

# **Hierarchical Gaussian Filters (HGFs) for behavioural modelling in computational psychiatry**

**Christoph Mathys**



**Methods and Primers for Computational Psychiatry and Neuroeconomics**  
Yale School of Medicine  
Psychiatry

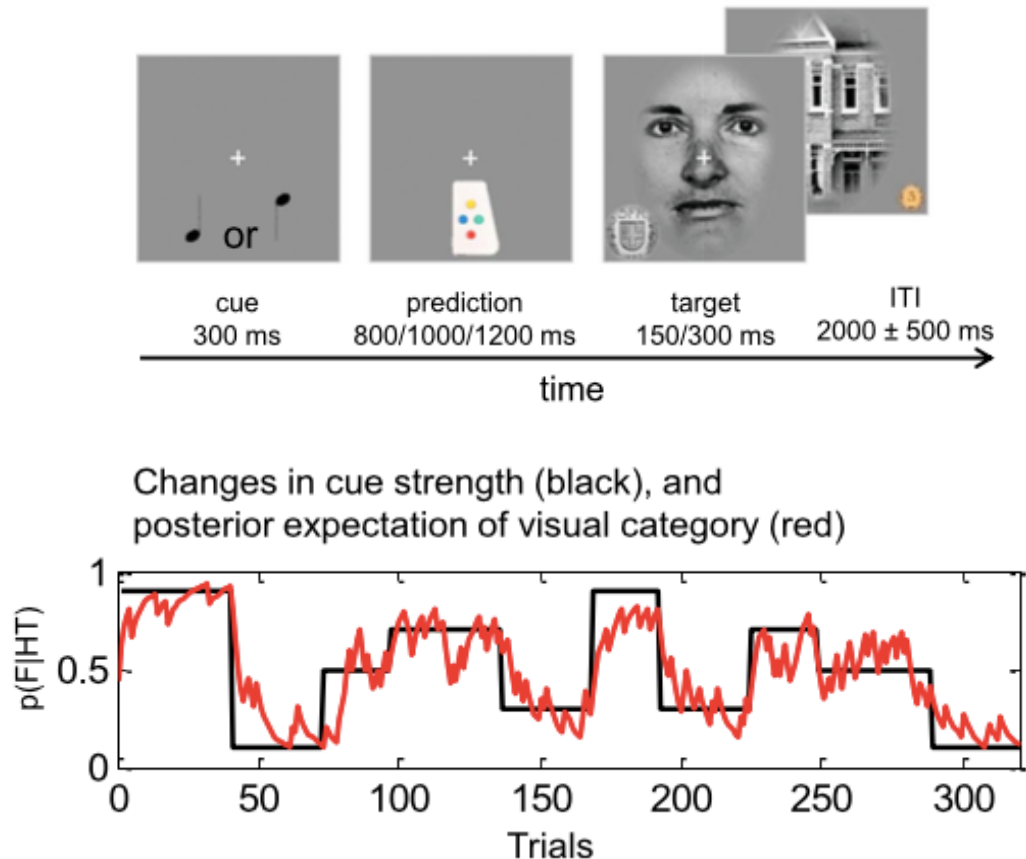
<https://medicine.yale.edu/psychiatry/map/>

4 February 2021

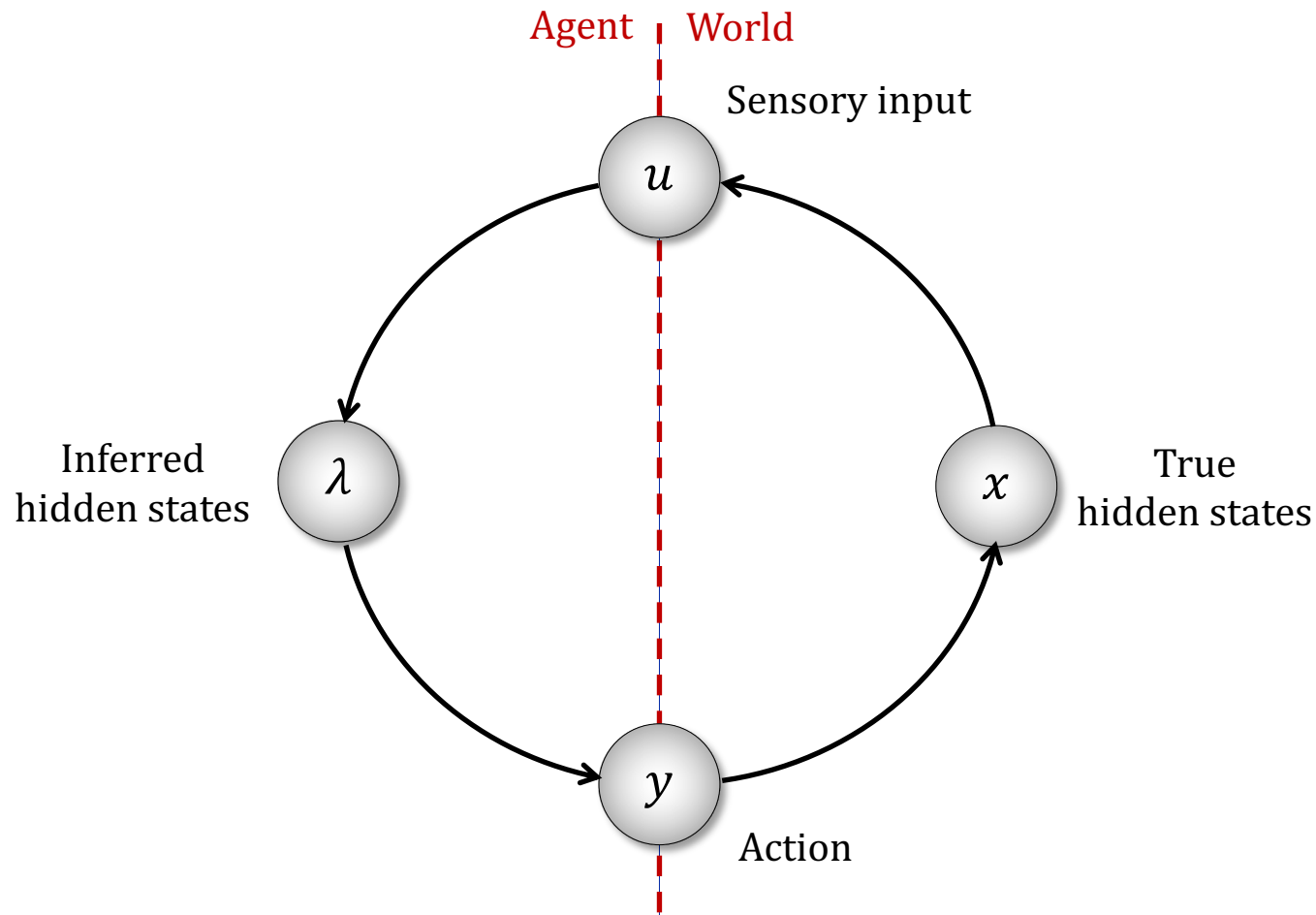
# Learning rate adaptation in dynamic environments:

## An important focus of research in computational psychiatry and computational neuroscience

Example: task of Iglesias et al., *Neuron*, (2013):



