

Uncovering the Internal Representations of Neural Networks with Regression Analysis

Hiroshi Nonaka

May 12, 2025

Abstract

With the rise of AI technologies in different social sectors, understanding the internal states of deep neural networks is becoming a critical research field that tackles the black box problem. One possible approach is probing, a regression analysis that takes internal activations as independent variables and regresses target information. This method quantifies how informative the representations are in relation to the dependent variables. I implemented probing for a simple three-layer neural network for binary image classifications with three regression models: linear probability models, logistic regression, and neural networks. The results revealed that all three layers possess equally rich information for the binary image labels. Through further statistical analysis, I discovered that the representation vectors from deeper layers are more multicollinear. By applying Principal Component Analysis (PCA) to mitigate multicollinearity, I also found out that the vectors essentially acquired an equally useful representation. Heteroskedasticity tests showed that the data was heteroskedastic at a high significance level, and the training steps did not affect it.

1 Introduction

1.1 Overview

Neural networks are essential building components of recent AI models such as ChatGPT [1]. Although neural network predictions often outperform linear regression models because of their flexible representation capacity, their internal states are hard to interpret compared to linear models. For instance, the common multivariate linear regression can be written in the following form:

$$\hat{Y} = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_n X_n$$

where $\hat{Y}$ is an estimated value of a dependent variable Y , and β_i are the coefficients of independent random variables X_i . This is easy to interpret. In contrast, neural networks with m layers are often represented as:

$$\hat{Y} = (T_l \circ f_{l-1} \circ T_{l-1} \circ \dots \circ f_1 \circ T_1)(\vec{X}) \quad (1)$$

where $\circ$ is the function composition operation, $\vec{X} = [X_1, X_2, \dots, X_n]$, T_i is a transformation such that $T_i(\vec{V}) = A_i \vec{V} + b_i$ for some input vector $\vec{V}$, matrix A_i and bias vector b_i , and f is some function (mostly non-linear) that slightly modify outputs from T_i . Repetitively applying $f \circ T$ gives rise to the non-linearity of neural network functions. We *train* neural networks so they learn patterns in a dataset through multiple training steps. During the training, the error between the output $\hat{Y}$ and the true value Y is calculated. This error function is called a *loss function*, and neural networks update the values of their matrices A_i and bias vectors b_i by taking the gradient of the loss function. This is how training optimizes neural networks to *fit* to the given dataset.

1.2 Motivation

Due to the multiple transformations and non-linear functions within neural networks, it becomes hard to interpret what the matrix and bias vector in each transformation represent compared to simple linear regression models, in which the coefficients denote how much each variable contributes to the estimated point $\hat{Y}$. The uninterpretability of neural networks is named the *black box problem* and has stimulated various research endeavors, including *explainable AI (XAI)* [6]. Given that AI systems are implemented in many different sectors today, understanding *how AI thinks* is one of the critical themes of ongoing research. I aim to tackle this problem through a popular approach called probing [5]. My research questions are mentioned in Section 4.

1.3 Approach

Probing applies regression analysis to measure how rich information an internal vector contains in relation to target data. Let an *activation* $\vec{Y}_i = [Y_{i,1}, Y_{i,2}, \dots, Y_{i,m}]$ represent the output vector of the i -th transformation T_i in Equation 1: $\vec{Y}_i = T_i(f_{i-1}(\vec{Y}_{i-1}))$. Note that the first activation is produced from the observation vector: $\vec{Y}_1 = T_1(\vec{X})$. In prob-

ing, regression models P , namely *probes*, regress some target information using the vector $\vec{Y}_i = [Y_{i,1}, Y_{i,2}, \dots, Y_{i,m}]$ as m independent variables. If P produces good predictions, this indicates that $\vec{Y}_i$ are good independent variables, or more specifically, contain rich information about the target.

2 Notations

Here, I outline notations that frequently appear in this paper. T_i indicates the i -th transformation layer in a neural network. f_i represents a function applied to the output from T_i . Most models use a single non-linear function throughout their pipeline, so I occasionally write f to denote the unique function. $\vec{X} = [X_1, X_2, \dots, X_n]$ is an n -dimensional observation and an input to a neural network. An activation $\vec{Y}_i = [Y_{i,1}, Y_{i,2}, \dots, Y_{i,m}]$ is an m -dimensional vector from the i -th layer in the neural network. Therefore, each observation $\vec{X}$ has corresponding activations $\vec{Y}_1, \vec{Y}_2$ through $\vec{Y}_l$ for each transformation in Equation 1. In my paper, P is a probe/function/regression model that takes activations and maps them to a probability p : $p = P(\vec{Y}_i)$. One example of P is logistic regression.

