

Computational Cognitive Science

Lecture 2: Model Building

Frank Keller

24 September 2026

School of Informatics
University of Edinburgh
keller@inf.ed.ac.uk

Slide credits: Chris Lucas, Benjamin Peters

Precision and Predictive Scope

Example: Reasoning about Causal Relations

Model-building Case Study

Random-walk Model

R Implementation

Predictions

Trial-to-trial variability

Parameters

Reading: Chapters 1.3, 1.4, and 2 of [Farrell and Lewandowsky \(2018\)](#).

Optional: [CCS R Cheat Sheet](#) for students familiar with Python or Java.

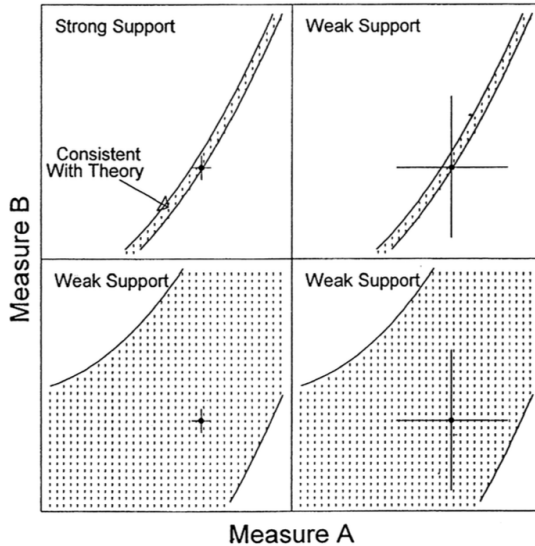
1. Precision and predictive scope
2. A model-building case study: Random-walk model of response times in decision-making
 - Turning a theory into a model
 - Qualitative predictions, intuitions, assumptions
 - Modifications and free parameters

Why should we translate a theory expressed as verbal statements into a model?

- See where theories are vague; make them more precise.
- Overcome disagreement about a theory or its implications: “There is a rich variety of misinterpretations of Quillian’s theory” ([Collins and Loftus 1975](#))
- Pin down one version of a theory from many alternatives

Precision and Predictive Scope

Precision and “predictive scope”



Precision and “predictive scope”

What does it mean for a model to be precise?

Specificity, or “predictive scope” is an important part of it.

Confidence matters too: Strong predictions; “falsifiability”

Example: Reasoning about Causal Relations

Suppose we are interested in how people make judgments about a causal relationship, based on evidence they accumulate over time.

More details about causal reasoning in Lectures 11 and 12.

Example

“There is a village where people frequently experience rashes. We want to know if the Dax plant causes or prevents rashes, or makes no difference at all.”

1. Villager 1 touched a Dax plant. Villager 1 had a rash. (1,1)
2. Villager 2 did not touch a Dax plant. No rash. (0,0)

...

79. Villager 79 did not touch a Dax plant. Rash. (0,1)
80. Villager 80 touched a Dax plant. No rash. (1,0)

Example

Two experimental conditions:

- **Generate-first (GF):** Early evidence favors plant-causes-rash, late evidence favors plant-prevents-rash.
- **Prevent-first (PF):** Late evidence favors plant-causes-rash, early evidence favors plant-prevents-rash.

After seeing 80 events, people make a judgment about the underlying causal relationship on a -100 to 100 scale.

- -100 : The plant definitely prevents rashes.
- 0 : No relationship.
- 100 : The plant definitely causes rashes.

“My theory predicts. . .

1. . . .judgments differ between the conditions.” (Alice)
2. . . .judgments will be higher in the GF condition.” (Bob)
3. . . .the mean will be 25 in GF and -25 in PF.” (Claire)

If our mean judgments are 20 in GF ($\text{stErr}=7$) and -30 in PF ($\text{stErr}=10$), whose theory should we prefer?

