
Dynamic Reaction Time Modeling with Deep Learning

JD Pruett
Symbolic Systems
Stanford University

1 Introduction

Cognitive neuroscience has traditionally relied on so-called "handcrafted" models of cognition. These models are intuitive and interpretable, but often lack explanatory power of human behavior.

Deep learning is emerging as a distinct methodological paradigm from the handcrafted model approach. If a computational model can exhibit patterns similar to human cognition, this suggests a way that humans and computers might process a task similarly. Since computational models are easier to perturb and analyze than the human mind, realistic computational models of human behavior provide a promising new medium to understand cognitive processes.

In this paper, we apply a deep-learning framework to model reaction times on a sequence of cognitive tasks specifically designed to investigate the interplay between cognitive "flexibility" and "stability."

Cognitive flexibility and stability are both adaptive behaviors particularly relevant to the study of cognitive control and clinical psychology. Stability helps in maintaining current goals and shielding them from distractions, while flexibility allows for switching between tasks or adjusting to new information.

For example, individuals with attention-deficit/hyperactivity disorder often exhibit high variability in reaction times and task performance, indicating a challenge in maintaining stability (Nassar & Troiani, 2021). On the other hand, individuals with Autism Spectrum Disorder often show a preference for stability over flexibility, leading to difficulties in adapting to changes (Nack & Yu-Chin, 2023). This rigidity can result in struggles with social interactions, transitions, and responding to unexpected events.

We implement two deep-learning architectures: a Long Short-Term Memory (LSTM) Recurrent Neural Network (RNN) and a Dynamic Variational Autoencoder (DVAE). Our analysis includes exploring the latent-space representations of the DVAE and examining subject-specific tradeoffs between cognitive flexibility and stability. Both models take as input a sequence of trials from human subjects performing a cognitive task. Each trial is characterized by features indicating whether it demands cognitive stability or flexibility. The output is the predicted human reaction time for each trial.

2 Background and Related Work

The concept of a "stability-flexibility tradeoff" has been a central tenet in neuroscience, suggesting that improvements in cognitive flexibility come at the cost of stability and vice versa (Allport et al., 1994; Goschke, 2013; Braem & Egner, 2018). Essentially, the mind's ability to switch between tasks and adapt to changing environmental stimuli trades off with its ability to maintain focus on a single task amidst distractions.

Last year, researchers at Stanford University presented Task-DyVa (Jaffe et al., 2023), which reproduced subject-specific sequences of response times with high temporal precision. An analysis of the model’s dynamics revealed that the different regions of latent space represented the two kinds of tasks - those that required stability and others that required flexibility. The model’s ability to differentiate between tasks in its latent space reflects the stability-flexibility tradeoff in cognitive control. This tradeoff suggests that individuals must balance maintaining current task goals (stability) with adapting to changing demands (flexibility), often leading to slower reaction times when switching between tasks.

However, recent work by Geddert & Egner (2022) challenged this long-standing notion. They presented evidence that the cognitive processes mediating stable and flexible tasks might be *independent*. So while previous models dictated that improvements in stability came at the expense of reduced flexibility, their data did not support this tradeoff.

Importantly, Task-DyVa was built using a large dataset comprised of thousands of hours of gameplay from Lumosity, a brain-training app. This is substantially more data than can be reasonably collected in traditional cognitive task batteries.

In this paper, we use a deep-learning approach to contribute to the cognitive stability-flexibility tradeoff debate. Specifically, we built a set of deep learning models optimized for the Geddert & Egner (2022) cognitive data. If reaction times are longer on so-called "switch" trials, this would be support of a tradeoff.

This paper makes two important advancements: First, we provide the first-ever test of whether a deep learning model can accurately recreate human-like reaction time patterns where stability and flexibility are *independent*. Second, by training these models on data specifically designed for studying cognition—rather than large, repurposed datasets—we assess whether deep learning can effectively capture the complex dynamics of human cognitive behavior using data volumes that align with the typical scale of cognitive testing batteries.