3 Literature Review

In this section, I aim to introduce some prior works on probing and simultaneously deepen the audience’s understanding of its method through concrete examples.

3.1 Machine Learning in Regression Analysis

Before discussing probing techniques, it is worth noting that machine learning algorithms are closely tied to regression analysis in many fields. Decision trees [10] and their families, such as random forest and gradient boosting, are popular regression algorithms for their accuracy and interpretability. Ridge and Lasso regressions [13], regularized linear regression models, are also important linear models in machine learning. These machine learning models are applied in many different areas, such as political sciences [7], chemistry [4], and economics [15].

3.2 Probing Neural Networks

Gurnee and Tegmark [8] applied Ridge regression and neural nets to measure the spatiotemporal representations within large language models (LLMs). LLMs are neural-network-based AI models that deal with language, such as ChatGPT. For spatial information, they gave LLMs prompts asking for the latitude and longitude of a specific place in the world, and for temporal information, the LLMs were asked about historic dates related to well-known books, figures, and et cetera. Activations produced by those stimuli were collected and used as independent variables to regress the latitudes and longitudes of the places or the associated years to the historic events. The researchers discovered that the linear model

can predict the dependent variables with high accuracy, concluding that the spatio-temporal representations within LLMs are linearly related to the target values.

Multiple research works have applied probing to investigate representations of neural networks in various modalities, like text and image. McCoy and Leshinskaya [12] applied regression analysis to investigate how an LLM understands subject-object relations. Specifically, they focused on the *has-color* relation; for instance, the subject is *bananas*, and the relation between the subject and object is *have the color of*, then the expected object is *yellow*. They gave prompt queries asking the color of a particular subject, collected activations from the fifth layer of the LLM ($\vec{Y}_5$), and regressed the final output $\vec{Y}_l$.

Alain and Bengio [3] introduced linear classifier probes as a tool to understand the information captured in intermediate layers of neural networks. They collected activations from each layer of their neural network for image classification. Results revealed that the error of probes decreased as they used activations from deeper (later) layers as independent variables. Rashtchian et al. [17] probed image encoders of foundation models (AI models processing multimodal data) to quantify how the encoders generalize image representations. They collected activation vectors by giving the same images, but with different visual effects, and discovered that the representations that the encoders obtain from these differentially edited images converge well. Ilharco et al. [9] focused on the multimodal relations between contextual text representations with corresponding visual features by regressing image features based on activations from text-only LLMs. Results indicated that language representations are strong signals for predicting visual features.

Akhondzadeh et al. [2] leveraged linear, Bayesian, and pairwise probing to investigate the relations between the representations of graph-based models with important chemical properties like functional groups and odors. Ngo and Kim [14] trained probes to align language representations in LLMs with sound representations from audio models.

4 Research Directions and Questions

Despite the considerable number of prior studies, few have focused on the statistical analysis, such as heteroskedasticity test and multicollinearity, of representations along with probing. In this work, I aim to (1) mildly replicate existing research on probing techniques and (2) apply statistical analysis to internal representations for further understanding than probing. The research questions in this work are (1) which activation outputs $\vec{Y}_i$ contain rich information, and (2) how the amount of training affects representation learning.

5 Data Collection

5.1 Training and Architecture

To collect activation vectors, I first trained my neural network with 2,000 images of cats and dogs from the CIFAR10 dataset [11]. The model classifies whether a given image is likely to be a cat or a dog. The training pipeline is as follows: (1) give the model an image and have

it predict probability p that the image is a cat; (2) calculate binary cross-entropy between p and its true label as an error score; (3) repeat step (1) and (2) for 1,600 training images; (4) calculate a validation score with 400 unseen data; and (5) repeat step (1) through (4) for 100 times with the same training and validation data. Each repetition of steps (1) through (4) is called an *epoch*.

The model consists of a *convolution* part and a neural net part. Convolution layers compress a $64 \times 64 \times 3$ (height, width, and channels) image into $13 \times 13 \times 16$. Note that the convolution layers are also trainable.