The value of precision

- Vague $\prec$ inaccurate $\prec$ accurate
- Easier to gather evidence for/against precise models (“falsification”)
- Among models whose predictions are consistent with the data, more precise predictions provide stronger evidence when confirmed
- We will be more precise about precision in future lectures

Model-building Case Study

Response times in forced-choice tasks

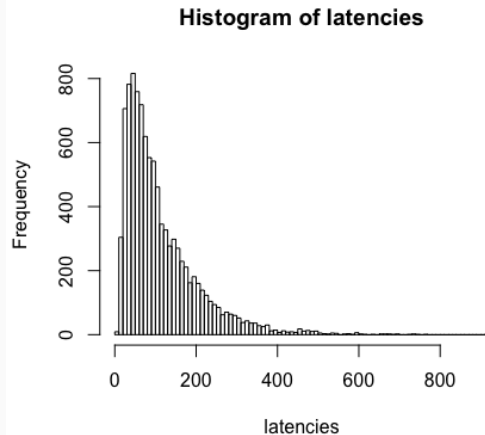
How do we trade off between speed and accuracy?

- “Do I know that person? Should I wave?”
- “Rock or shadow?” while running downhill at night
- “Are these lines pointing left or right?”



The relationship between task difficulty, accuracy, and response time can be used to test models of decision-making. If you want to try the experiment, use [Pablo Leon-Vilagra's replication](#).

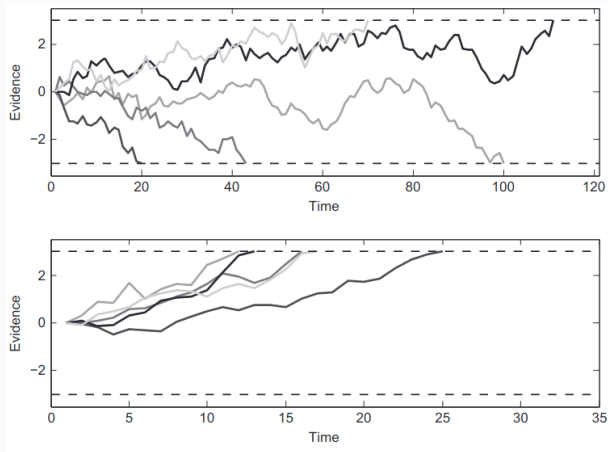
Response times in forced-choice tasks



There are fast errors (speed-accuracy trade-off) and slow errors.

Random-walk model

Idea: People sequentially accumulate evidence for one decision or another, and decide when the evidence exceeds a threshold. This process noisily and additively combines information from stimuli.



R Implementation

We assume the following parameters with psychological interpretations:

<code>drift</code>	Average evidence supplied per sample
<code>sdrw</code>	Within-trial evidence noise
<code>criterion</code>	Amount of evidence required before responding

R Implementation

Variable initialization:

```
nreps <- 10000 # The number of complete simulations
nsamples <- 2000 # The max number of samples/steps

drift <- 0.0 # Non-informative stimulus
sdrw <- 0.3 # Standard deviation of the random walk
criterion <- 3 # Fixed response threshold

latencies <- rep(0,nreps) # An empty vector to start
responses <- rep(0,nreps) # An empty vector to start
evidence <- matrix(0,nreps,nsamples+1) # Empty matrix
```

R Implementation

The main loop:

```
for (i in c(1:nreps)) {  
  evidence[i,] <- cumsum(c(0,rnorm(nsamples,drift,sdrw)))  
  # The first point that exceeds the threshold  
  p <- which(abs(evidence[i,]) > criterion)[1]  
  # The sign tells us whether the response is left/right  
  responses[i] <- sign(evidence[i,p])  
  # Latency: How many time steps did it take?  
  latencies[i] <- p  
}
```

At each time step, the movement is a random sample with mean equal to the drift. We take the cumulative sum of the individual movements. Everything is repeated `nreps` times

R Implementation

The first line within the loop on the previous slide can be unpacked:

```
evidence[i,] <- cumsum(c(0,rnorm(nsamples,drift,sdrw)))
```

