

Convex order and preservation of convexity for Bayesian posterior updates

Yuqiong Wang
yuqw@umich.edu

Joint work with Erhan Bayraktar

Department of Mathematics
University of Michigan

Workshop on Optimal Stopping and Its Applications, Manchester

September 14, 2026

- 1 Motivation and problem setup
- 2 Main results
- 3 Examples and consequences

A motivating example

- A coin has unknown probability $\Theta \in (0, 1)$. Two observers, Alice and Bob observe $X_1, X_2 \dots$ and try to predict the next outcome. Each wrong prediction is penalized by 1.
- Conditioning on Θ , predicting T would have penalization $P(X_{n+1} = H|\Theta) = \Theta$, while predicting H has $1 - \Theta$.
- After n tosses, $\Pi_n = E[\Theta | \mathcal{F}_n]$ is the predictive probability of heads. The total penalization is $g(\Pi_n) = \min\{\Pi_n, 1 - \Pi_n\}$.
- Assume one is allowed to observe at most one more time with cost c , before making the prediction.

The one-step stopping problem is

$$V^{(1)}(n, z) = \min \{g(z), c + E_{n,z} [g(\Pi_{n+1})]\}.$$

Assume a Uniform prior $\text{Beta}(1, 1)$ for Θ . Alice observes one H and one T; Bob observes 10 heads and 10 tails:

$$\mu_2^A = \text{Beta}(2, 2), \quad \mu_{20}^B = \text{Beta}(11, 11), \quad \Pi_2^A = \Pi_{20}^B = \frac{1}{2}.$$

Question: who is better off (has a lower cost)?

A motivating example

- If Alice and Bob each do one more toss, each outcome has probability $1/2$:

$$\Pi_3^A \in \left\{ \frac{2}{5}, \frac{3}{5} \right\}, \quad \Pi_{21}^B \in \left\{ \frac{11}{23}, \frac{12}{23} \right\}.$$

- Thus Alice is better off: $\Pi_3^A \mid \Pi_2^A = \frac{1}{2} \geq_{\text{cx}} \Pi_{21}^B \mid \Pi_{20}^B = \frac{1}{2}$.
- For all $c \geq 0$,

$$V^{(1)}\left(2, \frac{1}{2}\right) = \min\left(\frac{1}{2}, c + \frac{2}{5}\right) \leq \min\left(\frac{1}{2}, c + \frac{11}{23}\right) = V^{(1)}\left(20, \frac{1}{2}\right).$$

- If $1/46 < c < 1/10$, Alice buys the observation and Bob stops.

Questions:

- Isn't more information better? By Jensen, $\Pi_m \leq_{\text{cx}} \Pi_n$ for $m < n$. But we consider the conditional one: coming back to the same estimate itself is evidence.
- Is Alice's advantage a consequence of Beta/Bernoulli conjugacy, or does it reflect a general property?

A motivating example

In this work, we show that it is a general property, for one-parameter exponential families, arbitrary priors, and any non-decreasing statistic T .

In particular, we show

- Along a posterior level curve, future updates decrease in convex order as observation time increases. (cf. [Ekström-Vaicenavicius](#), [Ekström-Karatzas-Vaicenavicius](#), [Ekström-Wang](#))
- Posterior transition operators preserve convexity. (cf. [Janson-Tysk](#), [Ekström-Tysk](#))

Together, these yield time monotonicity of the value function for a broad class of stopping problems.

- Conditioning on $\Theta = u$, observations X_i 's are i.i.d with density $p_u(x) \propto e^{ux-B(u)}$.
- For $Y_n = X_1 + \cdots + X_n$ and an arbitrary prior μ , define the posterior

$$\mu_{n,Y_n}(du) \propto e^{uY_n-nB(u)}\mu(du).$$

- We interpolate this posterior distribution continuously:

$$\mu_{t,y}(du) \propto e^{uy-tB(u)}\mu(du), \quad t \geq 0.$$

- For any non-decreasing T , define

$$\Pi_t^T = \mathbb{E}[T(\Theta) \mid \mathcal{F}_t], \quad \Lambda_t(y) = \int T d\mu_{t,y}(du) =: z.$$

- Since $\partial_y \Lambda_t(y) = \text{Cov}_{\mu_{t,y}}(T(\Theta), \Theta) > 0$, a state z determines a unique posterior $\mu_{t,z} := \mu_{t,y(t,z)}$.

Some possible choices of T :

- **Estimation:** $T(u) = u$. Then Π_n^T is the posterior mean.
- **Sequential testing:** $T(u) = \mathbf{1}_{(\theta_0, \infty)}(u)$. Then $\Pi_n^T = \mathbb{P}(\Theta > \theta_0 \mid \mathcal{F}_n)$.
- **Forecasting:** $T(u) = B'(u)$. Then $\Pi_n^T = \mathbb{E}[X_{n+1} \mid \mathcal{F}_n]$.