# The general framework for modelling agents navigating dynamic environments

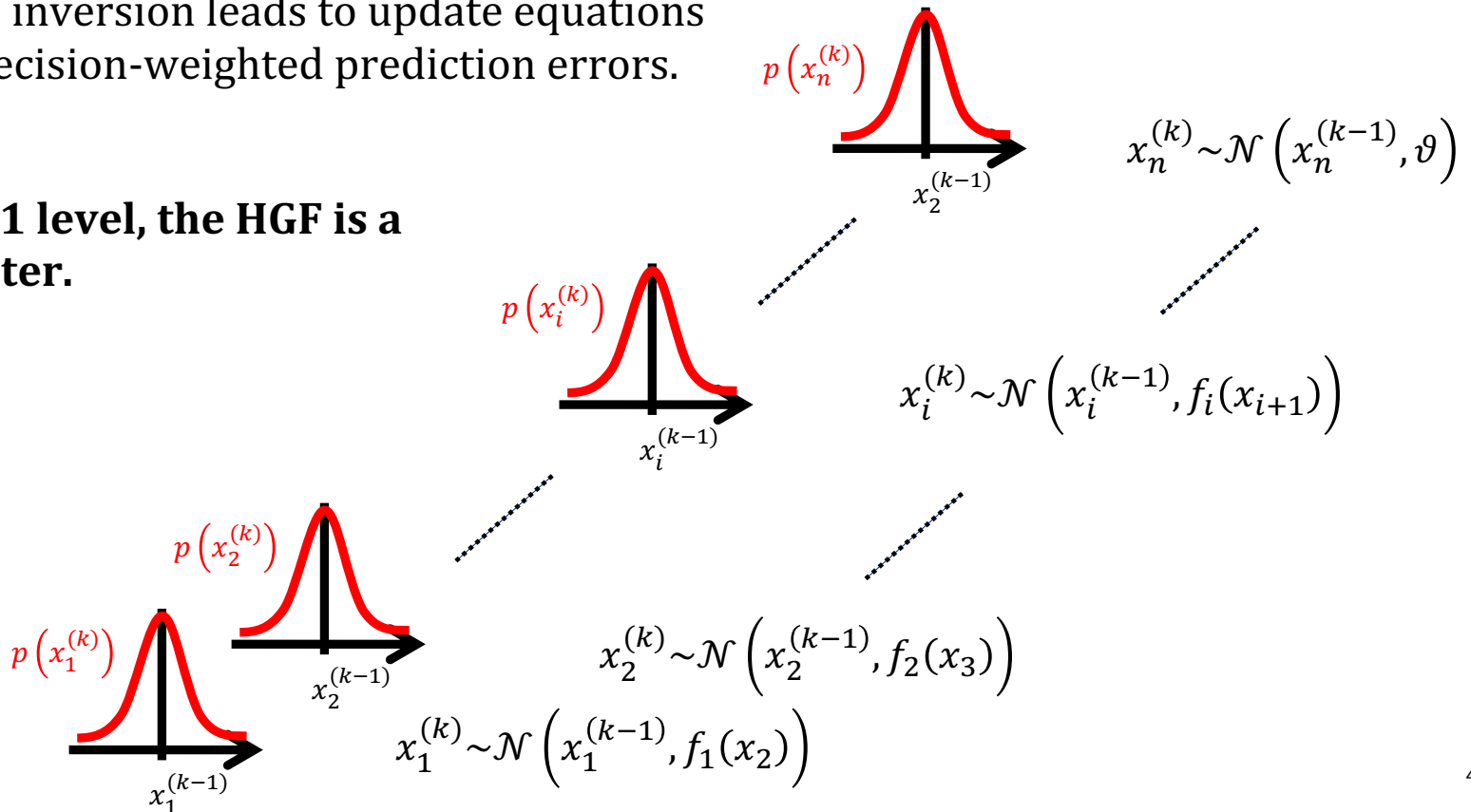


# Today's focus: hierarchical Gaussian filters (HGFs, Mathys et al., 2011; 2014)

HGFs provide a generic solution to the problem of adapting one's learning rate in a dynamic environment.

Variational inversion leads to update equations that are precision-weighted prediction errors.

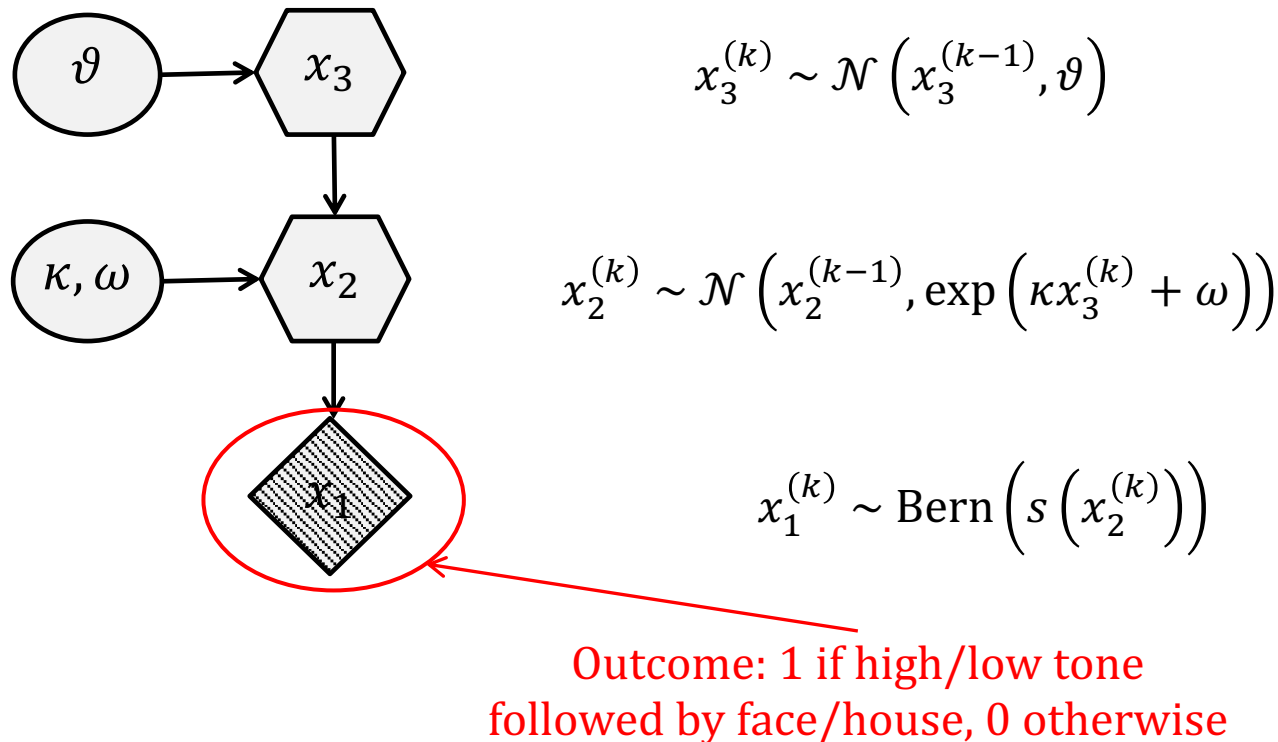
**With only 1 level, the HGF is a Kalman filter.**





# Generative model for trial outcomes in the Iglesias et al. task:

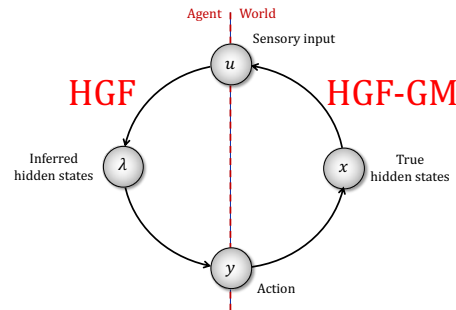
## 3-level HGF for binary observations



Mathys et al., 2011; Iglesias et al., 2013; Vossel et al., 2014a; Hauser et al., 2014; Diaconescu et al., 2014; Vossel et al., 2014b; ...

# Variational inversion and update equations

- Important distinction: **generative model** (HGF-GM) vs its inversion, **inference model** (HGF proper)



- Inversion of HGF-GM proceeds by introducing a mean field approximation and fitting quadratic approximations to the resulting variational energies (Mathys et al., 2011).
- This leads to **simple one-step update equations** (HGF proper) for the sufficient statistics (mean and precision) of the approximate Gaussian posteriors of the states  $x_i$ .
- The updates of the means have the same structure as value updates in Rescorla-Wagner learning:

$$\Delta\mu_i \propto \frac{\hat{\pi}_{i-1}}{\pi_i} \delta_{i-1}$$

Precisions determine learning rate
Prediction error

- The updates are **precision-weighted prediction errors**.

## Updates at the first level

At the outcome level (i.e., at the very bottom of the hierarchy), we have

$$u^{(k)} \sim \mathcal{N} \left( x_1^{(k)}, \hat{\pi}_u^{-1} \right)$$

This gives us the following update for our belief on  $x_1$  (our quantity of interest):

$$\pi_1^{(k)} = \hat{\pi}_1^{(k)} + \hat{\pi}_u$$

$$\mu_1^{(k)} = \mu_1^{(k-1)} + \frac{\hat{\pi}_u}{\pi_1^{(k)}} \left( u^{(k)} - \mu_1^{(k-1)} \right)$$

The familiar structure again – but now with a learning rate that is responsive to all kinds of uncertainty, including environmental (unexpected) uncertainty.

