

Asymptotic comparison of nonparametric jump rate estimators for one-dimensional piecewise-deterministic Markov processes

Romain Azaïs and Solune Denis

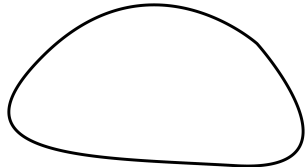
July 1st 2025 | Paris

Piecewise-deterministic Markov process (PDMP)

- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP

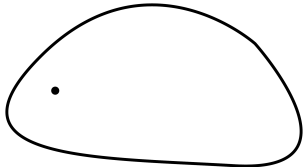
Piecewise-deterministic Markov process (PDMP)

- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP



Piecewise-deterministic Markov process (PDMP)

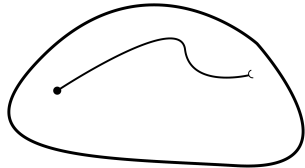
- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP



- Initial condition: $X_0 = x$

Piecewise-deterministic Markov process (PDMP)

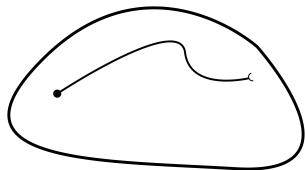
- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP



- Initial condition: $X_0 = x$
- Deterministic motion: $\forall 0 \leq t < T_1, X_t = \Phi(t|x)$

Piecewise-deterministic Markov process (PDMP)

- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP

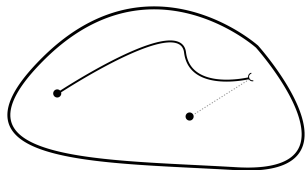


- Initial condition: $X_0 = x$
- Deterministic motion: $\forall 0 \leq t < T_1, X_t = \Phi(t|x)$
- Random time:

$$\mathbb{P}(T_1 > t) = \exp\left(-\int_0^t \lambda(\Phi(s|x)) ds\right) \mathbb{1}_{\{0 \leq t < t^+(x)\}}$$

Piecewise-deterministic Markov process (PDMP)

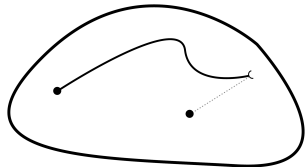
- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP



- Initial condition: $X_0 = x$
- Deterministic motion: $\forall 0 \leq t < T_1, X_t = \Phi(t|x)$
- Random time:
$$\mathbb{P}(T_1 > t) = \exp\left(-\int_0^t \lambda(\Phi(s|x)) ds\right) \mathbb{1}_{\{0 \leq t < t^+(x)\}}$$
- Random jump:
$$\mathbb{E}[\varphi(X_{T_1}) | \Phi(T_1|x)] = \int \varphi(u) Q(du | \Phi(T_1|x))$$

Piecewise-deterministic Markov process (PDMP)

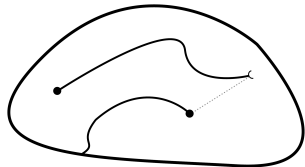
- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP



- $X_{T_1} = x$

Piecewise-deterministic Markov process (PDMP)

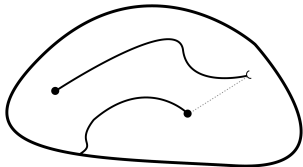
- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP



- $X_{T_1} = x$
- Deterministic motion: $\forall 0 \leq t < S_2, X_{T_1+t} = \Phi(t|x)$

Piecewise-deterministic Markov process (PDMP)

- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP

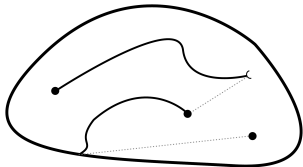


- $X_{T_1} = x$
- Deterministic motion: $\forall 0 \leq t < S_2, X_{T_1+t} = \Phi(t|x)$
- Random time:

$$\mathbb{P}(S_2 > t) = \exp\left(-\int_0^t \lambda(\Phi(s|x)) ds\right) \mathbb{1}_{\{0 \leq t < t^+(x)\}}$$

Piecewise-deterministic Markov process (PDMP)

- Introduced in the 80's by Davis
- The class is “wide enough to include as special cases virtually all the non-diffusion models of applied probability”
- MC, CT-MC, M/G/1 queue, G/G/1 queue $\subset$ PDMP



