



Introduction to Computational Cognitive Modeling

Ti-Fen Pan

UC Berkeley

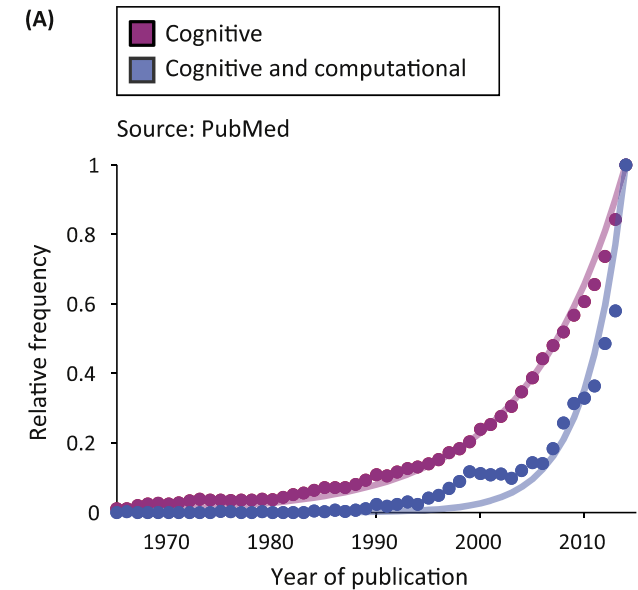
SBI summer school, July 27-31, 2026

Slide credits to Anne Collins



Computational modeling is useful

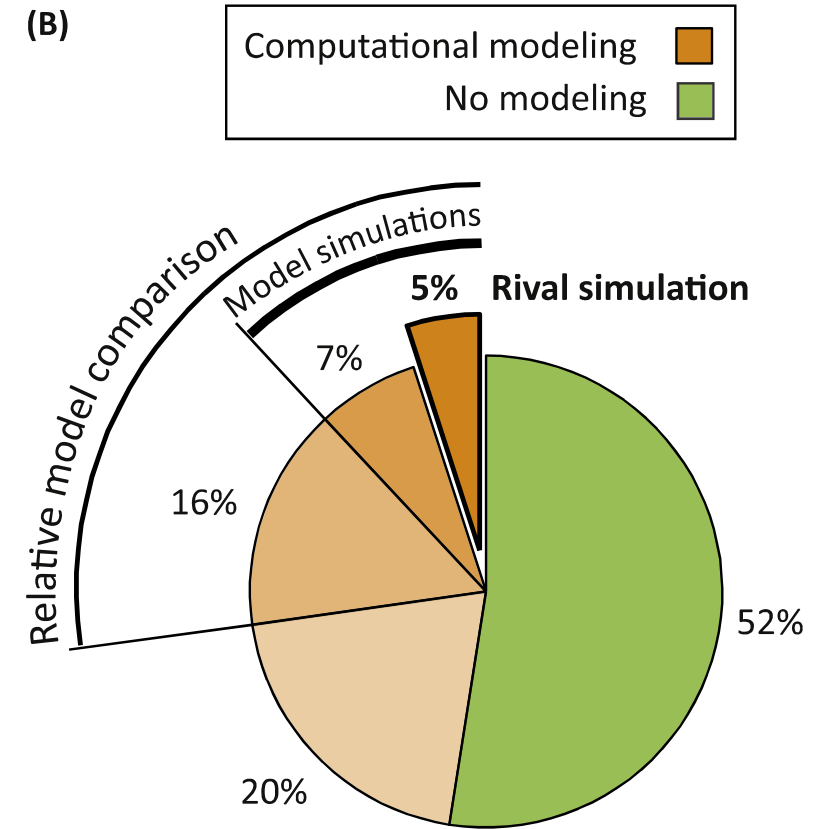
- Precise, testable quantitative predictions
- Descriptive, but also generative/mechanistic explanations
 - Cognitive scientists and cognitive neuroscientists
- Relate neural function to mechanisms giving rise to behavior
 - Cognitive and behavioral neuroscientists
- Summarize individual differences into interpretable variables
 - Development, computational psychiatry



