきかを指示し……、考慮すべき変数を決定し……、モデルに投入されるデータの信頼性を確保し、モデルの自動的な数学的处理を促進する特定の「取っ手 (Knobs)」を調整することである。この作業は、……困難で複雑である。しかし、作業がどれほど困難なものであろうと、プログラマーはアルゴリズムの基本的任務、すなわち、アルゴリズムの予測を形作るロジックを決定することには参加しない。つまり、機械は予測において各変数がどれほど重要であるべきか、また各変数が予測にどのように影響するかを独立して決定するのである。

3. プログラマーは機械学習アルゴリズムが特定の予測を行う理由を説明できない。ほとんどの場合、アルゴリズムは

「ブラックボックス」であり、そのルールは人間が理解できないほど複雑である。また、ルールの集合が膨大すぎて、人間が実際上処理できない場合もある。この本質的な不透明性と、プログラマーがルールを記述していないという事実は、プログラマーが通常その動作原理を理解していないことを意味する。

4. 機械学習アルゴリズムの出力には、プログラマーに帰すことのできない誤りが常に存在する。機械学習アルゴリズムは確率を計算して予測を行う。確率の本質上、アルゴリズムは結果を100%正確に予測できず、結果として、少なくとも一部の予測は必然的に誤りうる。確かに、プログラマーは従来型のコードを用いて予測を行うこともできる——これもまた必然的に誤りを生む予測である。しかし従来型のコードでは、プログラマーがルールを記述するため、予測が誤っていてもアルゴリズムはプログラマーが書いたルールを忠実に実行したに過ぎず、誤りを直接プログラマーに帰すことができる。機械学習では、プログラマーはアルゴリズムのルールを記述せず、完全に説明することできない⁽³⁹⁾。

以上のように、推薦システムにおける機械学習モデルの導入が意味するのは、当然ながら、コンテンツの順位付け等に対して人間の関与が限定的になるということである。また、アテンション・エコノミーをビジネスモデルとするソーシャルメディアプラットフォームの場合、このAIが、必然的にユーザーのエンゲージメントの最大化を目的として与えられていることにも注意が必要である。AIは、かかる指示を受け、自ら推薦したコンテンツに対するユーザーの反応を細かく観察し、その成果に応じてまた次の推薦を調整するという、終わらなき調整ゲームを自ら独立して行い、「人間の行動のための最適化ループ (optimization loop for human behavior)」に作り出すと指摘される⁽⁴⁰⁾。実際、この「ループ」の効果は非常に高く、機械学習によるAI推薦システムの導入により、ユーザーのエンゲージメント(動画に対する視聴時間)は三〇〜五〇%も増加したという⁽⁴¹⁾。

3 検討

(1) AI推薦システムの操作性

では、このようなAI推薦システムによる情報出力は、修正一条により保護されるのか、保護されずともどの程度の保護を受けるのか。ピーターズは、ストラウスの説得性原則や自律性に関する議論を踏まえて、それが人間による編集とは根本的に異なると主張している。彼が両者の質的相違をいう際にまず注目するのは、AI推薦システムによる情報出力の「操作性」である。ピーターズによれば、「AI推薦システムは人間の編集者よりも聞き手を操作でき」、⁽⁴²⁾「戦略的目標達成のために人間を操作する能力を一貫して占めている」というのである。もとより操作は、理性的プロセスに訴える「説得」とは異なり、聞き手の自律性を害するががゆえに、もし彼の理解が正しければ、それは修正一条の価値に反するものということになる。以下、他の論者の指摘も適宜参照しながら、ピーターズの議論をみていきたい。

まずピーターズは、ササーら (Daniel Susser et al.) の研究に依拠しつつ、操作性を、「意思決定の脆弱性を標的にし、それを利用することで、〔①〕意図的かつ〔②〕秘密裏に〔誰かの〕意思決定に影響を与えること」と⁽⁴⁴⁾定義し、AI推薦システムによる出力は①および②を満たすという。①(意図)についてピーターズは、コンテンツに対するユーザーのエンゲージメントを増大させれば報酬が得られるよう設定されたAI推薦システムは、「人間の振る舞いを変化させることで得られる報酬を常に追求しようとするため、常に「人間の意思決定に」影響を与える意図をもっている」と⁽⁴⁵⁾言いうると指摘する。また②(秘密性)についてピーターズは、推薦システムの挙動を理解するために必要な情報の多くがプラットフォームにより公開されていないこと、機械学習モデルは「ブラックボックス」とも言われるように、その意思決定過程に関する人間の理解を一般に拒むことなどを挙げ、ユーザーはAI推薦システムによりどのような影響を受けているかを具体的に知ることができないがゆえに、操