The neural network has three layers: T_1, T_2 and T_3 . Each T_i maps $f(\vec{Y}_{i-1})$ or $\vec{X}$ to a vector of 60 dimensions. Therefore, $T_1 : \mathbb{R}^{2704=13 \cdot 13 \cdot 16} \rightarrow \mathbb{R}^{60}$, $T_2 : \mathbb{R}^{60} \rightarrow \mathbb{R}^{60}$, and $T_3 : \mathbb{R}^{60} \rightarrow \mathbb{R}^{60}$. A non-linear function f called *Rectified Linear Unit* follows T_1, T_2 , and T_3 . One final layer T_4 transforms the 60-dimensional vector from the previous transformation T_3 to a scalar. The logistic function σ turns the scalar into probability p .

The overall model architecture is outlined below:

$$p = (\sigma \circ T_4 \circ f \circ T_3 \circ f \circ T_2 \circ f \circ T_1 \circ conv)(\vec{X})$$

where *conv* is the convolution part, and $\vec{X}$ is an image vector.

5.2 Data Requirement

After each *epoch* during training, I gave the same 2,000 images to the network and collected $\vec{Y}_1, \vec{Y}_2$, and $\vec{Y}_3$ for *each image*. Therefore, I collected 600,000(= 2,000 images $\times$ 3 layers $\times$ 100 epochs) activations in total. I restricted the output dimension of T_1, T_2 , and T_3 to 60 in order for $\vec{Y}_i$ to have the same dimension size, which, later in probing, prevents the variability of regression performance due to the different number of independent variables. In short, the data I collected is cross-sectional data consisting of 300 .csv files (3 files for each of the three layers, multiplied by 100 epochs) with a table of 60 independent variables ($\vec{Y}_i$) and 2,000 observations for each file (refer to Figure 1 for details). I did not implement any control or instrument variables because of the nature of my research. The details of data preprocessing are explained in Section 6.2.

5.3 Assumptions and Diagnostics

For statistical analysis, I focused on two metrics, heteroskedasticity and multicollinearity. For heteroskedasticity, I ran the Breusch-Pagan test to calculate p-values for the null hypothesis of homoskedasticity. The left plot in Figure 2 shows that the entire data is heteroskedastic at a very high significance level (p-value $\lll 0.001$). Only the first layer experiences rapid increases of p-value, but this trend is inconsistent.

VIF of the second and third layers are also extremely high. VIFs of the third layer activations converges around 3 as the epoch approaches the end of training. Therefore, the activations from the first layer are mildly multicollinear, and those of the second and third layers have high multicollinearity. However, I can observe the downward trend of VIF as the amount