Past computational models primarily include recurrent neural networks (RNNs) and dynamical systems approaches. Musslick et al. (2018) used RNNs to show how a meta-control parameter influences cognitive stability and flexibility by manipulating network activation gains. Dynamical approaches, as demonstrated by Jaffe et al. (2024), conceptualize control signals as attractors, where control constraints affect attractor depth, enhancing flexibility but impairing stability.

While these models support the traditional tradeoff view, emerging models like the Dual-Dimension Framework (Nack & Yu-Chin, 2024) advocate for independent regulation of cognitive processes. Our work contributes to this debate by leveraging deep learning on cognition-specific data, aiming to further explore the potential independence of cognitive stability and flexibility.

3 Dataset and Features

3.1 Cognitive Task Structure

We will be using data from the Geddert & Egner study [2], publicly available at <https://github.com/rmgeddert/stability-flexibility-tradeoff>. The raw data consists of 65 subjects performing a series of trials, totaling $\sim 40,000$ trials in the experiment across all subjects.

Participants in the experiment completed a classic task-switching paradigm designed to independently assess cognitive stability and flexibility. On each trial, participants categorized a single-digit number (1–9, excluding 5) according to either its parity (odd/even) or its magnitude (greater/less than 5). The task was determined by the color of a rectangular frame surrounding the number, with red and blue frames randomly assigned to indicate the parity and magnitude tasks across participants. Responses were made using two keys ("Z" and "M"), with the category-response mapping randomized between participants (e.g., "Z" for odd numbers and "M" for even numbers or vice versa).

The stimuli were designed to vary in congruency, depending on whether the response required for the parity task matched the response required for the magnitude task. Trials were labeled congruent if both tasks required the same keypress (e.g., the number 3, if "Z" was mapped to both "odd" and "less than 5") and incongruent if the tasks required different keypresses (e.g., the number 8, if "Z" was mapped to "odd" and "less than 5"). The congruency effect, defined as the difference in response time and accuracy between incongruent and congruent trials, served as a measure of cognitive stability.

Smaller congruency effects reflected greater stability, characterized by a better-shielded task focus that minimized interference from task-irrelevant stimulus features.

The task structure also included a task-switching component, with trials categorized as switch trials if the task changed from the previous trial (e.g., a parity task followed by a magnitude task) or repeat trials if the task remained the same. The switch cost, defined as the difference in response time and accuracy between switch and repeat trials, provided a measure of cognitive flexibility. Smaller switch costs reflected a greater readiness to switch tasks, indicative of a more flexible cognitive state. The relative frequencies of congruent/incongruent and switch/repeat trials were independently manipulated between blocks (25% or 75% for each condition).

4 Methods

We experimented with various deep learning models and compared it to a baseline linear regression model. The models we implemented included a classic Neural Network, a Recurrent Neural Network, and a Dynamical Variational Autoencoder.

4.1 Classic Neural Network

For our classic NN approach, we used a 6-layer neural network with ReLU nonlinear activations between each layer except for the output layer which outputs a single numerical value. The model's input layer accepts six features representing task-specific information such as the current and previous trial's task cues, task types, stimulus identities, and the task-switching event. The five hidden layers comprise of 64, 128, 64, 32, and 16 neurons, respectively. The loss function used during training was the mean squared error loss for regression between the predicted and true response times. During training, the model uses the Adam Optimizer for its gradient descent, running with a learning rate of 0.0001 for 1000 epochs. We also clipped gradients to a maximum norm of 1.0.

4.1.1 Recurrent Neural Network

Our RNN method processed 10 timesteps of the cognitive experiments for each subject at a time in order to encode any immediate and longer-term reaction dependencies from successive trials. The input layer takes in five features representing the stimulus, task cue, switch and congruency category, as well as accuracy of the user response. It comprises four hidden LSTM layers with 128, 64, 32, and 16 neurons, respectively, each utilizing ReLU activation functions. The RNN also includes a dropout layer after the first hidden layer of 128 neurons to increase regularization and discourage overfitting. The final output layer is a TimeDistributed Dense layer that outputs a single value predicting the response time for each of the timesteps in the input sequence of 10 trials. The model was trained with the Adam optimizer (learning rate = 0.0001) and mean squared error as the loss function, over 100 epochs until convergence.

