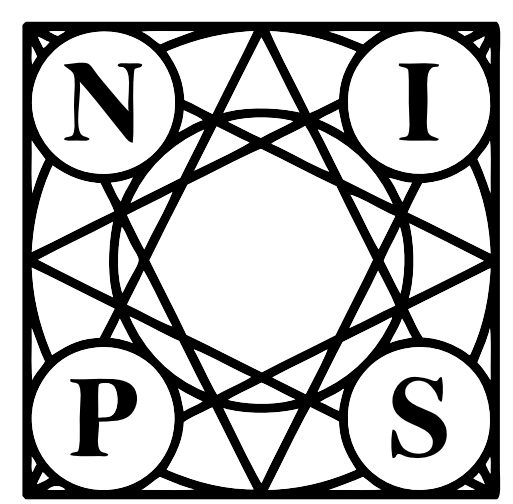
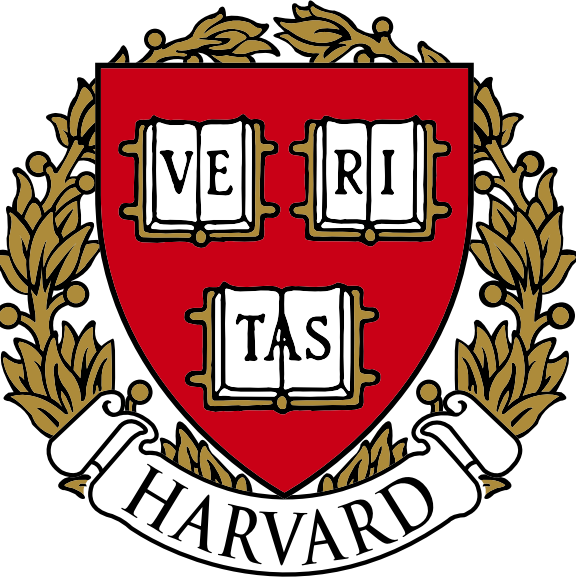




STOCHASTIC EXPECTATION PROPAGATION

Yingzhen Li¹, José Miguel Hernández-Lobato² and Richard E. Turner¹
¹University of Cambridge, ²Harvard University
{yl494, ret26}@cam.ac.uk, jmh@seas.harvard.edu



FROM EP TO STOCHASTIC EP

Goal: approximate the true posterior $q(\theta) \approx p(\theta|\mathcal{D})$

$$p(\theta|\mathcal{D}) \propto p_0(\theta) p(x_1|\theta) p(x_2|\theta) p(x_3|\theta) \approx q(\theta) \propto p_0(\theta) f_1(\theta) f_2(\theta) f_3(\theta)$$

idealised

$$p(\theta|\mathcal{D}) \propto p_0(\theta) p(x_1|\theta) p(x_2|\theta) p(x_3|\theta) \approx q^{new}(\theta) \propto p_0(\theta) p(x_1|\theta) p(x_2|\theta) f_3^{new}(\theta)$$

EP

$$\tilde{p}(\theta) \propto p_0(\theta) f_1(\theta) f_2(\theta) p(x_3|\theta) \approx q_{EP}^{new}(\theta) \propto p_0(\theta) f_1(\theta) f_2(\theta) f_3^{new}(\theta)$$

SEP

$$p(\theta|\mathcal{D}) \propto p_0(\theta) f(\theta)^{N-1} p(x_3|\theta) \approx q_{SEP}^{new}(\theta) \propto p_0(\theta) f^{new}(\theta)^N$$

memory $\mathcal{O}(ND^2)$
(Gaussian approx.)

memory $\mathcal{O}(D^2)$

AN ALTERNATIVE VIEW OF SEP

Equivalent factorisation: $\overline{p(x|\theta)} = (\prod_n p(x_n|\theta))^{1/N}$

$$p(\theta|\mathcal{D}) \propto p_0(\theta) p(x_1|\theta) p(x_2|\theta) p(x_3|\theta)$$

EP on the equivalent factorisation will converge to $f_n(\theta) = f(\theta)$ that captures the **average affect** of likelihood terms on posterior!

