

Control Strategies for the Fokker-Planck Equation

Tobias Breiten, **Karl Kunisch** and Laurent Pfeiffer

Cortona

June, 2016



Motivation

- ▶ dragged Brownian particle following Langevin equation

$$dX_s = Y(s)ds, \quad dY_s = -\beta Y_s ds + F(X, s)ds + \sqrt{2\beta kT/m} dB(s)$$

particle confined by potential V : force $F(X, s) = -\nabla V(X, s)$

$$\beta \gg, \quad t = \frac{s}{\beta}, \quad \nu = \frac{kT}{m},$$

- ▶ Smoluchowski equation

$$dX_t = -\nabla V(X, t) + \sqrt{2\nu} dB_t$$

- ▶ atomic force microscopy, single molecule pulling, optical trapping
- ▶ collimate light from laser into aperture of microscope objective
- ▶ control by optical tweezer $V(X_t, t) = W(X_t) + u(X_t, t)$.

The Fokker-Planck equation

Consider **probability distribution function**

$$\rho(x, t)dx = \mathbb{P}[X_t \in [x, x + dx)]$$

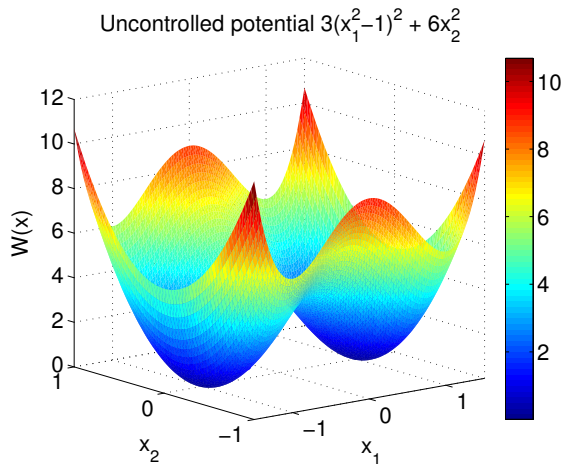
Fokker-Planck equation

$$\begin{aligned}\frac{\partial \rho}{\partial t} &= \nu \Delta \rho + \nabla \cdot (\rho \nabla V) && \text{in } \Omega \times (0, \infty), \\ 0 &= (\nu \nabla \rho + \rho \nabla V) \cdot \vec{n} && \text{on } \Gamma \times (0, \infty), \\ \rho(x, 0) &= \rho_0(x) && \text{in } \Omega,\end{aligned}$$

- ▶ $\Omega \subset \mathbb{R}^n$ bounded open set with boundary $\Gamma = \partial\Omega$,
- ▶ ρ_0 initial probability distribution with $\int_{\Omega} \rho_0(x)dx = 1$.

[HARTMANN ET AL.'13]

The Fokker-Planck equation



The Fokker-Planck equation

The Fokker-Planck equation

Control of the Fokker-Planck equation

$$\frac{\partial \rho}{\partial t} = \nu \Delta \rho + \nabla \cdot (\rho \nabla V) \quad \text{in } \Omega \times (0, \infty),$$

$$0 = (\nu \nabla \rho + \rho \nabla V) \cdot \vec{n} \quad \text{on } \Gamma \times (0, \infty),$$

$$\rho(x, 0) = \rho_0(x) \quad \text{in } \Omega,$$

- ▶ system converges to **stationary distribution** $\rho_\infty(x)$,
(Boltzmann distribution)
- ▶ particles have to cross **energy barrier** between potential wells,
 $\leadsto$ may be inadequately slow, $\sim \exp(\Delta W / \nu)$
- ▶ can we **control** the potential $V(x, t) = W(x) + \alpha(x)u(t)$,
- ▶ if we can, what are “good” **choices for** α ?

Assumptions:

- ▶ $W \in W^{2, \max(2, n+\varepsilon)}(\Omega), \varepsilon > 0$
- ▶ $\alpha \in W^{2, \max(2, n+\varepsilon)}(\Omega), \varepsilon > 0$ with $\nabla \alpha \cdot \vec{n} = 0$ on Γ

Control of the Fokker-Planck equation

$$\frac{\partial \rho}{\partial t} = \nu \Delta \rho + \nabla \cdot (\rho \nabla V) \quad \text{in } \Omega \times (0, \infty),$$

$$0 = (\nu \nabla \rho + \rho \nabla V) \cdot \vec{n} \quad \text{on } \Gamma \times (0, \infty),$$

$$\rho(x, 0) = \rho_0(x) \quad \text{in } \Omega,$$

- ▶ system converges to **stationary distribution** $\rho_\infty(x)$,
(Boltzmann distribution)
- ▶ particles have to cross **energy barrier** between potential wells,
 $\leadsto$ may be inadequately slow, $\sim \exp(\Delta W / \nu)$
- ▶ can we **control** the potential $V(x, t) = W(x) + \alpha(x)u(t)$,
- ▶ if we can, what are “good” **choices for** α ?