- 1 Motivation and problem setup
- 2 Main results
- 3 Examples and consequences

Lemma

Let T and $q \geq 0$ be non-decreasing, let G be concave, and suppose $\text{Cov}_\mu(T, G) = 0$. Then

$$\text{Cov}_{q\mu / \int q d\mu}(T, G) \leq 0.$$

The inequality becomes reversed when G is convex.

- Concavity means that the slope of G decreases as the parameter increases.
- The condition $\text{Cov}_\mu(T, G) = 0$ says that the contributions from smaller and larger parameter values initially balance.
- An increasing q puts relatively more weight on the larger u where G' is smaller.

What $q(u)$ do we choose

Fix a posterior $\mu_{t,z}$ and a threshold $a \in \mathbb{R}$. After the next k observations, define

$$q_a(u) = \mathbb{P}_u(\Pi_{t+k}^T > a).$$

- We evaluate at $\mu_{t,z}$,

$$\mathbb{E}_{t,z} [(\Pi_{t+k}^T - a)^+] = \int (T(u) - a) q_a(u) \mu_{t,z}(du).$$

- We use $(z - a)^+$ because
 - For distributions α and β with the same expected value,

$$\alpha \leq_{\text{cx}} \beta \iff \int (x - a)^+ \alpha(dx) \leq \int (x - a)^+ \beta(dx) \quad \text{for every } a \in \mathbb{R}.$$

- Together with affine functions they generate every convex function.

Main result 1: convex order comparison

Hold the state fixed: $\Lambda_t(y_t) = z$. Differentiating along the level curve,

$$\frac{d}{dt} \mu_{t,z}(du) = G_t(u) \mu_{t,z}(du), \quad G_t(u) = u \frac{\partial}{\partial t} y(t, z) - B(u) - \text{constant}.$$

- G_t is concave
- $\mu_{t,z}$ is a probability measure: $\frac{d}{dt} \int d\mu_{t,z} = \int G_t d\mu_{t,z} = 0$.
- We stay on the z -level: $\frac{d}{dt} \int T d\mu_{t,z} = \int T G_t d\mu_{t,z} = 0$. Hence $\text{Cov}_{\mu_{t,z}}(T, G_t) = 0$.

With q_a frozen at the later time, the Corollary gives

$$\frac{d}{dt} \int (T(u) - a) q_a(u) \mu_{t,z}(du) \leq 0 \implies \text{every } \mathbb{E}[(\Pi_s - a)^+ | \Pi_t = z] \text{ decreases in } t, \quad t < s.$$

Theorem (1)

For $m < n$, $k \geq 1$, and a state z attainable throughout $[m, n]$,

$$\Pi_{m+k}^T \mid \{\Pi_m^T = z\} \geq_{\text{cx}} \Pi_{n+k}^T \mid \{\Pi_n^T = z\}.$$

Corollary: $t \mapsto \text{Cov}_{\mu_{t,z}}(T(\Theta), \Theta)$ is non-increasing.

Main result 2: preservation of convexity

- Fix t and differentiate the posterior twice with respect to z :

$$\frac{\partial^2}{\partial z^2} \mu_{t,z}(\mathrm{d}u) = Q_z(u) \mu_{t,z}(\mathrm{d}u),$$

where Q_z is quadratic and convex.

- The same two constraints differentiated twice gives $\int Q_z \mathrm{d}\mu_{t,z} = \int TQ_z \mathrm{d}\mu_{t,z} = 0$.

Theorem (2)

Fix t and $k \geq 1$. If f is convex, then $z \mapsto \mathbb{E}_{t,z}[f(\Pi_{t+k}^T)]$ is convex.

- For a concave payoff g , consider

$$V(n, z) = \inf_{\tau} \mathbb{E}_{n,z} \left[g(\Pi_{n+\tau}^T) + c\tau \right],$$

$$V_{j+1}(n, z) = \min \left\{ g(z), c + \mathbb{E}_{n,z} \left[V_j(n+1, \Pi_{n+1}^T) \right] \right\}.$$

- Theorem 2 keeps every iterate concave, so Theorem 1 can be applied to it.
- $V(n, \cdot)$ is concave and $V(m, z) \leq V(n, z)$ for $m \leq n$.