4.1.2 Dynamical Variational Autoencoder

Dynamic Variational Autoencoders (DVAEs) are an extension of the standard variational autoencoder (VAE) architecture, specifically designed to model temporal sequences and dynamic data. They merge the generative capabilities of VAEs with time-series modeling to capture the evolving nature of sequential data.

The scope of this paper does not permit a comprehensive discussion of VAEs, but the primary components are an encoder and a decoder. The encoder maps high-dimensional input data to a lower-dimensional latent space, where the data is represented as a probabilistic distribution. The decoder maps these latent variables back to the high-dimensional space, reconstructing the original input data.

Our training set consists of temporal sequences of cognitive tasks, where the response on each trial is influenced by previous cognitive states and responses. DVAEs are well-suited for modeling such sequences, capturing how cognitive states evolve over time. The probabilistic nature of the latent space in DVAEs enables a compact representation of cognitive states, making it possible to distinguish between those requiring stability and those requiring flexibility. Additionally, DVAEs' generative capability allows for the synthesis of new sequences that reflect learned patterns of cognitive behavior,

enabling simulations and predictions of human cognitive responses across different task conditions, and offering valuable insights into the stability-flexibility tradeoff.

Our DVAE architecture uses a multi-layer LSTM network to capture the complex interdependencies between successive trials. The Encoder network is implemented as a multi-layer LSTM that maps input sequences to a probabilistic latent space. Specifically, the encoder processes input tensors through a two-layer LSTM with a hidden dimension of 128 and a dropout rate of 0.2 to mitigate overfitting. The network employs linear projection layers to generate both a mean vector and a log-variance vector in the latent space. To ensure consistent processing, the encoder initializes zero-state hidden and cell states for the LSTM, allowing for robust handling of variable-length input sequences.

The Decoder network complements the encoder by reconstructing the original input from the sampled latent representation. It employs an identical LSTM configuration of two layers with 128 hidden units and a 0.2 dropout rate, culminating in a linear output layer that projects the final LSTM hidden state to the desired output dimension.

This approach enables the model to learn how factors such as task-switching, stimulus congruency, and previous trial accuracy contribute to variations in reaction time. The reparameterization technique allows for stochastic sampling of these latent representations, providing insight into the probabilistic nature of cognitive processing.

The model’s objective function integrates two critical components: a Mean Squared Error (MSE) reconstruction loss and a Kullback-Leibler (KL) divergence regularization term. This combined loss function encourages the learned latent distribution to approximate a standard Gaussian prior while minimizing reconstruction error. Key hyperparameters were carefully selected, including a latent dimension of 64, four LSTM layers, and an Adam optimizer with a learning rate of 0.001.

The training procedure follows a standard variational autoencoder approach, alternating between training and validation phases across 100 epochs. During each iteration, the model minimizes the combined reconstruction and KL divergence loss, systematically tracking both training and validation performance. This approach enables the model to learn a compressed, probabilistic representation of input sequences while maintaining the capability to reconstruct the original human subject reaction time with minimal information loss.

5 Results/Discussion

5.1 Quantitative Results

For each of our deep learning experiments, we used a train-test split of 90-10 for tuning and evaluating our models. We left out a third dev/test set split from our original data due to the limited number of samples (less than 40,000 total trials). We observed a few quantitative regression metrics from each model including the final mean squared error between the model’s predicted response times and the actual user response times for the test set and the correlation coefficient between the two sets of response times. Table 1 presents these metrics from our various experiments. During training, we also adjusted our learning rate and number of epochs based on the training and validation loss progression and with early stopping based on convergence of the average losses. Figure 1 shows the loss progression during training for the RNN model; we limited the optimization to only 100 epochs due to its convergence.