Palminteri et al, 2017

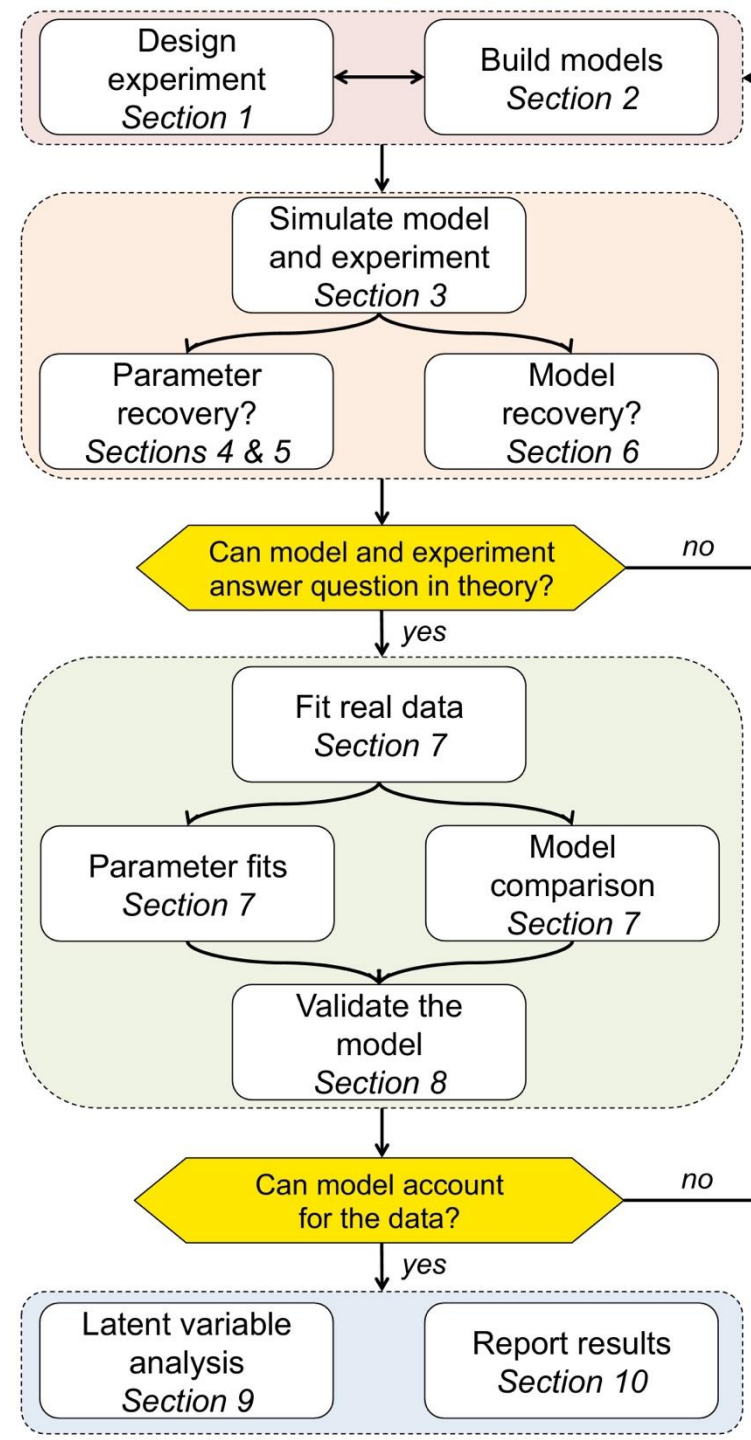
Computational modeling can be dangerous

- It's important to understand the basics very well and all the potential pitfalls (the goal of this talk!).
- More bad than good examples in the literature. Be careful!

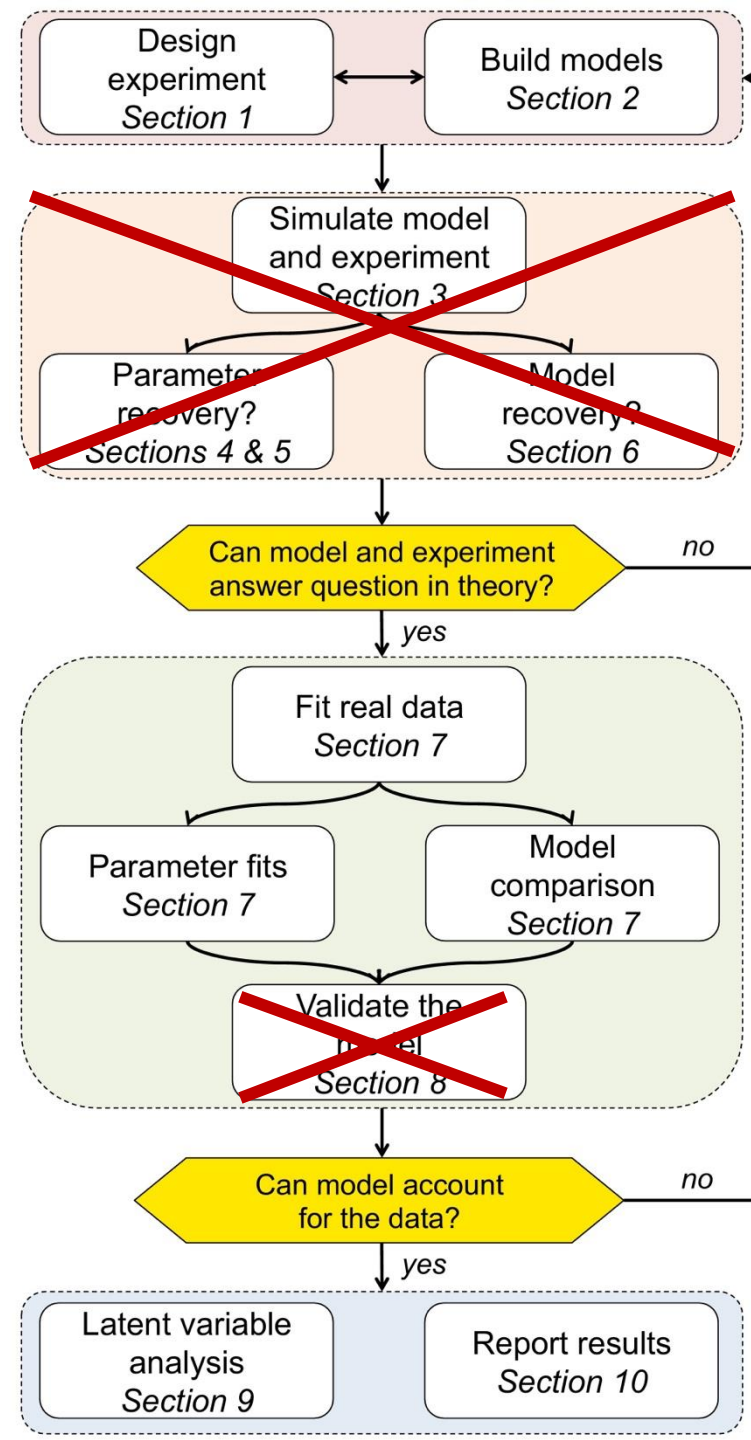


Palminteri et al, 2017

Life cycle of a modeling project



Life cycle of a modeling project



Steps

1. Data analysis
2. Model simulation
 - Role of model parameters
 - Setting up models representing competing hypotheses
3. Parameter estimation (model fitting)
4. Model comparison
5. Model validation
6. Improving models

Data analysis and model simulation

- First step: data analysis. Know your data set!
 - What questions are you trying to ask?
 - What questions can you answer in a model-independent way?
 - What questions is your protocol not designed to answer?
 - How do participants behave?

