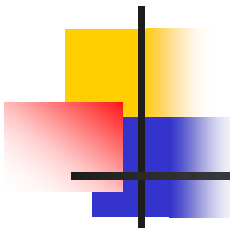


Bayesian Method for Repeated Threshold Estimation

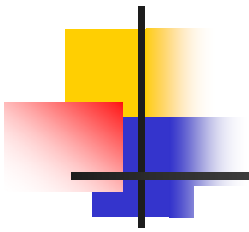
Alexander Petrov

Department of Psychology
Ohio State University

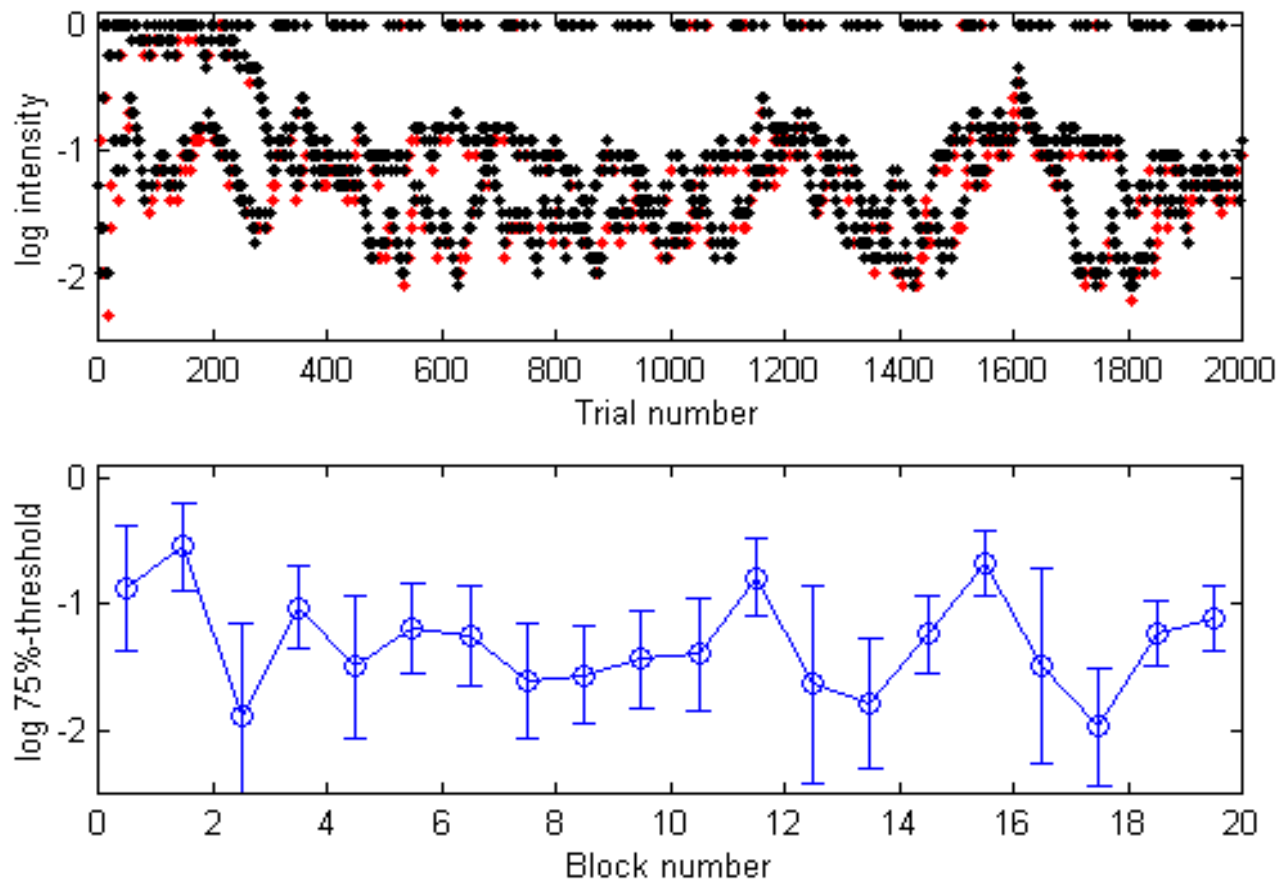


Motivation: Perceptual Learning

- Non-stationary thresholds
- Dynamics of learning is important
- Must use naïve observers
- Low motivation → high lapsing rates
- Slow learning → many sessions
- Large volume of low-quality binary data



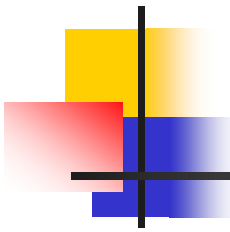
Objective: Data Reduction



Isn't This a Solved Problem?

- Up/down (Levitt, 1970)
- PEST (Taylor & Creelman, 1967)
- BEST PEST (Pentland, 1980)
- QUEST (Watson & Pelli, 1979)
- ML-Test (Harvey, 1986)
- Ideal (Pelli, 1987)
- YAAP (Treutwein, 1989)
- and many others...