Assumptions:

- ▶ $W \in W^{2, \max(2, n+\varepsilon)}(\Omega), \varepsilon > 0$
- ▶ $\alpha \in W^{2, \max(2, n+\varepsilon)}(\Omega), \varepsilon > 0$ with $\nabla \alpha \cdot \vec{n} = 0$ on Γ

Outline

Motivation

The Fokker Planck equation

A Riccati-based local feedback law

A Lyapunov-based global feedback law

A numerical study

Solutions to the Fokker-Planck equation

For $T > 0$ we call ρ a **(variational) solution** on $(0, T)$ if

$$\rho \in W(0, T) = L^2(0, T; H^1(\Omega)) \cap H^1(0, T; (H^1(\Omega))^*)$$

and for a.e. $t \in (0, T)$

$$\langle \rho_t(t), v \rangle + \langle \nu \nabla \rho(t) + \rho(t) \nabla W, \nabla v \rangle + u(t) \langle \rho(t) \nabla \alpha, \nabla v \rangle = 0 \quad \forall v \in H^1(\Omega)$$

$$\rho(0) = \rho_0.$$

Proposition

(i) $u \in L^2(0, T), \rho_0 \in L^2(\Omega) \Rightarrow \exists!$ solution $\rho \in W(0, T)$

(ii) If moreover $\Delta \alpha \in L^\infty(\Omega)$

$$\Rightarrow \rho_t \in L^2(0, T; L^2(\Omega)), \rho \in C([0, T]; H^1(\Omega))$$

Proposition

Let $u \in L^2(0, T)$ and $\rho_0 \in L^2(\Omega)$.

(i) For every $t \in [0, T]$ we have $\int_{\Omega} \rho(t) \, dx = \int_{\Omega} \rho_0 \, dx$.

(ii) If $\rho_0 \geq 0$ a.e. on Ω , then $\rho(x, t) \geq 0 \, \forall t > 0$ and a.e. on Ω .

Solutions to the Fokker-Planck equation

For $T > 0$ we call ρ a **(variational) solution** on $(0, T)$ if

$$\rho \in W(0, T) = L^2(0, T; H^1(\Omega)) \cap H^1(0, T; (H^1(\Omega))^*)$$

and for a.e. $t \in (0, T)$

$$\langle \rho_t(t), v \rangle + \langle \nu \nabla \rho(t) + \rho(t) \nabla W, \nabla v \rangle + u(t) \langle \rho(t) \nabla \alpha, \nabla v \rangle = 0 \quad \forall v \in H^1(\Omega)$$

$$\rho(0) = \rho_0.$$

Proposition

(i) $u \in L^2(0, T), \rho_0 \in L^2(\Omega) \Rightarrow \exists!$ solution $\rho \in W(0, T)$

(ii) If moreover $\Delta \alpha \in L^\infty(\Omega)$

$$\Rightarrow \rho_t \in L^2(0, T; L^2(\Omega)), \rho \in C([0, T]; H^1(\Omega))$$

Proposition

Let $u \in L^2(0, T)$ and $\rho_0 \in L^2(\Omega)$.

(i) For every $t \in [0, T]$ we have $\int_{\Omega} \rho(t) \, dx = \int_{\Omega} \rho_0 \, dx$.

(ii) If $\rho_0 \geq 0$ a.e. on Ω , then $\rho(x, t) \geq 0 \, \forall t > 0$ and a.e. on Ω .

A bilinear control system

Consider the **bilinear control system**

$$\begin{aligned}\dot{\rho}(t) &= \mathcal{A}\rho(t) + u(t)\mathcal{N}\rho(t), \\ \rho(0) &= \rho_0,\end{aligned}$$

where the operators $\mathcal{A}$ and $\mathcal{N}$ are defined as follows

$$\mathcal{A}: \mathcal{D}(\mathcal{A}) \subset L^2(\Omega) \rightarrow L^2(\Omega),$$

$$\mathcal{D}(\mathcal{A}) = \{ \rho \in H^2(\Omega) \mid (\nu \nabla \rho + \rho \nabla W) \cdot \vec{n} = 0 \text{ on } \Gamma \},$$

$$\mathcal{A}\rho = \nu \Delta \rho + \nabla \cdot (\rho \nabla W),$$

$$\mathcal{N}: H^1(\Omega) \rightarrow L^2(\Omega), \quad \mathcal{N}\rho = \nabla \cdot (\rho \nabla \alpha).$$

A bilinear control system

Consider the **bilinear control system**

$$\begin{aligned}\dot{\rho}(t) &= \mathcal{A}\rho(t) + u(t)\mathcal{N}\rho(t), \\ \rho(0) &= \rho_0,\end{aligned}$$

its $L^2(\Omega)$ -adjoints are given by

$$\begin{aligned}\mathcal{A}^* &: \mathcal{D}(\mathcal{A}^*) \subset L^2(\Omega) \rightarrow L^2(\Omega), \\ \mathcal{D}(\mathcal{A}^*) &= \{ \varphi \in H^2(\Omega) \mid (\nu \nabla \varphi) \cdot \vec{n} = 0 \text{ on } \Gamma \}, \\ \mathcal{A}^* \phi &= \nu \Delta \varphi - \nabla W \cdot \nabla \varphi, \\ \mathcal{N}^* &: H^1(\Omega) \rightarrow L^2(\Omega), \quad \mathcal{N}^* \varphi = -\nabla \varphi \cdot \nabla \alpha.\end{aligned}$$

The symmetrized operator

cp. [RISKEN, '96]

We introduce $\Phi(x) = \log(\nu) + \frac{W(x)}{\nu}$ and the operator

$$\begin{aligned}\mathcal{A}_s &: \mathcal{D}(\mathcal{A}_s) \subset L^2(\Omega) \rightarrow L^2(\Omega), \\ \mathcal{D}(\mathcal{A}_s) &= \left\{ \varrho \in H^2(\Omega) \left| (\nu \nabla \varrho + \frac{1}{2} \varrho \nabla W) \cdot \vec{n} = 0 \text{ on } \Gamma \right. \right\}, \\ \mathcal{A}_s &= e^{\frac{\Phi}{2}} \mathcal{A} e^{-\frac{\Phi}{2}}.\end{aligned}$$