作性の定義における「秘密性」要件も満たすと述べている。⁽⁴⁶⁾

仮にこのように、AI推薦システムによる出力が「操作性」の前記定義の①・②を満たすとすれば、重要になるのは、当該出力が人間の意思決定に実際に影響を与えているかであろう。ピーターズがここで強調しているのは、当該出力の「個別化 (personalization)」である。人間の編集者も、確かに情報の受け手に影響を与えることを意図してコンテンツを選ぼうとするかもしれないが、この編集者は、特定個人の属性や脆弱性を詳細に把握する能力をもたず、不特定多数者やせいぜい特定集団の一般的傾向を利用することしかできない。他方でAI推薦システムは、大量の個人データを通じたプロファイリングにより、特定個人の属性や脆弱性を把握し、コンテンツを個別化してフォードできるといふ⁽⁴⁷⁾。またこのことは、特定個人に既に備わっている脆弱性を利用できるだけでなく、脆弱性を新たに作出することも可能であるとされる。この点は、既にカロ (Ryan Calo) が、プラットフォームは個人データやAIを用いて、「かつて想像もできなかった規模で消費者がどのように合理的意思決定から逸脱するのを見通し、それにつけ込むことができる」とし、いまやカモを待つのではなく、自らカモを作り出せるようになったと指摘していることとも符合する。加えてピーターズは、AI推薦システムはユーザーの選好や思考そのものを変化させることも可能であると指摘する。彼は、「アルゴリズムは、将来、人々をより影響を受けやすく、より高いレートでマネタイズできるようにすべく、その行動を変えるように人々を操作することを学習してきた⁽⁴⁹⁾」とするラッセル (Stuart Russell) の研究などを参照し、「AIは、エンゲージメントを最大化する最善の方法が、コンテンツを我々の選好に合わせるのではなく、AIが提示するコンテンツに合わせて我々の選好の方を変えることなのだと考えるようになってきている⁽⁵⁰⁾」と主張するのである。

マークス (Mason Marks) も、プラットフォームの操作的権力に着目した論文「生の支配力 (biosupremacy)」⁽⁵¹⁾において同様の指摘を行っている。マークスは、プラットフォームはユーザーのデータを収集・分析してその選

考や意識を理解しようとする「感知ネット (sensing net)」と、ユーザーの行動に影響を与えようとする「コントロールネット (control net)」の両方を持ち、両者を絶えず循環させることで「プラットフォーム」企業の定立した規範 (norms) に自らの行動を適合させるよう人々を誘導している」⁽⁵²⁾と指摘している。マークスの言葉を借りれば、「感知ネットが行動の推論や予測というかたちで知識を生成し、然る後に、コントロールネットがこの知識を利用することで、集団や個人を操作する」⁽⁵³⁾。

こうした「操作」が実際に効果を発揮していることは、機械学習モデルの導入によりユーザーのエンゲージメントが三〇〜五〇%も増加したというプラットフォーム自身の告白からも、また近年社会問題になっているソーシャルメディア依存の深刻化などからも明らかであるように思われる。なお、オーストラリア国立科学機関は、その研究を通じて、「人間の意思決定プロセスにおける脆弱性」につけ込むようデザインされたAIが、約七〇%の確率で人間を意図した選択に誘導できることを明らかにしている。⁽⁵⁴⁾ この研究の責任者によると、かかる結果は「機械が人間との相互作用を通じて人間の意思決定を誘導することを学習しうることを示している」⁽⁵⁵⁾という。

(2) 受け手の自律性

このようなAI推薦システムによる情報出力が憲法上の価値を有しない理由について、先述のとおりピーターズは、第一に、事実に関する虚偽言明と同様、それが「標的になる者に気づかれることなくその隠れた脆弱性につけ込み、理性的な意思決定を迂回して行動に影響を与えようとする」⁽⁵⁶⁾点を挙げる。つまり、個人の自律性を害するために修正一条の価値をもたないというのである。ここで彼は、「AIは自らの内部目標に奉仕するようにデザインされた個別化した情報のループへとユーザーを誘導することで、彼らが自由に理性を働かせて判断を下し、独自の結論に至ることを妨げる」と説明したうえで、同システムによって受け手の「自律性 (autonomy)」が自律性 (automaticity) —— 意識的な注意を必要としない習慣的で予測認知的 (precognitive) な振る舞い —— に置