- $X_{T_1} = x$
- Deterministic motion: $\forall 0 \leq t < S_2, X_{T_1+t} = \Phi(t|x)$
- Random time:

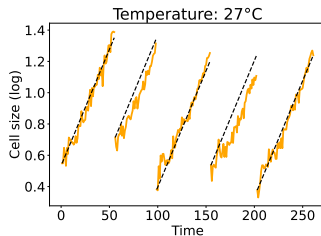
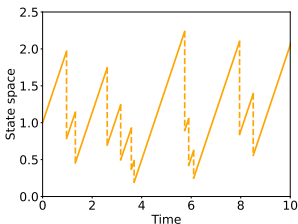
$$\mathbb{P}(S_2 > t) = \exp\left(-\int_0^t \lambda(\Phi(s|x)) ds\right) \mathbb{1}_{\{0 \leq t < t^+(x)\}}$$

- Random jump:

$$\mathbb{E}[\varphi(X_{T_2}) | \Phi(S_2|x)] = \int \varphi(u) Q(du | \Phi(S_2|x))$$

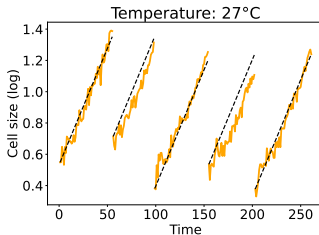
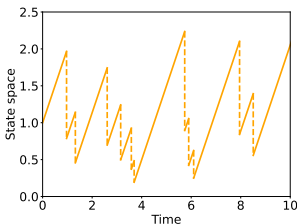
Typical PDMP in dimension 1: growth-fragmentation

- Growth: $\Phi(\cdot|x) \nearrow$, e.g. $\Phi(t|x) = x + \theta t$
- Fragmentation: $Q((0, x)|x) = 1$, e.g. $Q(\cdot|x) = \delta_{\{h(x)\}}$



Typical PDMP in dimension 1: growth-fragmentation

- Growth: $\Phi(\cdot|x) \nearrow$, e.g. $\Phi(t|x) = x + \theta t$
- Fragmentation: $Q((0, x)|x) = 1$, e.g. $Q(\cdot|x) = \delta_{\{h(x)\}}$



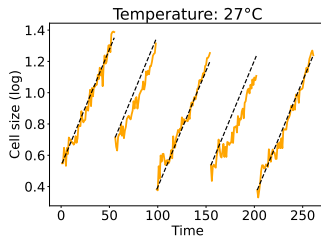
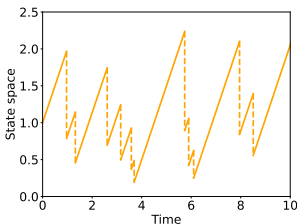
Statistical framework: one-dimensional PDMP

- Φ is known
- First n jumps (exactly) observed: $(Z_k, S_{k+1})_{0 \leq k \leq n}$
- Ergodicity condition: asymptotic statistics

Question: nonparametric estimation of λ ?

Typical PDMP in dimension 1: growth-fragmentation

- Growth: $\Phi(\cdot|x) \nearrow$, e.g. $\Phi(t|x) = x + \theta t$
- Fragmentation: $Q((0, x)|x) = 1$, e.g. $Q(\cdot|x) = \delta_{\{h(x)\}}$



Statistical framework: one-dimensional PDMP

- Φ is known
- First n jumps (exactly) observed: $(Z_k, S_{k+1})_{0 \leq k \leq n}$
- Ergodicity condition: asymptotic statistics

Question: nonparametric estimation of λ ?

Invariant measures

$$X_t \rightarrow \mu^{\text{cr}} \quad Z_n \rightarrow \mu$$

$$Z_n^- = \Phi(S_n, Z_{n-1}) \rightarrow \mu^-$$

Strategies from the literature (1/2)

Fujii (2013)



- $\Phi(t|x) = x + \int_0^t \Delta(\Phi(s|x))ds$
- Estimation of μ^{ct} through the local time, in relation with μ^- and λ
- μ^- estimated by kernel methods (in the (true) time of the process)
- Uniform convergence in probability

Strategies from the literature (1/2)

Fujii (2013)