EP (on the equivalent factorisation)

$$p(\theta|\mathcal{D}) \propto p_0(\theta) f(\theta)^{N-1} \overline{p(x|\theta)} \approx q^{new}(\theta) \propto p_0(\theta) f^{new}(\theta)^N$$

SEP

$$p(\theta|\mathcal{D}) \propto p_0(\theta) f(\theta)^{N-1} p(x_3|\theta) \approx q_{SEP}^{new}(\theta) \propto p_0(\theta) f^{new}(\theta)^N$$

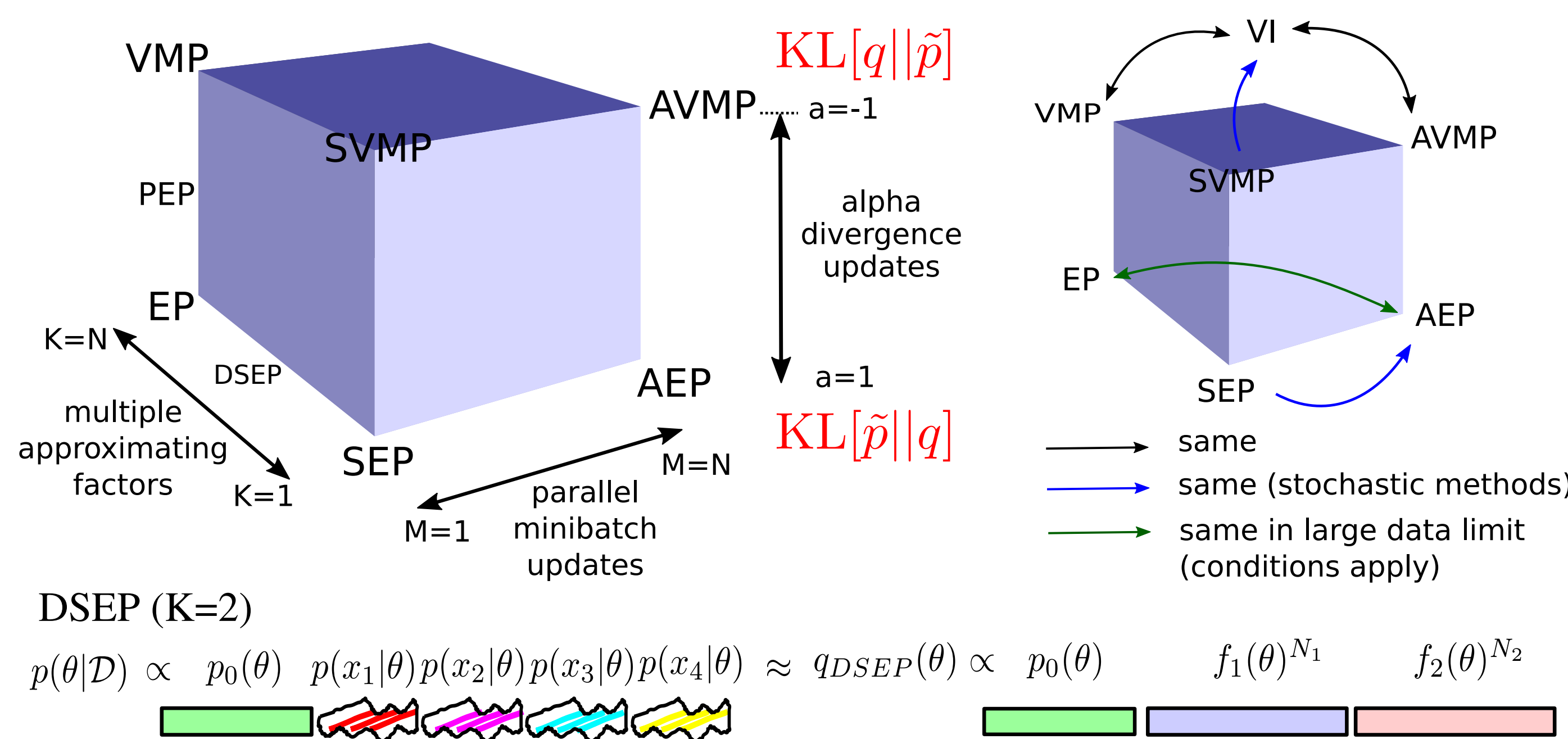
Work in progress:
guarantees of lower bound
analysis of approximation error