き換えられていく⁽⁵⁷⁾とするコーエン (Julie Cohen) の言葉を引いている。

この「自動性」は、ストラウスが「誘導」の典型とした「軽率な行動 (ill-considered action)」を惹起する言明」と密接に関連しているため、以下、右言明について一瞥しておきたい。

ストラウスは、この「軽率な行動を惹起する言明」が憲法上の価値を有しないことは、ブランドイス判事 (Brandeis, J.) が定式化した「明白かつ現在の危険」でも示唆されているという。その定式は以下のとおりである。

「邪悪な助言 (counsels) に対する適切な救済策は、善意の助言である。……言論から生じる危険は、その悪影響が差し迫り、十分な議論の機会を得る前に (Before there is opportunity for full discussion) 現実化する恐れがある場合を除き、明白かつ現在の危険と見なすことはできない。議論を通じて、虚偽 (の思想) や誤謬を暴き、教育のプロセスを通じて邪悪さを回避する時間があるならば、適用すべき救済策はさらなる言論なのであって、強制された沈黙ではない⁽⁵⁸⁾」(強調は筆者。「」内はストラウスの脚注を踏まえて筆者が補完したもの)。

ストラウスによれば、「明白かつ現在の危険」に関するこのブランドイスの定式は、理性に基づく「十分な時間」を与えずに直接行動を惹起するような「言論」は規制しうることを示しているという。右定式は、「邪悪な目的を『十分な議論の機会を得る前に』実現させることで軽率な行動を惹起する言論に対しては、これを保護の対象としないことを明らかにしている⁽⁶⁰⁾」というのである。かような理解に基づくならば、AI 推薦システムの出力が、コーエンのいう自動的あるいは反射的なユーザーの行動を絶えず促すのであれば、それはブランドイスの「明白かつ現在の危険」の定式からみても憲法上の価値を有しないことになる。

実際、ソーシャルメディアが採用する情報技術は、「脳活動の」潜在的な部分に意図的に働きかけようとして「おり、「思考・判断を司る」大脳新皮質ではなくて「情動を司る」大脳辺縁系、小脳、もつといえは脳幹部、脊髄に近い橋など、いわゆる「爬虫類の脳」といわれているようなところに特化して刺激が与えられて、それに反応しているということが起きている」（強調は筆者）⁽⁶¹⁾という認知神経学者の指摘や、人間の思考のあり方に関する心理学者カーネマン（Daniel Kahneman）の二重過程理論（Dual Process theory）⁽⁶²⁾を踏まえて、AI推薦システム等の技術は、「システム1」という直感的で処理速度の速い思考モードにトリガーをかけ、「システム2」という論理的・内省的で処理速度の遅い——認知的リソースを多く消費する——思考モードを抑え込むように発展してきているとの指摘があることを踏まえれば、AI推薦システムの出力は、「自律性」ではなく「自動性」と結びつくもので、ストラウスによって憲法上の価値を否定される「軽率な行動を惹起する言明」とも関連しているように思われる。

なおピーターズは、こうした議論に対してプラットフォームは、AI推薦システムはユーザーが欲するもの、ユーザーにとって価値のあるものを提供しているのであって、個人の自律性を害するものではないと反論するかもしれないという。ピーターズは、このようなプラットフォームの主張に再反論を加えるにあたり、「人間＝餌」の関係に関する以下の例を挙げている。

「人間が餌を使って犬を訓練する場合、犬は、報酬を得るために喜んで自らの振る舞いを変える。確かに、餌は犬が最も欲し、必要とするものである。しかし、餌が犬にとって価値があるという事実は、人間が「犬に対して」支配権をもち、犬の欲求や必要性を利用して自らの目的を達成しようとしているという事実を何ら変えるものではない」⁽⁶⁴⁾。

ピーターズによれば、AI 推薦システムが出力する「餌」を我々が「欲する」としても、それがプラットフォームによる「操作」であり、我々の自律性ないしは尊厳を蔑ろにしているという事実は変わらない。

またピーターズは、AI 出力が憲法上の価値を有しない第二の理由として、それが理性的な思考（システム2）を妨げることで、表現の自由のもう一つの価値である民主的自己統治の前提をも掘り崩すことを挙げている。かつてブランドイス判事は、建国者は「熟慮の力（*deliberative forces*）」が統治における恣意性を打ち破るものと信じた」（強調は筆者）⁽⁶⁵⁾と述べたが、仮にこうして、「理性と『熟慮の力』の行使こそが民主的自己決定の基礎」⁽⁶⁶⁾だと考えるならば、その行使を妨げる「操作」は、民主的自己統治や思想の自由市場の前提を揺るがすものとして憲法上の価値を否定されるだろう。リドスキー（Lyrisa Barnett Lidsky）も、同様に以下のように指摘している。