# The learning rate ('Kalman gain') in the HGF

Unpacking the learning rate, we see:

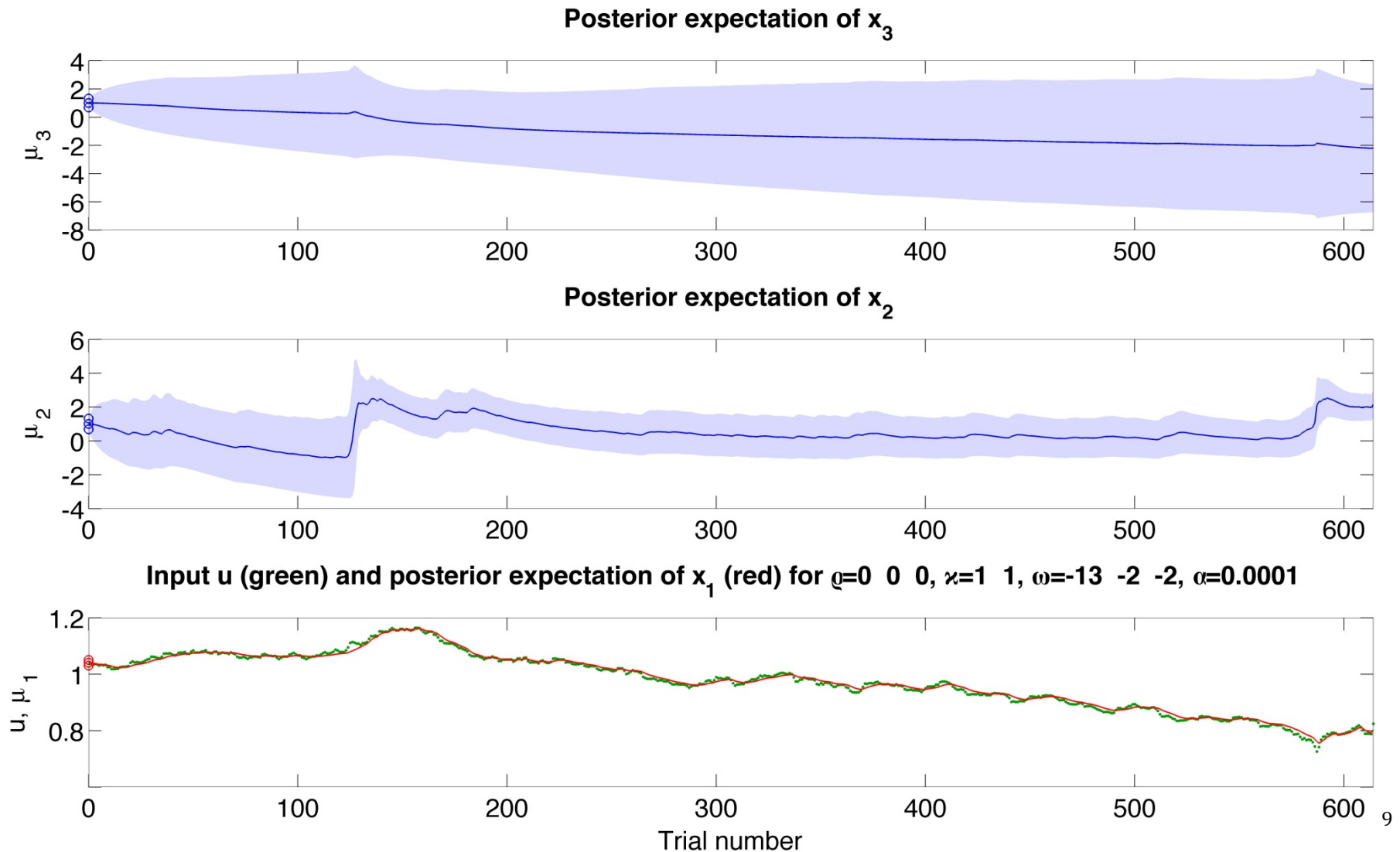
$$\frac{\hat{\pi}_u}{\pi_1^{(k)}} = \frac{\hat{\pi}_u}{\hat{\pi}_1^{(k)} + \hat{\pi}_u} = \frac{\hat{\pi}_u}{\frac{1}{\sigma_1^{(k-1)} + \exp(\kappa_1 \mu_2^{(k-1)} + \omega_1)} + \hat{\pi}_u}$$

outcome uncertainty

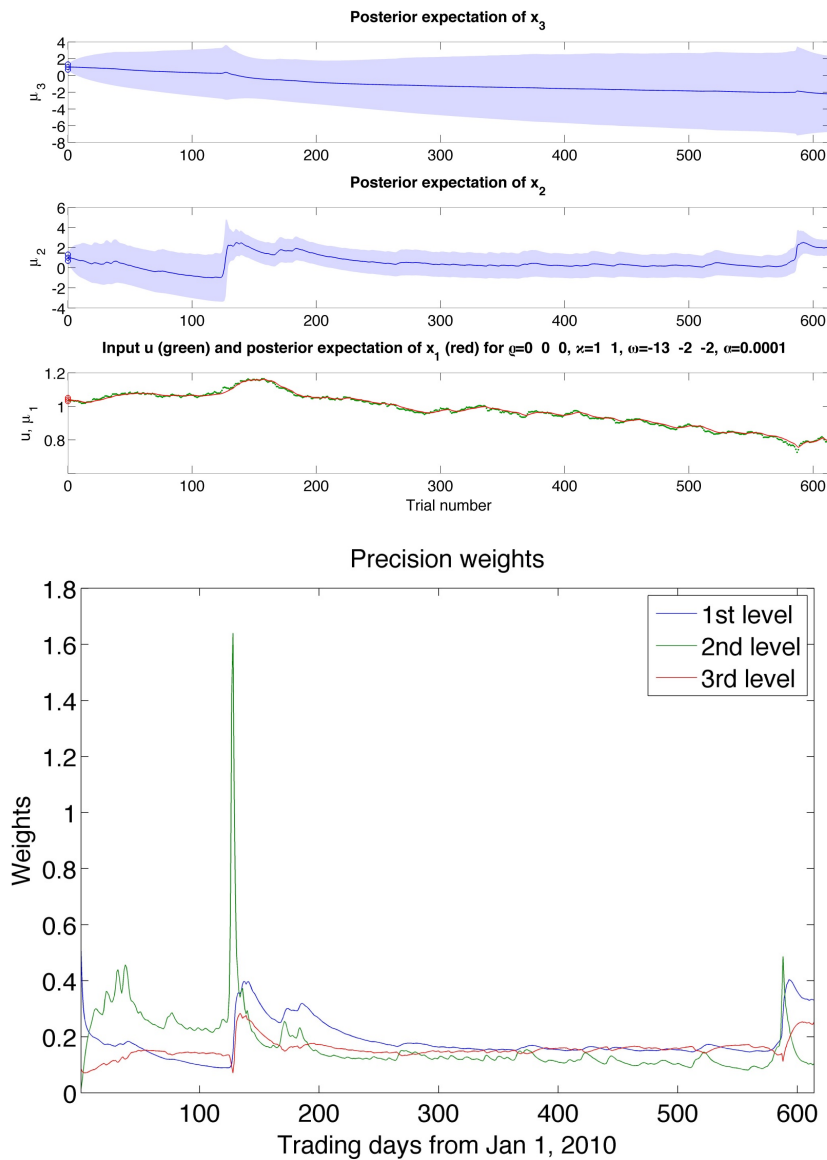
informational uncertainty

environmental uncertainty  
(instead of the constant  $\vartheta$  in the Kalman filter)