Table 1: Results of our experiments

Model	Test Loss	Correlation Coefficient
Linear Regression	0.053	0.18
Classic NN	0.045	0.28
RNN	0.046	0.27
DVAE	0.048	0.27

From the results above, we can see that all our experiments achieved roughly the same level of performance in predicting the user’s response times. The low correlation between our models’ predictions and the actual data could be due to the difficulty of encoding the user response times

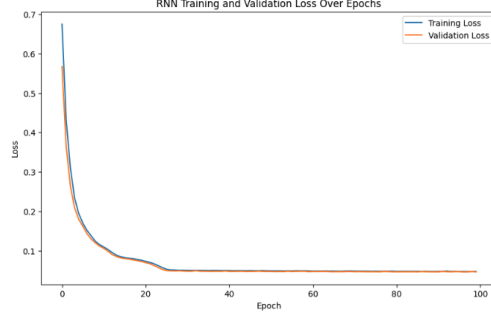


Figure 1: Recurrent Neural Network Training Loss

based off of just the task stimuli features available from the Geddert data. However, out of the models we implemented, the classic NN seemed to have the best performance with the lowest test loss of 0.045 and correlation coefficient of 0.28. A reason for this could be that only using information from the previous task and the current task was sufficient for our model to achieve its best performance and the simpler nature of the classic NN compared to the RNN and DVAE could have actually worked better due to the limited features available.

5.2 Qualitative Results

As outlined in the introduction, this paper aimed to achieve two primary objectives: first, to determine whether a deep learning model could accurately predict human reaction times using a dataset typical of cognitive task studies, and second, to explore the characteristics of the latent representation when the conventional stability-flexibility tradeoff is not present.

To address these objectives, we conducted several qualitative analyses. First, we investigated whether the model demonstrated a tradeoff between stability and flexibility, represented by the switch cost and congruency effect (as defined in the methods). Figure 2 shows the scatter plots for the two metrics where each data point represents one subject and the x-axis is the actual data and the y-axis is the predicted model outputs. Although our models faced difficulties in fully encoding the input features to predict accurate response times, there does seem to be a slight positive correlation between the actual and predicted switch costs and congruency effects as shown in the plots.

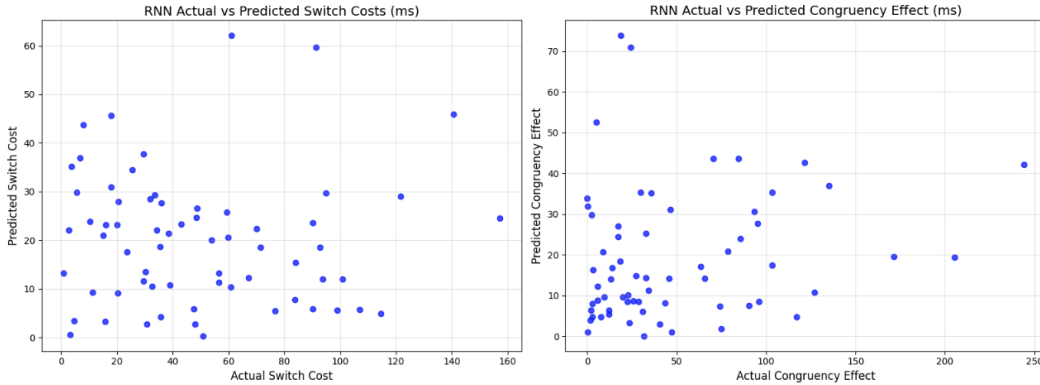


Figure 2: Per Subject Switch Costs and Congruency Effects

Next, we examined the latent space representation of the DVAE. As described in the methods, the structure of the DVAE is designed to learn latent representations of different task features. Figure 3 visualizes the latent representation of various task types, where different combinations of task features assess stability and flexibility independently.

Consistent with the findings of Jaffe et al. (2024), the latent representations of different feature sets—corresponding to distinct adaptive behaviors—are separated in latent space, particularly for

Latent state trajectories by switch type and stimulus congruency

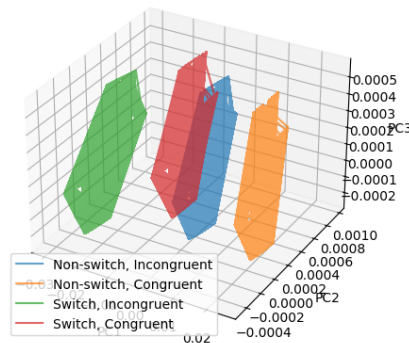


Figure 3: Latent Representations of Different Tasks from DVAE

switch versus non-switch trials. However, given the model’s relatively low accuracy in predicting reaction times, we caution against over-interpreting this as evidence supporting the stability-flexibility tradeoff. More accurate predictions would provide stronger support for this hypothesis.