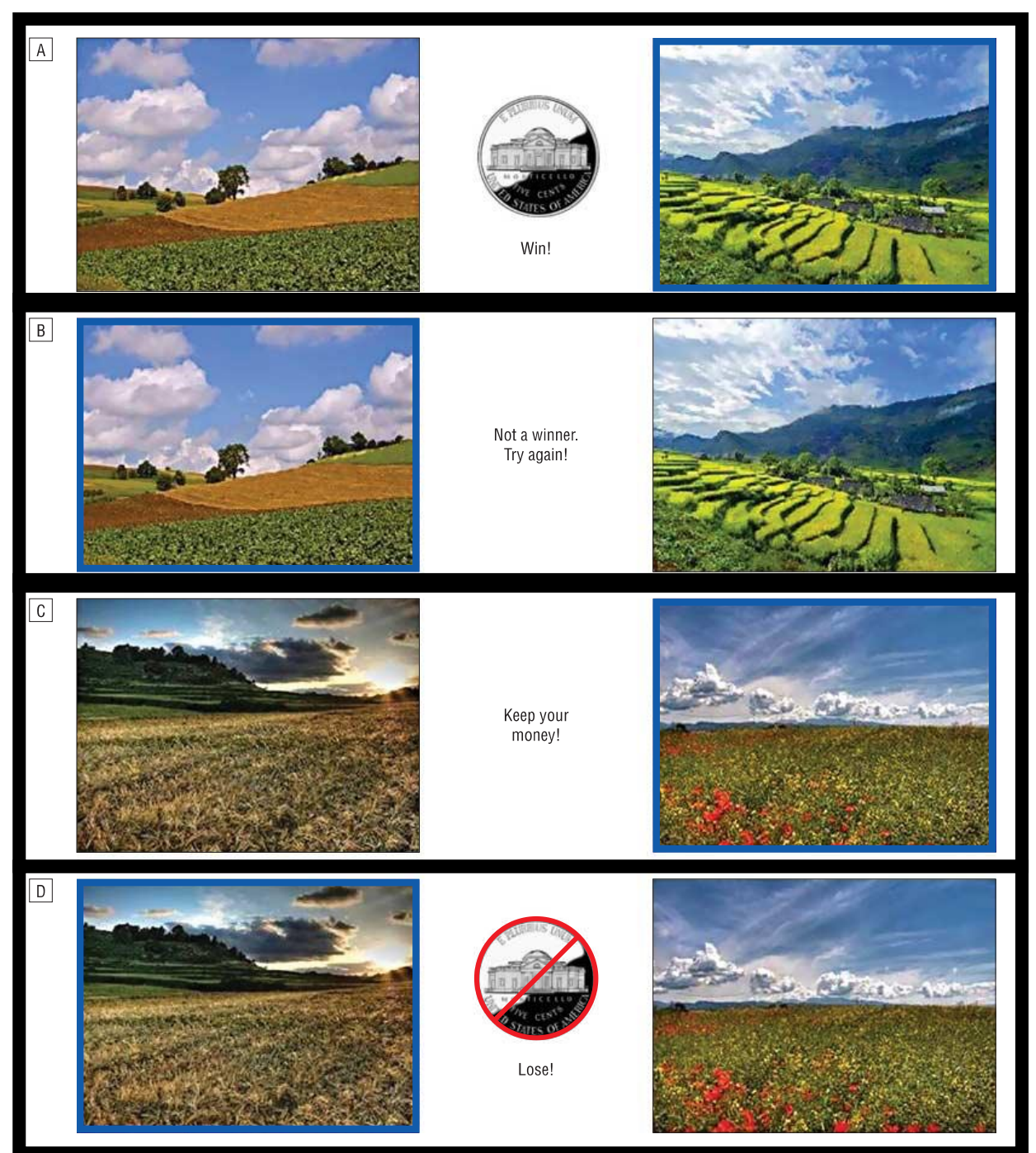
RL example

Two uncertainty levels

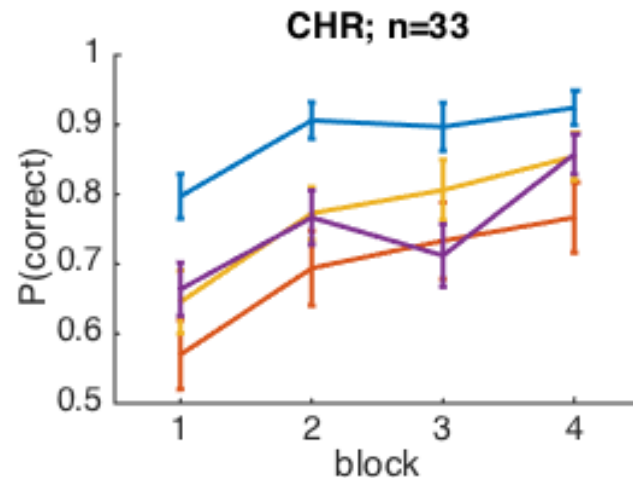
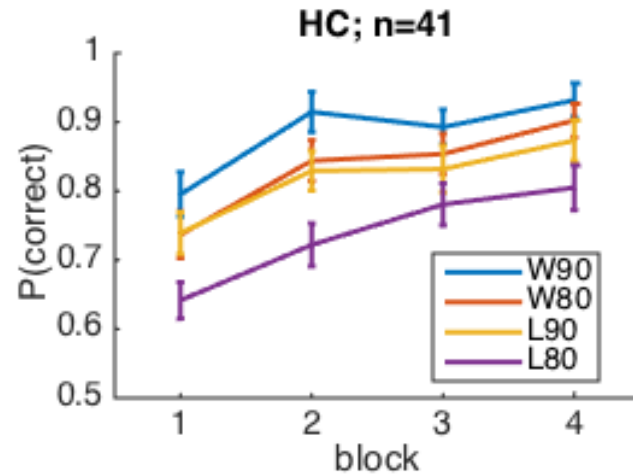
	90%	80%
Win:		
	AB	CD
\$0.05 vs. 0		
Lose		
	EF	GH
0 vs. -\$0.05		