intractable

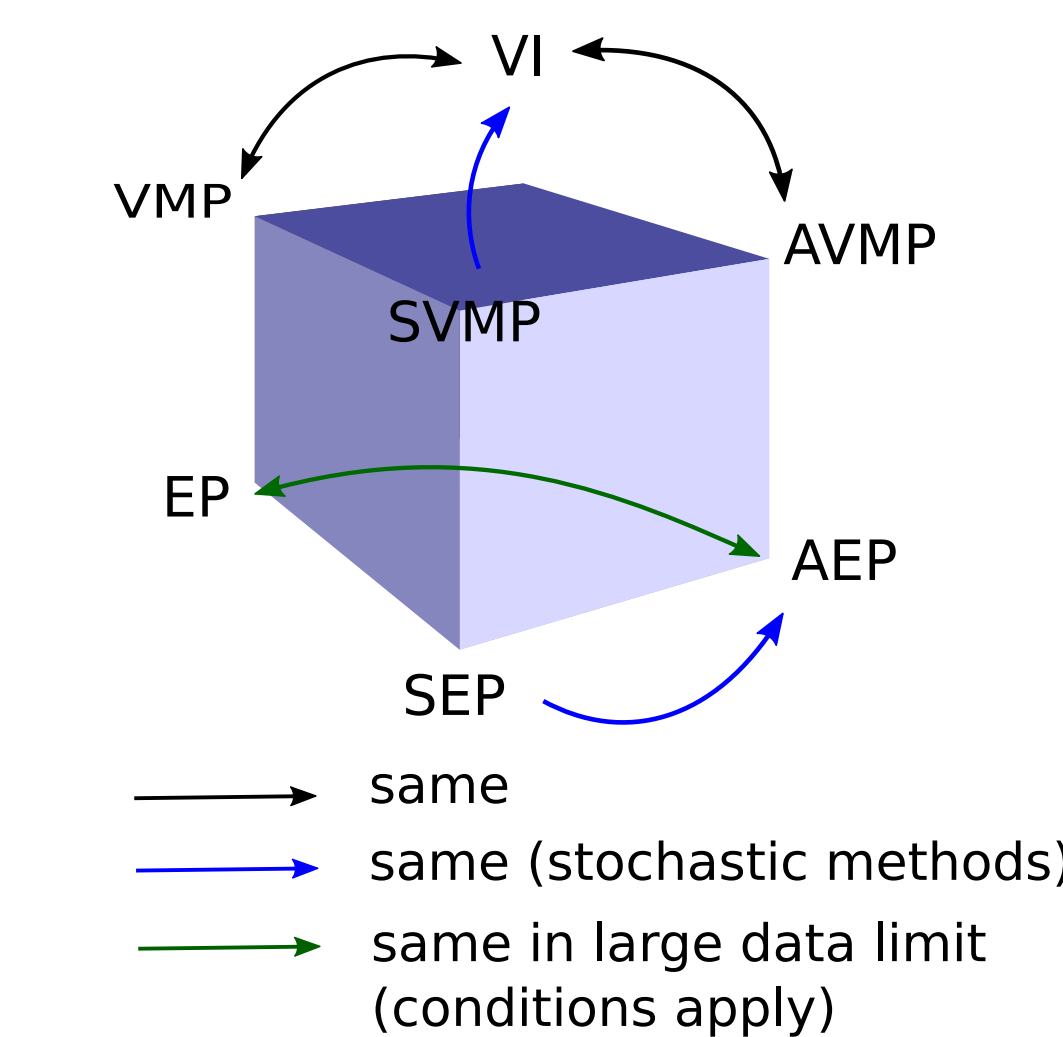
stochastic approximation

RELATED ALGORITHMS

A) Relationships between algorithms



B) Fixed points properties



SEP FOR LATENT VARIABLE MODELS

Model: $p(\mathbf{x}_n|\mathbf{h}_n, \theta), \mathbf{h}_n \sim p_0(\mathbf{h}_n)$ i.i.d., posterior