Lemma

- (i) $\mathcal{A}_s = \mathcal{A}_s^*$, $\sigma(\mathcal{A}_s) = \sigma_p(\mathcal{A}_s) = \sigma(\mathcal{A}) \subset \overline{\mathbb{R}}_-$, eigenfunctions $\{\varrho_i\}_{i=0}^\infty$ of $\mathcal{A}_s$ form a complete orthogonal set
- (ii) ψ_i eigenfunction of $\mathcal{A} \Leftrightarrow e^{\frac{\Phi}{2}} \psi_i$ eigenfunction of $\mathcal{A}_s$
 ψ_i eigenfunction of $\mathcal{A} \Leftrightarrow e^{\Phi} \psi_i$ is an eigenfunction of $\mathcal{A}^*$
- (iii) $\rho_\infty = e^{-\Phi}$ is an eigenfunction of $\mathcal{A}$ with $\mathcal{A} \rho_\infty = 0$

A shifted problem

Instead of ρ , consider the shifted state $y := \rho - \rho_\infty$

$$\begin{aligned}\dot{y}(t) &= \mathcal{A}y(t) + u(t)\mathcal{N}y(t) + \mathcal{B}u(t), \\ y(0) &= \rho_0 - \rho_\infty,\end{aligned}$$

with $\mathcal{B} = \mathcal{N}\rho_\infty$.

The control operator $\mathcal{B}$ and its adjoint are defined as

$$\begin{aligned}\mathcal{B}: \mathbb{R} &\rightarrow L^2(\Omega), \quad \mathcal{B}c = c\mathcal{N}\rho_\infty = c\nabla \cdot (\rho_\infty \nabla \alpha), \\ \mathcal{B}^*: L^2(\Omega) &\rightarrow \mathbb{R}, \quad \mathcal{B}^*v = \langle \mathcal{N}\rho_\infty, v \rangle.\end{aligned}$$

Some useful projections

Projection $\mathcal{P}$ onto $\mathbb{1}^\perp$ along ρ_∞

$$\mathcal{P}: L^2(\Omega) \rightarrow L^2(\Omega), \quad \mathcal{P}y = y - \int_{\Omega} y \, dx \, \rho_\infty,$$

$$\mathcal{Y}_{\mathcal{P}} = \text{im}(\mathcal{P}) = \left\{ v \in L^2(\Omega): \int_{\Omega} v \, dx = 0 \right\}, \quad \ker(\mathcal{P}) = \text{span} \{ \rho_\infty \}.$$

Complementary projection $\mathcal{Q}$

$$\mathcal{Q}: L^2(\Omega) \rightarrow L^2(\Omega), \quad \mathcal{Q}y = (I - \mathcal{P})y = \int_{\Omega} y \, dx \, \rho_\infty,$$

$$\text{im}(\mathcal{Q}) = \ker(\mathcal{P}), \quad \ker(\mathcal{Q}) = \text{im}(\mathcal{P}).$$

Orthogonal projection on $\mathcal{Y}_{\mathcal{P}}$:

$$\Theta: L^2(\Omega) \rightarrow L^2(\Omega), \quad \Theta y = y - \frac{1}{|\Omega|} \int_{\Omega} y \, dx \, \mathbb{1},$$

$$\text{im}(\Theta) = \text{im}(\mathcal{P}) = \mathcal{Y}_{\mathcal{P}}, \quad \ker(\Theta) = \{ \mathbb{1} \}.$$

Some useful projections

Projection $\mathcal{P}$ onto $\mathbb{1}^\perp$ along ρ_∞

$$\mathcal{P}: L^2(\Omega) \rightarrow L^2(\Omega), \quad \mathcal{P}y = y - \int_{\Omega} y \, dx \, \rho_\infty,$$

$$\mathcal{Y}_{\mathcal{P}} = \text{im}(\mathcal{P}) = \left\{ v \in L^2(\Omega): \int_{\Omega} v \, dx = 0 \right\}, \quad \ker(\mathcal{P}) = \text{span} \{ \rho_\infty \}.$$

Complementary projection $\mathcal{Q}$

$$\mathcal{Q}: L^2(\Omega) \rightarrow L^2(\Omega), \quad \mathcal{Q}y = (I - \mathcal{P})y = \int_{\Omega} y \, dx \, \rho_\infty,$$

$$\text{im}(\mathcal{Q}) = \ker(\mathcal{P}), \quad \ker(\mathcal{Q}) = \text{im}(\mathcal{P}).$$

Orthogonal projection on $\mathcal{Y}_{\mathcal{P}}$:

$$\Theta: L^2(\Omega) \rightarrow L^2(\Omega), \quad \Theta y = y - \frac{1}{|\Omega|} \int_{\Omega} y \, dx \, \mathbb{1},$$

$$\text{im}(\Theta) = \text{im}(\mathcal{P}) = \mathcal{Y}_{\mathcal{P}}, \quad \ker(\Theta) = \{ \mathbb{1} \}.$$