- $\Phi(t|x) = x + \int_0^t \Delta(\Phi(s|x))ds$
- Estimation of μ^{CT} through the local time, in relation with μ^- and λ
- μ^- estimated by kernel methods (in the (true) time of the process)
- Uniform convergence in probability

Azaïs and Muller-Gueudin (2016)



- Multidimensional PDMP with forced jumps
- $\lambda(\Phi(t|x))$ estimated as the ratio of two invariant measures
 $\leadsto$ collection of consistent estimators of $\lambda(x)$
- Selection of the estimator with minimal asymptotic variance
- Pointwise almost sure convergence and central limit theorem

Strategies from the literature (2/2)

Krell (2016)



- $Q(\cdot|x) = \delta_{\{h(x)\}}$ assumed to be known
- $h \circ \Phi(\cdot|x)$ as a change of variable $\leadsto$ relation between λ and μ
- μ estimated by kernel methods
- Convergence of the mean integrated square error to 0

Strategies from the literature (2/2)

Krell (2016)



- $Q(\cdot|x) = \delta_{\{h(x)\}}$ assumed to be known
- $h \circ \Phi(\cdot|x)$ as a change of variable $\leadsto$ relation between λ and μ
- μ estimated by kernel methods
- Convergence of the mean integrated square error to 0

Krell and Schmisser (2021)



- General Q , not assumed to be known
- $\Phi(\cdot|x)$ as a change of variable $\leadsto$ relation between λ and μ^-
- μ^- estimated by projection methods
- Estimator nearly minimax (up to a $\log^2 n$ factor)

Observation 1

- Different statistical frameworks
- Different estimation methods
- Different convergence results

~> Estimation strategies are difficult to compare

Observation 1

- Different statistical frameworks
- Different estimation methods
- Different convergence results

↪ Estimation strategies are difficult to compare

Observation 2

Fujii (2013) and Krell and Schmisser (2021) estimate λ via

$$\lambda(x) = \frac{\Delta(x)\mu^-(x)}{\mathbb{P}_\mu(Z_0 \leq x < Z_1)}$$

where $\Delta(x) = \partial_t \Phi(0|x)$

Standardization (1/2)

- General one-dimensional PDMP
- Invariant measures estimated by (non-recursive) kernel methods
- Asymptotics in the number of jumps observed
- Goal: central limit theorem with explicit variance

Standardization (1/2)

- General one-dimensional PDMP
- Invariant measures estimated by (non-recursive) kernel methods
- Asymptotics in the number of jumps observed
- Goal: central limit theorem with explicit variance

Krell (2016): $Q(\cdot|x) = \delta_{\{h(x)\}}$ where h is assumed to be known

$$\hat{\lambda}_n^{\clubsuit}(x) = h'(x)\Delta(x) \times \frac{\frac{1}{n} \sum_{i=0}^{n-1} K_{h_n}(Z_i - h(x))}{\frac{1}{n} \sum_{i=0}^{n-1} \mathbb{1}_{\{Z_i \leq x\}} \mathbb{1}_{\{Z_{i+1} \geq h(x)\}}}$$

Standardization (1/2)

- General one-dimensional PDMP
- Invariant measures estimated by (non-recursive) kernel methods
- Asymptotics in the number of jumps observed
- Goal: central limit theorem with explicit variance

Krell (2016): $Q(\cdot|x) = \delta_{\{h(x)\}}$ where h is assumed to be known

$$\hat{\lambda}_n^{\text{🌸}}(x) = h'(x)\Delta(x) \times \frac{\frac{1}{n} \sum_{i=0}^{n-1} K_{h_n}(Z_i - h(x))}{\frac{1}{n} \sum_{i=0}^{n-1} \mathbb{1}_{\{Z_i \leq x\}} \mathbb{1}_{\{Z_{i+1} \geq h(x)\}}}$$

Fujii (2013) and Krell and Schmisser (2021)

$$\hat{\lambda}_n^{\text{💎}}(x) = \Delta(x) \times \frac{\frac{1}{n} \sum_{i=1}^n K_{h_n}(Z_i^- - x)}{\frac{1}{n} \sum_{i=0}^{n-1} \mathbb{1}_{\{Z_i \leq x < Z_{i+1}^-\}}}$$

Standardization (2/2)

