

Stochastic Calculus:
IV. Hidden Markov Models and Optimal Dynamic
Contests

Lones Smith

March 24, 2026

Learning about a Hidden Markov Model

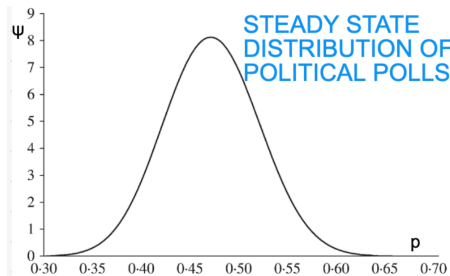
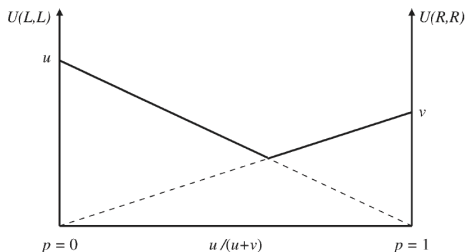
- ▶ Keppo, Davydov, and Smith (*REStud*, 2008): “Optimal Electoral Timing: Exercise Wisely and You May Live Longer”
- ▶ An unobservable political state $\theta = L, R$ at time t
- ▶ Party ℓ, r is best in the unobserved political state $\theta = L, R$
- ▶ State switches from θ in $[t, t + \Delta t]$ with chance $\lambda_\theta \Delta t > 0$
- ▶ Observation process (X_t) obeys $dX_t = \theta_t dt + \gamma dW_t$
- ▶ **Lemma:** The *political slant* $p_t = P[\theta_t = R]$ obeys

$$dp_t = a(b - p_t)dt + \boxed{\sigma p_t(1 - p_t)d\bar{W}(t)} \quad (\heartsuit)$$

where $a = \lambda_L + \lambda_H$, $b = \lambda_L/a$, $\sigma = (R - L)/\gamma$.

- ▶ Volatility vanishes near $p = 0, 1$, and so $0 < p_t < 1$ for all t
- ▶ Since the state is constantly shifting, a steady-state exists
- ▶ Expected time for θ_t to switch from L to R and back to L equals $(1/\lambda_L) + (1/\lambda_R) = 1/(ab)(1 - b) \approx 2.5$ years

A Model of Political Polls



- ▶ Political preferences captures how you weight benefits of matching the party ℓ, r to the state L, R
- ▶ If u, v are independently & identically distributed across voters, with exponential density, the political slant $P(t)$ is the share of voters for party R in any election at time t .
- ▶ If $dP(t) = \mu(P)dP + \sigma(P)dW$ has a stationary density ψ , then $[\sigma(p)\psi(p)]'' - 2[\mu(p)\psi(p)]' = 0$ (Fokker-Planck equation)

$$\psi(p) \propto \frac{\exp\left(-\frac{2a}{\sigma^2} \left(\frac{1-b}{1-p} + \frac{b}{p}\right)\right) \left(\frac{p}{1-p}\right)^{2a(2b-1)/\sigma^2}}{p^2(1-p)^2}$$

Predictive Enforcement" (Yeon-Koo Che Theory Seminar)

- ▶ Let's modify Che's theory seminar paper to escape bang bang control and allow richer policies and states
- ▶ Diminishing returns to policing: y_t cop cars $\Rightarrow$ crime $y_t - \gamma y_t^2$
- ▶ Che assumes a state $\theta = L, H$ switching back and forth as in my election paper. Here, p is the current chance of state H .
 - ▶ Che: any crime reveals state $p = 1$. Then $dp_t = a(b - p_t)dt$
 - ▶ Let's assume slowly revealed Gaussian info as in ($\heartsuit$)
 - ▶ Or, let $p \in \mathbb{R}$ be the state, obeying a mean-reverting Ornstein-Uhlenbeck process: $dp = a(b - p_t)dt + \boxed{\sigma dW_t}$
- ▶ (y_t) minimizes $V(p_0) = \int_0^\infty (p_t[y_t - \gamma y_t^2] + c(1 - y_t)) e^{-rt} dt$

$$\text{FOC: } 0 = p_t[1 - 2\gamma y_t] - c \Rightarrow y_t^* = \frac{1 - c/p_t}{2\gamma}$$