Some useful projections cont'd

The $L^2(\Omega)$ adjoint of $\mathcal{P}$ is the **projection $\mathcal{P}^*$ onto $\rho_\infty^\perp$ along $\mathbb{1}$**

$$\mathcal{P}^*: L^2(\Omega) \rightarrow L^2(\Omega), \quad \mathcal{P}^*y = y - \int_{\Omega} \rho_\infty y \, dx \, \mathbb{1},$$

$$\text{im}(\mathcal{P}^*) = \left\{ v \in L^2(\Omega): \int_{\Omega} \rho_\infty v \, dx = 0 \right\}, \quad \ker(\mathcal{P}^*) = \{\mathbb{1}\}.$$

The **complementary projection $\mathcal{Q}^*$** is given by

$$\mathcal{Q}^*: L^2(\Omega) \rightarrow L^2(\Omega), \quad \mathcal{Q}^*y = \int_{\Omega} \rho_\infty y \, dx \, \mathbb{1},$$

$$\text{im}(\mathcal{Q}^*) = \ker(\mathcal{P}^*), \quad \ker(\mathcal{Q}^*) = \text{im}(\mathcal{P}^*).$$

Decoupling the problem

By means of $\mathcal{P}$ and $\mathcal{Q}$ we can **decompose** our **state space**

$$\mathcal{Y} = L^2(\Omega) = \text{im}(\mathcal{P}) \oplus \text{im}(\mathcal{Q}) =: \mathcal{Y}_{\mathcal{P}} \oplus \mathcal{Y}_{\mathcal{Q}},$$

$$y = y_{\mathcal{P}} + y_{\mathcal{Q}} = \mathcal{P}y + \mathcal{Q}y, \quad y \in L^2(\Omega).$$

$\leadsto$ applying $\mathcal{P}$ and $\mathcal{Q}$ yields

$$\begin{pmatrix} \dot{y}_{\mathcal{P}} \\ \dot{y}_{\mathcal{Q}} \end{pmatrix} = \begin{pmatrix} \mathcal{P}\mathcal{A} & \mathcal{P}\mathcal{A} \\ \mathcal{Q}\mathcal{A} & \mathcal{Q}\mathcal{A} \end{pmatrix} \begin{pmatrix} y_{\mathcal{P}} \\ y_{\mathcal{Q}} \end{pmatrix} + u \begin{pmatrix} \mathcal{P}\mathcal{N} & \mathcal{P}\mathcal{N} \\ \mathcal{Q}\mathcal{N} & \mathcal{Q}\mathcal{N} \end{pmatrix} \begin{pmatrix} y_{\mathcal{P}} \\ y_{\mathcal{Q}} \end{pmatrix} + u \begin{pmatrix} \mathcal{P}\mathcal{B} \\ \mathcal{Q}\mathcal{B} \end{pmatrix}.$$

$\leadsto$ simplifies to

$$\begin{pmatrix} \dot{y}_{\mathcal{P}} \\ \dot{y}_{\mathcal{Q}} \end{pmatrix} = \begin{pmatrix} \mathcal{P}\mathcal{A} & 0 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} y_{\mathcal{P}} \\ y_{\mathcal{Q}} \end{pmatrix} + u \begin{pmatrix} \mathcal{P}\mathcal{N} & \mathcal{P}\mathcal{N} \\ 0 & 0 \end{pmatrix} \begin{pmatrix} y_{\mathcal{P}} \\ y_{\mathcal{Q}} \end{pmatrix} + u \begin{pmatrix} \mathcal{P}\mathcal{B} \\ 0 \end{pmatrix}.$$

$\leadsto$ simplifies to

$$\begin{aligned} \dot{y}_{\mathcal{P}} &= \hat{\mathcal{A}}y_{\mathcal{P}} + u\hat{\mathcal{N}}y_{\mathcal{P}} + \hat{\mathcal{B}}u, & y_{\mathcal{P}}(0) &= \mathcal{P}\rho_0, \\ y_{\mathcal{Q}}(t) &= \mathcal{Q}\rho_0 - \rho_{\infty} = 0, & t &\geq 0. \end{aligned}$$

Decoupling the problem

By means of $\mathcal{P}$ and $\mathcal{Q}$ we can **decompose** our **state space**

$$\mathcal{Y} = L^2(\Omega) = \text{im}(\mathcal{P}) \oplus \text{im}(\mathcal{Q}) =: \mathcal{Y}_{\mathcal{P}} \oplus \mathcal{Y}_{\mathcal{Q}},$$

$$y = y_{\mathcal{P}} + y_{\mathcal{Q}} = \mathcal{P}y + \mathcal{Q}y, \quad y \in L^2(\Omega).$$

$\leadsto$ applying $\mathcal{P}$ and $\mathcal{Q}$ yields

$$\begin{pmatrix} \dot{y}_{\mathcal{P}} \\ \dot{y}_{\mathcal{Q}} \end{pmatrix} = \begin{pmatrix} \mathcal{P}\mathcal{A} & \mathcal{P}\mathcal{A} \\ \mathcal{Q}\mathcal{A} & \mathcal{Q}\mathcal{A} \end{pmatrix} \begin{pmatrix} y_{\mathcal{P}} \\ y_{\mathcal{Q}} \end{pmatrix} + u \begin{pmatrix} \mathcal{P}\mathcal{N} & \mathcal{P}\mathcal{N} \\ \mathcal{Q}\mathcal{N} & \mathcal{Q}\mathcal{N} \end{pmatrix} \begin{pmatrix} y_{\mathcal{P}} \\ y_{\mathcal{Q}} \end{pmatrix} + u \begin{pmatrix} \mathcal{P}\mathcal{B} \\ \mathcal{Q}\mathcal{B} \end{pmatrix}.$$