- All conditions interleaved
- 4 blocks of $4 \times 10 = 40$ trials
- 41 youths at risk for schizophrenia
- 33 matched controls

Waltz et al 2012



Preliminary behavioral results, modeling targets



- **Question:** Do we learn differently from **seeking gains vs. avoiding losses**? Do youths at risk for schizophrenia learn differently from matched controls?
- Behavior:
 - Learning from **losses** is slower than learning from **gains**
 - Effect of **uncertainty** on learning is different in the gain and loss domains for High Risk individuals
 - → two modeling targets

2. model simulation

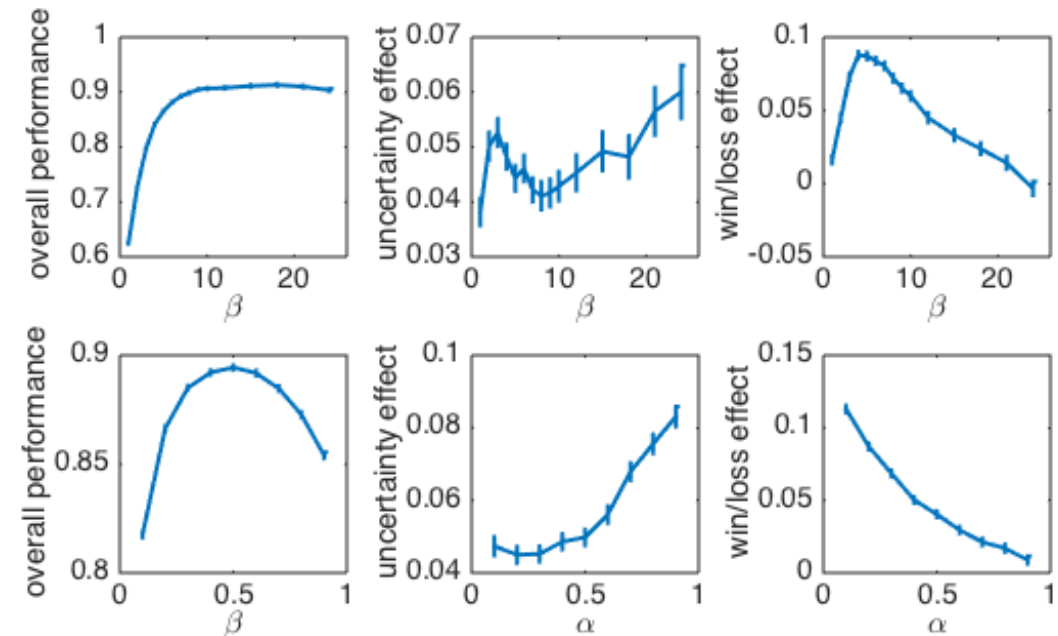
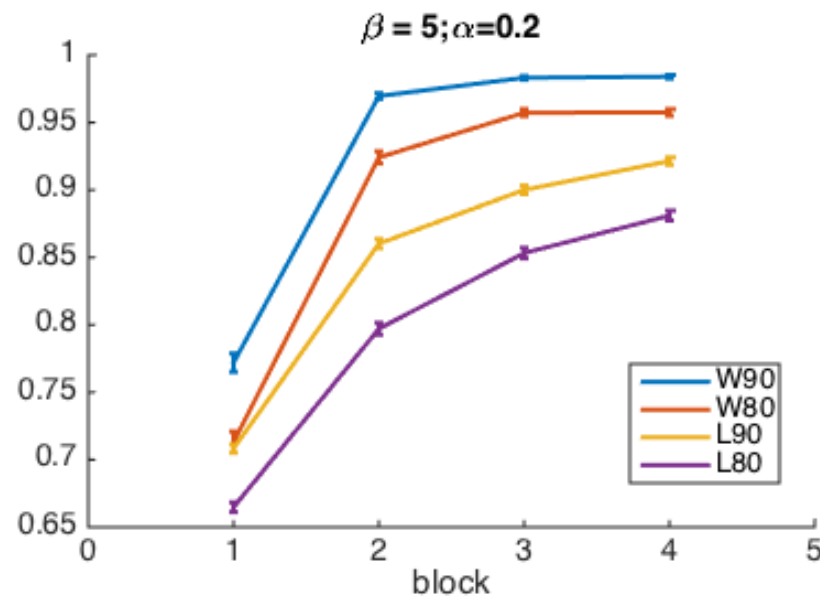
- A: model building
 - Develop a model that represents your main hypothesis
 - Develop competing models

$$\begin{aligned} Q_{t+1} &\leftarrow Q_t + \alpha \text{ RPE} \\ \text{Choice} &\sim \exp(\beta Q) \end{aligned}$$

- B: model simulation
 - Simulate your models **in the exact experimental environment**
 - Compare qualitative predictions of different models
 - Know the influence of all model parameters on model behavior

2. model simulation

- B: model simulation
 - Simulate your models **in the exact experimental environment**
 - Compare qualitative predictions of different models
 - Know the influence of all model parameters on model behavior



Steps

1. Data analysis
2. Model simulation
 - Role of model parameters
 - Setting up models representing competing hypotheses
3. Parameter estimation (model fitting)
4. Model comparison
5. Model validation
6. Improving models

3. Relate model to data – model fitting

- Bayes Rule!

$$P(\vartheta|D, M) \propto P(D|\theta, M)P(\theta|M)$$

Posterior probability
of the parameters

Likelihood of the data

Prior on model parameters

- What we want to know:
 - How likely is it that the data was generated with parameters (α^*, β^*) ?
- What we can compute:
 - How likely is the data given model assumptions and parameters (α^*, β^*) ?