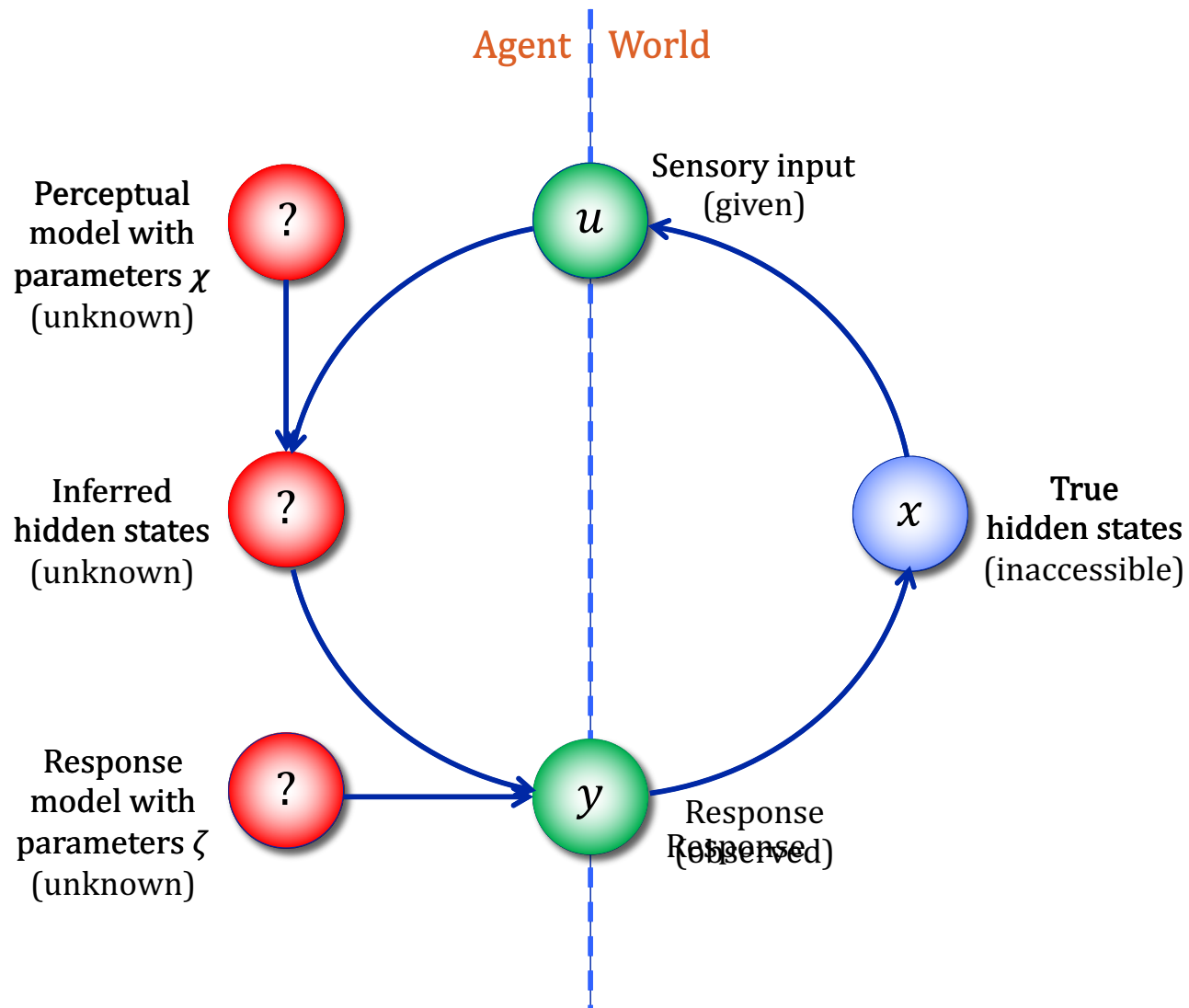
# 3-level HGF for continuous observations



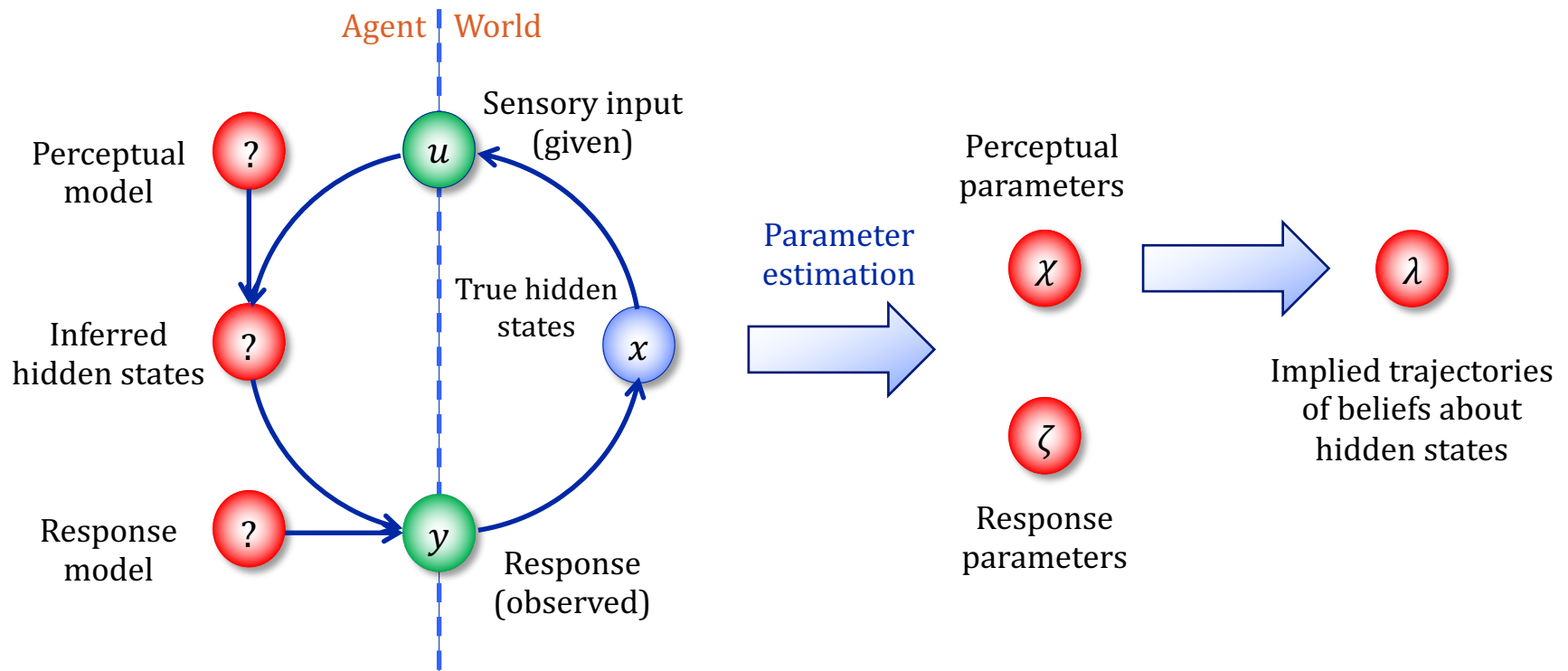
# Example of precision weight trajectory



# Application to experimental data



# Application to experimental data: parameter estimation





# Generative model, perceptual model, and decision model

- The *generative model* describes (probabilistically) how the states  $x_i$  evolve in time:  $x_i^{(k-1)} \rightarrow x_i^{(k)}$
- Variationally inverting the generative gives us the *perceptual model* (also called *inference model* or *recognition model*)
- The perceptual model describes (deterministically, via update equations) how *beliefs*  $\{\mu_i, \pi_i\}$  about states evolve in time:  $\{\mu_i^{(k-1)}, \pi_i^{(k-1)}\} \rightarrow \{\mu_i^{(k)}, \pi_i^{(k)}\}$
- The decision model (also called observation model or response model) describes (probabilistically) how beliefs are translated into observed actions:  
$$\left\{ \mu_i^{(k)}, \pi_i^{(k)} \right\}_{i=1, \dots, l} \rightarrow y^{(k)}$$
- In the process of applying the HGF to data, we only need the perceptual and decision models. The generative model has already done its work: it has supplied the update equations of the perceptual model
- Both the perceptual and decision models have parameters that can be estimated individually for each dataset. Doing so requires defining priors for these parameters.