Figure 1: .csv file of the first layer at epoch 1

Index	label	scalar_0	scalar_1	scalar_2	scalar_3	scalar_4	scalar_5	scalar_6	scalar_7	scalar_8	scalar_9	scalar_10	scalar_11	scalar_12	scalar_13	scalar_14	scalar_15	scalar_16	scalar_17	scalar_18	scalar_19	scalar_20	
0	1	-0.186014677003330	0.2389158295817	-0.254874179135000	0.244486938715000	-0.242653036142807	0.230876789450000	0.4005233458114800	-0.408026086995000	0.2477684338007	-0.333888179178000	-0.3780733815440010	-0.273218717277000	0.476710078866666	0.243087772208200	0.04308800675231130	0.248586678888200	0.11174142808160700					
1	1	0.25808456166179000	0.025826697882300	-0.214566180811760	0.716748191024878	-0.310436381602400	0.372375320911000	0.377535777416000	-0.329318824188400	0.377535777416000	-0.329318824188400	-0.329318824188400	-0.329318824188400	0.377535777416000	0.377535777416000	0.377535777416000	0.377535777416000	0.377535777416000	0.377535777416000	0.377535777416000	0.377535777416000	0.377535777416000	
2	0	0.0717278348438000	0.1503884873427800	-0.194000258202300	0.260202548880400	-0.2345180511474100	-0.143801170548000	0.3027186810681320	0.3207148828119000	0.0967667028000	-0.248029677640250	-0.328447181894700	-0.327100178812180	-0.281471848694400	1.103716683800000	0.210738888451410	-0.882416379134480	-1.17054817058000	0.8086544532820				
3	0	0.071826111932470	-0.2476215533915340	-0.715419148621700	0.448309148895370	-0.15716022091410	-1.11731930035400	-0.42826808387100	0.716775374271600	0.7967056577700	-0.359447181894700	-0.327100178812180	-0.281471848694400	1.103716683800000	0.210738888451410	-0.882416379134480	-1.17054817058000	0.8086544532820					
4	0	0.0515423640678800	-0.198058838079400	-0.547487326490000	0.03252561886315100	-0.5470682739816100	-0.03252561886315100	-0.221713205478600	0.30819747401120	0.446671071210	-0.371784037888300	-0.330734508585800	-0.084855508200007	0.5867200319871830	-0.0517941448841000	-0.1273588176881870	-0.380248677090450	0.293778841728800					
5	0	0.41749128695327700	-0.1074807547712330	-0.2443363817028500	0.2309180271825100	-0.282836854021000	0.04294512754578000	0.1788515481787000	-0.2905688887619	0.3385546426582	-0.280232323716900	-0.334108198835400	-0.188248888832400	0.0152158023284110	-0.053776788328170	0.075475820407010	-0.044428874851950	0.198702515232180					
6	0	0.673058627732880	0.8909254336759	-0.4583843454501900	0.958937174280000	-0.4012291545813450	0.108477032000000	-0.190286981291120	-0.142630251969340	0.51411989008	-0.00391407814400	-0.710341451937180	-0.178845241587000	-0.253871703885820	-0.404289911628700	0.023252841123100	-0.173927010589300	-0.471388800854940					
7	0	0.273815713288810	0.020719071922400	-0.238044877612500	-0.020719071922400	-0.102891118835400	0.462026513188800	0.238250197195000	0.17887890071200	0.4666728871400	-0.288471037454000	0.5301256431984700	0.095757673448800	0.287817708682000	0.0230260554044900	0.028883866557000	0.424048853457800	0.424048853457800	0.424048853457800	0.424048853457800	0.424048853457800	0.424048853457800	
8	0	0.73738458623340	0.071705648586820	-0.2096869591480200	0.80781063420200	-0.2632918617725700	0.289287711584700	0.708426888414100	-0.320714882811900	0.1524368144471	-0.32426381045100	-0.4280336478300	-0.27888943083190	-0.27888943083190	-0.27888943083190	-0.27888943083190	-0.27888943083190	-0.27888943083190	-0.27888943083190	-0.27888943083190	-0.27888943083190	-0.27888943083190	
9	1	0.307724401908400	0.087288382758400	-0.523533847475810	0.0368778773078910	-0.40472818778917	-0.377273488987710	0.012584884181000	-0.34687032758370	0.3346230312001	-0.3247415462200	-0.36550878622890	-0.171108114714400	-0.221089322234800	-0.4748821028840	0.65617148916400	-0.2811941025040	0.0814711718430					
10	0	0.198830815171000	0.077428323290000	-0.5136203033587340	-0.392000195030240	-0.178305244885200	0.6402141132011	-0.380254647925100	0.13133478948328	0.3354423031200	-0.3247415462200	-0.36550878622890	-0.171108114714400	-0.221089322234800	-0.4748821028840	0.65617148916400	-0.2811941025040	0.0814711718430					
11	1	0.388492038788800	0.079721259486200	-0.27158684714760	0.04432839131100	-0.294028128789100	-0.18677258487470	0.0029716235206840	-0.350677028182960	0.025444072770	-0.3247415462200	-0.36550878622890	-0.171108114714400	-0.221089322234800	-0.4748821028840	0.65617148916400	-0.2811941025040	0.0814711718430					
12	1	0.388492038788800	0.079721259486200	-0.27158684714760	0.04432839131100	-0.294028128789100	-0.18677258487470	0.0029716235206840	-0.350677028182960	0.025444072770	-0.3247415462200	-0.36550878622890	-0.171108114714400	-0.221089322234800	-0.4748821028840	0.65617148916400	-0.2811941025040	0.0814711718430					
13	1	0.225305102968000	0.08642714841780	-0.488910088758670	0.382292628388800	-0.12277723934600	0.648035821609000	0.134487228654100	-0.588773023044100	0.4954888380504	-0.451858412454000	-0.38714984471880	-0.27111825345700	0.102237881202800	0.023918467133380	0.280713305119880	0.7828182183320	0.86858877028840					