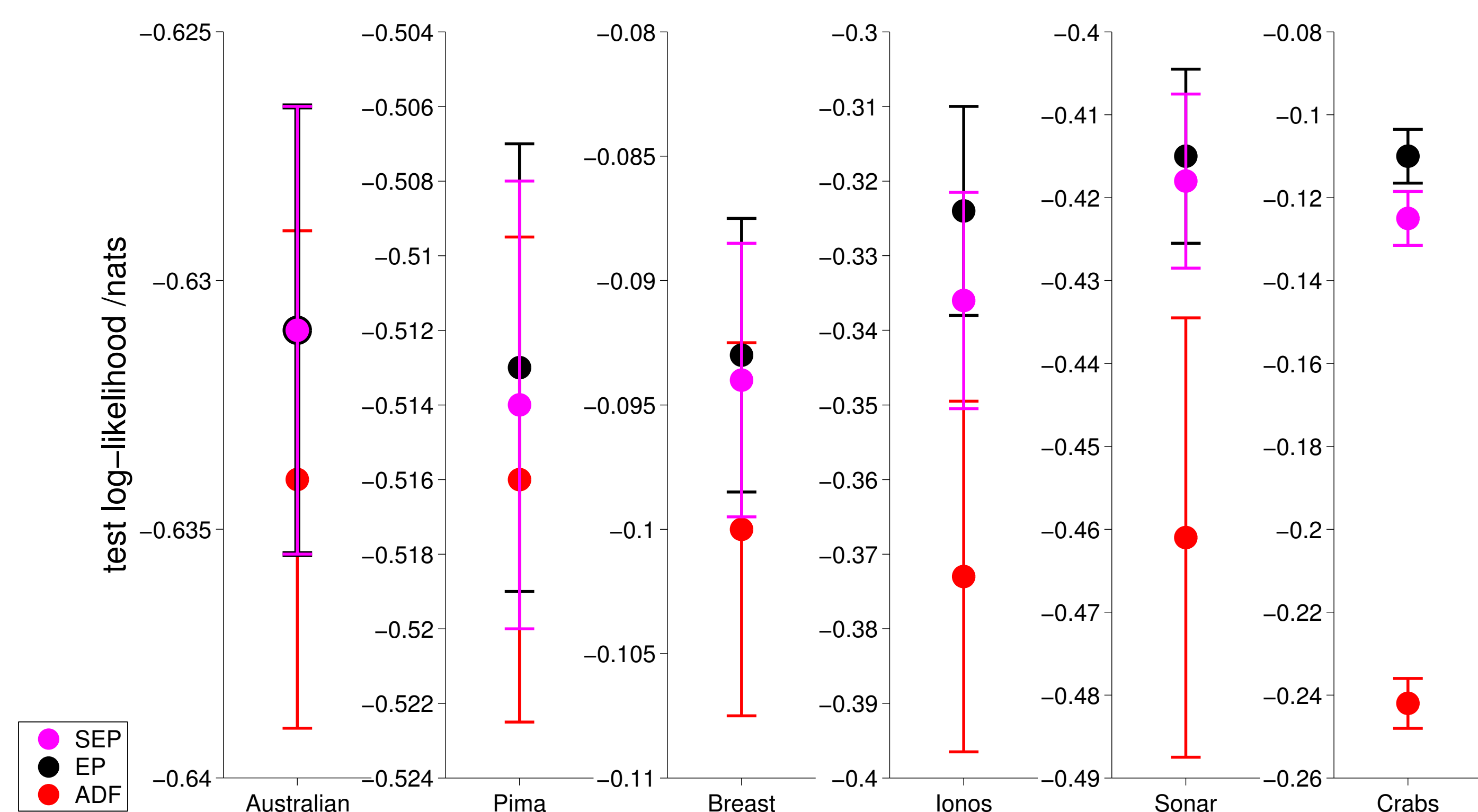
$$p(\theta, \{\mathbf{h}_n\}|\mathcal{D}) \approx q(\theta, \{\mathbf{h}_n\}) \propto p_0(\theta) f(\theta)^N \prod_{n=1}^N g_n(\mathbf{h}_n)$$