「民主主義は、市民に集団的運命を選択する権利を認める。その選択が強制または操作されたものであるならば、その選択は無意味である。実際、この基本原理に反する民主的な決定は、端的に正統性をもたない」⁽⁶⁷⁾

（3）送りの自律——言論の確実性原則

ピーターズは、人間による編集とAIによる「編集」との差異、また後者の憲法的価値を検討するうえで、（2）で述べた受け手の自律性への影響だけでなく、送りの自律性への影響をも考慮すべきと説く。

ピーターズは、その際にまず、連邦最高裁が、編集権について「新聞に掲載する素材の選択、紙面のサイズや内容に関する決定、公共問題や公職者の扱い……は編集上の統制と判断の行使を構成する」（強調は筆者）⁽⁶⁸⁾とし、「この重要なプロセスへの政府の規制」が修正一条違反となりうると述べたり、編集権を「自らの言論を統制する自律性（*autonomy to control one's own speech*）」と表現したりしてきたことなどを確認する。彼は、このような文言に着目することで、連邦最高裁が編集権を保護してきたのは、理性に従って自らのメッセージをコントロールする送りの自律性に価値を認めてきたからなのだという。⁽⁷¹⁾ そのうえで彼は、もしAI推薦システムの出力

を人間が実質的にコントロールできないのなら、かかる出力は——これまで連邦最高裁が編集権行使として保護してきたものと全く異なり——憲法上の価値を有しないと結論付けるのである。

こうした行論においてポイントになるのは、AIの出力に対して人間のコントロールが実際にどこまで及ぶかである。既に2において、オースティンとレヴィの見解を示したが、ピーターズも同様に、プラットフォームは「[AI推薦システムの]モデル構造に対して広範で、一般的なコントロール(broad, generalized control)を行使している」ものの、「各ユーザーにどのコンテンツが推薦されるのか、またAIがどのようにそのコンテンツを推薦するかを把握していない」とし、「AIが特定の状況下でどのような判断を下すのかを知らない」と述べる。⁽⁷²⁾

また、「機械学習システムは人間の完全なコントロール下にないという事実がますます明らかになっている」とする専門家の見解に依拠しつつ、より具体的に、「AIは人間の介入から独立して、特定した内部目標を追求すること、人間が定義したエンゲージメント目標を達成する方法を学習する」(強調は筆者)⁽⁷⁴⁾など指摘している。こうしてピーターズは、「編集判断にAIを用いる場合、プラットフォームはもはや編集判断に際して自らの理性や判断を用いているとは言え[ず]」、「人間の意思決定プロセスと比較して、AIへの委任は、プラットフォームが自らのメッセージに対して行使する自律性を低下させ、したがって修正一条による保護の価値を損なう」と述べるのである。⁽⁷⁵⁾

オースティンとレヴィの説く「言論の確実性(speech certainty)」原則も、同様の結論を導く。彼らのいう言論の確実性原則とは、歴史的に修正一条は、発話者がその発話の時点で、自らが発話する内容を確実に理解していた「言論」のみを保護してきたと考える見解である。⁽⁷⁶⁾彼らによれば、これを規範論的に支えているのも、やはり「自律性」だといふ。⁽⁷⁷⁾発話者が発話時点でその内容を確実に特定できる言論のみが保護されてきたのは、修正一条が、発言するかしないかの選択に対する個人の自律性を保障することを(少なくとも一つの)目的としてき

たからだ、⁽⁷⁸⁾というわけである。

オースティンとレヴィは、詳細な判例分析により言論の確実性原則を裏付け、また自律性原理から規範論的にもこれを基礎づけたうえで、当該原則を AI 推薦システムへと適用する。容易に想像しうるように、彼らによれば、「従来型アルゴリズムはプログラマーが記述したルールに忠実に従い、その出力を決定するのに対して、機械学習アルゴリズムは自らルールを記述する⁽⁷⁹⁾」がゆえに、従来型アルゴリズム (Facebook が機械学習を広く採用する前に使っていた二〇一〇年頃の EdgeRank など) の場合は別論、機械学習については、「発話者であるプログラマーは、出力が生成される時点でその内容を確認にすることができない⁽⁸⁰⁾」。あるいはまた、プログラマーは「アルゴリズムが公開するコンテンツが、プラットフォーム「自ら」が公開を意図したものと一致することを確信できない⁽⁸¹⁾」。かくしてオースティンとレヴィは、言論としての確実性を欠く AI 推薦システムの出力を、修正一条が保護してきた「言論」からは逸脱したものと結論付けるのである。