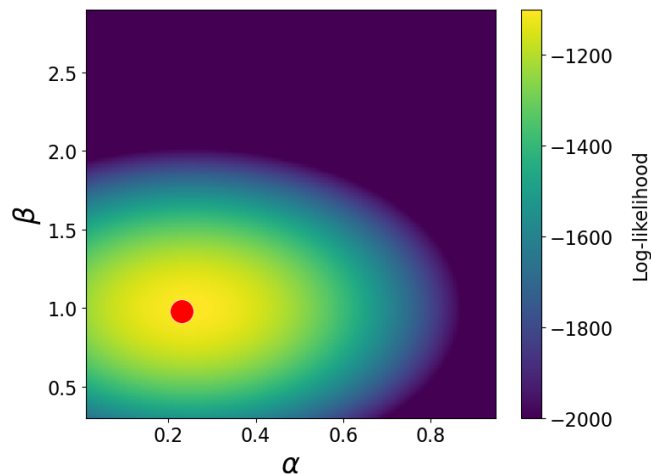
The theory –a simple trick

- Likelihood is a tiny number: product of many numbers < 1

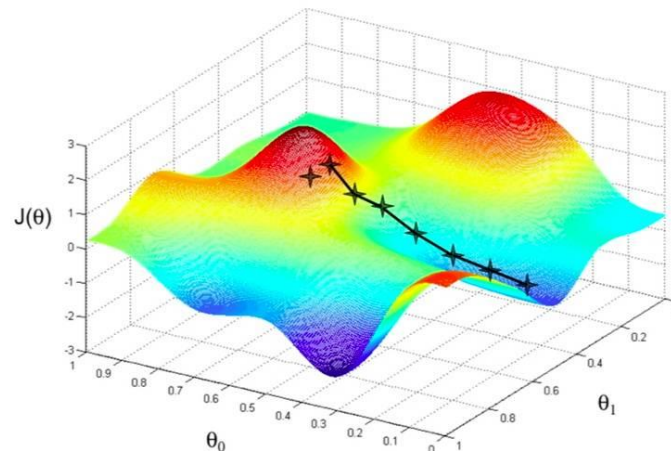
$$LLH = \log(LH) = \log\left(\prod P(a_t|H_{t-1})\right) = \sum \log(P(a_t|H_{t-1}))$$

Example!

Grid Search



Gradient Descent



```
>> probabilities = rand(1,1000);  
>> prod(probabilities)
```

```
ans =
```

```
0
```

```
>> sum(log(probabilities))
```

```
ans =
```

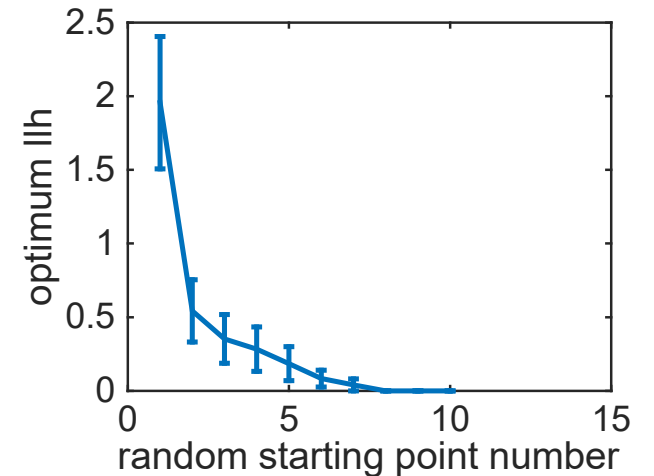
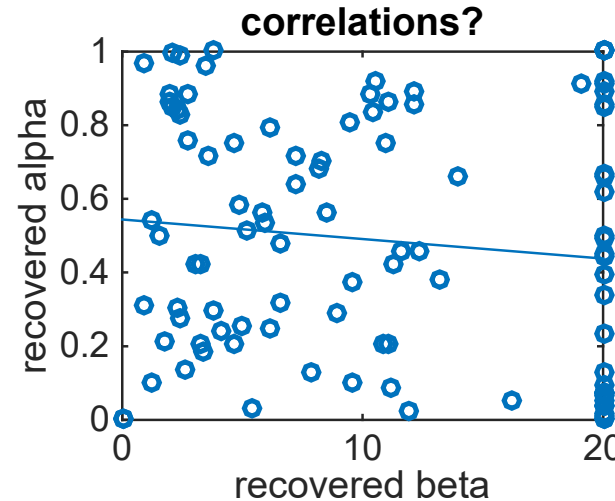
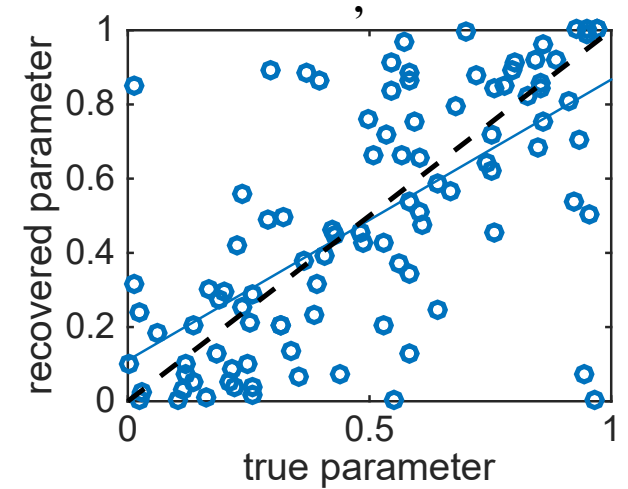
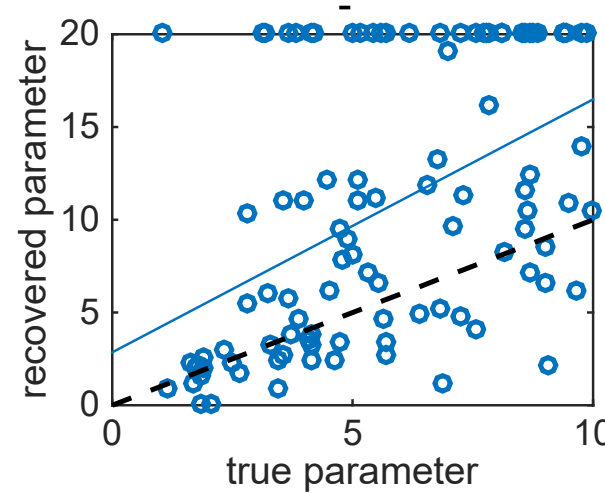
```
-965.7124
```

Parameter recovery: Generate and recover!