Update $f(\theta)g_n(\mathbf{h}_n)$ such that

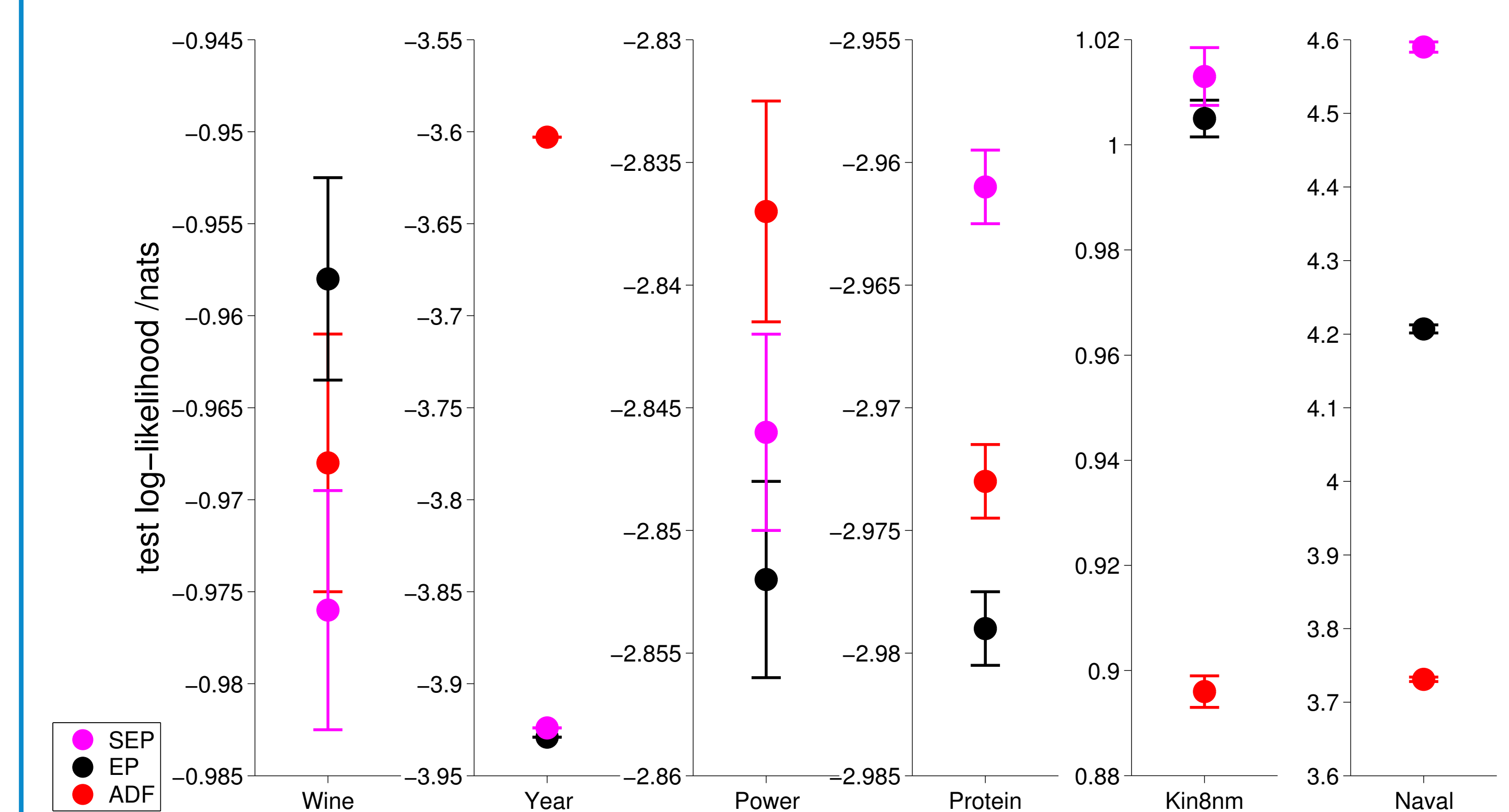
$$p_0(\theta) f(\theta)^{N-1} \prod_{m \neq n} g_m(\mathbf{h}_m) p(\mathbf{x}_n|\mathbf{h}_n, \theta) p_0(\mathbf{h}_n) \approx p_0(\theta) f(\theta)^{N-1} \prod_{m \neq n} g_m(\mathbf{h}_m) f^{new}(\theta) g_n^{new}(\mathbf{h}_n)$$