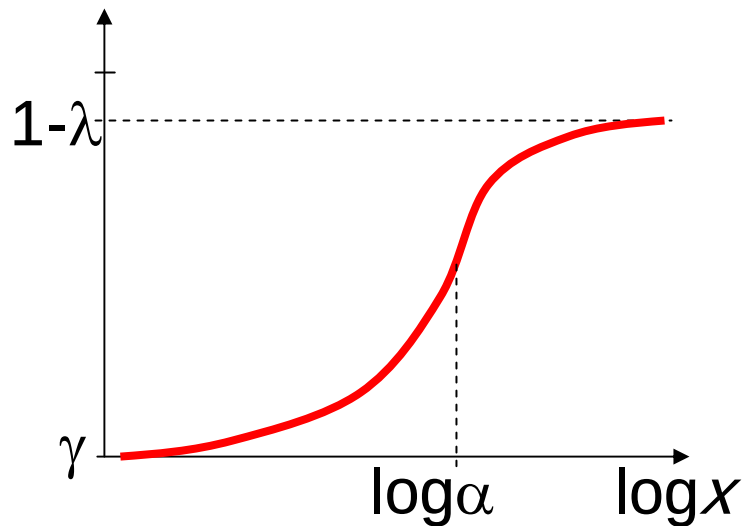
We Solve a Different Problem

- Standard methods:
 - Adaptive stimulus placement
 - Stopping criterion
 - Threshold estimation
- Our method:
 - Threshold estimation
 - Integrate information across blocks

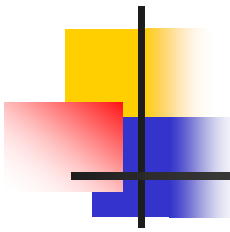
Weibull Psychometric Function

$$W(x; \alpha, \beta) = 1 - \exp(-\exp((\log x - \log \alpha)\beta))$$

$$P(x; \alpha, \beta, \gamma, \lambda) = \gamma + (1 - \gamma - \lambda)W(x; \alpha, \beta)$$



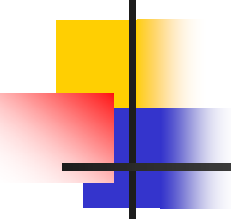
- Threshold $\log \alpha$
- Slope β
- Guessing rate γ
- Lapsing rate λ



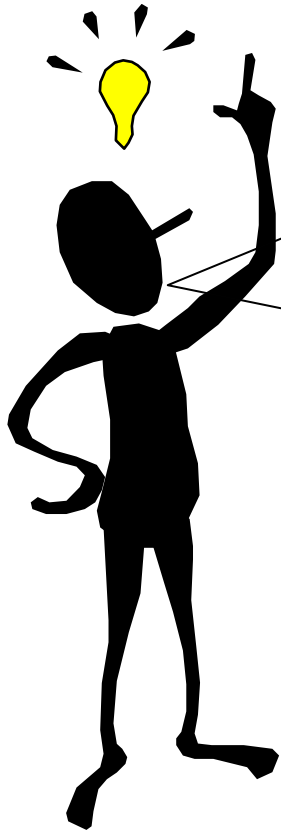
Two Kinds of Parameters

- Threshold $\log \alpha$
 - Slope β
 - Guessing rate γ
 - Lapsing rate λ
- Parameters of interest θ
- Nuisance parameters ϕ

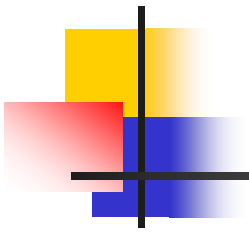
The nuisance parameters are harder to estimate but change more slowly than the threshold parameter.



Get the Best of Both Worlds



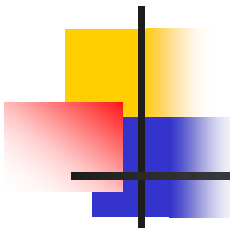
Use long data sequences to constrain the nuisance parameters; use short sequences to estimate the thresholds.



Joint Posterior of θ_k, ϕ

$$p(\theta_k, \phi \mid \mathbf{y}_k; \mathbf{y}_1 \cdots \mathbf{y}_{k-1}, \mathbf{y}_{k+1} \cdots \mathbf{y}_n) =$$
$$\underbrace{p(\mathbf{y}_k \mid \theta_k, \phi)}_{\text{Likelihood of current data}} \underbrace{p(\theta_k)}_{\text{Priors}} \underbrace{p(\phi) \prod_{i \neq k} \int p(\theta_i) p(\mathbf{y}_i \mid \theta_i, \phi) d\theta_i}_{\text{Information about } \phi \text{ extracted from the other data sets}}$$

Modified prior for the current block



Two-Pass Algorithm

- Pass 1: for each block i , calculate

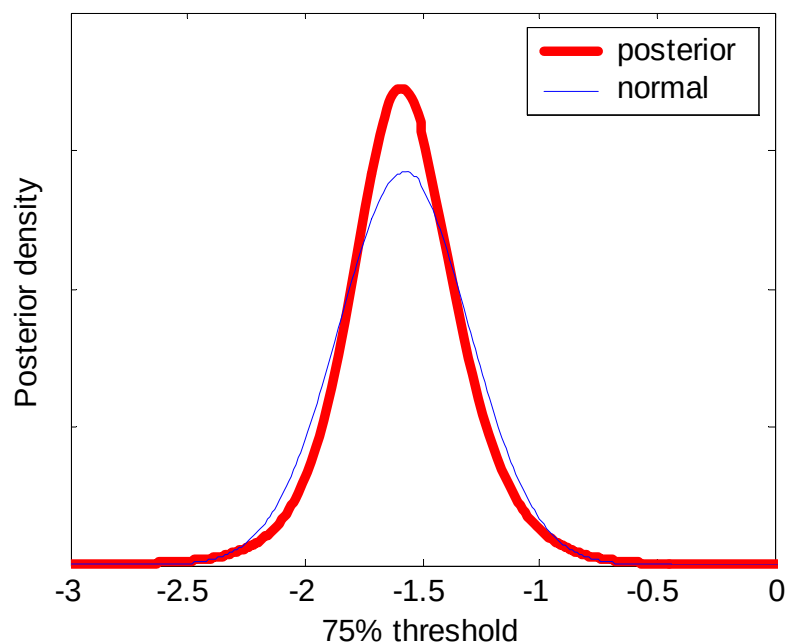
$$p(\phi | \mathbf{y}_i) = \int p(\theta) p(\mathbf{y}_i | \theta, \phi) d\theta$$