Azaïs and Muller-Gueudin (2016)

$$\widehat{\lambda \circ \Phi}_n(t|x) = \frac{\frac{1}{n} \sum_{i=0}^{n-1} K_{h_n^\bullet}(Z_i - x) K_{h_n^\bullet}(S_{i+1} - t)}{\frac{1}{n} \sum_{i=0}^{n-1} K_{h_n^\bullet}(Z_i - x) \mathbb{1}_{\{S_{i+1} > t\}}}$$

$$\overset{\circ}{\lambda}_n^{\spadesuit}(x) = \widehat{\lambda \circ \Phi}_n(\tau_x(\xi)|\xi) \quad \text{with} \quad \xi = \arg \max_{C_x} \mu(\cdot) G(\tau_x(\cdot)|\cdot)$$

$$\widehat{\lambda}_n^{\spadesuit}(x) = \widehat{\lambda \circ \Phi}_n(\tau_x(\xi)|\xi) \quad \text{with} \quad \xi = \arg \max_{C_x} \sum_{i=0}^{n-1} K_{h_n^\bullet}(Z_i - \cdot) \mathbb{1}_{\{S_{i+1} > \tau_x(\cdot)\}}$$

Baseline

- Almost sure convergence: $\hat{\lambda}_n^{\clubsuit}(x)$, $\hat{\lambda}_n^{\diamondsuit}(x)$ and $\hat{\lambda}_n^{\spadesuit}(x)$ tend to $\lambda(x)$
- Central limit theorem for $\hat{\lambda}_n^{\spadesuit}(x)$

$$\sqrt{nh_n^s h_n^t} \left[\hat{\lambda}_n^{\spadesuit}(x) - \lambda(x) \right] \rightarrow \mathcal{N}(0, \sigma_{\spadesuit}^2(x))$$

where

$$\sigma_{\spadesuit}^2(x) = \frac{\tau^4 \lambda(x)}{\max_{\xi \in \mathcal{C}_x} \mu(\xi) G(\tau_x(\xi) | \xi)}$$

Baseline

- Almost sure convergence: $\hat{\lambda}_n^{\clubsuit}(x)$, $\hat{\lambda}_n^{\diamondsuit}(x)$ and $\check{\lambda}_n^{\spadesuit}(x)$ tend to $\lambda(x)$
- Central limit theorem for $\check{\lambda}_n^{\spadesuit}(x)$

$$\sqrt{nh_n^s h_n^t} \left[\check{\lambda}_n^{\spadesuit}(x) - \lambda(x) \right] \rightarrow \mathcal{N}(0, \sigma_{\spadesuit}^2(x))$$

where

$$\sigma_{\spadesuit}^2(x) = \frac{\tau^4 \lambda(x)}{\max_{\xi \in \mathcal{C}_x} \mu(\xi) G(\tau_x(\xi) | \xi)}$$

Main theorem: asymptotic normality

Under ergodicity and regularity conditions,

$$\begin{aligned} \sqrt{nh_n} \left[\hat{\lambda}_n^{\clubsuit}(x) - \lambda(x) \right] &\rightarrow \mathcal{N}(0, \sigma_{\clubsuit}^2(x)) \\ \sqrt{nh_n} \left[\hat{\lambda}_n^{\diamondsuit}(x) - \lambda(x) \right] &\rightarrow \mathcal{N}(0, \sigma_{\diamondsuit}^2(x)) \end{aligned}$$

where

$$\sigma_{\clubsuit}^2(x) = \frac{\tau^2 \lambda(x)^2 h'(x)}{\mu^-(x)} \quad \text{and} \quad \sigma_{\diamondsuit}^2(x) = \frac{\tau^2 \lambda(x)^2}{\mu^-(x)}$$

$$\hat{\lambda}_n^{\clubsuit}(x) \text{ vs. } \hat{\lambda}_n^{\blacklozenge}(x)$$

- Same rate: $\sqrt{nh_n}$
- $\sigma^2_{\clubsuit}(x) = h'(x)\sigma^2_{\blacklozenge}(x)$

$$\hat{\lambda}_n^{\clubsuit}(x) \text{ vs. } \hat{\lambda}_n^{\blacklozenge}(x)$$