No need to maintain g_n if updates can be computed fast!

BAYESIAN PROBIT CLASSIFICATION



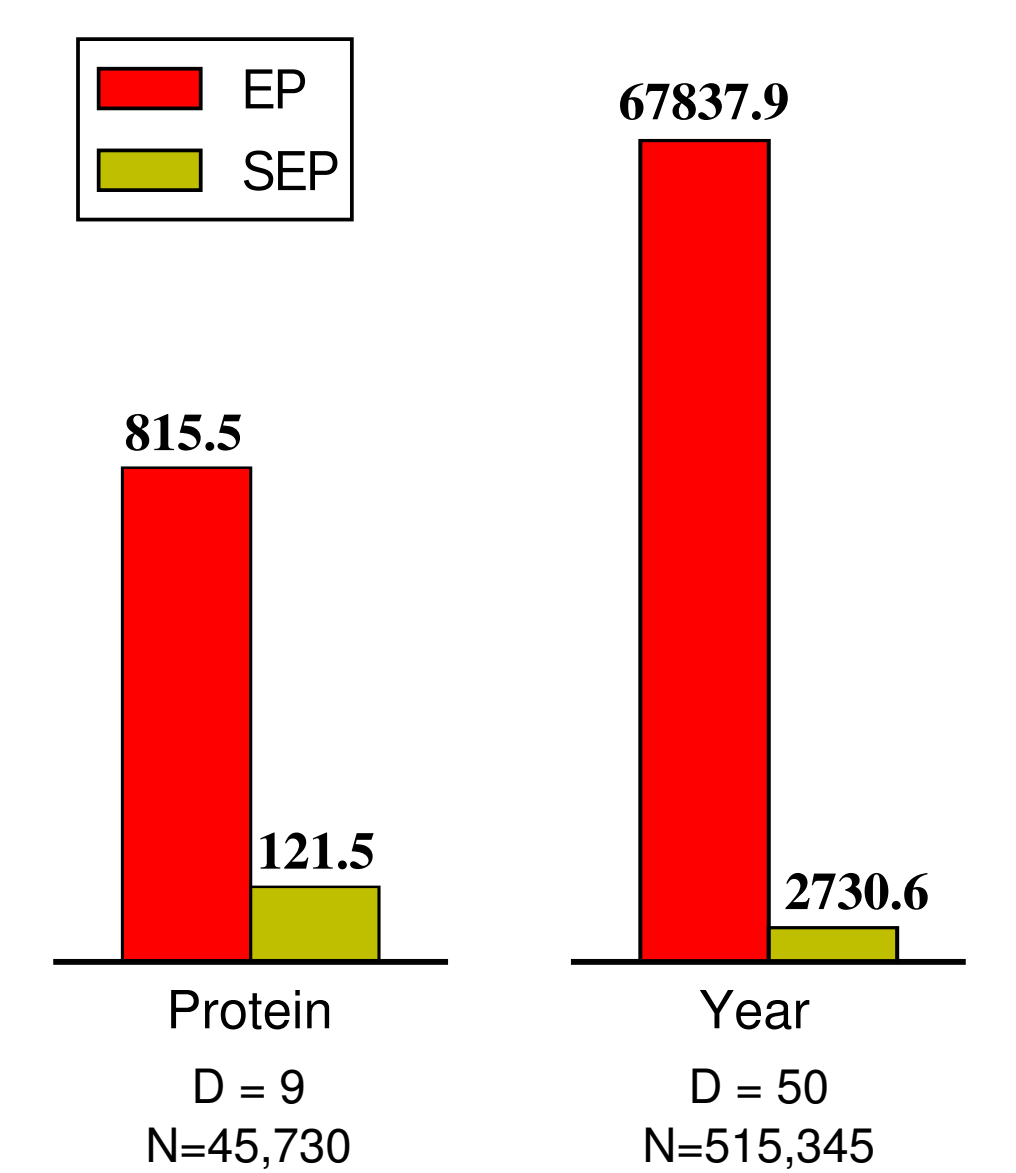
BAYESIAN NEURAL NETWORK



- Prior:** $p_0(\lambda|\gamma) = \text{Gam}(\alpha_0^\lambda, \beta_0^\lambda)$, $p_0(\mathbf{w}|\lambda) = \prod_{ijl} \mathcal{N}(w_{ij,l}; 0, \lambda^{-1})$
- Likelihood:** $p(\gamma) = \text{Gam}(\alpha_0^\gamma, \beta_0^\gamma)$, $p(y|\mathbf{x}, \mathbf{w}, \gamma) = \mathcal{N}(y; f(\mathbf{x}; \mathbf{w}), \gamma^{-1})$
- $f(\mathbf{x}; \mathbf{w})$ is the output of a neural network with ReLU units.
- Inference:** PBP^a works better when under-estimating uncertainty.

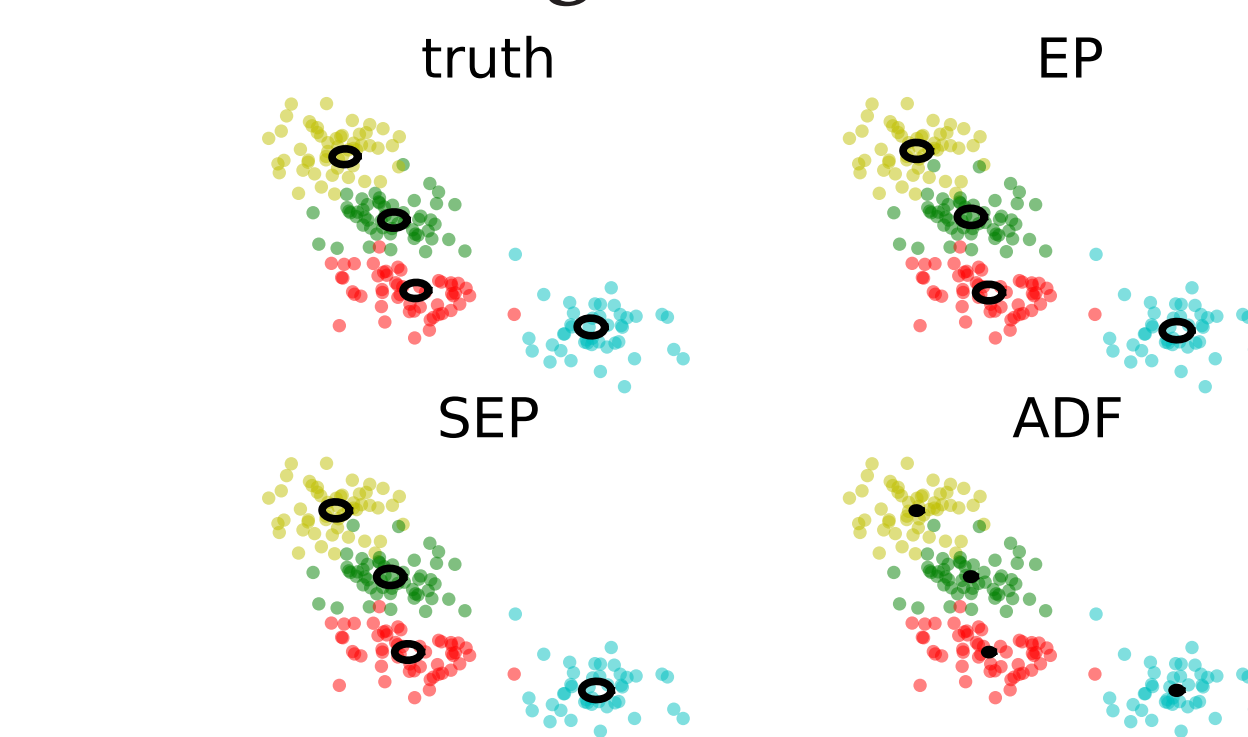
^aHernández-Lobato J M, Adams R P. Probabilistic Backpropagation for Scalable Learning of Bayesian Neural Networks. In ICML 2015.