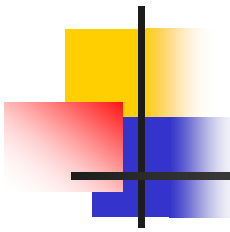
- Pass 2: for each block k , calculate

$$p(\theta_k, \phi | \mathbf{y}_k) = p(\mathbf{y}_k | \theta_k, \phi) p(\theta_k) p(\phi) \prod_{i \neq k} \int p(\phi | \mathbf{y}_i)$$

Posterior Thresholds

$$p(T_k) = \int P^{-1}(.75; \theta_k, \phi) p(\theta_k, \phi | \mathbf{y}_k) d\theta_k d\phi$$





Some Details

- Vaguely informative priors:

$$p(\log \alpha) \propto N(\mu_{\alpha}, \sigma_{\alpha})$$

$$p(\beta) \propto N(\mu_{\beta}, \sigma_{\beta})$$

$$p(\lambda) \propto \text{Beta}(a_{\lambda}, b_{\lambda})$$

- Implemented on a grid: $\log \alpha \times \beta \times \lambda$
- Assume $\gamma = .5$ for 2AFC data
- MATLAB software available at <http://alexpetrov.com/softw/>

Simulation 1: Stationary

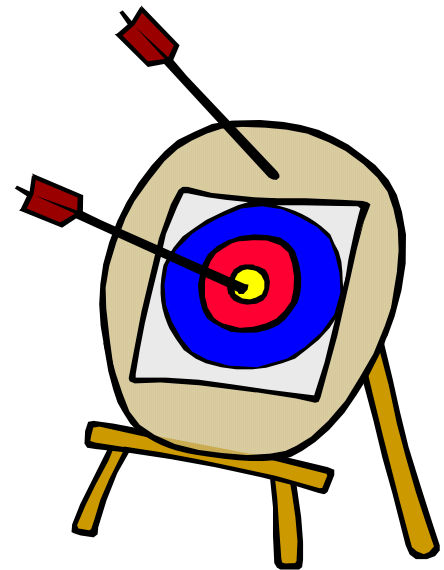
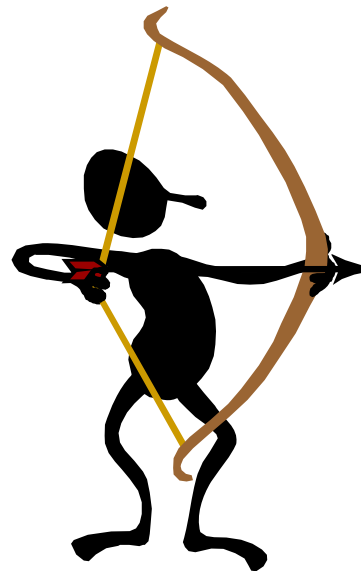
$$\log \alpha = -1.204 = \text{const}$$

$$\beta = 1.5$$

$$\lambda = .10$$

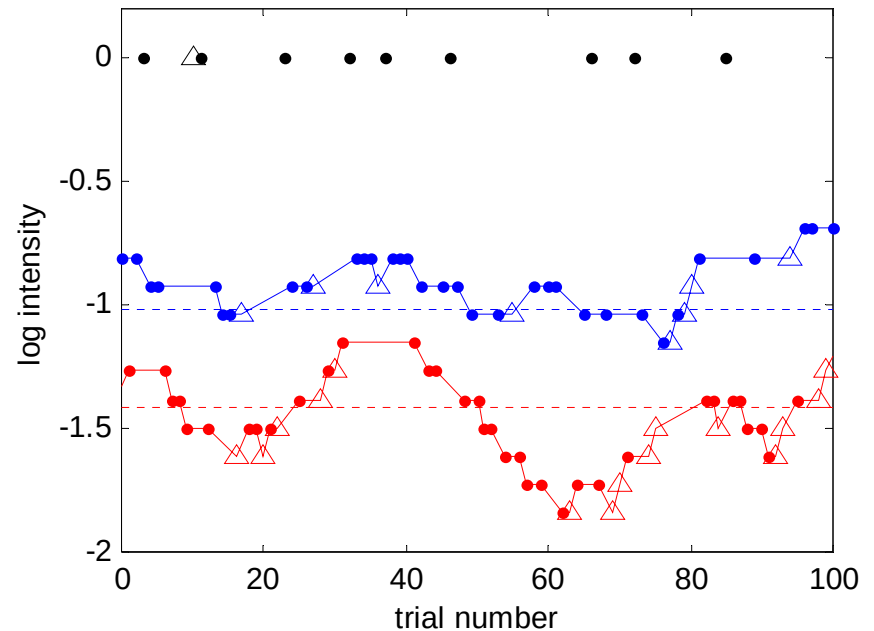


$$T_{75} = -1.217$$



Stimulus Placement

- 2 interleaved staircases
- 100 trials/block
 - 10 catch
 - 40 x 3down/1up
 - 50 x 2down/1up
- 100 runs of 12 blocks each