# Parameter estimation with the HGF Toolbox

- Available at

[www.translationalneuromodeling.org/tapas](http://www.translationalneuromodeling.org/tapas) or

[translationalneuromodeling.github.io/tapas](http://translationalneuromodeling.github.io/tapas)

- Start with README, manual, and interactive demo

- Modular, 

```
est2 = tapas_fitModel(sim2.y,...  
                      usdchf,...  
                      'tapas_hgf_config',...  
                      'tapas_gaussian_obs_config',...  
                      'tapas_quasineutron_optim_config');
```

Parameter estimates for the perceptual model:

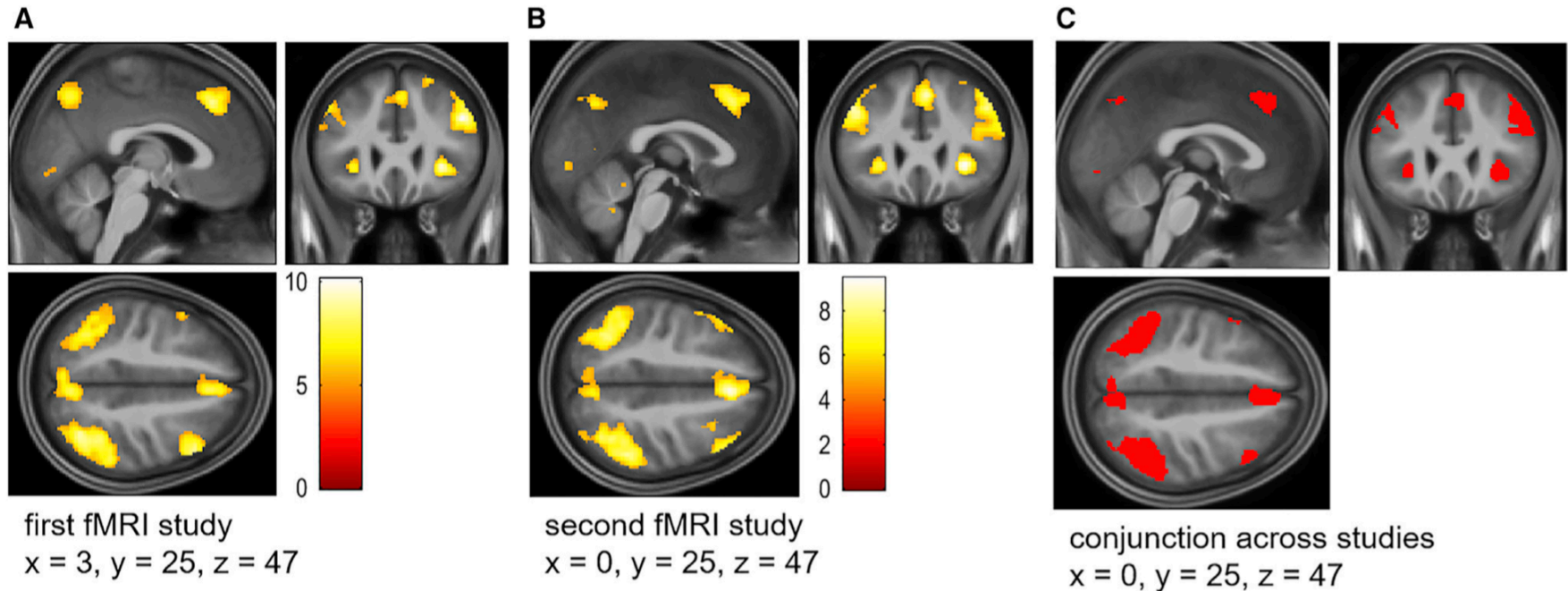
```
mu_0: [1.0352 1]  
sa_0: [3.7101e-05 0.0996]  
rho: [0 0]  
ka: 1  
om: [-12.8680 -1.8689]  
pi_u: 9.8449e+03
```

Parameter estimates for the observation model:

```
ze: 2.3970e-05
```

# Parameter estimation with the HGF Toolbox

## Activations by Precision-Weighted Visual Outcome Prediction Error $\varepsilon_2$

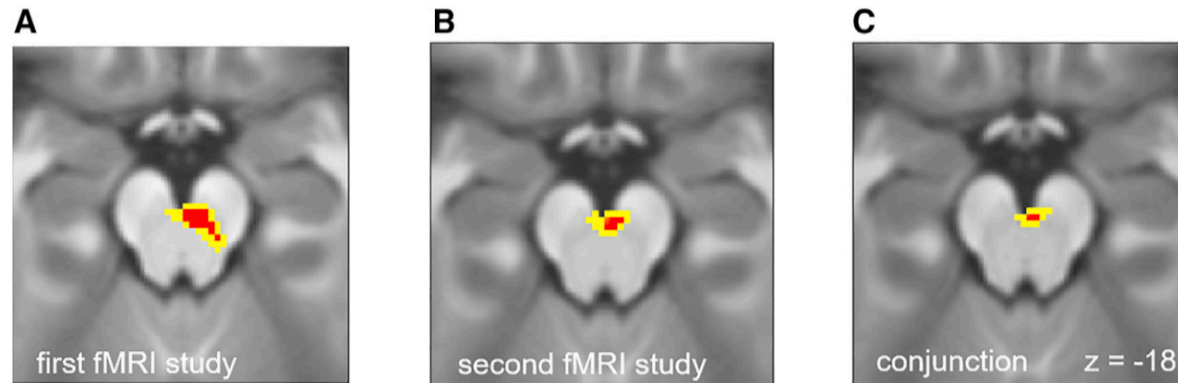


**Figure 2. Whole-Brain Activations by  $\varepsilon_2$**

Activations by precision-weighted prediction errors about visual stimulus outcome,  $\varepsilon_2$ , in the first fMRI study (A) and the second fMRI study (B). Both activation maps are shown at a threshold of  $p < 0.05$ , FWE peak-level corrected for multiple comparisons across the whole brain. To highlight replication across studies, (C) shows the results of a “logical AND” conjunction, illustrating voxels that were significantly activated in both studies.

Iglesias et al, 2019 (Correction to Iglesias et al., 2013)

# Parameter estimation with the HGF Toolbox



## Figure 3. Midbrain Activation by $\varepsilon_2$

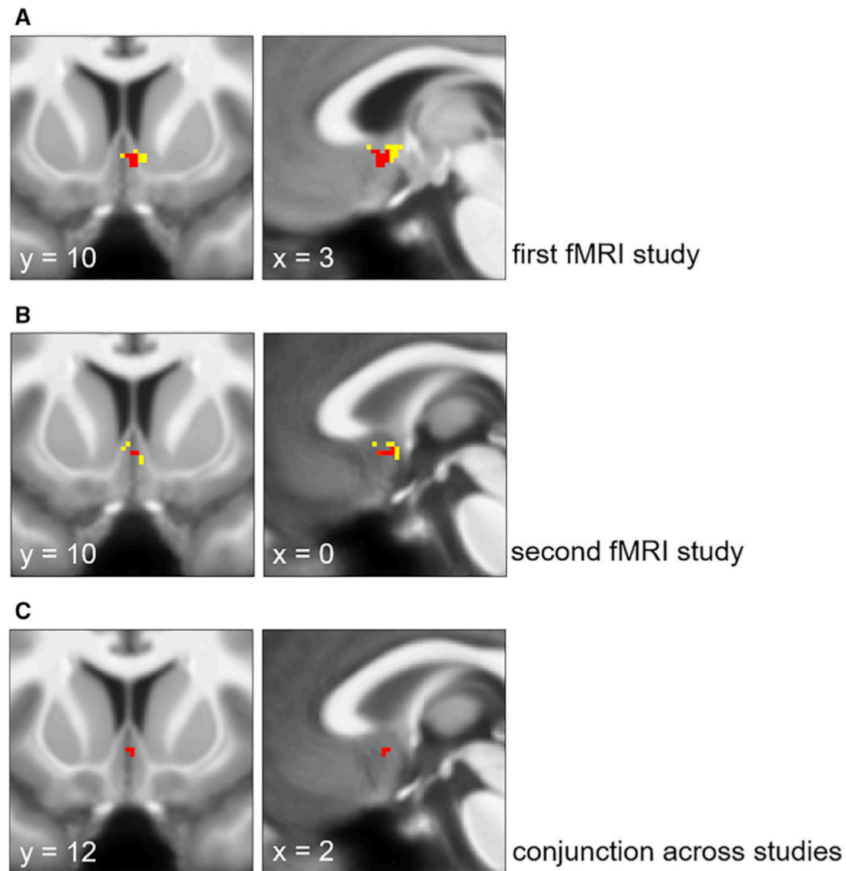
Activation of the dopaminergic VTA/SN by precision-weighted prediction errors about visual outcome,  $\varepsilon_2$ . The activation at  $p < 0.05$  FWE peak-level corrected for the volume of our anatomical mask (comprising both dopaminergic and cholinergic brain structures: VTA/SN, PPT/LDT, and basal forebrain) is shown in red. The activation thresholded at  $p < 0.001$  uncorrected is shown in yellow.