14	1	0.20848478118880	-0.08888204112800	-0.028573210813820	-0.42844327784670	-0.43844282888880	0.308487108723780	-0.44738715887087	-0.1801182274810	0.1194477877805	-0.3715445403827750	-0.78014884810700	-0.35688545705790	-1.088140344000	-0.415001888810400	1.128844188847000	0.33884282027480	-0.70148844884840					
15	1	0.263711167579100	0.084547919682800	-0.11818274474180	0.494646278619400	-0.42072118129270	0.078807098888700	-0.424283514310000	0.078807098888700	-0.424283514310000	0.078807098888700	-0.424283514310000	0.078807098888700	-0.424283514310000	0.078807098888700	-0.424283514310000	0.078807098888700	-0.424283514310000					
16	0	1.386548743579100	0.325438357519100	-0.241548710889100	-0.32471313688900	-0.32471313688900	1.08071733709400	1.7586188848740	-0.41758888308354	0.494646278619400	-0.42072118129270	0.078807098888700	-0.424283514310000	0.078807098888700	-0.424283514310000	0.078807098888700	-0.424283514310000	0.078807098888700	-0.424283514310000				
17	1	0.102160811217000	0.1787578343100	-0.336804164455000	0.0941328309238	-0.1180533587481400	-0.0718677547342800	-0.34623552226600	0.410318613850	-0.188484884780	-0.102160811217000	-0.102160811217000	-0.102160811217000	-0.102160811217000	-0.102160811217000	-0.102160811217000	-0.102160811217000	-0.102160811217000					
18	0	0.18348574838887	0.251031575871280	-0.4004348471707200	0.242788486778100	-0.291822943354800	-0.172295259111822	0.003068848888888	-0.244884085100700	0.3818889243382	-0.244884085100700	-0.244884085100700	-0.244884085100700	-0.244884085100700	-0.244884085100700	-0.244884085100700	-0.244884085100700	-0.244884085100700					
19	1	0.306921100854400	0.29257914547410	-0.37620511024470	0.7104008269200	-0.266114524312840	0.53428415381870	0.073763454820900	-0.27194217077400	0.477759232600	-0.27194217077400	-0.27194217077400	-0.27194217077400	-0.27194217077400	-0.27194217077400	-0.27194217077400	-0.27194217077400	-0.27194217077400					
20	1	0.284469054547000	0.18241378237100	-0.02570655595480	0.9805891183800	-0.186847426271410	0.7827551918200	1.0487142848100	0.0487142848100	0.0487142848100	-0.0487142848100	-0.0487142848100	-0.0487142848100	-0.0487142848100	-0.0487142848100	-0.0487142848100	-0.0487142848100	-0.0487142848100					
21	1	0.087488518210800	-0.282071688234600	-0.451688951945600	1.731008014281100	-0.58438218243200	1.7822548619380	1.54297104876000	0.588714782728880	0.4474588483450	-0.4474588483450	-0.4474588483450	-0.4474588483450	-0.4474588483450	-0.4474588483450	-0.4474588483450	-0.4474588483450	-0.4474588483450					
22	1	0.137581083788800	0.1936363644820	-0.159995179163800	0.022296115760340	-0.23303108120700	0.24845828112026	0.21825526217710	-0.177841188236700	0.1520894893374	-0.1520894893374	-0.1520894893374	-0.1520894893374	-0.1520894893374	-0.1520894893374	-0.1520894893374	-0.1520894893374	-0.1520894893374					
23	1	0.020174768888800	0.10941324545430	-0.06331497744500	0.55872712141450	-0.85185432784780	0.090262304158170	0.67618168434000	-0.13378167862025	-0.13378167862025	-0.13378167862025	-0.13378167862025	-0.13378167862025	-0.13378167862025	-0.13378167862025	-0.13378167862025	-0.13378167862025	-0.13378167862025					
24	0	0.24274425089100	0.443891227468800	-0.285551142820570	0.495886776120000	-0.1914424830283100	0.111680257528490	0.107870170884940	0.254112061161400	0.443891227468800	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570					
25	0	0.514742424242400	0.75768683185840	-0.1973584184074000	0.102868868687100	-0.328587174738700	-0.102868868687100	-0.08060242286800	0.60020073723540	0.443891227468800	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570	-0.285551142820570					
26	1	0.054778472242350	-0.258218748549800	-0.595651844138100	0.387735545823230	-0.606438878174900	0.4585485580120	0.10400544488800	-0.5641518187800	0.03332788474071	-0.37786881701240	-0.40521542316440	-0.38845789134370	-0.27888131058860	0.299267915568400	-0.10184732187100	-0.84058881964110	0.7654938973250					
27	0	0.118844323257800	-0.024338779320900	-0.51308180500310	-0.31167819202030	-0.46108891308380	0.54418815340010	-0.28388084103790	-0.48989760370000	0.0287877732814	-0.38877888184200	-0.30278180876310	-0.171108114714400	0.0882828413877	-0.058810278880000	-0.05881027888000	-0.05881027888000	-0.05881027888000					
28	0	0.118844323257800	-0.024338779320900	-0.51308180500310	-0.31167819202030	-0.46108891308380	0.54418815340010	-0.28388084103790	-0.48989760370000	0.0287877732814	-0.38877888184200	-0.30278180876310	-0.171108114714400	0.0882828413877	-0.058810278880000	-0.05881027888000	-0.05881027888000	-0.05881027888000					
29	0	0.166388888888888	-0.234888888888888	-0.234888888888888	0.335555555555555	-0.234888888888888	0.038488888888888	0.105888888888888	0.105888888888888	0.105888888888888	-0.105888888888888	-0.105888888888888	-0.10										