$\leadsto$ simplifies to

$$\begin{pmatrix} \dot{y}_{\mathcal{P}} \\ \dot{y}_{\mathcal{Q}} \end{pmatrix} = \begin{pmatrix} \mathcal{P}\mathcal{A} & 0 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} y_{\mathcal{P}} \\ y_{\mathcal{Q}} \end{pmatrix} + u \begin{pmatrix} \mathcal{P}\mathcal{N} & \mathcal{P}\mathcal{N} \\ 0 & 0 \end{pmatrix} \begin{pmatrix} y_{\mathcal{P}} \\ y_{\mathcal{Q}} \end{pmatrix} + u \begin{pmatrix} \mathcal{P}\mathcal{B} \\ 0 \end{pmatrix}.$$

$\leadsto$ simplifies to

$$\begin{aligned} \dot{y}_{\mathcal{P}} &= \hat{\mathcal{A}}y_{\mathcal{P}} + u\hat{\mathcal{N}}y_{\mathcal{P}} + \hat{\mathcal{B}}u, \quad y_{\mathcal{P}}(0) = \mathcal{P}\rho_0, \\ y_{\mathcal{Q}}(t) &= \mathcal{Q}\rho_0 - \rho_{\infty} = 0, \quad t \geq 0. \end{aligned}$$

Outline

Motivation

The Fokker Planck equation

A Riccati-based local feedback law

A Lyapunov-based global feedback law

A numerical study

An LQR problem for the linearized system

For $\delta > 0$ let us focus on the **linearized system**

$$\dot{y}_{\mathcal{P}} = (\hat{\mathcal{A}} + \delta I)y_{\mathcal{P}}(t) + \hat{\mathcal{B}}u, \quad y_{\mathcal{P}}(0) = \mathcal{P}\rho_0,$$

together with the **quadratic cost functional**

$$J(y_{\mathcal{P}}, u) = \frac{1}{2} \int_0^\infty \langle y_{\mathcal{P}}(t), \mathcal{M}y_{\mathcal{P}}(t) \rangle_{L^2(\Omega)} dt + \frac{1}{2} \int_0^\infty |u(t)|^2 dt,$$

where $\mathcal{M} \in \mathcal{L}(\mathcal{Y}_{\mathcal{P}})$ is a self-adjoint nonnegative operator on $\mathcal{Y}_{\mathcal{P}}$.

Riccati-based feedback law: $u = -\hat{\mathcal{B}}^* \hat{\Pi} y_{\mathcal{P}}$

The choice of α

Recall that we have $\widehat{\mathcal{B}} = \mathcal{B} = \mathcal{N}\rho_\infty = \nabla \cdot (\rho_\infty \nabla \alpha)$
 $\leadsto$ what is a **good** choice for α ?

Let φ_i be **eigenfunctions of $\mathcal{A}^*$** associated to the eigenvalues

$$-\delta \leq \lambda_q \leq \dots \leq \lambda_2 < 0 = \lambda_1.$$

Determine α as solution to the **elliptic equation**

$$\begin{aligned} \nabla \cdot (\rho_\infty \nabla \alpha) &= \mathcal{P} \sum_{i=2}^q e^{-\Phi} \varphi_i && \text{in } \Omega, \\ (\rho_\infty \nabla \alpha) \cdot \vec{n} &= 0 && \text{on } \Gamma. \end{aligned}$$

$\Rightarrow \exists!$ solution $\alpha \in W^{2, \max(2, n+\varepsilon)}(\Omega)/\mathbb{R}$

The ∞ -dimensional Hautus test

Consider now the operator **Riccati equation**

$$(\hat{\mathcal{A}} + \delta I)^* \hat{\Pi} + \hat{\Pi}(\hat{\mathcal{A}} + \delta I) - \hat{\Pi} \hat{\mathcal{B}} \hat{\mathcal{B}}^* \hat{\Pi} + \mathcal{M} = 0.$$

The pair $(\hat{\mathcal{A}}, \hat{\mathcal{B}})$ is **δ -stabilizable** if

$$\ker(\lambda I - \hat{\mathcal{A}}^*) \cap \ker(\hat{\mathcal{B}}^*) = \{0\} \quad \text{for } \lambda \in \overline{\mathbb{C}}_{-\delta} \cap \sigma(\hat{\mathcal{A}}^*).$$

Assume that (λ_j, ϕ_j) , $j \in \{2, \dots, q\}$ is an eigenpair of $\hat{\mathcal{A}}^*$.

$$\hat{\mathcal{B}}^* \phi_j = \langle \hat{\mathcal{B}}, \phi_j \rangle = \left\langle \mathcal{P} \sum_{i=2}^q e^{-\Phi} \varphi_i, \phi_j \right\rangle = \sum_{i=2}^q \langle e^{-\Phi} \varphi_i, \varphi_j \rangle = \|e^{-\frac{\Phi}{2}} \varphi_j\|^2$$

e.g. [CURTAIN/ZWART'95]

The ∞ -dimensional Hautus test

Consider now the operator **Riccati equation**

$$(\hat{\mathcal{A}} + \delta I)^* \hat{\Pi} + \hat{\Pi}(\hat{\mathcal{A}} + \delta I) - \hat{\Pi} \hat{\mathcal{B}} \hat{\mathcal{B}}^* \hat{\Pi} + \mathcal{M} = 0.$$