以上、受け手および送り手の自律性を修正一条の重要な価値とみる種々の議論を概観してきたが、それらによれば、AI 推薦システムによる出力は、少なくとも人間による編集と同等の憲法的価値をもたず、したがって、少なくともそれと同等の保護を受けないということになる。

4 示唆

これまで紹介してきたアメリカの議論は、日本の判例の読解にも一定の示唆を与える。周知のとおり日本の最高裁は、Google の検索結果の削除請求に関する事案で、⁽⁸²⁾自動化されたプログラムに基づく「検索結果の提供は検索事業者 [Google] 自身による表現行為という側面を有する」(強調は筆者) と述べるにとどめた。

最高裁が、かようにアルゴリズムによる「言論」と憲法二二条が保障する表現の自由との関係性をあえて曖昧

にしたのは、ピーターズらの議論を踏まえれば、アルゴリズムの動作に対する人間のコントロールの限定性や、人間による編集とアルゴリズムによる編集との本質的相違——「送り手の自律性」にかかわる問題——を意識していたようにも解されよう。もともと二〇一七年に出された本判決は、検索結果提供行為が違法となるか否かは、「当該「プライバシーに属する」事実を公表されない法的利益と当該URL等情報を検索結果として提供する理由に関する諸事情を比較衡量して判断すべきもので、その結果、当該事実を公表されない法的利益が優越することが明らかな場合には、検索事業者に対して、当該URL等情報を検索結果から削除することを求めることができる」（強調は筆者）と述べていた。無論、この「明らか」要件は、アルゴリズムを通じたプラットフォームの「言論」の重要性を示すものであり、本判決は、アルゴリズムによる「言論」の憲法上の位置付けに曖昧さを残しつつも、実質的にはその憲法的価値を承認していたと考えることができる。

しかし、その後最高裁は、Twitter（現X）の投稿記事の削除請求事件で、いま述べた「明らか」要件を外して削除の可否を判断した（二〇二二年）⁽⁸³⁾。もちろん、本件における「明らか」要件の不使用は、「Twitterというソーシャルメディアの機能・利用実態と、Googleが提供する検索機能——「情報流通の基盤として」⁽⁸⁵⁾」大きな役割——との相違によって説明可能であり、実際そのような理解が一般的であるが、本稿のこれまでの検討からは、これとは異なる見方もありうるように思われる。ストラウス、ピーターズ、さらにはオースティンとレヴィの議論を踏まえると、機械学習の精度が上がったTwitterのAI推薦システムは、人間によるコントロールが困難となることで送り手の自律性がこれまで以上に制約されるとともに（プログラマーは、出力の生成段階でその内容を確実に知ることができないために、言論確実性の原則からも「言論」該当性が疑われる）、受け手のエンゲージメント奪取の方法をAIが学習、洗練させることで、受け手に対する操作可能性がこれまで以上に高まる——受け手の自律性も害される——とも言える。こう考えれば、TwitterのAI推薦システムによる出力は、Google

の(当時の)検索プログラムによる出力と比べて、人間の統制が及びにくく、かつ操作可能性が高いがゆえに、「言論(speech)」としての憲法上の価値が低く、対立利益との衡量面で「明らか」要件が外されたと解する余地もあるだろう(同じ理は、アルゴリズムがさらに高度化・複雑化した対話型生成AIによる出力にも当てはまる)。

もっとも、これらの二つの最高裁判決・決定は、連邦最高裁のMoody判決と同様、AIによる「言論」の憲法上の位置付けについてあえて多くを語らず、技術の進歩に対して敬讓を払ったようにも見える。しかし、生成AIによる出力を含め、AIによる「言論」が我々に多くの便益をもたらす一方、様々な課題を惹起しているとすれば、「なぜ表現の自由が憲法上保障されなければならないのか」という目的や価値にまで遡り、その憲法上の位置付けを探究する営みは不可避であるように思われる。

四 終わりに

以上、本稿は、ストラウスの説得性原則とその後継の議論を参照しながら、事実に関する虚偽言明とAI推薦システムによる出力の憲法上の位置付けについて若干の考察を加えてきた。

現在の情報空間を「デユオニソス劇場」に喩えるケリーの見立てが情緒的に過ぎるとしても、インターネットの発展、さらにはAIの劇的進化が、かかる空間の性質を大きく変容させていることは疑いようがない。そこでは——これまでの人類が想定してこなかった情報出力を含む——雑多な「言論」が溢れる。このときに必要になるのは、表現の自由の価値や目的にまで遡って、憲法上本来的に保護される「言論」の意味を改めて検討してみることではないだろうか。⁽⁸⁶⁾