target labels. We clipped any $p > 1$ as 1 and $p < 0$ as 0.

The second model is the logistic regression model. Logistic regression captures the natural shape of the cumulative Gaussian distribution functions and is a popular model for binary classification.

$$p = \frac{1}{1 + \exp(-(\beta_0 + \beta_1 Y_{i,1} + \beta_2 Y_{i,2} + \dots + \beta_{60} Y_{i,60}))}$$

Finally, I introduce simple three-layer neural networks. Neural networks serve as a good indicator of non-linearity in activations. If the accuracy of linear probability models and logistic regression is low while the neural networks regress well, there is a non-linear relationship between the activations and true labels. I applied the three models to each of the 300 .csv files and recorded the accuracy. I split 2,000 observations into training and test datasets.

6.2 Data Preprocessing

Most values of my data range from 0.3 and -0.3, which resulted in the poor convergence of the regression models. Therefore, I standardized the values. Standardization is given by $x_{std} = \frac{x - \mu}{\sigma}$ where μ is the mean of the independent variable X , which sample x belongs to, and σ is the standard deviation of the variable. In the follow-up research mentioned later, I also performed principal component analysis (PCA) to reduce the dimensions of the observations from 60 to 10. PCA essentially removes unnecessary, sparse representations and compresses data to lower dimensions while retaining most of the information contained in the original data.