5.2.1 Future Work

While our model made strides in exploring the stability-flexibility tradeoff, it struggled to predict reaction times with the accuracy necessary for a definitive analysis. This limitation points to the need for more advanced approaches in future studies. One promising direction involves creating individualized models for each subject, as demonstrated by Jaffe et al. (2024). By tailoring the model to each participant’s specific cognitive profile, we may improve prediction accuracy and capture subject-specific variations in cognitive performance more effectively. Furthermore, incorporating additional data, such as individual task histories or neural correlates, could enhance the model’s ability to account for the complex dynamics of cognitive processes. These refinements would provide a clearer understanding of the stability-flexibility tradeoff and help bridge the gap between computational models and human cognition in cognitive neuroscience.

6 Contributions

We had our third team member drop out before the final report was due and we changed our approach a lot from the milestone. So all of the final modeling was completed equally by JD and Shannon. JD wrote the baseline linear regression and DVAE and did the latent analysis. Shannon implemented the classic NN and RNN and did the analysis for those models. JD and Shannon both contributed to the data preprocessing and writing of the final report equally.

References

All code is available at: <https://github.com/jdpru/CS-230-DynaRT>

- [1] Jaffe, P. I., Poldrack, R. A., Schafer, R. J., & Bissett, P. G. (2023). Modelling human behaviour in cognitive tasks with latent dynamical systems. *Nature Human Behaviour*. <https://doi.org/10.1038/s41562-022-01510-8>
- [2] Gedder, R., & Egner, T. (2022). No need to choose: Independent regulation of cognitive stability and flexibility challenges the stability-flexibility trade-off. *Journal of Experimental Psychology: General*. <https://doi.org/10.1037/xge0001241>
- [3] Bi, Z., & Zhou, C. (2020). Understanding the computation of time using neural network models. *Proceedings of the National Academy of Sciences of the United States of America*, 117(19), 10530–10540. <https://doi.org/10.1073/pnas.1921609117>

- [4] Yen, P. Y., Kelley, M., Lopetegui, M., Rosado, A. L., Migliore, E. M., Chipps, E. M., & Buck, J. (2017). Understanding and Visualizing Multitasking and Task Switching Activities: A Time Motion Study to Capture Nursing Workflow. *AMIA Annual Symposium Proceedings*, 2016, 1264–1273.
- [5] Allport, D. & Styles, Elizabeth & Hsieh, Shulan. (1994). Shifting attentional set – Exploring the dynamic control of tasks. *Attention and Performance XV*. Vol. XV.
- [6] Braem, S., & Egner, T. (2018). Getting a grip on cognitive flexibility. *Current directions in psychological science*, 27(6), 470–476. <https://doi.org/10.1177/0963721418787475>
- [7] Goschke, T. (2013). Volition in action: Intentions, control dilemmas, and the dynamic regulation of cognitive control.
- [8] Yamashita, A., Rothlein, D., Kucyi, A. et al. Variable rather than extreme slow reaction times distinguish brain states during sustained attention. *Sci Rep* 11, 14883 (2021). <https://doi.org/10.1038/s41598-021-94161-0>
- [9] Nassar, M.R., Troiani, V. The stability flexibility tradeoff and the dark side of detail. *Cognitive, Affective, and Behavioral Neuroscience* 21, 607–623 (2021). <https://doi.org/10.3758/s13415-020-00848-8>
- [10] Madlon-Kay, S., Pearson, J., & Egner, T. (2024). Optimal decision-making under task uncertainty: a computational basis for cognitive stability versus flexibility. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 46. Retrieved from <https://escholarship.org/uc/item/2d40418r>
- [11] Nack, C. & Yu-Chin, C. (2023). Cognitive flexibility and stability at the task-set level: A dual-dimension framework. *advances.in/psychology*, 1(1), 1-28. <https://doi.org/10.56296/aip00007>