人間の自律性や理性を重視するストラウスの説得性原則やその後継理論にはもちろん批判もありうる。しかし、

表現の自由の原理論に正面から向き合い、そこから、あらゆる表現ないし情報出力が憲法上同等の価値をもつわけではないということを包み隠さず主張する学問的姿勢は、やはり注目に値しよう。「低価値表現」という言葉が憲法の教科書に堂々と書かれるように、そもそも憲法学は、あらゆる表現や情報出力が憲法二一条により等しく保障されるべきと解釈してきたわけではなかった。そうだとすれば、学問的に誠実なのは、表現ないし情報出力の間の序列や「差別」を消極的に受容し、見て見ぬふりをするのではなく、序列の存在を前提に、その理由を積極的に探究することであるようにも思われる。

思えば、日本の最高裁は、猥褻文書の頒布等の禁止（刑法一七五条）を合憲としたチャタレイ事件判決⁽⁸⁷⁾において、「猥褻文書は性欲を興奮、刺激し、人間をしてその動物的存在の面を明瞭に意識させる」、「それは人間の性に関する良心を麻痺させ、理性による制限を度外視し、奔放、無制限に振舞い、性道德、性秩序を無視することを誘発する危険を包蔵している」と述べていた。しかし日本の憲法学説には、「猥褻」の定義困難性やそれによる表現の萎縮効果等、ある種手続的な問題を指摘しておきながら、人間の動物的側面を「刺激」し、理性を迂回して「振舞い」を「誘発」する——と最高裁がいう——情報出力が、はたして「言論」なのかという実質的な問題には直接応答しないものも少なくなかった。⁽⁸⁸⁾

我々の理性に訴える「説得」に憲法上特別な意味を与え、それと「誘導」・「操作」とを区別するストラウスの説得性原則からは、人間の動物的側面を刺激して振舞いを誘発するこのような情報出力は、仮に予防的・政策的に憲法上保護されるべきだとしても（したがって刑法一七五条は違憲と解される余地はあるけれども）、本来的な「言論」とはみなされないと評されることになる。⁽⁸⁹⁾ こうした原理論は——説得性原則を批判するものも含めて——身元不明の種々の情報出力が溢れかえる時代においては必要であり、これを回避することは、かえって表現の自由の基底的価値を害するリスクがあるように思われる。例えば、事実に関する虚偽言明が、憲法が正面から

保護する「言論」なのか、憲法上予防的・政策的に保護される情報出力に過ぎないのかを理論上曖昧にしておくことは、リドスキーが指摘するように、「多数の市民が嘘、ごまかし、プロパガンダに基づき政策的決定を行う」よう誘い、「主権を国民の手から虚偽情報の提供者へと明け渡す」ことにもなるだろう。⁽⁹¹⁾ 骨の折れる仕事ではあるものの、この時代には表現の自由についても骨太の「総論」が必要ということである。⁽⁹²⁾

- (1) Jenima Kelly, *Twitter isn't the town square, it's the theatre*, FINANCIAL TIMES (May 5, 2022), at <https://www.ft.com/content/93c17a9e-8f65-48e1-8273-c7772d6f1819>.
- (2) 山本龍彦『フナハンモン・エロノミーの逆説』(KADOKAWA, 二〇一四年)。
- (3) See e.g., Tim Wu, *Machine Speech*, 161 U. PA. L. REV. 1495, 1507 (2013); Lawrence Lessig, *The First Amendment Dose Not Protect Replicants*, in SOCIAL MEDIA, FREEDOM OF SPEECH, AND THE FUTURE OF OUR DEMOCRACY 273 (Lee C. Bollinger & Geoffrey R. Stone, eds., 2022).
- (4) David A. Strauss, *Persuasion, Autonomy, and Freedom of Expression*, 91 COLUM. L. REV. 334 (1991).
- (5) Alec Peters, *Machine Manipulation: Why an AI Editor Does Not Serve First Amendment Value*, 95 U. COLO. L. REV. 307 (2024).
- (6) Strauss, *supra* note 4, at 337
- (7) *Id.* at 335.
- (8) *Id.*
- (9) *Id.*
- (10) *Id.* at 336.
- (11) Turner Broad. Sys., Inc. v. FCC, 512 U.S. 622, 641 (1994).
- (12) See e.g., C. Edwin Baker, *Autonomy and Free Speech*, 27 CONST. COMMENT. 251, 274 (2011); Mackenzie Austin & Max Levy, *Speech Certainty: Algorithmic Speech and the Limit of the First Amendment*, 77 STAN. L. REV. 1,