7 Results

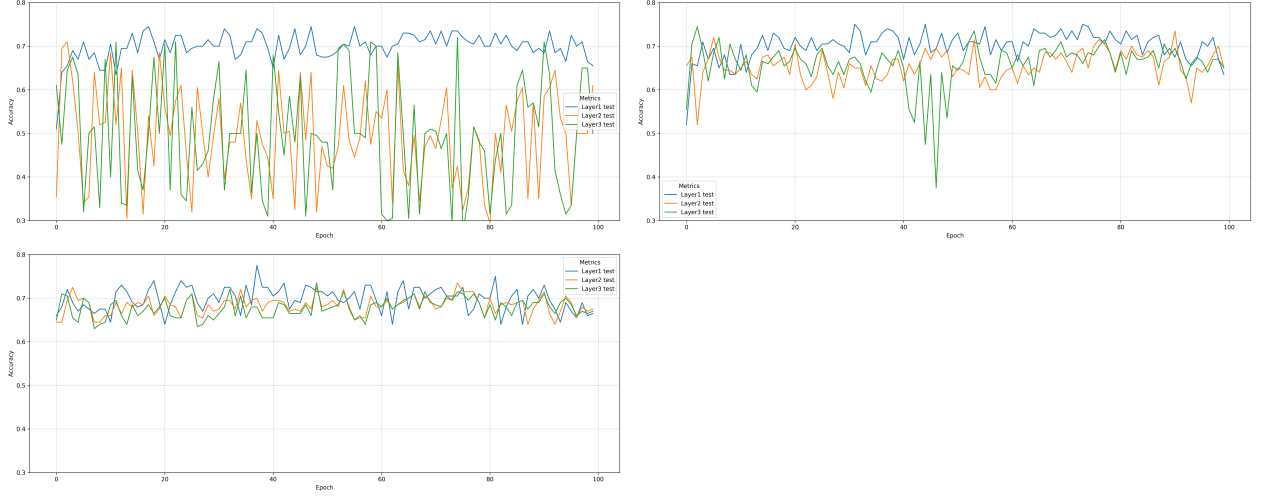
7.1 Probing Results

Figure 3 shows the results of probing. The linear probability models suffer from fluctuating accuracy. Logistic regression is more robust than linear probability models, but there still exists a large decline in accuracy. In contrast, neural networks regressed the true labels more stably. Moreover, the overall results show that the first layer contains more stable representations, while the activations from the second and third layers resulted in slightly lower or significantly lower accuracy in the logistic models and linear probability models. These graphs present counterintuitive results, as I initially speculated that representations would grow richer as the layer goes deeper, closer to the output layer. However, in Section 4, I identified the multicollinearity in the activation data and hypothesized that it degraded the accuracy of the probes.

7.2 Probing Results with PCA

Speculating that the high multicollinearity contributes to the low accuracy with the second and third layers, I performed PCA to reduce dimensionality. PCA mitigates multicollinearity by creating new variables that are linearly unrelated to each other [16]. Probing results

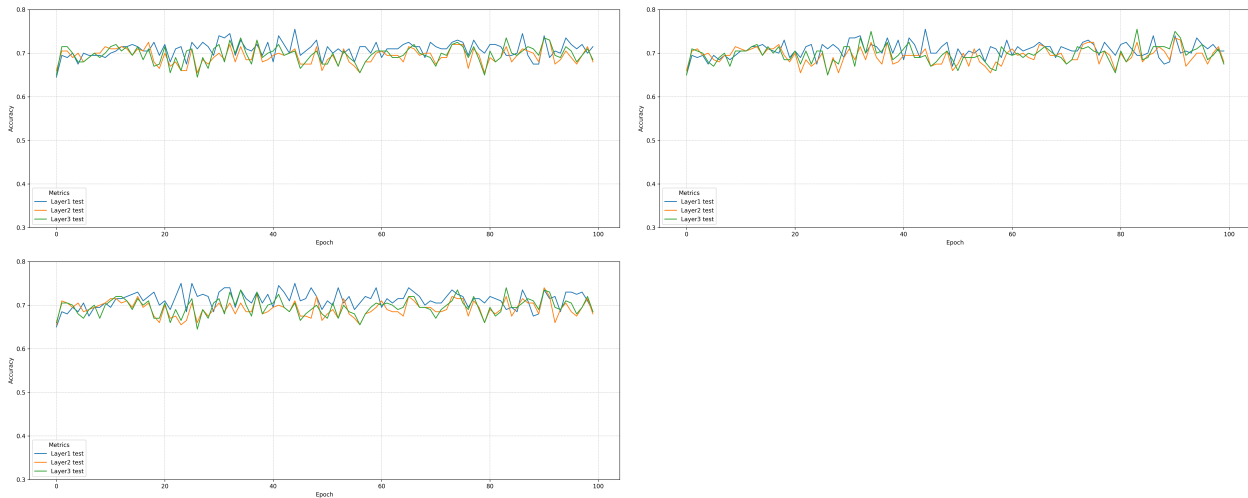
Figure 3: Accuracy of Probes on Test Data



The upper-left graph is the linear probability models, the upper-right is the logistic models, and the bottom-left is the neural networks. The blue lines represent the scores of probes on the activations from the first layer, the orange lines correspond to the second layer, and the green lines track that of the third layer.

after PCA are outlined in Figure 4. I can observe clear improvements in the accuracy scores across the three probes, especially in the linear probability models. Furthermore, the accuracy scores across different layers generally stay at almost the same level, indicating that activations from the first, second, and third layers possess equally rich information. Moreover, the stable accuracy scores of the linear probability models indicate that there is a solid linear relationship between the activations and true labels.

Figure 4: Accuracy of Probes on Test Data after PCA



The upper-left graph is the linear probability models, the upper-right is the logistic models, and the bottom-left is the neural networks.

8 Conclusion

In this study, I investigated how the internal representations of neural networks in binary image classification change depending on the depth of layers and training amount. I discovered that the layers later in the network pipeline possess representations with higher multicollinearity. However, the multicollinearity in every layer declines and converges to certain points of value as the amount of training increases. Further regression analysis on the data after PCA revealed that the representations from different layers possess the equivalent level of information about the true binary labels.

