

フィジカル AI を支える制御工学

東京理科大学 創域理工学部 電気電子情報工学科 教授 なかむら ひさかず 中村 文一

はじめに

AI（人工知能）が身近なものになり、みなさんの中にも、ChatGPT、Geminiなどに様々なことを質問している方がいるのではないだろうか？近年は、フィジカル AI という概念が注目を集めている。これは、身体を持った AI ということであるが、自動車やロボットのようなリアルワールドでの実システム（フィジカルシステム）に対しセンサを取り付け、制御対象の状態や制御対象を取り巻く環境を計測し、AI 利用して制御指令を生成し、実システムを制御するシステムである¹⁾。フィジカル AI で考える実システムは、「入力」を与えると「制御出力」が出てくる制御対象を扱う【図1】。例えば制御対象が自動車である場合、入力はアクセルの踏み込み量とハンドルのステアリング角度、制御出力は自動車の位置、となる。つまり、自動車のアクセルを踏むと、そのうち速度が上昇し、自動車の位置が変化する、というわけであり、自動車は、アクセルの踏み込みとステアリング角度を入れると、位置が出てくる制御対象、と考えられるのである。実際には自動車には多数のアクチュエータ（駆動装置）やセンサ（計測装置）が搭載されているが、これらをひとまとめのもの、と考えて扱うのである（扱ってよいかどうかはまた別の問題である）。



【図1】 制御対象

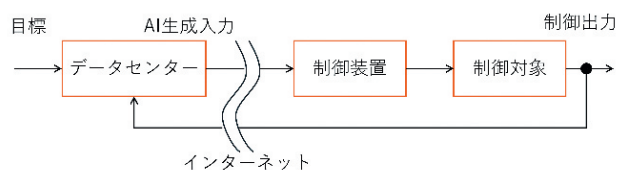
フィジカル AI

AI が身体を持つという、具体的な構造を【図2】に示した。人間がはじめに与える制御対象に対する所望の目標を、言語あるいは数値でデータセンターの AI に示し、AI は人間が与えた目標値とインターネットを介して接続されている制御対象からの制御出力をもとに入力を計算し、制御対象に送るという構造である。

これだけなら「え？単にネットワークでデータセンターと制御システムをくっつけばいいだけじゃないの？」と思うかもしれないが、フィジカル AI 特有の問題があり、うまくいかない。大きな問題は2つあり、1つ目は計算の問題、2つ目はハルシネーションの問題である。

まず前者の計算の問題について紹介しよう。われわれがいわゆる AI と呼んでいる大規模言語モデル（LLM）では、性能の競争から速度競争へと移り変わっているが、それでも複雑な問題に対して数秒以上かかることは珍しくない。一方、ロボットの力制御などでは最低1ミリ秒の間隔で処理が行われなければならない、技術的課題は非常に高い。また、制御対象が自動車などの場合、数秒の遅れがあった場合、もはや事故は避けられない。

後者のハルシネーションとは、AI が事実に基づかない情報を生成してしまう現象である。その多くは一見もっともらしい情報のように見えるため判別が難しく、AI の信頼性を損ねている。AI が、フィジカルシ



【図2】 フィジカル AI システム

システムに影響を与えないのであれば、直接的に人間が危険になることはない。ところが身体を持った瞬間に、誤った情報に基づいてロボットや自動車が人間を攻撃する可能性が否定できないのである。

このような重大な問題をどのように解決するのか？これにはフィジカル AI の基盤構造が重要である。フィジカル AI が現実世界で成立し続けるためには、以下の 4 つの基盤が同時に成立していなければならない²⁾。

1. L1：知能・計算の成立基盤：世界を理解し、未来を予測し、行動計画を生成する判断の中核。
2. L2：身体・感覚の成立基盤：賢さを現実世界に変換する基礎。アクチュエータ、センサ、力制御、安全制御など。
3. L3：学習加速の成立基盤：失敗を学習に変えるための仕組み。
4. L4：社会・需要の成立基盤：事故時の責任、説明可能性、規制・受容性など社会的ルール。

このなかで、L1 と L3 は、LLM がまさしく得意とする分野であり、L4 は社会的なルール設定の問題である。そのように考えると、フィジカル AI の技術的課題は、L2 に集約されると考えてよい。筆者が主に研究を行っている制御工学は、このまさに LLM が苦手の L2 の問題解決を研究の主題とする学術分野である。昔からの制御工学を知っている方には「？」となるかもしれないが、現在制御工学はフィジカル AI に対応させるために急速に変化している。