- ▶ No more bang-bang: *policing optimally smoothly rises in p !*
- ▶ V solves ODE (for volatility $\xi(p) = \sigma$ or $\xi(p) = \sigma p(1 - p)$)

$$rV(p) = (p_t y_t^*[1 - \gamma y_t^*] + c(1 - y_t^*)) + a(b - p)V'(p) + \frac{1}{2}\xi^2(p)V''(p)$$

$$\Rightarrow rV(p) = \left(p_t \frac{1 - c^2/p_t^2}{4\gamma} + c(1 - \frac{1 - c/p_t}{2\gamma}) \right) + a(b - p)V'(p) + \frac{1}{2}\xi^2(p)V''(p)$$

“Optimal Dynamic Contests”, Moscarini and Smith (2008)

► The Tug-of-War

- Two players spend costly effort to produce output.
- Noise renders the outcome stochastic at each moment in time.
- The first player to outperform the opponent by a given margin wins a prize (tennis game)
 - Optimal finish line? (points per game, games per set, sets?)
 - Optimal prize?
 - Optimal scoring, handicaps etc.?
- Unlike war of attrition, we allow for “back-and-forth”.

► Literature

- Harris and Vickers (1987): discrete time tug-of-war up: tie games are most intense; leaders try harder than followers; numerically solved
- Budd, Harris and Vickers (1993) return to the problem in continuous time. But they cannot solve it, and proceed to use asymptotic expansions and more numerical simulations.

The Model

- ▶ Players $i = 1, 2$ compete for a prize $\pi > 0$.
- ▶ The state of competition is a point z on the interval $[-L, L]$
 - ▶ Example: literally tug of war; tennis at "love"
 - ▶ For simplicity, call $L_1 = -L$ and $L_2 = L$.
 - ▶ The game stops when the state z first hits L or $-L$.
 - ▶ Player 2 wins if z_t hits L , and player 1 wins if it hits $-L$.
 - ▶ The winner gets $\pi > 0$ and the loser gets nothing.
 - ▶ Payoffs are undiscounted (zero interest rate $r = 0$).
- ▶ Chance and effort move the state: Given flow effort levels $n_2, n_1 \geq 0$,

$$dz = (n_2 - n_1)\mu dt + \sigma dW$$

- ▶ Omitted is any role for ability (creates an asymmetric game)
- ▶ Effort n incurs a flow cost $c(n) = c_0 n^2 / 2$
- ▶ We crucially ignore changing variance, which does happen:
 - ▶ late in a hockey game, pulling the goalie
 - ▶ playing a basketball defense that shuts down scoring

Symmetric Markov Perfect Nash Equilibrium

- ▶ The game ends at the **stopping time**

$$T = \inf \{t \geq 0 : z_t = L_1 \text{ or } z_t = L_2\}.$$

- ▶ A **Markov Perfect** Equilibria (MPE) is a SPE where strategies can only depend on the state
- ▶ MPE: pair of admissible Markov strategies $\{n_i(z)\}_{i=1,2}$ such that for all $z_0 \in (-L, L)$, for player $i = 1, 2$

$$n_i(z_0) = \arg \max_u \left\{ \pi \Pr(T_{L_i} < T_{L_{3-i}}) - E \left[\int_0^{T_L \vee T_{-L}} c_0 \frac{u^2}{2} dt \middle| z_0 \right] \right\}$$

- ▶ We look for **symmetric Markov Perfect Equilibrium** (SMPE)

$$n_2(z) = n_1(-z) \equiv n(z)$$

Equilibrium Analysis

- ▶ We characterize an SMPE (guess and verify logic).
- ▶ Assume value $\bar{V}_i(z)$. By symmetry

$$\bar{V}_2(z) = \bar{V}_1(-z) \equiv \bar{V}(z).$$

⇒ $\bar{V}(-L) = 0$ and $\bar{V}(L) = \pi$, and $\bar{V}$ obeys the HJB equation:

$$0 = \sup_u -c_0 \frac{1}{2} u^2 + [u - n(-z)] \mu \bar{V}'(z) + \frac{1}{2} \sigma^2 \bar{V}''(z)$$