The pair $(\hat{\mathcal{A}}, \hat{\mathcal{B}})$ is **δ -stabilizable** if

$$\ker(\lambda I - \hat{\mathcal{A}}^*) \cap \ker(\hat{\mathcal{B}}^*) = \{0\} \quad \text{for } \lambda \in \overline{\mathbb{C}}_{-\delta} \cap \sigma(\hat{\mathcal{A}}^*).$$

Assume that (λ_j, ϕ_j) , $j \in \{2, \dots, q\}$ is an eigenpair of $\hat{\mathcal{A}}^*$.

$$\hat{\mathcal{B}}^* \phi_j = \langle \hat{\mathcal{B}}, \phi_j \rangle = \left\langle \mathcal{P} \sum_{i=2}^q e^{-\Phi} \varphi_i, \phi_j \right\rangle = \sum_{i=2}^q \langle e^{-\Phi} \varphi_i, \varphi_j \rangle = \|e^{-\frac{\Phi}{2}} \varphi_j\|^2$$

e.g. [CURTAIN/ZWART'95]

The Riccati equation on $L^2(\Omega)$

Consider the two Riccati equations

$$(1) \quad (\mathcal{A} + \delta\mathcal{P})^*\Pi + \Pi(\mathcal{A} + \delta\mathcal{P}) - \Pi\mathcal{B}\mathcal{B}^*\Pi + \mathcal{P}^*\mathcal{M}\mathcal{P} = 0, \quad \Pi \in \mathcal{L}(\mathcal{Y}),$$

$$(2) \quad (\hat{\mathcal{A}}^* + \delta I)\hat{\Pi} + \hat{\Pi}(\hat{\mathcal{A}} + \delta I) - \hat{\Pi}\hat{\mathcal{B}}\hat{\mathcal{B}}^*\hat{\Pi} + \mathcal{M} = 0, \quad \hat{\Pi} \in \mathcal{L}(\mathcal{Y}_{\mathcal{P}}).$$

Lemma

- (i) If $\hat{\Pi} \in \mathcal{L}(\mathcal{Y}_{\mathcal{P}})$ is a solution to (2), then for all $\gamma \in \mathbb{R}$,
 $\Pi = \mathcal{P}^*\hat{\Pi}\mathcal{P} + \gamma\mathbb{1}\mathbb{1}^* \in \mathcal{L}(\mathcal{Y})$ is a solution to (1).
- (ii) If the Hautus criterion is satisfied, then a solution $\Pi \in \mathcal{L}(\mathcal{Y})$ to (1) exists and $\exists \gamma \in \mathbb{R}$ such that $\Pi = \mathcal{P}^*\hat{\Pi}\mathcal{P} + \gamma\mathbb{1}\mathbb{1}^*$ is a solution to (2) where $\hat{\Pi} := \Theta\Pi I_{\mathcal{P}} \in \mathcal{L}(\mathcal{Y}_{\mathcal{P}})$.

Local stabilization of the nonlinear system

What happens if we apply $u = -\hat{\mathcal{B}}^* \hat{\Pi} y_{\mathcal{P}}$ to the nonlinear system

$$\dot{y}_{\mathcal{P}} = \hat{\mathcal{A}} y_{\mathcal{P}} + u \hat{\mathcal{N}} y_{\mathcal{P}} + \hat{\mathcal{B}} u, \quad y_{\mathcal{P}}(0) = \mathcal{P} \rho_0.$$

Theorem

There exist $C, \tilde{C} > 0$ such that if $\|\rho_0\|_{L^2(\Omega)} \leq \frac{3}{16C^2\tilde{C}}$, then

$$\dot{y}_{\mathcal{P}} = (\hat{\mathcal{A}} - \hat{\mathcal{B}} \hat{\mathcal{B}}^* \hat{\Pi}) y_{\mathcal{P}} - (\hat{\mathcal{B}}^* \hat{\Pi} y_{\mathcal{P}}) \hat{\mathcal{N}} y_{\mathcal{P}}, \quad y_{\mathcal{P}}(0) = \mathcal{P} \rho_0,$$

admits a unique solution

$$y_{\mathcal{P}} \in W_{\mathcal{P}}^{\infty} := L^2(0, \infty; H^1(\Omega) \cap \mathcal{Y}_{\mathcal{P}}) \cap H^1(0, \infty; [H^1(\Omega) \cap \mathcal{Y}_{\mathcal{P}}]')$$

satisfying $\|e^{\delta t} y_{\mathcal{P}}\|_{W_{\mathcal{P}}^{\infty}} \leq \frac{1}{4C\tilde{C}}$. In particular, $\|y_{\mathcal{P}}\|_{L^2(\Omega)} \leq \tilde{C} e^{-\delta t} \|\rho_0\|_{L^2(\Omega)}$.

Compare e.g.: M. Badra, I. Lasiecka, J.-P. Raymond, R. Triggiani,...

Outline

Motivation

The Fokker Planck equation

A Riccati-based local feedback law

A Lyapunov-based global feedback law

A numerical study

An alternative choice for α

Let (λ_2, ψ_2) denote the second eigenpair of $\hat{\mathcal{A}}$.