Validate this procedure!

1. Simulate the model with a reasonable range of parameters to create ***“fake” (synthetic) data***
2. Fit the “fake data”
3. Compare the recovered parameters with the true parameters

n=1 block
10 starting points



Tips and tricks: Parameter fitting

- More data always leads to better parameter recovery
- The generate and recover procedure, on ***your*** experimental protocol, shows you whether you can recover parameters satisfactorily
- Carefully choose constraints on parameters.
- Beware of local minima (always run fitting multiple times with random initial conditions)
- **Once you are satisfied that you can recover parameters for a model, you can fit your real data!**

Steps

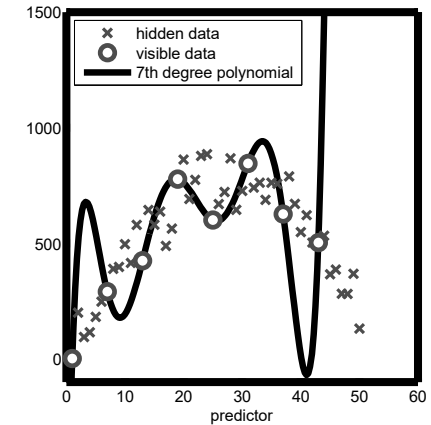
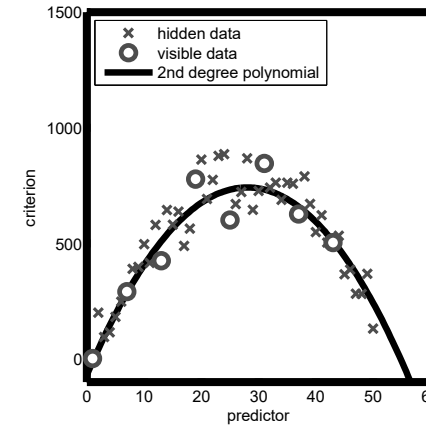
1. Data analysis
2. Model simulation
 - Role of model parameters
 - Setting up models representing competing hypotheses
3. Parameter estimation (model fitting)
- 4. Model comparison**
5. Model validation
6. Improving models

4. Model comparison: The question

- Which of two (or n) models fits the data best?
- A better question:
 - Which of the models generalizes better?
 - Which of the models will let me predict the participant's behavior best?

Bias-variance tradeoff / Occam's razor

- More is not always better!
 - Capture signal in the data, not noise
- Best solution: **out-of-sample fit measure**
 - how well your fit model predicts new data (generalization)



Measure fit out of sample. Why we rarely do it.

- Not enough data
- Not possible to separate data into independent data sets
 - for example, if trying to model a single learning block.

Approximate criteria for model comparison

- Principle:
 - measure model fit by likelihood, but penalize for model complexity.
- Various Information Criteria come from different mathematical theories, with different interpretations and guarantees
 - Bayesian Information Criterion (BIC)
 - Akaike Information Criterion (AIC)

Smaller is better! →

$$\text{AIC} = -2\text{llh} + 2np$$

← Fewer parameters is better

$$\text{BIC} = -2\text{llh} + np \cdot \log(n)$$

← More so with more trials

The diagram illustrates the Akaike Information Criterion (AIC) and Bayesian Information Criterion (BIC) formulas. It shows that both criteria consist of a log-likelihood term (-2llh) and a penalty term for model complexity. The AIC formula is $\text{AIC} = -2\text{llh} + 2np$, and the BIC formula is $\text{BIC} = -2\text{llh} + np \cdot \log(n)$. Arrows point from the text 'Smaller is better!' to both formulas, indicating that lower values are preferred. Another arrow points from the text 'Fewer parameters is better' to the $2np$ term in the AIC formula. A third arrow points from the text 'More so with more trials' to the $np \cdot \log(n)$ term in the BIC formula, highlighting that BIC imposes a stronger penalty for model complexity as the sample size n increases.

Model comparison – validating the procedure

- Which model comparison method should I use?
 - I.e. are my models **identifiable**? With my comparison method and the experimental protocol, can I distinguish between the models?
- Test it in the best-case scenario (where the true model is known: synthetic dataset). Create confusion matrices:
 - Generate data with models M1 and M2
 - Fit the generated data with model M1 and M2
 - Verify that data generated with M_i is best fit by model M_i

	M1's AIC better	M2's AIC is better
Data generated by M1	95%	5%
Data generated by M2	0	100%

What other factors could influence model identifiability?

- Model domain: range of parameters over which the model is simulated.
 - **Case study:** if range of beta is very low, choice will be mostly noisy, so other parameters will have little influence on the likelihood. This will influence model identifiability.
 - Many other such situations...
- In practice: what range of parameters should I pick to do my model recovery analyses?
 - **Before you have fit to real data: pick a range of parameters where the model behavior mimick expected experimental range.**
 - After you have fit to real data: use real fit parameters. This allows you to draw conclusions for your actual sample, but this will be less generalizable.

Tips and tricks: Model comparison

- Confusion matrices are computationally intensive.
 - Start small!
 - Small number of simulations/participants per model
 - Small number of starting points for fitting
 - Small number of trials (if possible)
- Model + parameter recovery
 - When doing model recovery, you're also doing parameter recovery (you're fitting M1 data with M1, M2 data with M2, etc...)
 - But in practice, to minimize computation cost, you can combine the two steps.
 - Downside: doing parameter recovery is faster (n instead of n^2), and you might not use all models in model recovery analysis – so it sometimes still makes sense to do the two steps separately.