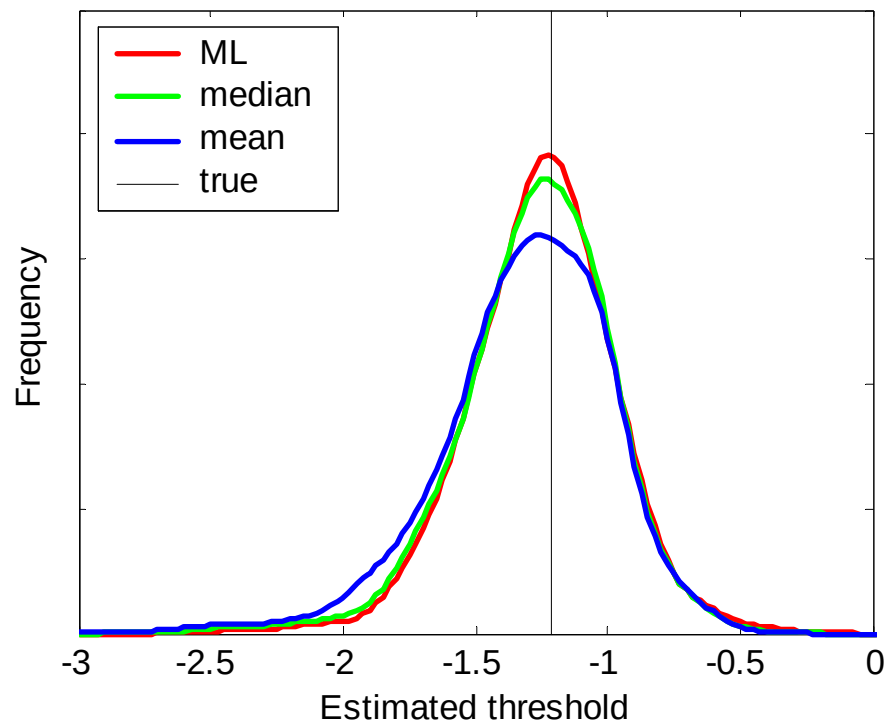


Threshold Estimators

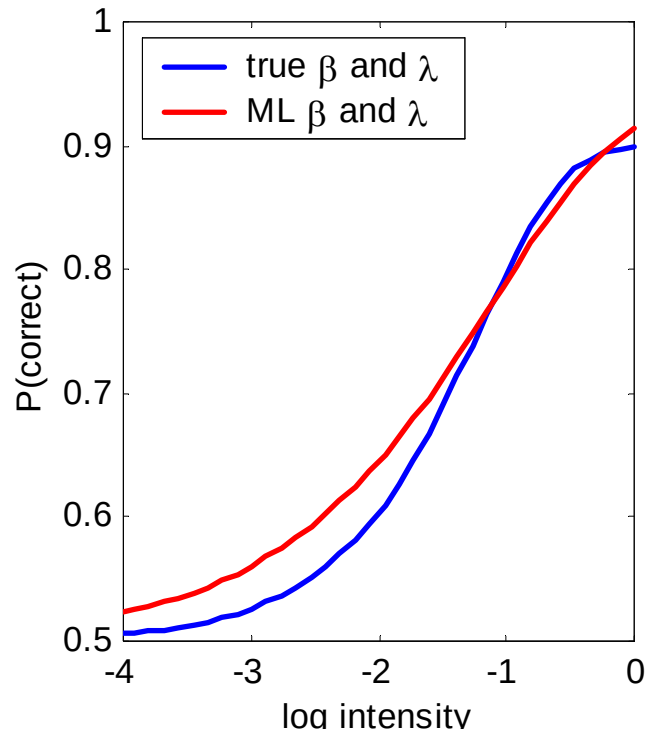
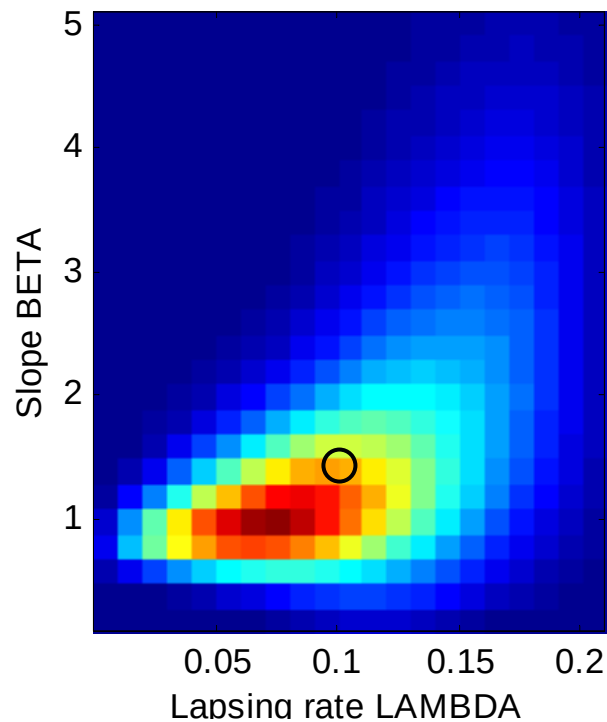
Estimator	Mean	Med	Std
ML	-1.24	-1.23	.27
Median	-1.26	-1.23	.28
Mean	-1.30	-1.27	.31
Std. dev.	0.41	0.36	.15