- 21-23 (2025).
- (13) Strauss, *supra* note 4, at 354.
- (14) *Id.*
- (15) *Id.* at 355.
- (16) Richard H. Fallon, Jr., *Two Senses of Autonomy*, 46 STAN. L. REV. 875, 877 (1984).
- (17) Strauss, *supra* note 4, at 355.
- (18) *Id.* at 339.
- (19) 418 U.S. 323 (1974).
- (20) *Gertz*, 418 U.S. at 340 (quoting New York Times Co. v. Sullivan, 376 U.S. 254, 270 (1964)).
- (21) Peters, *supra* note 5, at 328-329.
- (22) 最大判昭和四四年六月二五日刑集二三卷七号九七五頁。
- (23) David A. Strauss, *The Ubiquity of Prophylactic Rule*, 55 U. CHI. L. REV. 190 (1988). ストラウスの「予防的ルール」については、山本龍彦「違憲審査理論と権利論」大沢秀介編『東アジアにおけるアメリカ憲法』（慶應義塾大学出版会、二〇〇六年）。
- (24) 「子供SNS依存 米訴訟増」読売新聞二〇二五年一〇月三日一面、三面。
- (25) 最も重要なのは、EUのデジタルサービス法（DSA）であろう。概説として、生貝直人「EUデジタルサービス法とプラットフォームのガバナンス」法とコンピュータ四一号（二〇二三年）二七頁以下。
- (26) フロリダ州法は二〇二二年五月、テキサス州法は同年九月に成立した。どちらの州も知事は保守派である（この立法は、二〇二一年一月の連邦議会襲撃事件を受けたソーシャルメディアプラットフォームのトランプ支持派のアカウント停止等を背景にしている）。
- (27) 603 U.S. 707 (2024).
- (28) 同判決に関する日本語の優れた評釈として、上本翔大「コンテンツモデレーションと表現の自由」Moody v. NetChoice, LLC, 144 S. Ct. 2383 (2024)「世界人権問題研究センター研究紀要二〇号（二〇二五年）三七頁以下」。

- (29) *Moody*, 603 U.S. at 728.
- (30) *Id.* at 738.
- (31) *Id.*
- (32) 上本・前掲注(28)四七頁は、ケーガン判事が「SNS事業者によるコンテンツモデレーションを新聞社等の編集判断に匹敵する営為であると提えていることは、コンテンツモデレーションに対して手厚い保障を認めている立場である」を意味しつつ「知る」を指摘している。
- (33) *Moody*, 603 U.S. at 766 (Alito, J., concurring in the judgement).
- (34) *Id.* at 745 (Barrett, J., concurring).
- (35) *Id.* at 746.
- (36) *See e.g.*, Tim Wu, *The First Amendment is Out of Control*, NEW YORK TIMES (July 2, 2024), at <https://www.nytimes.com/2024/07/02/opinion/supreme-court-netchoice-free-speech.html>.
- (37) Peters, *supra* note 5, at 315.
- (38) Rachel Metz, *Facebook's top AI scientist says it's 'dust' without artificial intelligence*, CNN BUS (Dec. 5, 2018), at <https://edition.cnn.com/2018/12/05/tech/ai-facebook-lecun>.
- (39) Austin & Levy, *supra* note 12, at 42-43.
- (40) Peters, *supra* note 5, at 338-339.
- (41) *Id.* at 341.
- (42) *Id.* at 333.
- (43) *Id.*
- (44) Daniel Susser et al., *Technology, Autonomy, and Manipulation*, 8 INTERNET POL'Y REV. 1, 4 (2019).
- (45) Peters, *supra* note 5, at 335.
- (46) *Id.*
- (47) *Id.* at 339.