Determine α as solution to the elliptic equation

$$\begin{aligned}\nabla \cdot (\rho_\infty \nabla \alpha) &= \psi_2 & \text{in } \Omega, \\ (\rho_\infty \nabla \alpha) \cdot \vec{n} &= 0 & \text{on } \Gamma.\end{aligned}$$

Further choose $\mu > 0$ such that

$$\langle (\mu I - \hat{\mathcal{A}})y_{\mathcal{P}}, y_{\mathcal{P}} \rangle_{L^2(\Omega)} \geq \langle y_{\mathcal{P}}, y_{\mathcal{P}} \rangle_{H^1(\Omega)}, \quad \text{for all } y_{\mathcal{P}} \in \mathcal{D}(\hat{\mathcal{A}}).$$

Determine the solution Υ to the Lyapunov equation:

$$\langle \Upsilon y_{\mathcal{P}}, \hat{\mathcal{A}} z_{\mathcal{P}} \rangle + \langle \hat{\mathcal{A}} y_{\mathcal{P}}, \Upsilon z_{\mathcal{P}} \rangle = -2\mu \langle y_{\mathcal{P}}, z_{\mathcal{P}} \rangle, \quad y_{\mathcal{P}}, z_{\mathcal{P}} \in \mathcal{D}(\hat{\mathcal{A}}).$$

A global feedback law

$$\langle \Upsilon y_{\mathcal{P}}, \hat{\mathcal{A}} z_{\mathcal{P}} \rangle + \langle \hat{\mathcal{A}} y_{\mathcal{P}}, \Upsilon z_{\mathcal{P}} \rangle = -2\mu \langle y_{\mathcal{P}}, z_{\mathcal{P}} \rangle, \quad y_{\mathcal{P}}, z_{\mathcal{P}} \in \mathcal{D}(\hat{\mathcal{A}}).$$

Theorem

Let μ and Υ be as before. Consider the system

$$\dot{y}_{\mathcal{P}} = \hat{\mathcal{A}} y_{\mathcal{P}} + u \hat{\mathcal{N}} y_{\mathcal{P}} + \hat{\mathcal{B}} u, \quad y_{\mathcal{P}}(0) = \mathcal{P} \rho_0.$$

where the control u is defined by the feedback law

$$u = -\langle \hat{\mathcal{B}} + \hat{\mathcal{N}} y_{\mathcal{P}}, \Upsilon y_{\mathcal{P}} + y_{\mathcal{P}} \rangle.$$

Then $V(y_{\mathcal{P}}) := \langle y_{\mathcal{P}}, \Upsilon y_{\mathcal{P}} + y_{\mathcal{P}} \rangle$ is a **global Lyapunov function**.

Theorem

Let $\lambda_i, i = 2, 3, \dots$ denote the eigenvalues of $\hat{\mathcal{A}}$. Assume that

$$\tilde{\lambda}_2 := \lambda_2 - \|\psi_2\|^2 + \frac{\mu}{\lambda_2} \|\psi_2\|^2 \neq \lambda_j, \quad j = 3, \dots$$

Then for the **spectrum** of the **linearized closed loop operator** it holds that

$$\sigma(\hat{\mathcal{A}} - \hat{\mathcal{B}} \hat{\mathcal{B}}^* \Upsilon - \hat{\mathcal{B}} \hat{\mathcal{B}}^*) = \{\tilde{\lambda}_2\} \cup \{\lambda_j\}, \quad j \geq 3.$$

Outline

Motivation

The Fokker Planck equation

A Riccati-based local feedback law

A Lyapunov-based global feedback law

A numerical study

Finite-dimensional properties

Assume spatial discretization yields $A, N, \rho^d, e = \frac{1}{d} (1, \dots, 1)^T$.

- ▶ $\int_{\Omega} \rho(t) \, dx = \int_{\Omega} \rho_0 \, dx \rightsquigarrow e^T \rho^d(t) = e^T \rho_0^d = 1$
- ▶ $\rho_0 \geq 0 \Rightarrow \rho(x, t) \geq 0 \, \forall t \rightsquigarrow A$ is a Metzler matrix
- ▶ $A_s = DAD^{-1}$ with $D = \text{diag}(e^{\frac{\phi^d}{2}})$ is a symmetric matrix
- ▶ $A^T e = N^T e = 0 = A \rho_{\infty}^d$
- ▶ $\mu = \max(\varepsilon, \varepsilon + \frac{1}{2} \lambda_{\max}(\hat{A} + \hat{A}^T))$