1200 Monte Carlo estimates

True 75% threshold = -1.217



$\beta \times \lambda$ Distribution from Pass 1



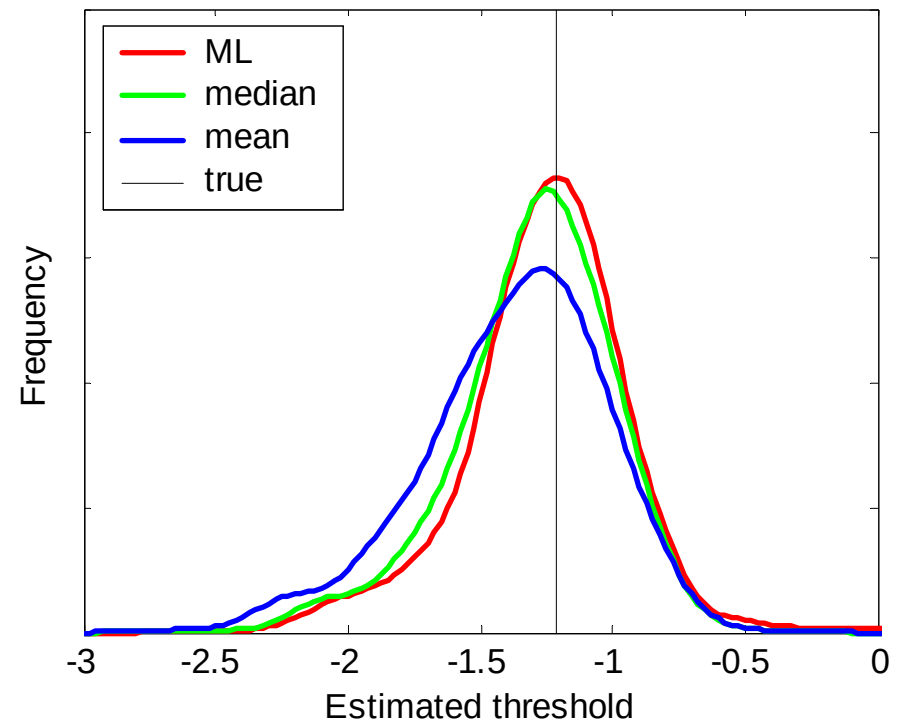
Catch Trials Are Worthwhile

Estimator	Mean	Med	Std
ML	-1.24	-1.22	.31
Median	-1.29	-1.26	.30
Mean	-1.36	-1.33	.34
Std. dev.	0.58	0.57	.16

1200 Monte Carlo estimates

No catch trials presented

True 75% threshold = -1.217

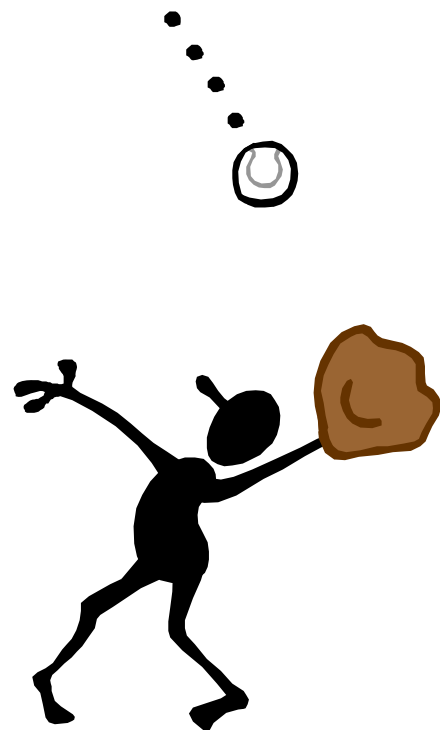
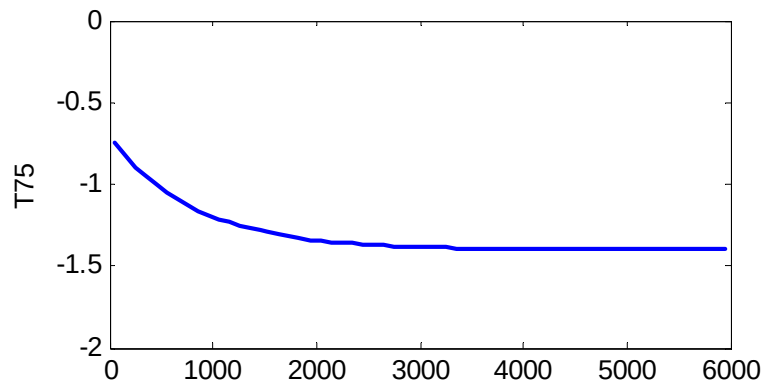