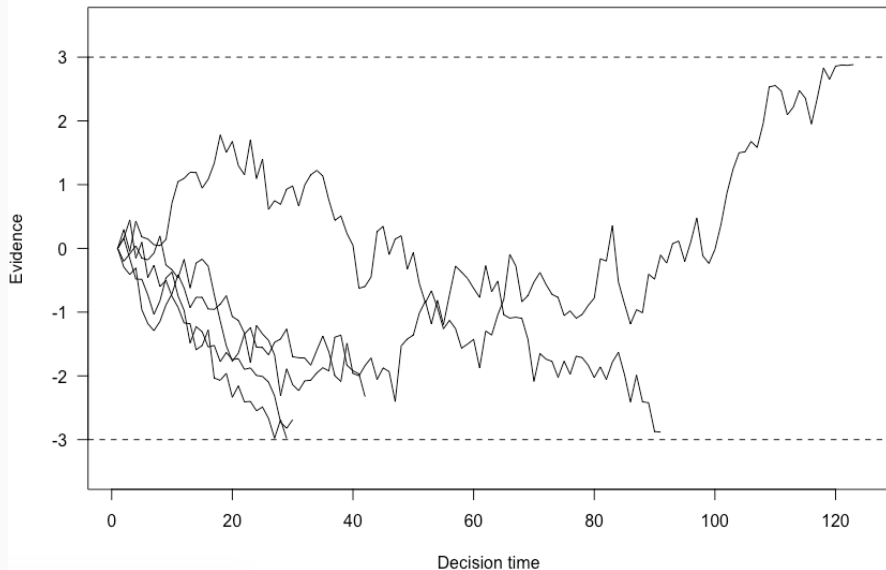
Is an efficient implementation of sequentially adding new pieces of evidence:

```
evidence[i,1] <- 0
for (j in c(2:nsamples+1)) {
  evidence[i,j] <- evidence[i,j-1] + rnorm(1,drift,sdrw)
}
```

Visualizing paths

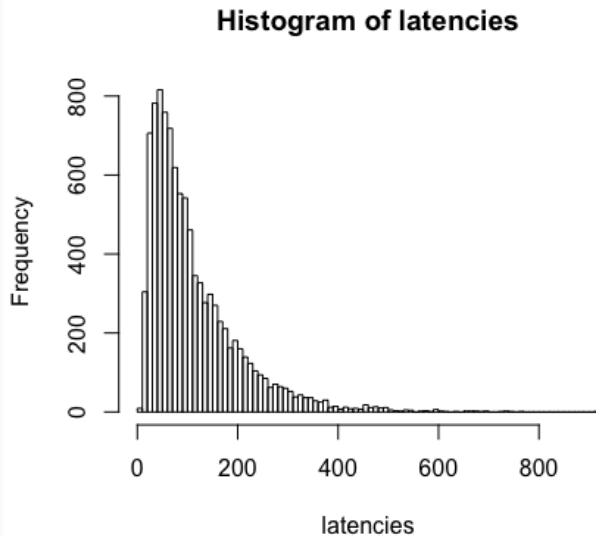
```
tbpn <- min(nreps,5) # Plot up to 5 lines
# Empty plot with axes and labels
plot(1:max(latencies[1:tbpn])+10,type="n",las=1,
     ylim=c(-criterion-.5,criterion+.5),
     ylab="Evidence",xlab="Decision time")
# The lines themselves
for (i in c(1:tbpn)) {
  lines(evidence[i,1:(latencies[i]-1)])
}
# Show the boundaries
abline(h=c(criterion,-criterion),lty="dashed")
```

Visualizing paths



The model predicts judgments and latencies.

```
hist(latencies,breaks = 100)
```



Suppose the stimulus is informative, e.g., lines tilt left. Do you think errors and correct responses will take the same amount of time?

Errors, latencies, and intuitions

This simple model predicts that correct and incorrect answers will take the same amount of time.

In reality, the distributions differ – including fast and slow errors.

Trial-to-trial variability

Idea: Drift and starting point can vary from trial to trial.

```
# t2tsd[1]: starting point standard deviation
# t2tsd[2]: drift standard deviation
t2tsd <- c(0.8,0.0)
drift <- 0.035
# ...
for (i in c(1:nreps)) {
  sp <- rnorm(1,0,t2tsd[1])
  dr <- rnorm(1,drift,t2tsd[2])
  # Prepend starting point to samples
  evidence[i,] <- cumsum(c(sp, rnorm(nsamples,dr,sdrw)))
  p <- which(abs(evidence[i,]) > criterion)[1]
  responses[i] <- sign(evidence[i,p])
  latencies[i] <- p
}
```

