

脳の様相と力学系

津田一郎（北大・理・数学）

§ 1. 身体性について

ここでは「身体性」を次のように定義する。

・内部記述と外部記述を整合させるようなインターフェイスの構築に伴う感覚を身体性という。

脳の問題において、内部記述とは、脳の中に埋め込まれた外部世界をその一部である脳がその働きそのものを観測し記述することであり、外部記述とは、それを脳の外にとりだし外部世界の側から記述することであるとする。簡単な例は、前庭動限反射である。空間の一点を凝視したまま頭を動かすと、眼は反対方向に反射的に動く。これを前庭動限反射という。これにより、脳は身体が運動していても静止物体をとらえることができるとされている。この時、網膜像は静止像となる。しかし、他方スモース・パシュートという眼の動きが知られている。これは、空間内をなめらかに動いている物体を頭を固定して追跡するときに、眼はその物体の運動方向になめらかに動くことをいう。これにより網膜像は静止像となる。従って、脳は網膜像の情報だけでは、物体が静止しているかどうか判断できない。少なくとも頭の動きと網膜像の二つの情報が必要である。ここで、空間内の静止物体を凝視しつつ頭を激しく左右に振ってみよう。物体が動いて見えるはずである。もはや脳は、その物体が静止しているとは判断できない。では何故、反射がこの状態で固定しなかったのだろうか。手で触れたとき物体が静止していると確認できたからか。恐らく、否であろう。感覚というものが、必ずしも外部世界を正確に写しとっているわけではない。しかし、脳は外部世界を知っているかのように反射を固定させている。外部世界に移された座標での記述と内的な座標での記述を整合させていると思われるのである。

サルの運動系に関係する高次視覚野に次のような神経細胞がみついている。サルが物体に近づいた時と、物体がサルに近づいた時の両方に反応する細胞がある。他方、いずれか一方の運動の場合にのみ反応する細胞がある。前者は相対運動にのみ反応しているので、このような内部記述ではコバリアンスのみが問題となる(Rössler)。後者は恐らく外部記述でありインバリアンスが獲得される。

こういった二つの記述系が脳内に構築されている。脳では内部記述の存在はむしろ自明であり、外部記述は非自明である。これらを整合させる機構をインターフェイスと呼ぶことにする。身体性とはインターフェイスを内部に構築することに伴う感覚のことだと考えようという提案である。

ここでは脳に関しての身体性を述べたが、上のように定義すると、これは脳の問題にとどまらない。神経節のみをもつ動物も、神経のない粘菌のような生物も、植物も、また生体内の高分子さえも身体性をもつだろう。

このような意味の身体性のモデルを構築したいという欲求から今までにいくつかのモデル構築の試みを行なった。

- ① 記憶と知覚あるいは記憶と情報処理の分離不可能性のモデル、
- ② 学習と想起あるいは情報の書き込みと読み出しの同時性のモデル、
- ③ 論理とダイナミックスの関係に関するモデル、
- ④ 自己記述のモデル、

などである。

以下、これらについて簡単に述べる。

§2. 記憶と知覚の分離不可能性

一般に記憶は貯えられたものと考えられており、神経回路網モデルにおいては特にこの意味での記憶を扱っている。その実現方法は、アトラクターによる記憶の表現である。しかし、記憶は常に現時点において知覚、認識の過程に伴って現われるので、論理的にはそれらの過程と記憶過程が分離して存在する必然性はない。貯えられた記憶をひっぱり出して認知活動に利用しているというのは皮相的な見方である。確かに記憶は貯えられるのであるが、独立な状態として表現でき情報処理そのものと区別できるといったものではないように思われる。

我々のモデルでは、記憶は、アトラクター痕跡として表現され、これらをつなぐ記憶の連鎖にカオスの遍歴を見い出したのだった。このアトラクター痕跡の可能な数学的表現の一つとしてミルナーアトラクターの弱崩壊の状態が考えられる。講演では、その意味で簡単だが自明ではないミルナーアトラクターモデルを紹介し、カオスの遍歴との関係（金子）を再吟味した。