- ▶ Greater effort u lowers the dividend and raises the drift
- ▶ The **equilibrium FOC** for the two players are

$$n_2(z) = n(z) = \frac{\mu}{c_0} \bar{V}'(z)$$

$$n_1(z) = n(-z) = -\frac{\mu}{c_0} \bar{V}'_1(z) = \frac{\mu}{c_0} \bar{V}'(-z).$$

- ▶ Plugging equilibrium FOCs back into the HJB equation yields a 2-boundary *Delayed Differential Equation* problem:

$$\bar{V}''(z) = \delta \bar{V}'(z) [2\bar{V}'(-z) - \bar{V}'(z)]$$

where $\delta \equiv \mu^2 / (c_0 \sigma^2)$

Simplifying the Equilibrium Equation

- ▶ Let the value of the race with finish line L is $\bar{V}_L(z) = \bar{V}_1(\frac{z}{L})$.

$$\Rightarrow n_L(z) = \frac{\mu}{c_0} \bar{V}_L'(z) = \frac{\mu}{c_0 L} \bar{V}_1'(\frac{z}{L}) = \frac{1}{L} n_1\left(\frac{z}{L}\right)$$

- ▶ We focus on game with $L = 1$. Then $V(y) \equiv \delta \bar{V}_1(y)$ obeys

$$V''(y) = V'(y)[2V'(-y) - V'(y)]$$

with boundary conditions $V(-1) = 0$ and $V(1) = \delta\pi$.

- ▶ Focus on game with $L = 1$ and recover values for $L \neq 1$ by scaling and parameters c_0, μ, σ collapsed into δ .
- ▶ Let

$$m(z) \equiv V'(z)$$

- ▶ Obtain a *first-order* Delayed Differential Equation (DDE)

$$m'(z) = m(z)[2m(-z) - m(z)]$$

subject to $\int_{-1}^1 m(z) dz = \delta\pi$

Equilibrium Predictions

- DDE equivalent to the 2nd order *Ordinary* DE

$$2m(z)m''(z) = 3(m'(z))^2 - 4m'(z)(m(z))^2 - 3(m(z))^4$$

1. Effort is positive.

Intuition: the marginal cost of zero effort is zero, since winning chance is zero.

2. Effort is first rising, and then maybe falling.

Logic: $m'(s) > 0$ somewhere, and $m'(y) = 0 \Rightarrow m''(y) < 0$.

3. The leader works harder than the follower.

Logic: $m'(0) = m(0)^2$ and so $m(s) > m(-s)$ for $s > 0$, because

$$m'(z) = m(z)[2m(-z) - m(z)] > 0$$

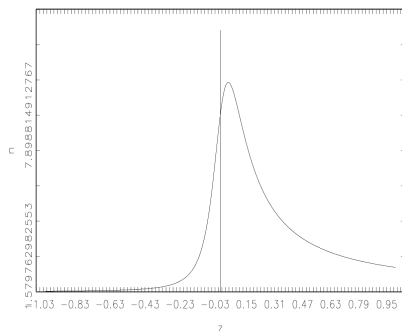
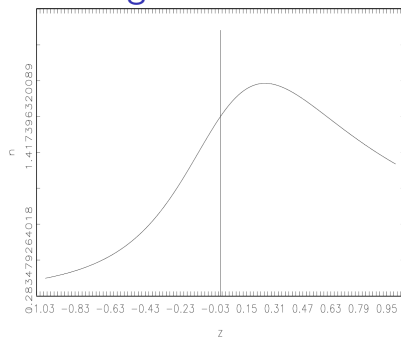
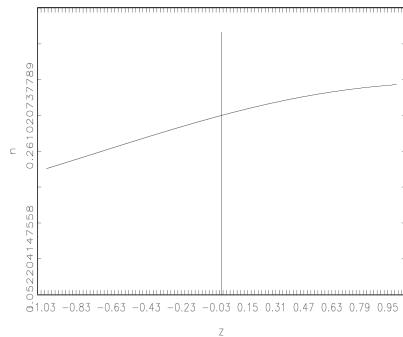
in a neighborhood of 0

4. Leader works at most twice as hard as the follower.

If ever $m(z) > 2m(-z)$ for some $z > 0$, then

$$m'(z) - m'(-z) = m(z)[2m(-z) - m(z)] - m(-z)[2m(z) - m(-z)] < 0$$

Equilibrium Predictions with Increasing Prizes



Optimal Effort Prize is Finite