- (48) Ryan Calo, *Digital Market Manipulation*, 82 GEO. WASH. L. REV. 995, 1018 (2014). カロの議論については、山本龍彦『〈超個人主義〉の逆説』（弘文堂、二〇一三年）三八、五三頁。
- (49) Robin Pomeroy, *Transcript: The Promises and Perils of AI - Stuart Russell on Radio Davos*, WORLD ECON. F. Jan. 6, 2022, at <https://www.weforum.org/stories/2022/01/artificial-intelligence-stuart-russell-radio-davos/>.
- (50) Peters, *supra* note 5, at 342.
- (51) Mason Marks, *Biosupremacy*, 55 UC. DAVIS L. REV. 513 (2021).
- (52) *Id.* at 519.
- (53) *Id.* at 518. 彼はこの力をフーコー (Michel Foucault) のいう「生権力 (biopower)」に擬える。 *Id.* at 524. ークスによれば、AI 推薦システムはコントロールネットワークの典型とされ、生権力の源泉ともされる。彼は、感知ネットワークとコントロールネットワークを通じて「人々の魂を掌握し、作り変える (refashion)」点で、プラットフォームはファシストとの共通点を有している点で近づく。 *See id.* at 548-549.
- (54) Jon Whittle, *AI Can Now Learn to Manipulate Human Behavior*, SCI. ALERT (Feb. 12, 2021), at <https://www.sciencelert.com/aican-now-learn-to-manipulate-human-behavior>.
- (55) *Id.*
- (56) Peters, *supra* note 5, at 345.
- (57) Julie E. Cohen, *Tailoring Election Regulation: The Platform Is the Frame*, 4 GEO. L. TECH. REV. 641, 652 (2020).
- (58) *Whiney v. California*, 274 U.S. 357, 375-77 (1927) (Brandeis, J., concurring).
- (59) Strauss, *supra* note 4, at 337 n.6.
- (60) *Id.* at 336.
- (61) 山本・前掲注(2)二〇二二〇三頁（下條信輔発言）。
- (62) ダニエル・カーネマン（村井章子訳）『ファスト&スロー（上、下）』（早川書房、二〇一四年）。
- (63) 山本・前掲注(48)二一九頁。

- (64) Peters, *supra* note 5, at 347.
- (65) *Whiney*, U.S. at 375 (Brandeis, J., concurring).
- (66) Peters, *supra* note 5, at 327.
- (67) Lyriisa Barnett Lidsky, *Nobody's Fools: The Rational Audience As First Amendment Ideal*, 2010 U. ILL. L. REV. 799, 840 (2010).
- (68) Miami Herald Pub. Co. v. Tornillo, 418 U.S. 241, 258 (1974).
- (69) *Id.*
- (70) Hurley v. Irish-Am. Gay, Lesbian & Bisexual Grp. Of Bos., 515 U.S. 557, 574 (1995).
- (71) ユーターズは「Tornillo 判決が『理性により公表すべきもの』と判断される内容を強制的に公表される行為に違憲である」(Tornillo, 418 U.S. at 256) 主張をなすよう使用している。Peters, *supra* note 5, at 319.
- (72) *Id.* at 349.
- (73) CHEN ET AL., PROCEEDING OF THE 2023 ACM CONFERENCE ON FAIRNESS, ACCOUNTABILITY, AND TRANSPARENCY, HARMS FROM INCREASINGLY AGENTIC ALGORITHMIC SYSTEMS 651, 652 (2023).
- (74) Peters, *supra* note 5, at 350.
- (75) *Id.* at 355.
- (76) Austin & Levy, *supra* note 12, at 5.
- (77) *Id.* at 20–23.
- (78) *Id.* at 22.
- (79) *Id.* at 40.
- (80) *Id.* at 69.
- (81) *Id.* at 6.
- (82) 最決平成二九年一月三二日民集七一巻一六三頁。
- (83) 最判令和四年六月二四日民集七六巻五号一一七〇頁。

- (84) 前掲注(82)。
- (85) 船所寛生「解説」法曹時報七六卷二号(二〇二四年)三〇五頁。
- (86) もちろん、こうした原理論はこれまでも行われてきた。代表的なものとして、奥平康弘『なぜ「表現の自由」か』(東京大学出版会、一九八八年〔新装版二〇一七年〕)。最近のものとして、梶原健佑「表現の自由の原理論」山本龍彦『横大道聡『憲法学の現在地』(日本評論社、二〇二〇年)一七九頁以下。
- (87) 最大判昭和三二年三月一三日刑集一一卷三号九九七頁。
- (88) もちろん、例外もある。ポルノグラフィーとの関係であるが、例えば、菅谷麻衣「言語と行為の臨界」法学政治学論究一〇三号(二〇一四年)六九頁。なお、本稿の議論を深化させるうえで、ストラウスの説得性原則と、例えばブラック判事の「言論／行為」論との関係に関する検討は不可避であろう。森脇敦史「ヒューゴ・ブラック」山本龍彦『大林啓吾編『アメリカ憲法の群像』(尚学社、二〇二〇年)一四五頁以下。
- (89) Strauss, *supra* note 23.
- (90) *Id.* at 345 n.35.
- (91) Lidsky, *supra* note 67, at 838-39.
- (92) 「総論」の必要性については、石川健治ほか「新・民法学を語る その1——石川健治先生をお招きして(下)」書斎の窓六八六号(二〇二三年)一五頁。