記憶痕跡の連鎖がカオスの遍歴になる様相が学習や知覚等の認知的過程と区別されないことは、我々の一連のモデル研究によって暗示されている。さらにこのことから、学習と想起、情報の読み出しと書き込みが同時に進行するというモデルをつくることができる。それはカオスの遍歴のモデルでもみられるが、特異連続でいたるところ微分不可能なアトラクター上でも認められる。

そこでは、縮小写像の空間にできたカントール集合上にカオスで読み出された情報が同時に書き込まれる過程が観測される。例えばカオスで生成された記号列 10^* はカントール集合のある部分集合のひとかたまりで整理され 101^* , 100^* はそのかたまりをより細かいスケールで分離する2つの部分のかたまりで表現される、といった具合である。これはある意味で学習、記憶、想起の力学系での表現であるが、ここでは、シナプス学習は仮定されていない。シナプス可塑性がなくともダイナミックスのみで学習、記憶、想起と類似の様相は表現可能となるのである。

このような過程は、脳のさまざまな領野で観測可能だと思われる。上のような情報表現（カントールコーディングと呼ぼう）は基本的に、安定な神経（回路網）が、カオスの

な神経（回路網）で駆動されている時に、縮小係数と拡大係数の比が適当に調節されていれば起こることだからである。臭い情報処理系の前梨状皮質や、記憶過程に関係することが古くから指摘されている海馬の CA1 か海馬傍回のいずれかでコントロールコーディングがみられる可能性を指摘することができる。

§3. 論理とダイナミックス — 自己記述のモデル

マッカロとピッツの研究以来、思考を神経モデルで表現する試みが行われてきた。他方、ブール以来、思考を論理と数学的形式で表現する試みが行われてきた。マッカロとピッツの試みはブールの研究に基づいている。しかし彼らは後にマッカロ・ピッツニューロンという名で呼ばれることになったニューロンモデルでそれを表現した。彼らは論理として古典論理、ダイナミックスとしてマッカロ・ピッツニューロンを用いることで計算能力としては万能チューリングマシンの能力を得た。

これに対し、我々は、論理として連続値論理（例としてウカシェビッチ論理）、ダイナミックスとしては動的関数写像（もしくはブートストラップ型神経回路）を用いることで、アナログ計算を実現しようと試みた。論理を仮に1つ固定して、命題の真理値を推論する過程を、結論を次のステップでの推論の前提と同一視することで再帰的な過程とみなす。これにより推論過程を力学系として扱うことができる。この過程には二つの解釈が可能である。一つの解釈では時間のキャリアーが、前提→結論にあり、もう一つの解釈では、結論→前提にある。論理が空間に一樣に作用するときは、この二つの解釈で力学に違いはないが、論理が部分空間ごとに非同期的に作用している場合や、それに類似の効果をもたらすような設定、例えば、ある一つの部分空間でのみ真理値が定数になるような命題の設定においては、どこに時間のキャリアーをもってくるかによって力学は異なったものになる。

我々はこのような系で、力学系を引き数としてもつ関数の写像を自己記述のモデルとして得た。この動的関数写像と類似のものは、ド・ラムの関数写像である。違いは、ド・ラム方程式が、縮小写像を座標変換とするのに対し、我々の方程式では、カオス写像が座標変換を与える点である。

次に我々は論理・力学変換の仕組を信念システムに対応づけた。この対応づけによって信念の階層の4段階をその上で実行する媒体のもつべき性質が示された。動的関数写像は4型の推論者が推論できる場を与えていると思われる。

しかし、こういった理論はあまりに形式的であるのでこれだけでは脳と心のミゾを埋めることはできない。このような形式論と並行して実験を行なう必要を強く感じた。そこでこの形式論に基づいた推論形式に関する実験を考案した。これは脳の実験家と協同して進めることになっている。