- Same rate: $\sqrt{nh_n}$
- $\sigma_{\clubsuit}^2(x) = h'(x)\sigma_{\blacklozenge}^2(x)$

Examples

- Linear fragmentation: $h(x) = \kappa x$, $0 < \kappa < 1$

Then $\sigma_{\clubsuit}^2(x) = \kappa\sigma_{\blacklozenge}^2(x)$ and $\hat{\lambda}_n^{\clubsuit}$ is uniformly better than $\hat{\lambda}_n^{\blacklozenge}$

$$\hat{\lambda}_n^{\clubsuit}(x) \text{ vs. } \hat{\lambda}_n^{\blacklozenge}(x)$$

- Same rate: $\sqrt{nh_n}$
- $\sigma^2_{\clubsuit}(x) = h'(x)\sigma^2_{\blacklozenge}(x)$

Examples

- Linear fragmentation: $h(x) = \kappa x$, $0 < \kappa < 1$

Then $\sigma^2_{\clubsuit}(x) = \kappa\sigma^2_{\blacklozenge}(x)$ and $\hat{\lambda}_n^{\clubsuit}$ is uniformly better than $\hat{\lambda}_n^{\blacklozenge}$

- $h(x) = x/(1 + \exp(-x))$ is sublinear
 - $h'(1) < 1 \rightsquigarrow \hat{\lambda}_n^{\clubsuit}(1)$ is better than $\hat{\lambda}_n^{\blacklozenge}(1)$
 - $h'(2) > 1 \rightsquigarrow \hat{\lambda}_n^{\blacklozenge}(2)$ is better than $\hat{\lambda}_n^{\clubsuit}(2)$

Comparison with $\mathring{\lambda}_n^\spadesuit(x)$ is trickier

- Rate of $\mathring{\lambda}_n^\spadesuit(x)$: $\sqrt{nh_n^s h_n^t}$
- Rate of $\hat{\lambda}_n^\clubsuit(x)$ and $\hat{\lambda}_n^\diamondsuit(x)$: $\sqrt{nh_n}$

Extremal cases

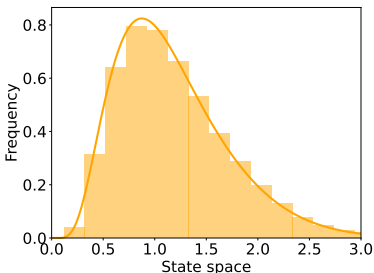
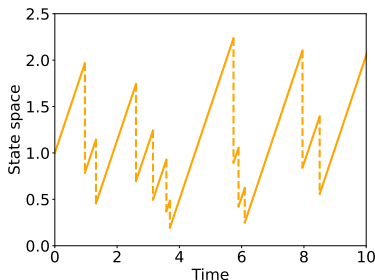
- $h_n = h_n^s$: $\mathring{\lambda}_n^\spadesuit$ is asymptotically disqualified
- $h_n = h_n^s h_n^t$: direct comparison of variances is valid

In practice, bandwidths are chosen from the available data

Growth-fragmentation model

$$\begin{cases} \lambda(x) &= x \\ \Phi(t|x) &= x + t \\ Q(dy|x) &= \delta_{\{\kappa x\}}(dy) \end{cases}$$

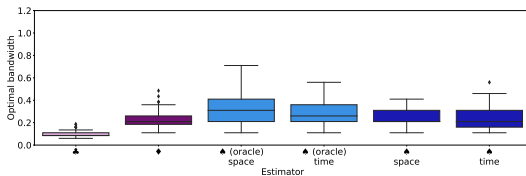
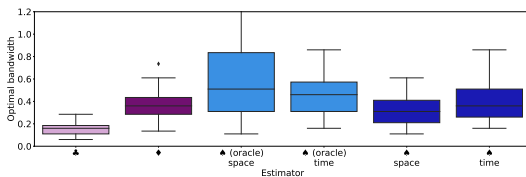
$$\mu^{\text{ct}}(x) = \frac{\sqrt{2/\pi}}{\prod_{n=0}^{+\infty} (1 - \kappa^{2n+1})} \sum_{n=0}^{+\infty} \frac{\kappa^{-2n}}{\prod_{k=1}^n (1 - \kappa^{-2k})} \exp(-\kappa^{-2n} x^2 / 2)$$