(A) Results from the first fMRI study. (B) Second fMRI study. (C) Conjunction (logical AND) across both studies.

Iglesias et al, 2019 (Correction to Iglesias et al., 2013)

# Parameter estimation with the HGF Toolbox

Activations by Precision-Weighted Prediction Error about Stimulus Probabilities  $\varepsilon_3$

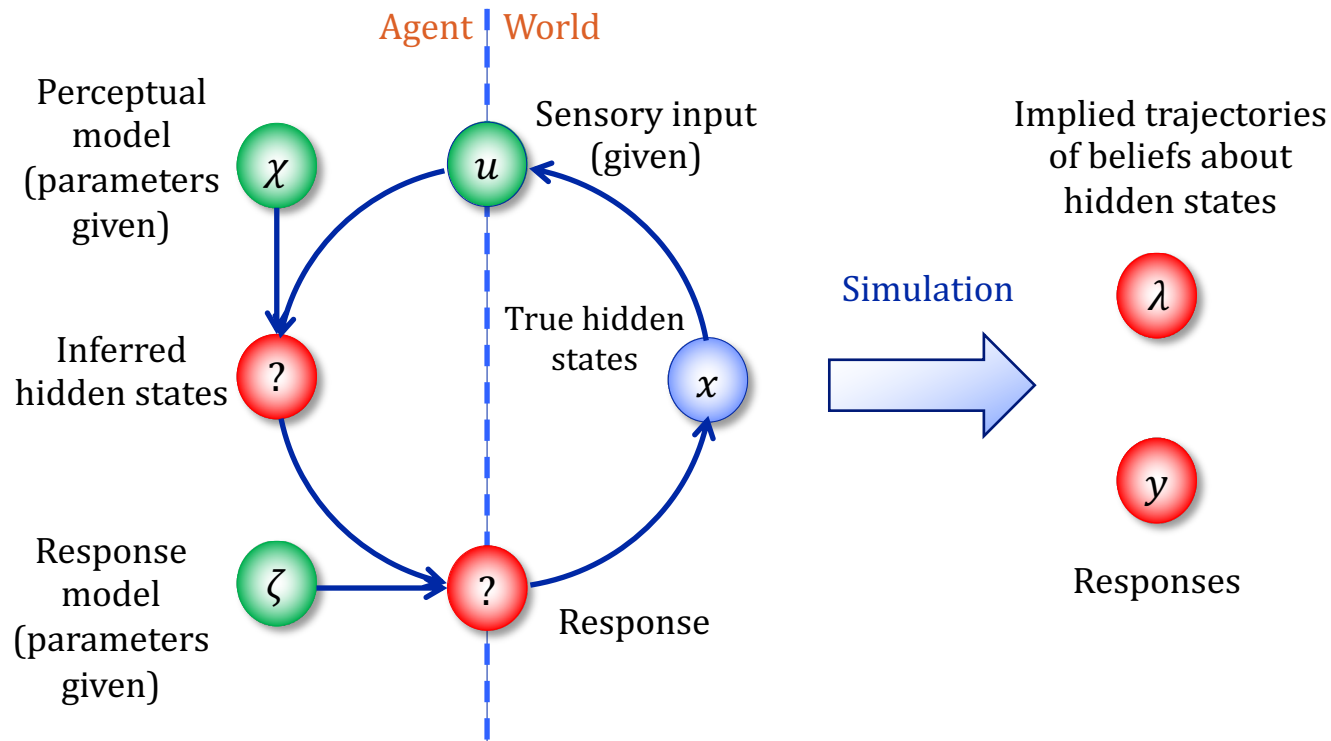


**Figure 6. Basal Forebrain Activations by  $\varepsilon_3$**

Activation of the basal forebrain by precision-weighted prediction error about stimulus probabilities  $\varepsilon_3$  within the anatomically defined mask. For visualization of the activation area, we overlay the results thresholded at  $p < 0.05$  FWE peak-level corrected for the entire anatomical mask (red) on the results thresholded at  $p < 0.001$  (yellow; the yellow cluster also survives  $p < 0.05$  FWE cluster-level correction for the entire anatomical mask). The anatomical mask comprised both dopaminergic and cholinergic brain structures: VTA/SN, PPT/LDT, and basal forebrain. (A) and (B) show results from the first (A: local maximum at  $x = 4$ ,  $y = 12$ ,  $z = -11$ ,  $t = 4.71$ ) and the second fMRI study (B: local maximum at  $x = 0$ ,  $y = 10$ ,  $z = -8$ ,  $t = 5.09$ ). (C) shows the conjunction analysis ("logical AND") across both studies. To ease visual comparison with Iglesias et al. (2013), the figure sections (x and y coordinates are indicated on each panel) are not located at the local maxima but correspond closely to those in Iglesias et al. (2013).

Iglesias et al, 2019  
(Correction to Iglesias et al.,  
2013)

# Generation of new data: simulation



## **HGF Toolbox demo**

# Remarks on parameter recovery

- *Parameter recovery* means the ability to estimate the ‘ground truth’ parameter values put into a simulation
- This is never possible to do exactly because simulation is stochastic
- The stochastic nature of the simulation means that the parameter value best able to explain the simulated data is not the same as that used in the simulation. By nature, the simulation only noisily represents the ‘ground truth’ parameter values.
- It is therefore misleading to call the value used in the simulation ‘true’. One could just as well argue that the estimated value is the ‘true’ one since it best explains the data.
- In order to get an idea of how exactly a parameter can be recovered, run many (on the order of tens to hundreds) simulations and look at the distribution of parameter estimates.
- **Estimating a parameter that isn’t well recovered from simulations can still make sense because adjusting its value can improve the model by leading to a better explanation of the data**