What **can** you conclude so far?

What can you **not** conclude so far?

Yes

- My procedure identifies M_x as most likely from $\{M_1, \dots, M_n\}$ to generate the data
- Best fit parameters for M_x are the most likely ones.

No

- M_x is a good model of my cognitive process of interest
- M_x is the best model out of $\{M_1, \dots, M_n\}$ to model my cognitive process.

Steps

1. Data analysis
2. Model simulation
 - Role of model parameters
 - Setting up models representing competing hypotheses
3. Parameter estimation (model fitting)
4. Model comparison
5. Model validation (after fitting your behavior data)
6. Improving models

Quantitative measures of model fit is not enough

Qualitative questions:

- Can this model produce the behavior patterns I'm expecting from this cognitive process, this experiment, this participant sample?
- Does the behavior of the model produce match participants' behavior?

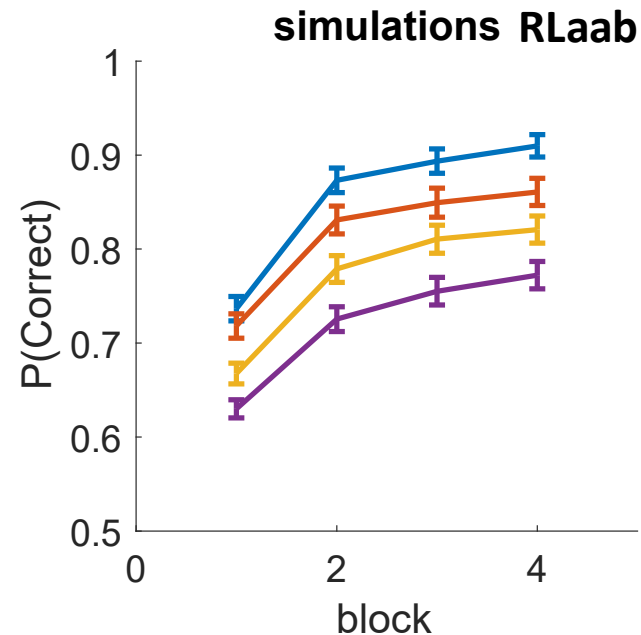
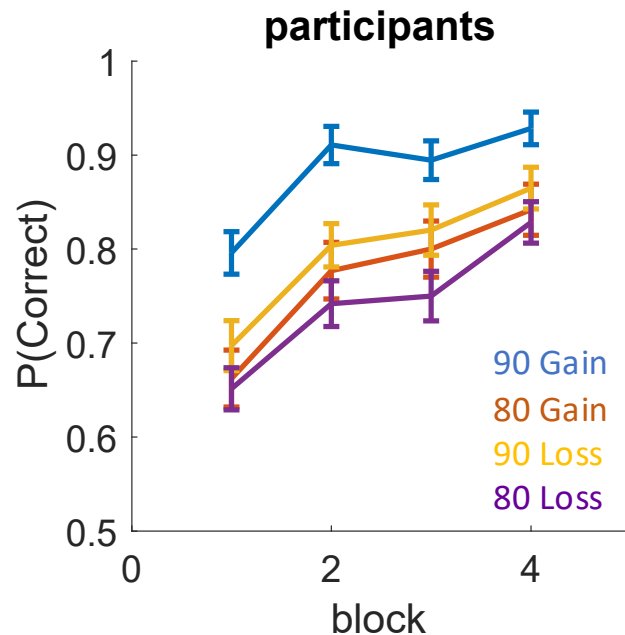
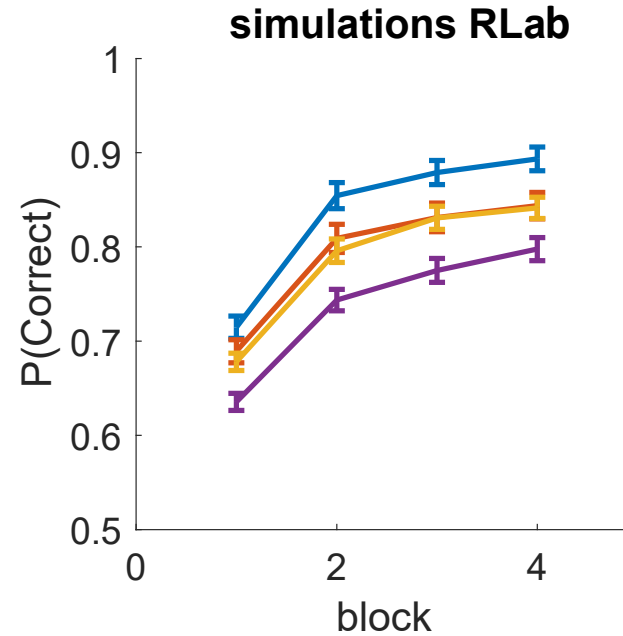
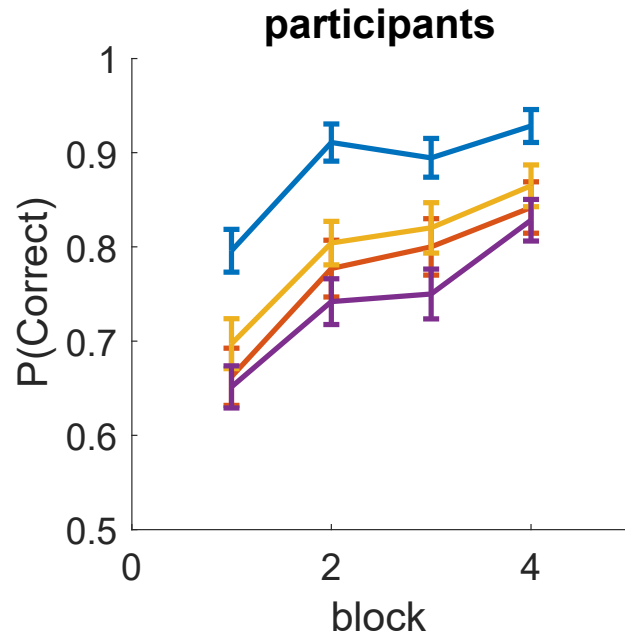
Validate your model!