Simulation 2: With Learning

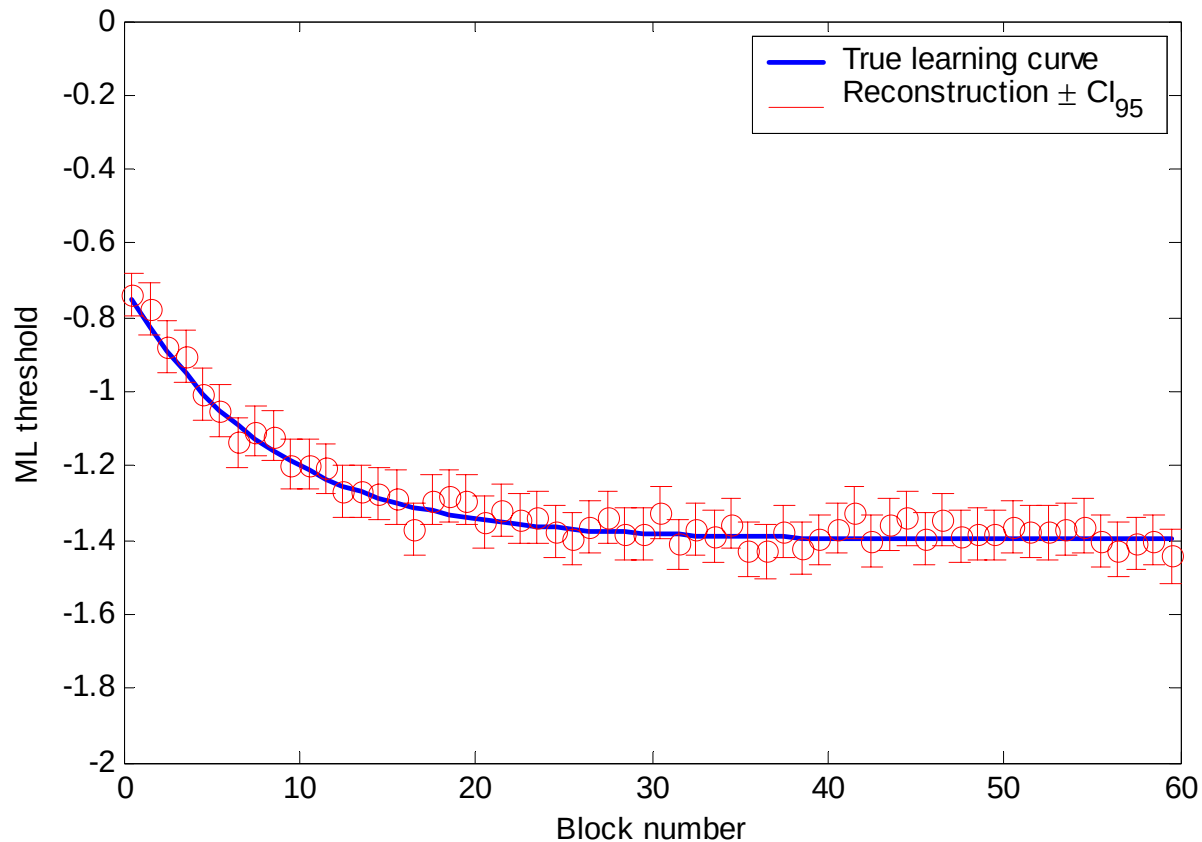
$$\log \alpha = -0.693 (e^{-t/800} - 2)$$