My research indicates that probing techniques are a suitable approach to understand *how AI thinks* by quantifying what kind of representations it acquires. The results also present answers to the research questions: (1) every layer contains equally rich information, and (2) the number of training steps does not affect accuracy or heteroskedasticity but makes multicollinearity converge.

This research contributes to the further understanding of deep neural network models not only through probing but also via statistical analysis. Limitations posed to this work are that I solely focused on neural networks for binary image classification. I also failed to replicate similar results with Alain and Bengio [3], who argued that deeper layers contain richer representations in binary image classification. Thus, further studies should be done to apply my analysis methods to neural networks with broader tasks and perform stricter replication studies.

References

- [1] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- [2] Mohammad Sadegh Akhondzadeh, Vijay Lingam, and Aleksandar Bojchevski. Probing graph representations. In Francisco Ruiz, Jennifer Dy, and Jan-Willem van de Meent, editors, *Proceedings of The 26th International Conference on Artificial Intelligence and Statistics*, volume 206 of *Proceedings of Machine Learning Research*, pages 11630–11649. PMLR, 25–27 Apr 2023.
- [3] Guillaume Alain and Yoshua Bengio. Understanding intermediate layers using linear classifier probes, 2018.
- [4] Vy Anh Tran, Le Thi Nhu Quynh, Thu-Thao Thi Vo, Phuc An Nguyen, Ta Ngoc Don, Yasser Vasseghian, Hung Phan, and Sang-Wha Lee. Experimental and computational investigation of a green knoevenagel condensation catalyzed by zeolitic imidazolate framework-8. *Environmental Research*, 204:112364, 2022.
- [5] Yonatan Belinkov. Probing classifiers: Promises, shortcomings, and advances. *Computational Linguistics*, 48(1):207–219, 04 2022.
- [6] Prashant Gohel, Priyanka Singh, and Manoranjan Mohanty. Explainable ai: current status and future directions, 2021.
- [7] Justin Grimmer, Margaret E. Roberts, and Brandon M. Stewart. Machine learning for social science: An agnostic approach. *Annual Review of Political Science*, 24, 2021.
- [8] Wes Gurnee and Max Tegmark. Language models represent space and time, 2024.
- [9] Gabriel Ilharco, Rowan Zellers, Ali Farhadi, and Hannaneh Hajishirzi. Probing contextual language models for common ground with visual representations. In Kristina Toutanova, Anna Rumshisky, Luke Zettlemoyer, Dilek Hakkani-Tur, Iz Beltagy, Steven Bethard, Ryan Cotterell, Tanmoy Chakraborty, and Yichao Zhou, editors, *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 5367–5377, Online, June 2021. Association for Computational Linguistics.
- [10] Yacine Izza, Alexey Ignatiev, and Joao Marques-Silva. On explaining decision trees, 2020.
- [11] Alex Krizhevsky. pages 32–33.
- [12] Michael B. McCoy and Anna Leshinskaya. Relational composition during attribute retrieval in GPT is not purely linear. In *NeurIPS 2024 Workshop on Compositional Learning: Perspectives, Methods, and Paths Forward*, 2024.
- [13] L.E. Melkumova and S.Ya. Shatskikh. Comparing ridge and lasso estimators for data analysis. *Procedia Engineering*, 201:746–755, 2017. 3rd International Conference “In-

formation Technology and Nanotechnology”, ITNT-2017, 25-27 April 2017, Samara, Russia.

- [14] Jerry Ngo and Yoon Kim. What do language models hear? probing for auditory representations in language models, 2024.
- [15] Saeed Nosratabadi, Amirhosein Mosavi, Puhong Duan, Pedram Ghamisi, Ferdinand Filip, Shahab S. Band, Uwe Reuter, Joao Gama, and Amir H. Gandomi. Data science in economics: Comprehensive review of advanced machine learning and deep learning methods. *Mathematics*, 8(10), 2020.
- [16] Lexi V Perez. Principal component analysis to address multicollinearity. *Whitman College: Walla Walla, WA, USA*, 99362, 2017.
- [17] Cyrus Rashtchian, Charles Herrmann, Chun-Sung Ferng, Ayan Chakrabarti, Dilip Krishnan, Deqing Sun, Da-Cheng Juan, and Andrew Tomkins. Probing the equivariance of image embeddings. In *NeurIPS 2023 Workshop on Distribution Shifts: New Frontiers with Foundation Models*, 2024.