How do I validate my model?

(also called *posterior predictive check*)

- Simulate the model on the **exact** experiment, with fit parameters
 - For each participant, with each participant's best fit parameters
 - Multiple times per participant, to average out randomness in choices
- Analyze the behavior of each simulation as if it was a participant
 - Create all the summary measurements you identified as important when you explored participants' behavior
 - Average measures across simulations for each participant, treat this as the validated participant's measure
- Analyze the behavior across participants, and compare to real data

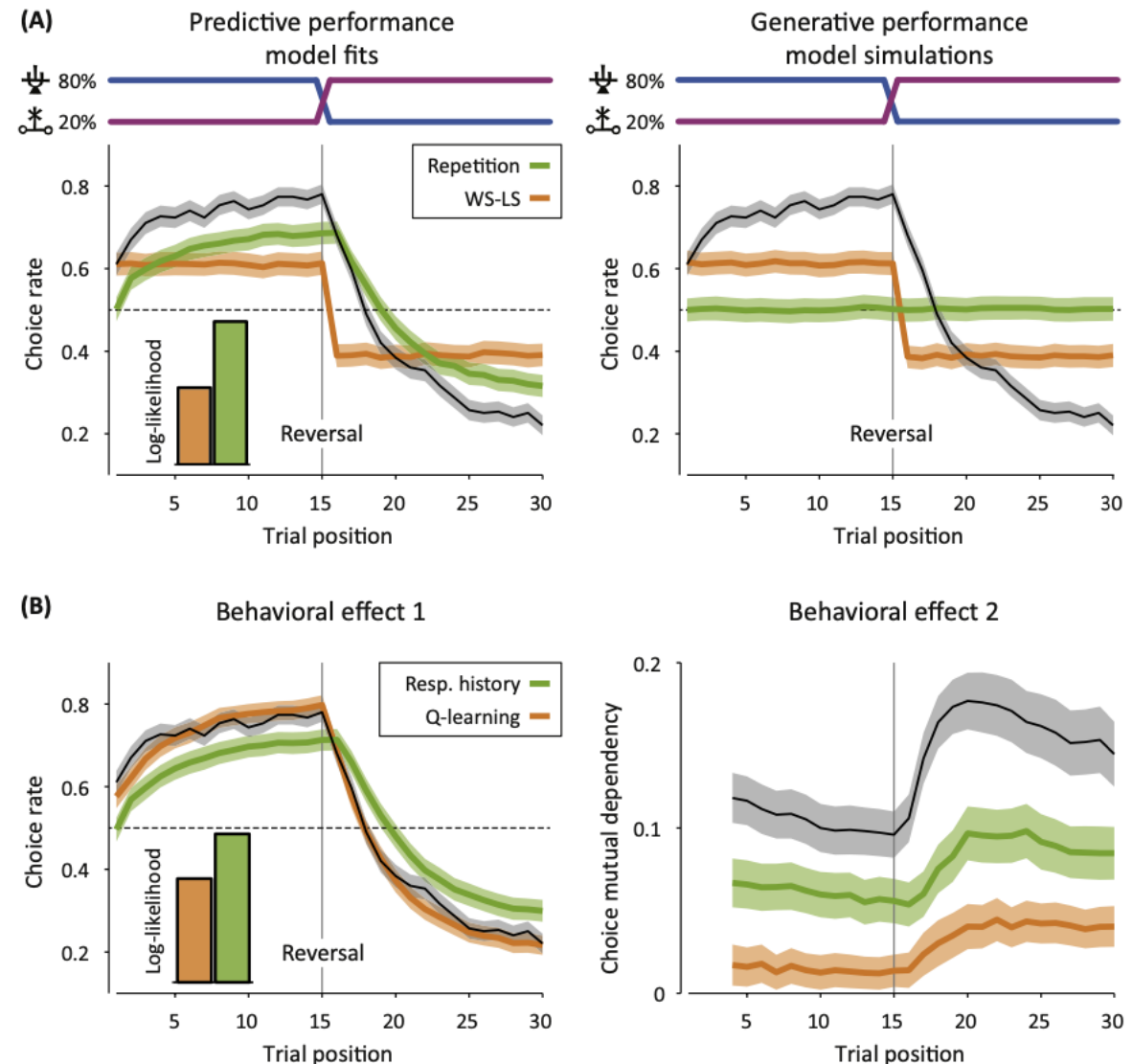
Model validation



- 100 simulations per participant, not looking at group effects yet.
- Good:
 - both capture Learning, Gain/Loss effects
- Not good:
 - underperformance
 - not capturing uncertainty *valence effects
 - RLaaB predicts 80%gain>90% loss learning. Wrong!

Model Predictions and Simulations Are Not Necessarily Related.

- Generating data is different from predicting data (A)
- Quantitative measures don't tell the whole story (A,B)



Simulation for validation, falsification, and model improvement

- Model fitting procedures don't ensure that the model can generate the observed data (*"predictive performance"*).
- Simulations **validate** this and can **falsify** other models.
- When they falsify a model, they indicate where the model can be **improved**.

Questions?



REVIEW ARTICLE



Ten simple rules for the computational modeling of behavioral data

Robert C Wilson^{1,2†*}, Anne GE Collins^{3,4†*}

Trends in Cognitive Sciences

CellPress

Opinion

The Importance of Falsification in Computational Cognitive Modeling

Stefano Palminteri,^{1,2,*,†} Valentin Wyart,^{1,2,*,†} and Etienne Koechlin^{1,2,*}