フィジカル AI を支える制御工学

制御工学は、19 世紀における蒸気機関の安全性、に端を発する長い歴史のある学問であり、制御出力をあらかじめ与えられた目標値に追従させるための制御器を与えるサーボ制御問題が主要な研究課題である。ここで制御対象は、数学的にモデル化され、その制御サイクルは非常に高速な問題を前提とする。制御工学では、長い間このサーボ制御問題、問題の解決に必須である安定化、について研究を行ってきた。従来の制御工学でも、アクチュエータ、センサ、力制御、ここまでは対応できた。しかし、そもそもワットの蒸気機関の安定性を保証することから始まった制御工学は、システムの安全保障に関してはほとんど議論されてこ

なかったのである。この問題に対し、近年、制御バリア関数、というポテンシャル関数を用いた安全保障制御理論が構築され、まだ一般には広まっていないが、産業界では急速に普及が進んでいる。

さて、制御工学分野で確立しつつある L2 基盤を、どのようにフィジカル AI に導入すればよいのだろうか？これについてもいくつかの研究が進められているが、特に重要なものが、カスケード型ハイブリッド AI と呼ばれる手法であり、システムの上流部に AI を使い、【図 2】のように制御対象が配置された場所に軽量な制御装置を配置する、というアプローチである。これにより、短い時間周期で「制御出力を制御」「安全保障」の 2 大問題を解決する。

制御工学の考え方

さて、それでは制御工学の考え方は、どのようなものであろうか？例えば、入力速度[km/h]、出力が位置[km]、であるような自動車の数学モデルを考えよう。これは、以下のような微分方程式で表現される。

$$\frac{dy}{dt}(t) = u(t) \quad (1)$$

ここで $y(t)$ [km] は制御出力である「自動車の位置」を表す変数であり、時々刻々変化するものであるから、時間 t [h] の関数として表現する。 $u(t)$ [km/h] はわれわれが自由に設計することのできる入力である。(1) 式は、位置の時間微分として、入力を与えられることを示した方程式である。

制御工学は、出力 $y(t)$ が別途与えた目標値 $r(t)$ に追従するように、入力 $u(t)$ を設計する、数学パズルの学問と考えてよい。純粋な数学とは異なり、物理法則を捻じ曲げて人間の都合のよくなるように、入力を設計するのである。具体的に、目標値が $r(t) = 30$ [km] の場合を考えてみよう。初期値が $x(0) = 0$ [km] の場合では、距離 = 速さ × 時間であるから 0.5[h] だけ $u(t) = 60$ [km/h] とすればよい。このような考え方をフィードフォワードという。一方、以下のような、時間に依存しない制御法も存在し、フィードバック制御と呼ばれる。

$$u(t) = -(y(t) - 30) \quad (2)$$

実用的にはこちらのほうが都合の良い場合も多く、自動車や産業用ロボットなど、多くの場面でこのような

フィードバック制御が用いられている。

制御バリア関数を用いた安全保障制御

ここでは、制御バリア関数を用いた安全保障理論を紹介しよう。これは Wieland により提案され⁴⁾、Ames らにより大きく発展した⁵⁾。AI や人間による、未来にどのような入力を与えられるかわからないものに対する安全保障理論は、筆者らの研究グループにより、世界で初めて与えられたと考えてよいと思う⁶⁾。筆者らは、2016 年ごろより人間を含んだシステムに対する安全保障に関する研究を行っているが、その当時はちょうど衝突被害軽減ブレーキを搭載した自動車の普及が始まったところであった。このころ人間の判断の誤りをシステムによってサポートするような理論の構築は、制御理論では少数の研究を除いて行われておらず、将来的な課題とされていた。一方、複雑化する先進運転支援システムの開発コストは増大が続いている、というニュースが日々流れており、人間の自由な操作下においても安全性を保障できる理論的なフレームワークの構築が、制御工学が今後も学術分野として継続するためには必須ではないかと考えたことが研究を開始する動機であった⁷⁾。