Bandwidth selection

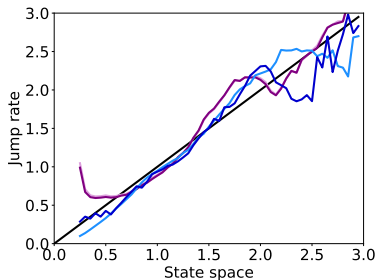
$$h^* = \arg \min_{h \in \mathbb{R}_+^*} \|\lambda - \hat{\lambda}\|_{\mathbb{L}_{[0.5, 2.5]}^2}$$

$n = 1000$

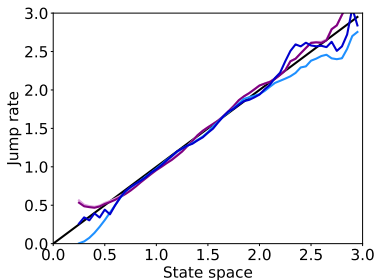


$n = 10\,000$

Comparison from one sample (not robust)

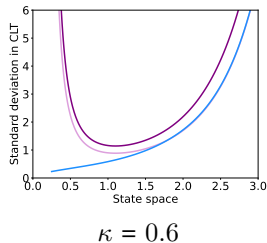
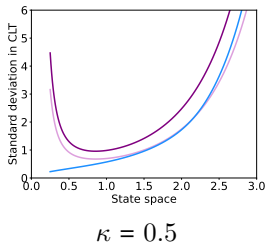
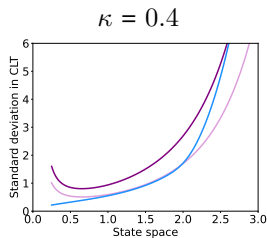
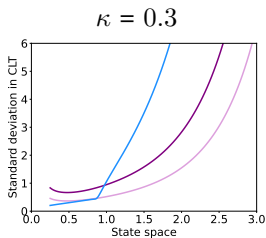


$n = 1000$



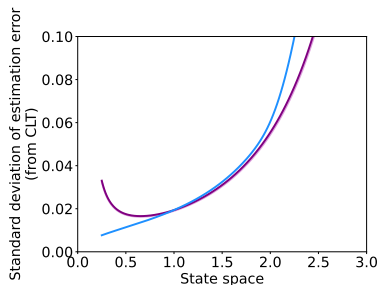
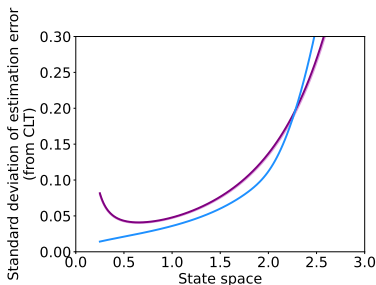
$n = 10\,000$

Variance from the central limit theorem (not normalized)



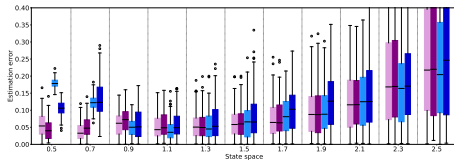
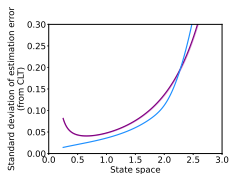
Expected asymptotic variance (normalized)

— $\kappa = 0.4$ —



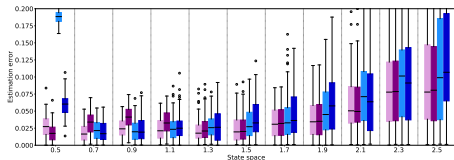
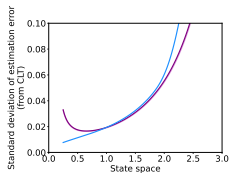
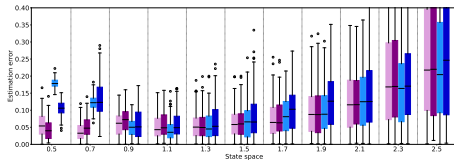
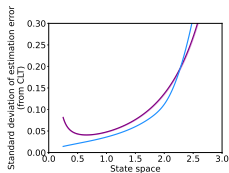
Expected asymptotic variance vs. estimation error