$$\beta = 1.5$$

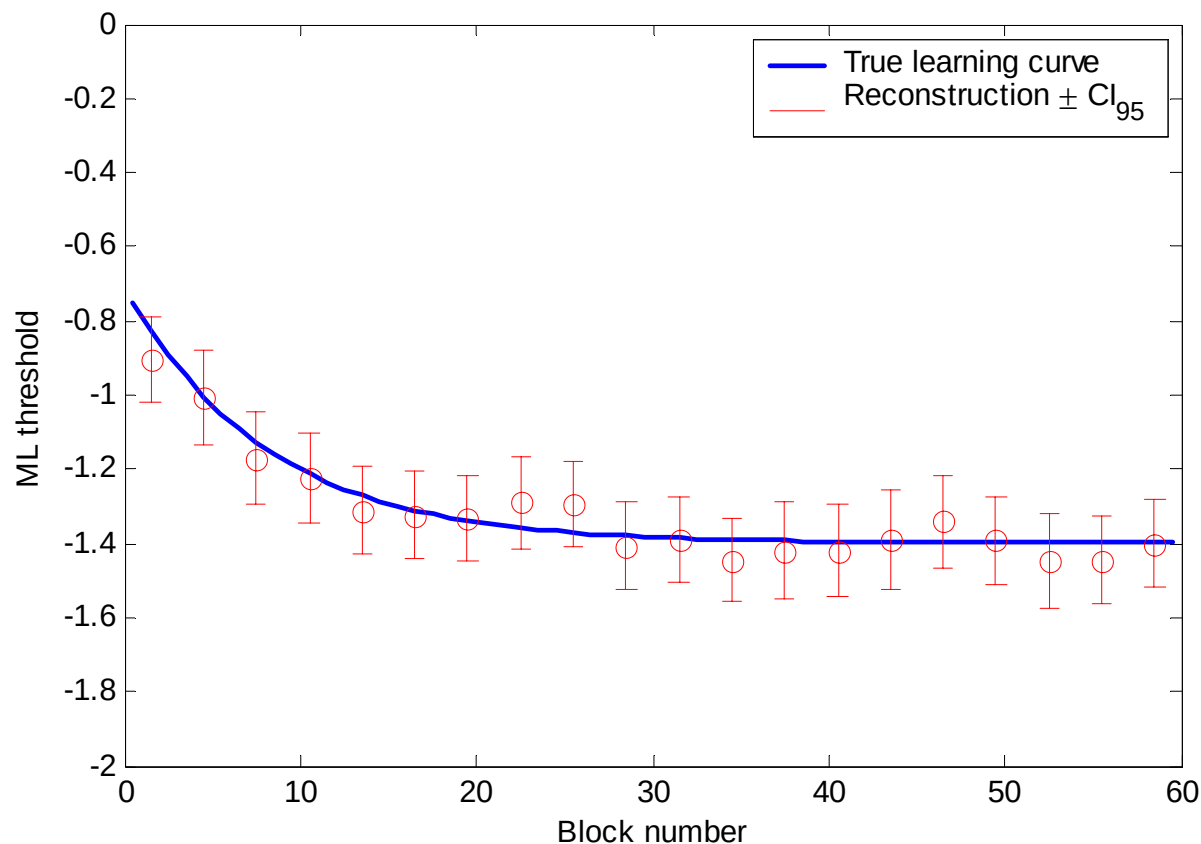
$$\lambda = .10$$

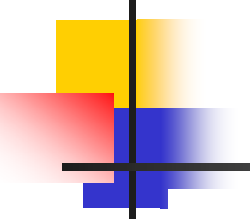


Group Learning Curve, N=100

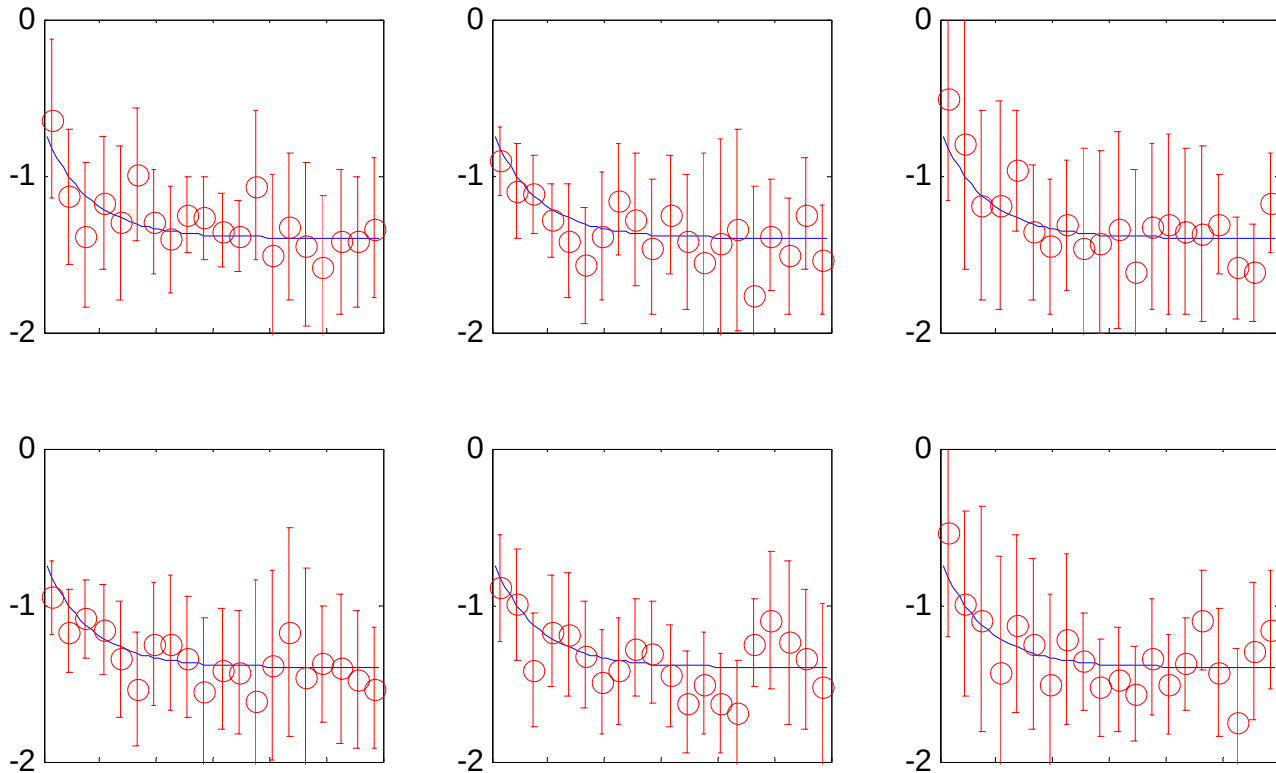


More Realistic Sample, N=10





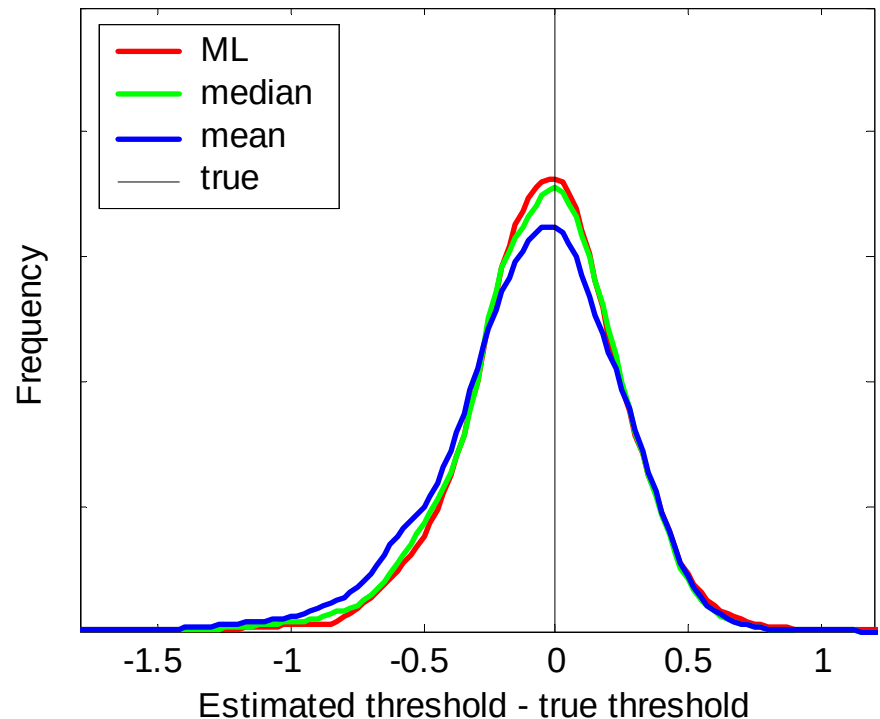
Individual Runs



The Method Performs Well

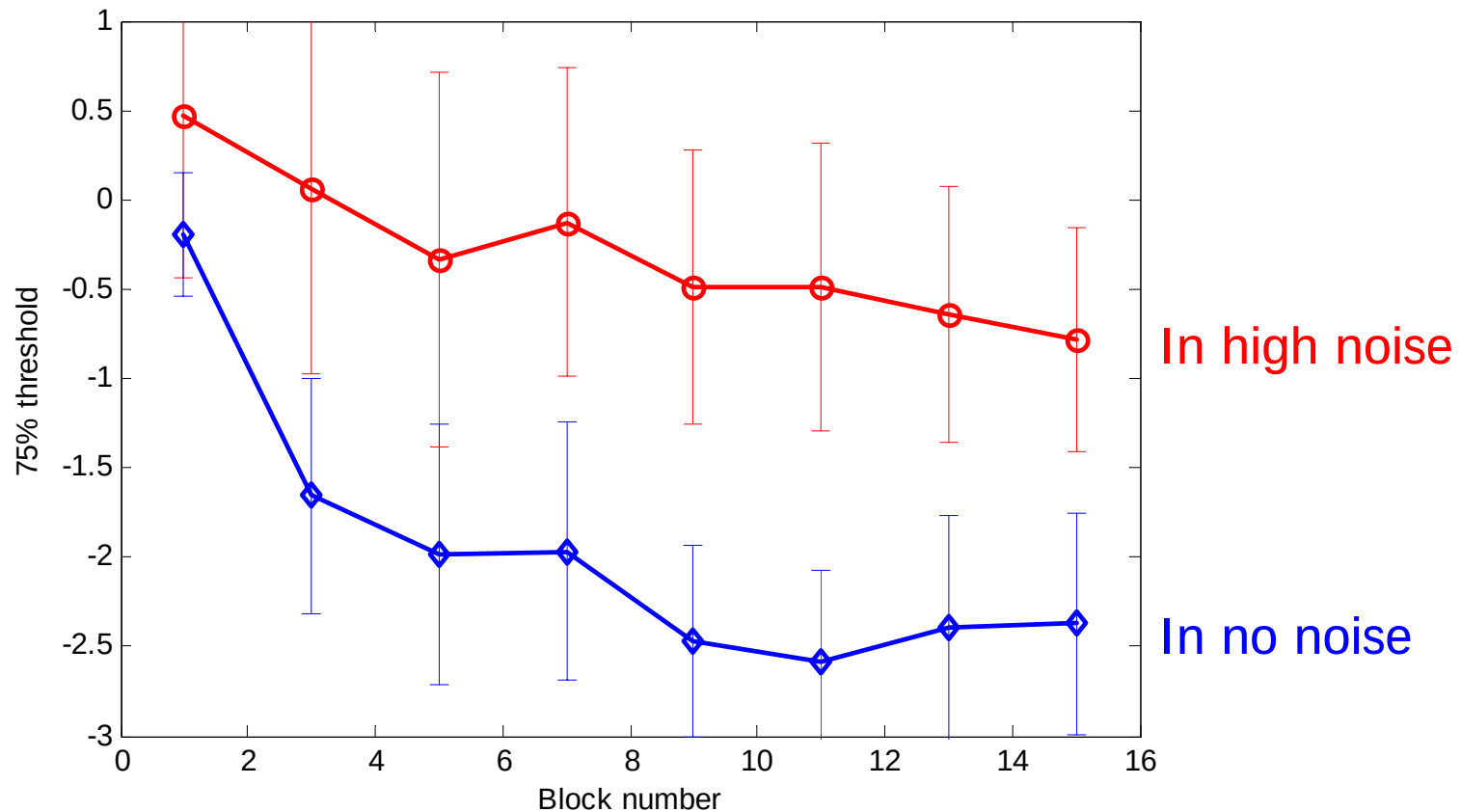
Estimator	Mean	Med	Std
ML	-0.03	-0.02	.28
Median	-0.05	-0.03	.29
Mean	-0.08	-0.05	.32
Std. dev.	0.42	0.39	.15

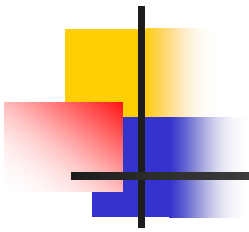
6000 Monte Carlo estimates
Similar to the stationary case
No systematic bias over time



Example: Actual Data, N=8

Jeter, Doshier, Petrov, & Lu (2005)

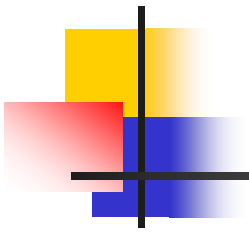




Future Work



- Sensitivity to priors?
- Compare with standard ML methods
- Individual differences
- Estimate slope in addition to threshold
- Non-stationary β and λ ?
- Recommended stimulus placement?
- Hierarchical models



The End