Memory savings (in MB)

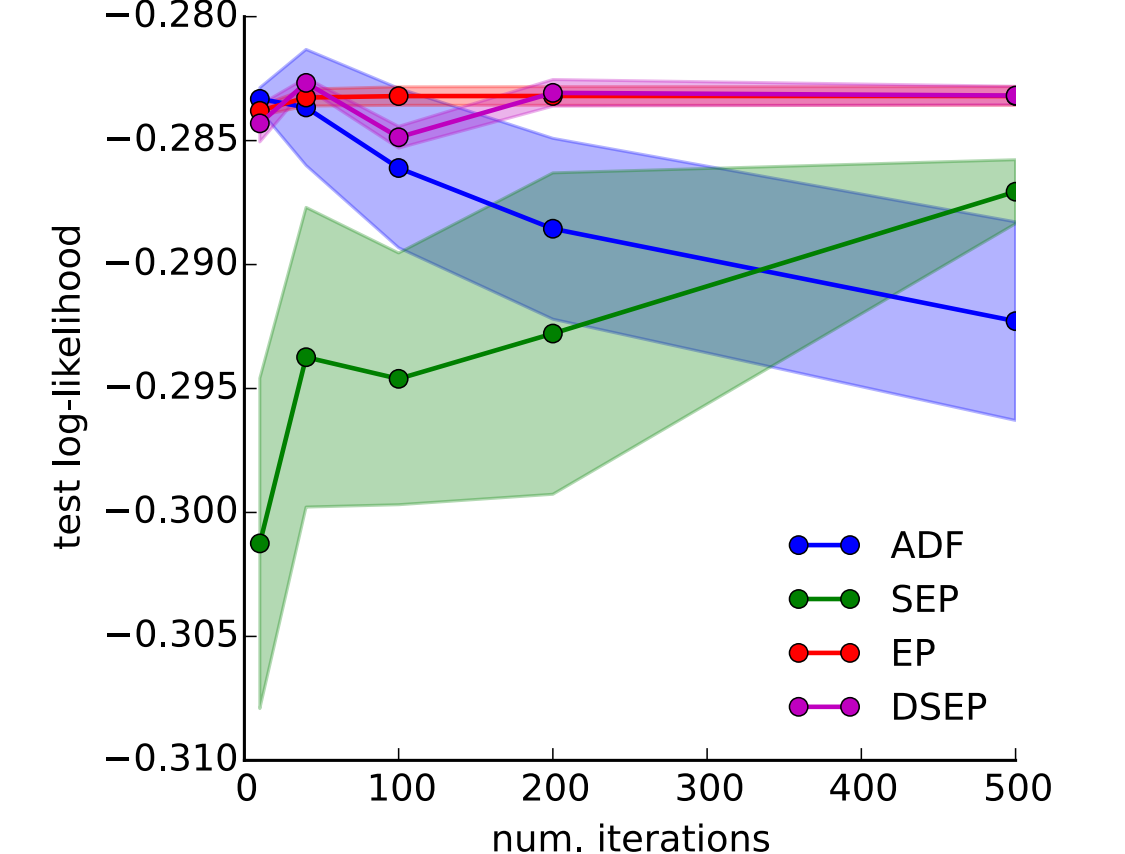


MORE EXAMPLES

1. Clustering with MoGs



2. Odd-vs-even digit classification



CONCLUSIONS

- We scaled EP to large datasets with nearly identical performances
- We extended stochastic EP and related it to variational Bayes
- SEP is well suited to “big model, big data” settings