- 1 Motivation and problem setup
- 2 Main results
- 3 Examples and consequences

Example 1: Bayesian sequential testing

Fix θ_0 and let $T(u) = \mathbf{1}_{(\theta_0, \infty)}(u)$, so $\Pi_n = P(\Theta > \theta_0 \mid \mathcal{F}_n)$. With unit loss for a wrong decision and cost c per observation,

$$V(n, \pi) = \inf_{\tau} \mathbb{E}_{n, \pi} [\Pi_{n+\tau} \wedge (1 - \Pi_{n+\tau}) + c\tau].$$

- Ekström-Wang impose the one-step comparison as Assumption 5.1 and conjectured for time-monotonicity.
- The corollary $t \mapsto \text{Cov}_{\mu_{t,z}}(T(\Theta), \Theta)$ non-increasing recovers their posterior concentration result, that the π -level curves spread out in t .

Example 2: investment under learning

Let $T(u) = u$ and $\hat{\Theta}_n = \mathbb{E}[\Theta \mid \mathcal{F}_n]$. The decision maker learns about an unknown profitability and chooses both when to stop and whether to invest, discounting at rate $r > 0$:

$$V(n, z) = \sup_{\tau \geq n} \mathbb{E} \left[e^{-r(\tau-n)} (\hat{\Theta}_\tau - K)^+ \mid \hat{\Theta}_n = z \right].$$

With $g(x) = (x - K)^+$.

- g is convex, so Theorem 2 and induction keep $x \mapsto V_j(n, x)$ convex;
- Theorem 1 then gives $V(n, z) \leq V(m, z)$ for $m < n$.

At the same posterior mean, the value of learning is larger to the one that has learned less.

Example (When T is not monotone)

For $T(p) = \mathbf{1}_{\{p=1/2\}}$ with prior $(\frac{1}{4}, \frac{1}{2}, \frac{1}{4})$ on $p \in \{\frac{1}{4}, \frac{1}{2}, \frac{3}{4}\}$,

$$\Pi_1^T \mid \{\Pi_0^T = \tfrac{1}{2}\} \sim \delta_{1/2}, \quad \Pi_2^T \mid \{\Pi_1^T = \tfrac{1}{2}\} \sim \tfrac{9}{16}\delta_{4/9} + \tfrac{7}{16}\delta_{4/7}.$$

T being monotone is essential.

Comment: the continuous-time analogue

In discrete time, after n observations,

$$\mu_{n,Y_n}(\mathrm{d}u) \propto e^{uY_n - nB(u)} \mu(\mathrm{d}u).$$

- For general exponential families, replacing n by a real number t initially gives only an analytic interpolation:

$$\mu_{t,y}(\mathrm{d}u) \propto e^{uy - tB(u)} \mu(\mathrm{d}u).$$

- If one member of the exponential family is infinitely divisible, the family admits a continuous-time realization by a Lévy process L .
- Thus t becomes genuine information time, and both theorems hold for all real times.

Examples: unknown Brownian drift, unknown Poisson intensity, gamma processes.

Not covered: the construction does not cover unknown volatility, arbitrary changes of the Lévy triplet, or multidimensional parameters.

Comment: from observation number to continuous time

Consider Brownian observations with unknown drift:

$$dY_t = \Theta dt + dW_t, \quad \Pi_t^T = \mathbb{E}[T(\Theta) \mid \mathcal{F}_t^Y].$$

Itô's formula gives

$$d\Pi_t^T = \partial_y \Lambda_t(Y_t) d\widehat{W}_t = \text{Cov}(T(\Theta), \Theta \mid \mathcal{F}_t^Y) d\widehat{W}_t.$$

- The conditional covariance is therefore the instantaneous volatility of the posterior statistic.
- For $T(u) = \mathbf{1}_{\{u > \theta_0\}}$, Ekström-Vaicenavicius show that this volatility decreases with t along each fixed posterior probability level.
- For $T(u) = u$, Ekström-Karatzas-Vaicenavicius establish the corresponding result for the posterior mean.

Example (3)

For $Y_t = \Theta t + W_t$, let

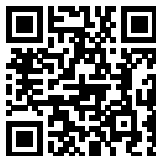
$$\Pi_t = \mathbb{P} \left(\sup_{r \leq h} (Y_{t+r} - Y_t) \geq a \mid \mathcal{F}_t^Y \right) = \mathbb{E}[T(\Theta) \mid \mathcal{F}_t^Y],$$

where

$$T(u) = \Phi \left(u\sqrt{h} - \frac{a}{\sqrt{h}} \right) + e^{2au} \Phi \left(-u\sqrt{h} - \frac{a}{\sqrt{h}} \right).$$

The person with more accumulated information has the less dispersed future revisions.

- Same estimates today but different history can give different value of information.
- Convexity survives Bayesian updating.
- These yield structural results for the value function of a large family of stopping problems. Convexity and time-monotonicity.
- We do not require conjugate priors. They hold for arbitrary priors and many standard observation models, both discrete time and continuous time.



Preprint: [arXiv:2609.05065](https://arxiv.org/abs/2609.05065)

Thank you!