制御バリア関数は、逆数型と零型の 2 種類が提案されており、現在では零型の採用例が多いが、両者は基本的に等価であるため、説明が簡単な逆数型を使って説明しよう。

前章とは少し違い、(1)式に人間の自由な入力 $u_h(t)$ が加わった以下のシステムを考えよう。

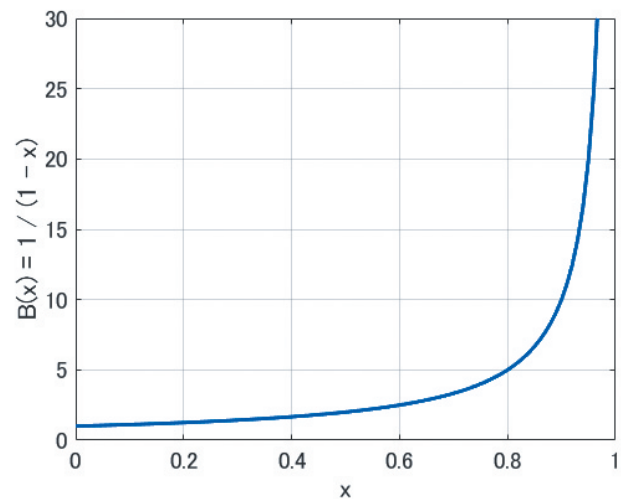
$$\frac{dy}{dt}(t) = u_h(t) + u(t) \quad (3)$$

ここで、1[km]先に故障車があり、事故を避けなければならない場合について考えよう。ここで、以下の関数 $B(x(t))$ は制御バリア関数とよばれる【図 3】。

$$B(x(t)) = \frac{1}{1-x(t)} \quad (4)$$

この関数は、障害物である故障車のある $x=1$ で無限大となる関数であり、自動車が安全である、というのは制御バリア関数 $B(x)$ が無限大にならない、ということの意味する。

関数が無限大にならないようにするためには、陰関数 $B(x(t))$ を陽関数 $B(t)$ のように考え、以下の時間微



【図 3】 制御バリア関数

分条件を満たせばよい、ということが安全保障制御の最も重要なポイントである。

$$\frac{dB}{dt}(t) \leq KB(x(t)) \quad (5)$$

ここで、 K は適当な正の定数であり、存在性が保証されさえすればよいが、本項では $K=1$ としよう。このとき、以下の入力を考えよう。

$$u(t) = \begin{cases} 0 & \left(\frac{u_h(t)}{x^2} \leq -\frac{1}{x} \right) \\ -u_h(t) - x & \left(\frac{u_h(t)}{x^2} > -\frac{1}{x} \right) \end{cases} \quad (6)$$

(6)式の入力は、2行で書けるほど簡単で、人間のどのような入力に対しても（将来どのような入力が入ってくるかわからなくても）(5)式を満足して自動車の安全性を常に保証し、最大限人間の入力を尊重し、しかも突然入力が介入することはない、という優れものである。

この入力 $u_h(t)$ を与えるのが人間ではなく AI であって、誤った指令が与えられても、理論的には事故を起こすことを防ぐことができることから、フィジカル AI において非常に重要な役割を果たすのである。

おわりに

制御工学、特に安全保障制御はフィジカル AI において欠かすことのできないツールになりつつある。その一方で、(5)式のシンプルな条件で、未来がわからない信号に対しても安全性を保証できることは、なんと 2010 年代までわかっていなかった。科学技術は非

常に進んでいるが、まだまだ人類が見つけていない簡単で便利な方法がたくさんあるのではないだろうか？



参考文献

- 1) 安井裕司：さいす学とフィジカル AI, 計測と制御, vol. 65, no. 3, pp. 179-185 (2026).
- 2) 田中道昭：フィジカル AI の時代 (前編), 富士通株式会社, <https://global.fujitsu/ja-jp/technology/key-technologies/news/ta-ces2026-report-1-20260108> (2026).
- 3) J. Seo, E. Kim and A. J. Choi: LLM-DWA: a hybrid path planning framework combining large language models with the dynamic window approach, Scientific Reports, vol. 16, 9898 (2026).
- 4) P. Wieland, F. Allgöwer: Constructive safety using control barrier functions, Proceedings of the 7th IFAC Symposium on Nonlinear Control System, pp. 462-467 (2007).
- 5) A. D. Ames, X. Xu, J. W. Grizzle, P. Tabuada: Control barrier function based quadratic programs for safety critical systems, IEEE Transactions on Automatic Control, vol. 62, no. 8, pp. 3861-3876 (2017).
- 6) 中村, 吉永, 小山, 江藤: 制御バリア関数を用いたヒューマンアシスト制御, 計測自動制御学会論文集, vol. 55, no. 5, pp. 353-361 (2019).
- 7) 中村文一: 制御バリア関数を用いたヒューマンアシスト制御, システム／制御／情報, vol. 65, no. 9, pp. 358-363 (2021).