

論文要旨

2023 年 1 月 10 日

※報告番号	甲第330号	氏 名	長瀬 准平
主論文題名			
深層ニューラルネットワークモデルの関数表現能力に関する数理解析			
内容の要旨			
深層学習は、近年情報科学の広い分野で活躍している機械学習技術であり、そのための学習モデルである深層ニューラルネットワークモデル（以下、DNN）も数多く提案されている。最も単純な DNN である多層パーセプトロンモデル（以下、MLP）は、アフィン関数と活性化関数が繰り返し合成されることによって設計される直列な接続のみをもつネットワークモデルであり、MLP を基本として様々な工夫や改良の加えられた DNN が提案されている。モデル設計上の工夫としては、直列以外の分岐や合流といった多様なネットワークの接続の利用、多様な非線型変換をもつ活性化関数の利用、畳み込み層などのようなアフィン関数のパラメータに関する制限などが例に挙げられる。一方で、DNN は実験によってのみ性能が保証されることも多く、説明性の低いブラックボックスな手法であるとされているため、その高い性能と広い応用、多くの研究に反して、信用性の問題や改良の難しさが課題である。また、既存のこれまでの機械学習手法と比較してパラメータ数が過剰に多く、従来の理論では学習は困難であると考えられるにも関わらず学習に成功しているなど、深層学習には理論的な疑問点も多い。これらの問題に対する理論的な研究として、MLP や残差スキップ接続付きモデルについての数理解析は近年盛んに進められているものの、DNN の設計は前述の通りとても多様なものとなっているために、それら DNN の統一的な解釈や設計に関する体系的な理論は未だに知られていない。 このような背景から、本論文では DNN の統一的な記述と評価を与えるための数理解析として、関数表現能力に着目した研究を行った。研究の方針として、学習モデルが設計可能な関数の集合を表現集合と呼んで定義し、異なる DNN が設計可能な関数をそれぞれ書き下すことによって関数の等価性と包含関係を調べることで、異なるモデルの関数表現能力の比較を行い、また、異なるいくつかのモデルのうち関数表現能力が等価であるものについての結果を得た。 （次頁に続く）			

論 文 要 旨

2023 年 1 月 10 日

※ 報告番号	第	号	氏 名	長瀬 准平
内容の要旨（続き）				
本論文の主結果は次の通りである．アフィン関数と ReLU 活性化関数を合成，加算，連結して設計される任意の関数が MLP によって設計可能であることが示された．このことから，DNN における接続の構造として，合成，加算，連結のようなネットワークの分岐や合流によって表される任意の接続の構造は，いずれも MLP のような直列な接続に帰着が可能である．次に，この結果をより一般化された活性化関数に拡張するために，ReLU 関数を一般化した区分線形関数を活性化関数に用いた MLP についての議論を行った．先行研究として，ReLU 活性化関数を用いた MLP が区分線形関数を設計可能であることが知られているが，本論文ではその逆の命題，区分線形関数を活性化関数とした MLP が ReLU 関数を設計可能であることを結果として示した．すなわち，両者を活性化関数として用いた MLP は層の深さと次元に関する適当な条件の下で等価な出力と表現集合をもつ．また，この結果から，区分線形関数を活性化関数とした MLP も任意の接続の構造を設計可能であることが導かれる．そして，上述のこれらの結果はいずれも構成的な証明によって導出されており，対応する両者のモデルの出力を保つようなモデルのパラメータの対応が存在することがわかる． 従来の DNN の設計と検討は，特定のデータセットにおける実験によって得られた経験的な性能評価に基づいていることが一般的であり，モデルの性能の評価を統一的に行うことはできていなかった．本論文の結果は，設計可能な関数の集合という観点から，異なるいくつかのモデルに対して等価な設計という基準を与えるものである．等価な設計の存在によって，既に設計されて高い性能が認められているモデルを別の形式のモデルで設計し直したり，特定のモデルに対して既に知られている理論的な結果を別の形式のモデルにも拡張したりすることが可能になる．すなわち，これまで特定のモデルや特定の活性化関数に対して行われていた理論的な結果について，対応する別のモデルへ置き換えてより広い結果を得られると考えられる．従来の万能近似性などの議論と異なり，本論文の結果はパラメータ数に関する極限操作を行わず，有限の範囲内でモデル同士の比較を行うことから，実際に実装されて応用されているモデルに対しても上記の結果を適用することができ，対応する異なるモデルへと置き換えることができる．各モデルの学習の性能として収束性や汎化性能などが今後明らかとなれば，本論文の結果を用いて学習中にそれぞれの状況に適した等価なモデルに置き換えて効果的に学習を進めることが可能となる．また，学習の終了後に計算コストやモデルの容量において優れたモデルに置き換えて運用するといった応用も可能であると考えられる．				