$n = 1000$



Expected asymptotic variance vs. estimation error

$n = 1000$



$n = 10\,000$

Real data analysis (1/4)

— Data description —

Tanouchi et al. (2017): cell growth and division data (E. coli)

Temperature		Lineages	Measurements	Division events
25°C	Total:	65	307 999	4 485
	Average:	—	4738.44	68.67
27°C	Total:	54	202 086	3 726
	Average:	—	3742.33	69
37°C	Total:	160	364 920	11 040
	Average:	—	2280.75	69

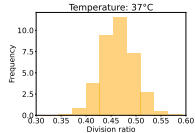
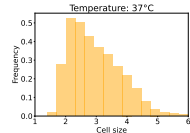
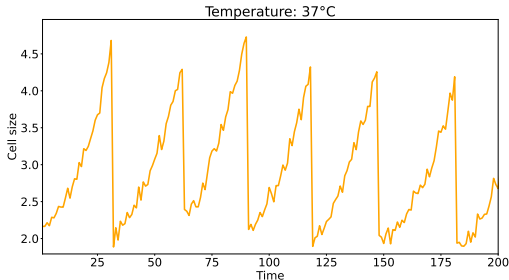
Real data analysis (1/4)

— Data description —

Tanouchi et al. (2017): cell growth and division data (E. coli)

SCIENTIFIC
DATA

Temperature		Lineages	Measurements	Division events
25°C	Total:	65	307 999	4 485
	Average:	—	4738.44	68.67
27°C	Total:	54	202 086	3 726
	Average:	—	3742.33	69
37°C	Total:	160	364 920	11 040
	Average:	—	2280.75	69



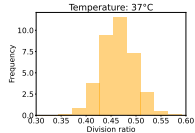
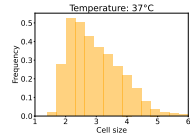
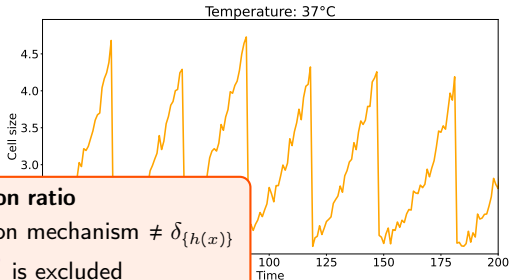
Real data analysis (1/4)

— Data description —

Tanouchi et al. (2017): cell growth and division data (E. coli)

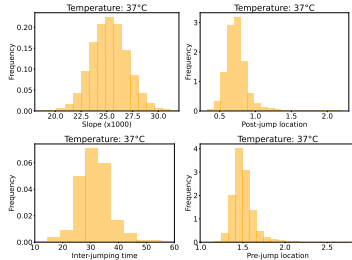
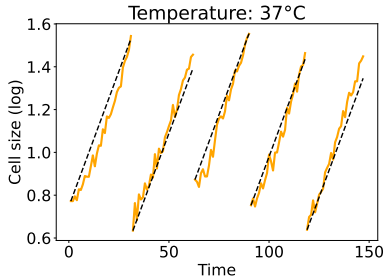
SCIENTIFIC
DATA

Temperature		Lineages	Measurements	Division events
25°C	Total:	65	307 999	4 485
	Average:	—	4738.44	68.67
27°C	Total:	54	202 086	3 726
	Average:	—	3742.33	69
37°C	Total:	160	364 920	11 040
	Average:	—	2280.75	69



Real data analysis (2/4)

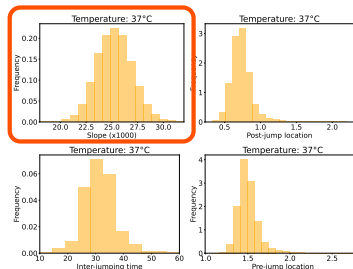
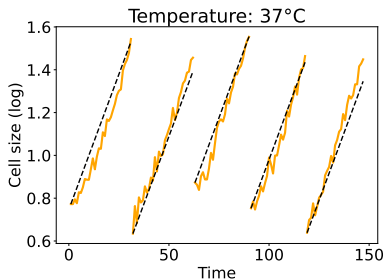
— Modeling —



- Cell growth looks linear in log scale
- Slope estimated for each of the cells

Real data analysis (2/4)