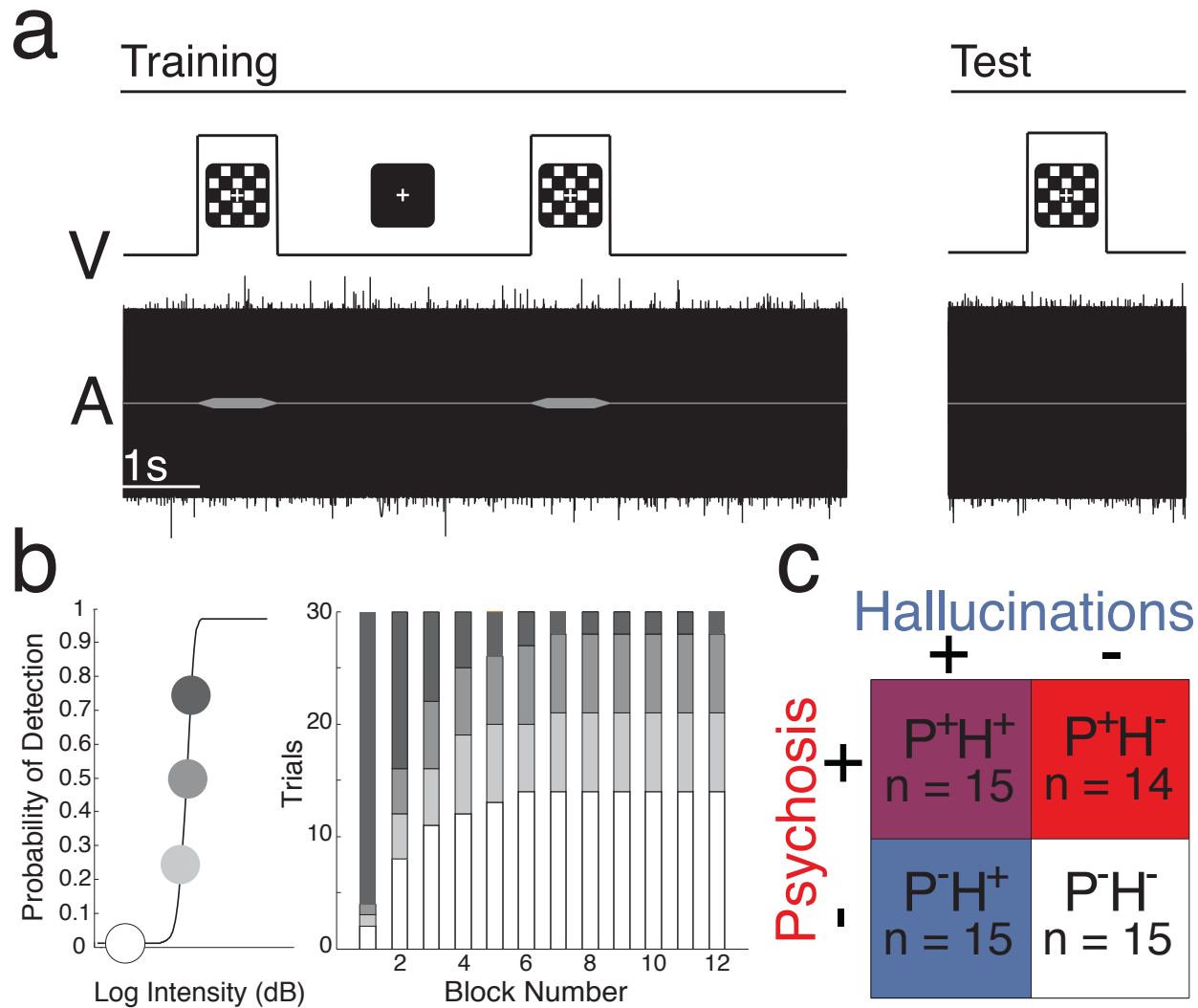
## Remarks on model comparison / model selection

- There is a range of scores that help in choosing a well-performing model: AIC (Akaike information criterion), BIC (Bayesian information criterion), Bayes factors, LME (log-model evidence), free energy, etc.
- Each model gets a particular score (which is on its own uninterpretable!)
- The difference in score between models is what counts
- However, model selection is not straightforward. AIC and BIC penalize complexity based on simple heuristics, which may not reflect complexity accurately. LME is better on that count, but is very sensitive to the modeller's choice of priors.
- **Three important considerations:**
  1. **Does the model allow me to answer my question of interest?**
  2. **Does the *prior predictive* distribution of observations make sense?**
  3. **Does the *posterior predictive* distribution of observations make sense?**

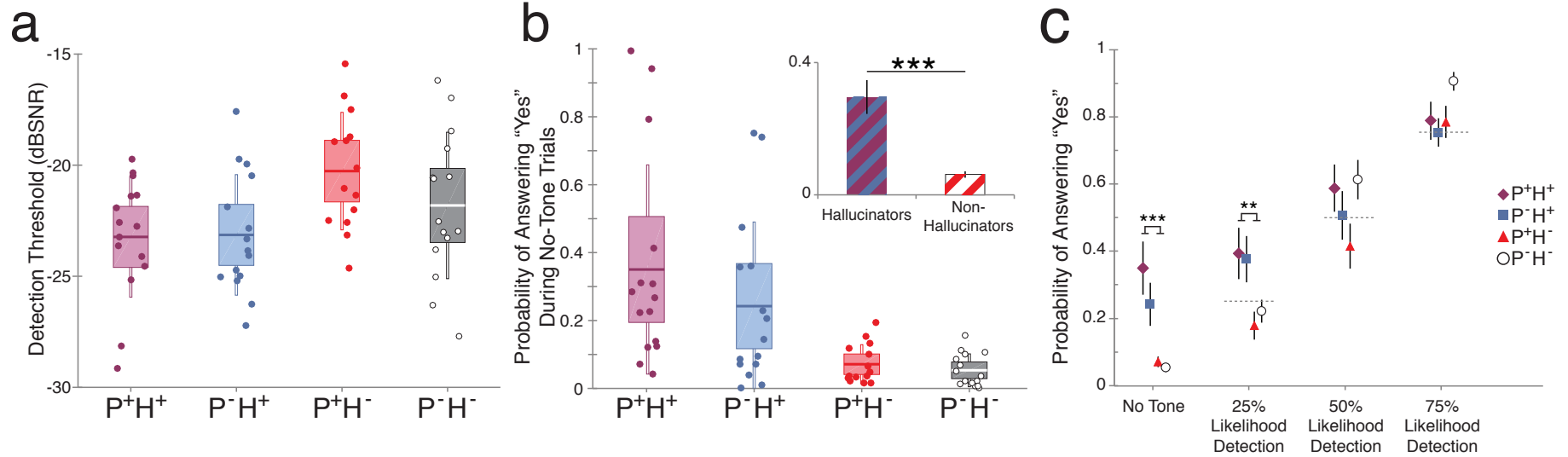
When the answer to all three is yes, the model is fine.

**Thanks**

# Conditioned hallucinations (Powers, Mathys, & Corlett, Science, 2017)



# Conditioned hallucinations



# Conditioned hallucinations

- Belief/percept formation model specifically created for this task
- Probability that the subject will respond “yes” to detection on a given trial:

$$P(\text{"yes"}|\textit{belief}) = \textit{sigmoid}(\textit{belief})$$

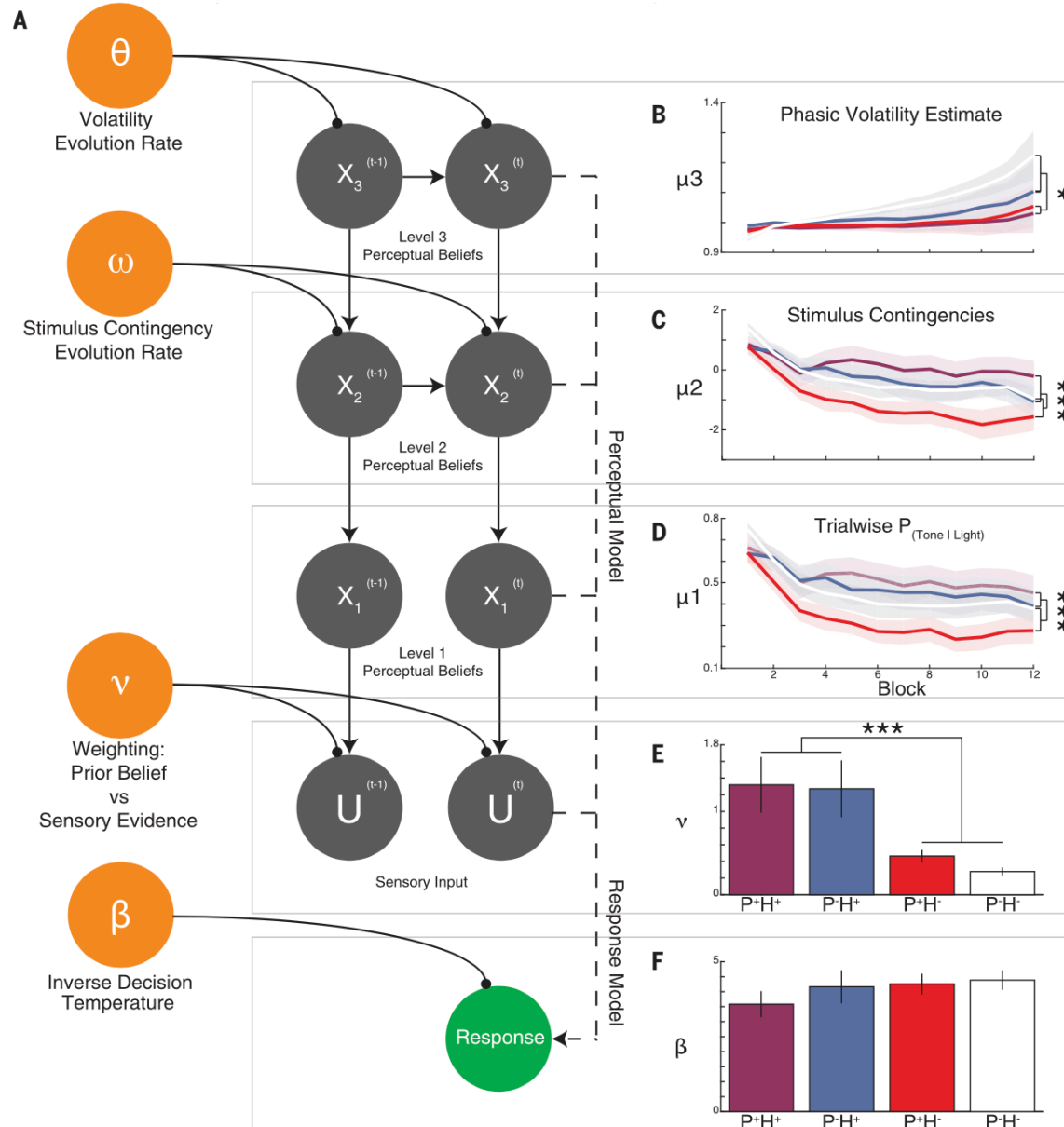
- *belief* is formalized as the Bayesian posterior mean of a beta distribution