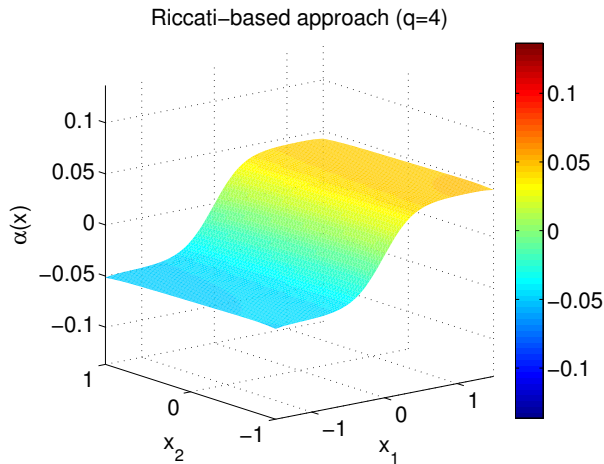
A two dimensional double well potential

As a numerical example, let us consider

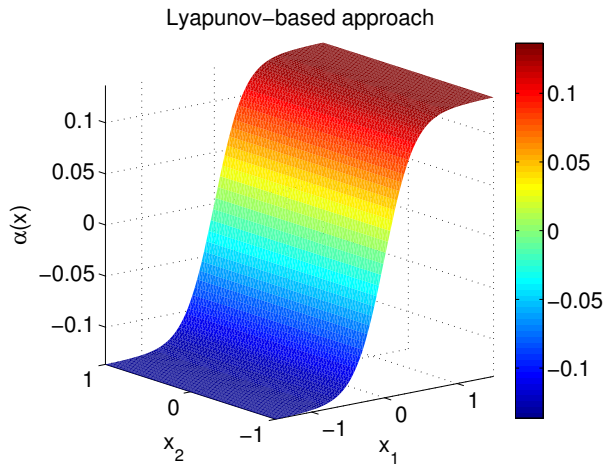
$$\begin{aligned}\frac{\partial \rho}{\partial t} &= \nu \Delta \rho + \nabla \cdot (\rho \nabla W) + u \nabla \cdot (\rho \nabla \alpha) && \text{in } \Omega \times (0, \infty), \\ 0 &= (\nu \nabla \rho + \rho \nabla W) \cdot \vec{n} && \text{on } \Gamma \times (0, \infty), \\ \rho(x, 0) &= \rho_0(x) && \text{in } \Omega.\end{aligned}$$

- ▶ $\Omega = (-1.5, 1.5) \times (-1.1)$
- ▶ $W(x) = 3(x_1^2 - 1)^2 + 6x_2^2$
- ▶ **finite differences** with $n_x \cdot n_y = 96 \cdot 64 = 6144$ grid points
- ▶ **upwind scheme** for convective terms

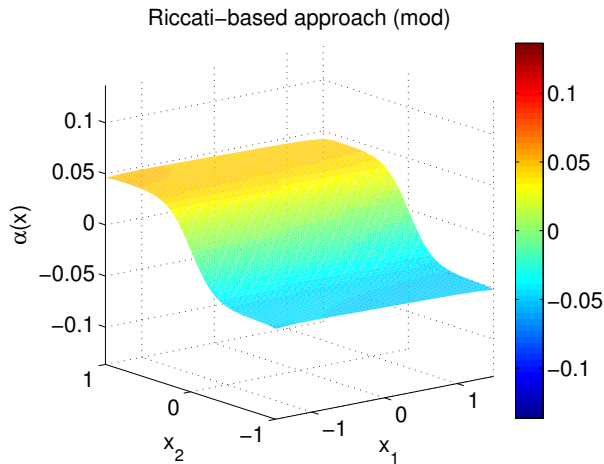
A comparison of different α



A comparison of different α

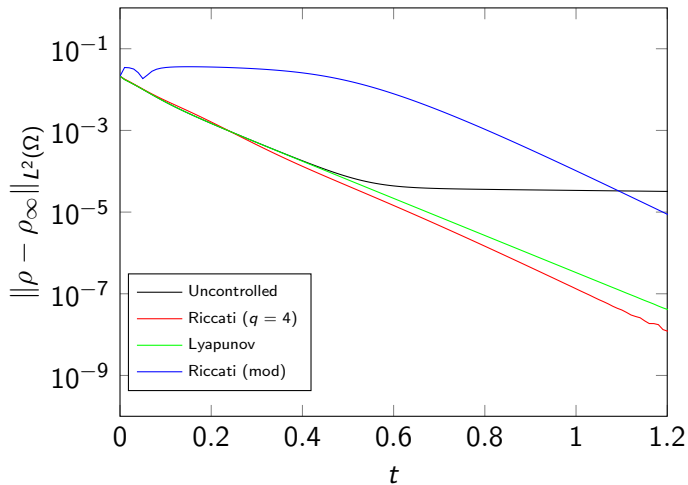


A comparison of different α



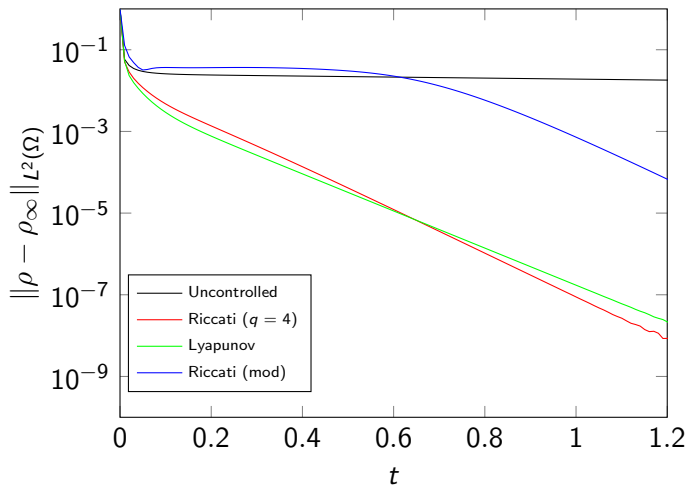
Some numerical results

Some numerical results



Some numerical results

Some numerical results



Conclusion

- ▶ Fokker-Planck equation yields a bilinear control system.
- ▶ Decoupled the system by spectral projections.
- ▶ Riccati-based feedback law $\Rightarrow$ local stabilization.
- ▶ Lyapunov-based feedback law $\Rightarrow$ global stabilization .
- ▶ Numerically efficient approaches?
- ▶ Optimal feedback law for bilinear system?

Conclusion

- ▶ Fokker-Planck equation yields a bilinear control system.
- ▶ Decoupled the system by spectral projections.
- ▶ Riccati-based feedback law $\Rightarrow$ local stabilization.
- ▶ Lyapunov-based feedback law $\Rightarrow$ global stabilization .
- ▶ Numerically efficient approaches?
- ▶ Optimal feedback law for bilinear system?

Thank you for your attention!