— Modeling —



- Cell growth looks linear in log scale
- Slope estimated for each of the cells $\leadsto$ pretty sharp distribution
- Common slope under each temperature condition: $\Phi(t|x) = x + \theta t$

Temperature	25°C	27°C	37°C
θ	0.012	0.014	0.025

Real data analysis (3/4)

— Estimation procedures —

- $\hat{\lambda}_n^{\spadesuit}$ can not be evaluated
- Argument selection simplified for $\hat{\lambda}_n^{\spadesuit} : \hat{\lambda}_n^{\spadesuit}(x) = \widehat{\lambda \circ \Phi_n}(\tau | x - \theta\tau)$

Temperature	25°C	27°C	37°C
τ	66.6	52.4	31.6

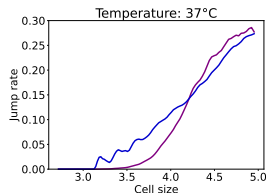
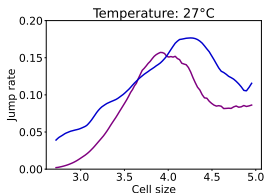
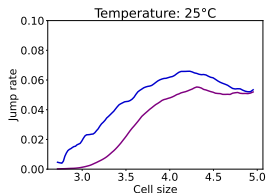
- Bandwidth selection

Temperature	Sample size	♦	♠ (space)	♠ (time)
25°C	4 485	0.05	0.06	4
27°C	3 726	0.07	0.08	8
37°C	11 040	0.02	0.03	3

Real data analysis (4/4)

— Estimation results —

$$\hat{\lambda}_n^{\blacklozenge}(\exp(x)) \text{ vs. } \hat{\lambda}_n^{\clubsuit}(\exp(x))$$



Take-home message

- Asymptotic variances can not be ordered

Take-home message

- Asymptotic variances can not be ordered
- $\hat{\lambda}_n^{\clubsuit}$ and $\hat{\lambda}_n^{\spadesuit}$ are the most affected by bias
- $\hat{\lambda}_n^{\blacklozenge}$ is more general than $\hat{\lambda}_n^{\clubsuit}$ and achieves pretty similar results

Take-home message

- Asymptotic variances can not be ordered
- $\hat{\lambda}_n^{\clubsuit}$ and $\hat{\lambda}_n^{\spadesuit}$ are the most affected by bias
- $\hat{\lambda}_n^{\diamondsuit}$ is more general than $\hat{\lambda}_n^{\clubsuit}$ and achieves pretty similar results
- Bandwidth parameters of $\hat{\lambda}_n^{\clubsuit}$ and $\hat{\lambda}_n^{\spadesuit}$ are difficult to tune

Take-home message

- Asymptotic variances can not be ordered
- $\hat{\lambda}_n^{\clubsuit}$ and $\hat{\lambda}_n^{\spadesuit}$ are the most affected by bias
- $\hat{\lambda}_n^{\blacklozenge}$ is more general than $\hat{\lambda}_n^{\clubsuit}$ and achieves pretty similar results
- Bandwidth parameters of $\hat{\lambda}_n^{\clubsuit}$ and $\hat{\lambda}_n^{\spadesuit}$ are difficult to tune

→ $\hat{\lambda}_n^{\blacklozenge}$ from Fujii (2013) and Krell and Schmisser (2021) tends to be the best estimator of λ in dimension 1!

Take-home message

- Asymptotic variances can not be ordered
- $\hat{\lambda}_n^{\clubsuit}$ and $\hat{\lambda}_n^{\spadesuit}$ are the most affected by bias
- $\hat{\lambda}_n^{\diamondsuit}$ is more general than $\hat{\lambda}_n^{\heartsuit}$ and achieves pretty similar results
- Bandwidth parameters of $\hat{\lambda}_n^{\clubsuit}$ and $\hat{\lambda}_n^{\spadesuit}$ are difficult to tune

→ $\hat{\lambda}_n^{\diamondsuit}$ from Fujii (2013) and Krell and Schmisser (2021) tends to be the best estimator of λ in dimension 1!

Azaïs and Denis (2025)

 Asymptotic analysis and practical evaluation of jump rate estimators in piecewise-deterministic Markov processes

Thank you for your attention!