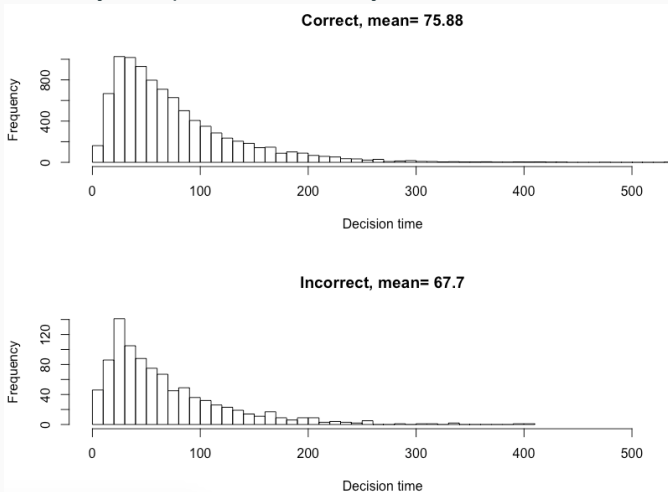
Trial-to-trial variability

Let's look:

```
par(mfrow = c(2, 1))
toprt <- latencies[responses>0]
topprop <- length(toprt)/nreps
hist(toprt,xlab="Decision time", xlim=c(0,max(latencies)),
     main=paste("Correct, mean=", signif(mean(toprt),4)),
     breaks=50)
botrt <- latencies[responses<0]
hist(botrt,xlab="Decision time", xlim=c(0,max(latencies)),
     main=paste("Incorrect, mean=", signif(mean(botrt),4)),
     breaks=50)
```

Trial-to-trial variability

- Starting-point variability can produce relatively fast errors
- Drift-rate variability can produce relatively slow errors

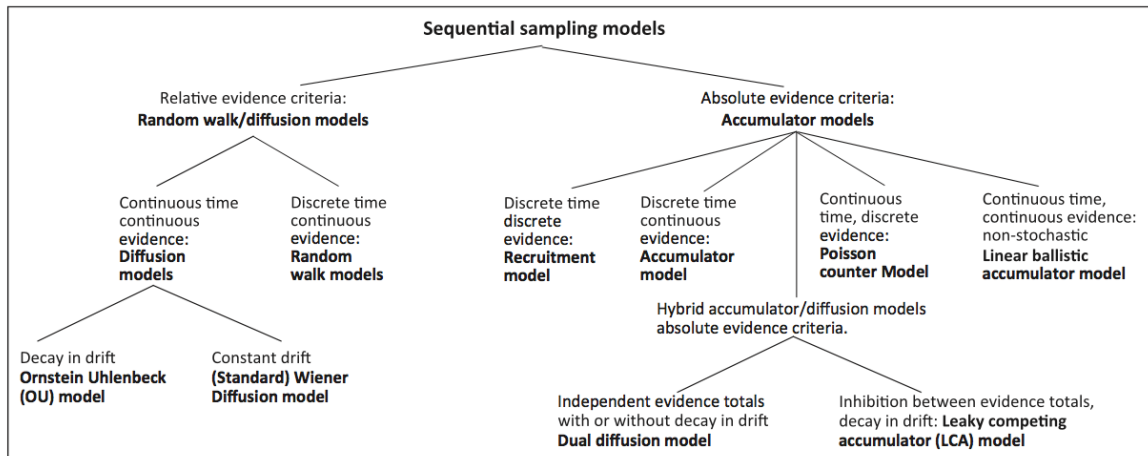


Assumptions

What assumptions are we making? Here are a few:

1. Fixed time between samples
2. Noise+drift distribution
 - Constant vs. changing between trials
 - Constant vs. decaying drift within trials
3. Starting state
 - Constant vs. changing
 - Independence of successive trials
4. Absolute vs. relative evidence

Assumptions



Parameters

We have introduced parameters that make the model more flexible; it can now capture both fast and slow errors.

- The model is flexible enough to capture more real patterns in human judgment.
- Greater flexibility allows the model to capture more patterns, but also increases complexity and reduces how strongly the model constrains the data.

For next time

How do we decide what free parameters to use?

- Estimating free parameters
- Pitfalls in estimating free parameters

How do we quantify the goodness of a model's fit to data?

- Necessary for parameter estimation
- Choices of discrepancy functions

- Collins, Allan M. and Elizabeth F. Loftus. 1975. A spreading-activation theory of semantic processing. *Psychological Review* 82(6):407–428.
- Farrell, Simon and Stephan Lewandowsky. 2018. *Computational Modeling of Cognition and Behavior*. Cambridge University Press, Cambridge.