$$\textit{belief} = \textit{prior} + \frac{1}{1 + \nu} (\textit{input} - \textit{prior})$$

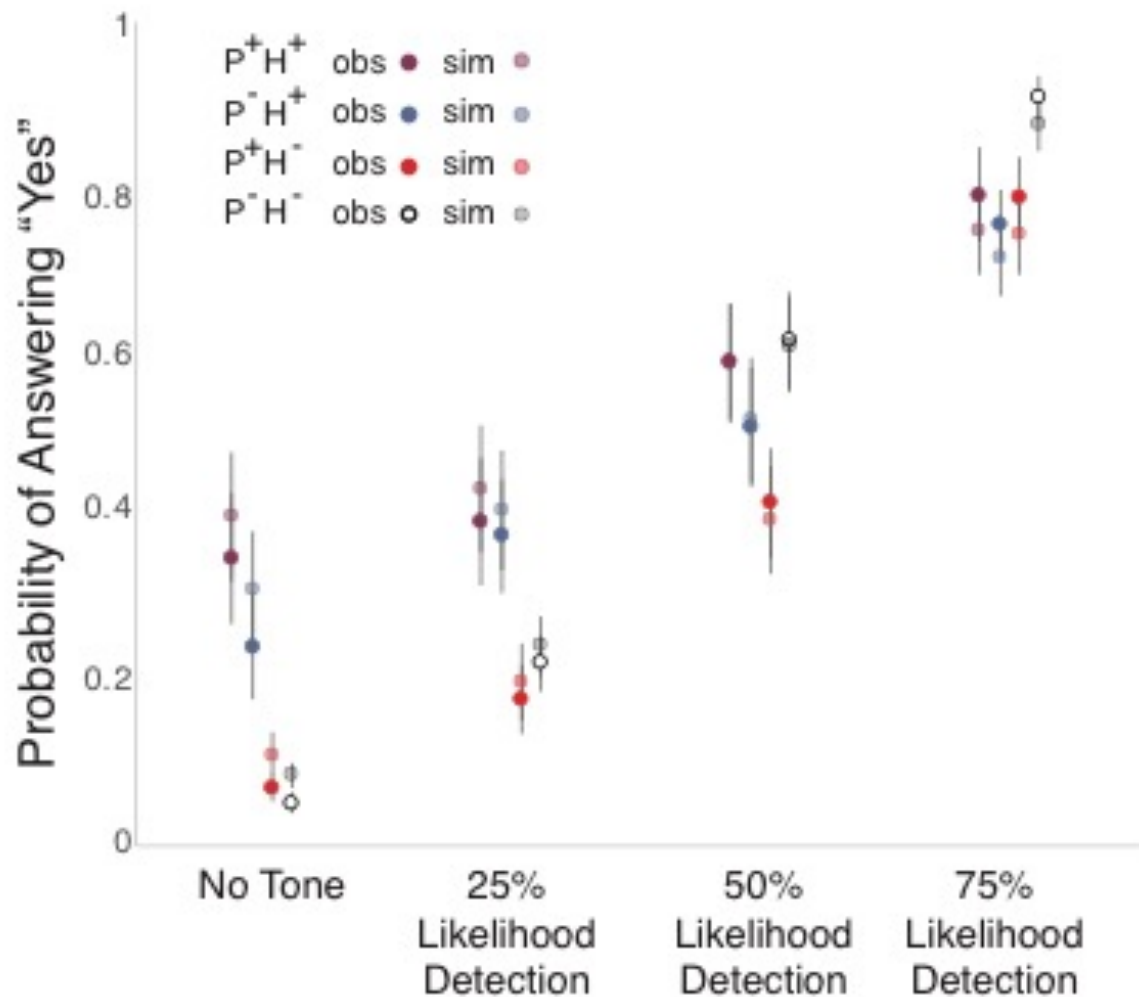
where

- *input* is given by experimental design: the true positive rate of the tone presented without light at each trial: 25%, 50%, or 75%
- *prior* is the prior from learning using the HGF:  $\hat{\mu}_1$
- $\nu$  is a subject-specific parameter indicating the relative weight of the prior compared to the input.
- For  $\nu = 1$ , prior and input have equal weight; for  $\nu > 1$  the prior has more weight than the input; and for  $\nu < 1$  the input has more weight than the prior.

# Conditioned hallucinations



# Conditioned hallucinations



# Conditioned hallucinations

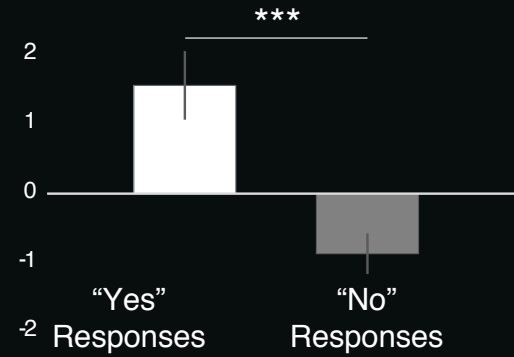
a

Thresholding, Tone-Responsive Regions



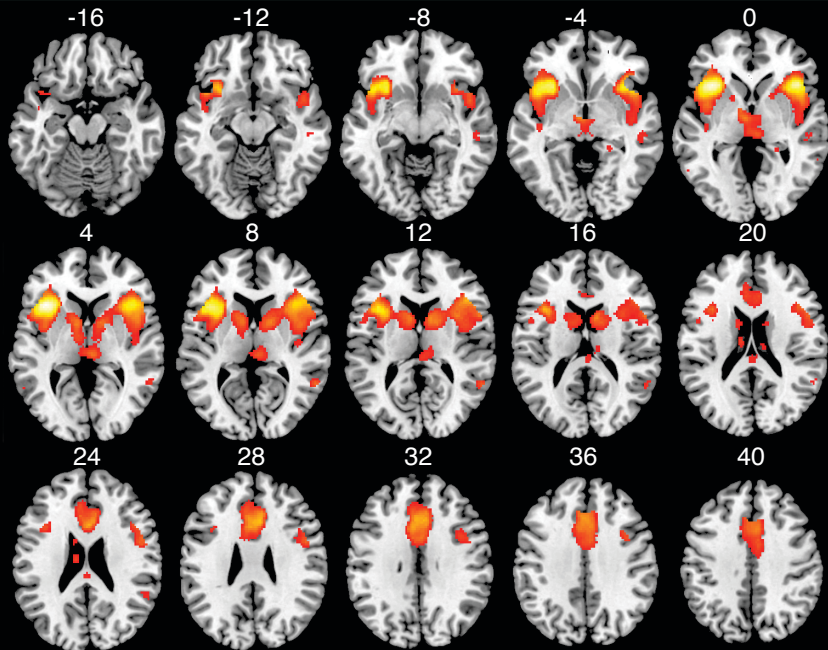
b

First Eigenvariate,  
No-Tone Test Trials



c

Test, "Yes" > "No" Responses, No-Tone Trials



d

Areas Active During AVH Symptom Capture

