

GBEES-GPU: An efficient parallel GPU algorithm for high-dimensional nonlinear uncertainty propagation

Benjamin L. Hanson^{a,c}, Carlos Rubio^b, Adrián García-Gutiérrez^b, Thomas Bewley^a

^a*Department of Mechanical and Aerospace Engineering, University of California San Diego, 9500 Gilman Dr, La Jolla, 92093, CA, USA*

^b*Department of Aerospace Engineering, Universidad de León, Av. Facultad de Veterinaria, 25, León, 24004, Spain*

^c*Corresponding author, Email: blhanson@ucsd.edu, Tel: +1 (913) 626-5877,*

Abstract

Eulerian nonlinear uncertainty propagation methods often suffer from finite domain limitations and computational inefficiencies. A recent approach to this class of algorithm, Grid-based Bayesian Estimation Exploiting Sparsity, addresses the first challenge by dynamically allocating a discretized grid in regions of phase space where probability is non-negligible. However, the design of the original algorithm causes the second challenge to persist in high-dimensional systems. This paper presents an architectural optimization of the algorithm for CPU implementation, followed by its adaptation to the CUDA framework for GPU execution. The algorithm is validated for correct convergence and accuracy, with performance evaluated across multiple GPUs. Tests include propagating a three-dimensional probability distribution subject to the Lorenz '63 model and a six-dimensional probability distribution subject to the Lorenz '96 model. The results imply that the improvements made result in a speedup of over 1000 times compared to the original implementation.

Keywords:

Eulerian uncertainty propagation, Corner Transport Upwind, Dynamic gridding, Hashtables, CUDA, GPU

1. Introduction

Nonlinear uncertainty propagation methods can generally be classified as Kalman, Lagrangian, and/or Eulerian. The Kalman approach approximates uncertainty as Gaussian, or as a mixture of Gaussians. For linear dynamics and measurement models, Gaussian uncertainty remains Gaussian globally, but in the presence of nonlinearities, Kalman methods are suboptimal [1]. Analytical linearizations, in the case of the Extended Kalman Filter (EKF), and statistical linearizations, in the cases of the Unscented Kalman Filter (UKF) [2] and the Ensemble Kalman Filter (EnKF) [3], are utilized to more accurately represent nonlinearities, but both tend to diverge when Gaussian measurement corrections are relatively infrequent. Alternatively, Gaussian Mixture Models (GMMs) represent non-Gaussian uncertainty as a weighted superposition of Gaussians [4]. Certain GMM methods use splitting procedures triggered by entropy flags to increase the number of components in the superposition as true uncertainty becomes more non-Gaussian [5].

Langrangian methods perform Sequential Monte Carlo (SMC) estimation on point mass representations of probability densities. Ensemble members are randomly drawn from an *a priori* distribution and then time-marched up to a measurement correction epoch via the true dynamics model. The simplest SMC method, Monte Carlo (MC) integration, attributes equal weight to each point [6]; this method prevents measurement corrections, limiting its applicability to uncertainty prediction. More sophisticated approaches assign weight based on an importance sampling distribution; these are known as Sequential Importance Sampling (SIS) methods, or, more commonly, particle filters [7]. For particle filters, measurement corrections are incorporated via weight adjustments and resampling procedures are utilized to avoid particle degeneracy [8, 9]. Hybrid Kalman/Lagrangian methods attribute Gaussian kernels to each particle and update the associated moments according to the GMM formulation [10].

The Eulerian approach considers and updates probability at fixed points in space. For stochastic processes dominated by deterministic forces, the time-evolution of the full probability density function is a hyperbolic partial differential equation (PDE). Considerable effort has been put forth by the fluid mechanics community towards numerically solving these types of PDEs via finite difference/volume methods [11] which can be divided into two categories: semidiscrete and fully discrete. Semidiscrete methods, like the Essentially Non-Oscillatory (ENO) and Weighted Essentially Non-Oscillatory

(WENO) schemes [12], use adaptive stencils to create smooth interpolations of probability across discontinuities, then use Runge-Kutta time stepping to march the system of semidiscrete equations. Fully discrete methods use fluxes defined at grid cell interfaces to update discretized probability at grid cell centers. The Lax-Wendroff and Godunov methods are two, first-order accurate examples [13, 14], but higher-order corrections are necessary to achieve second-order accuracy, and flux limiters ensure these methods are total variation diminishing [15].

Of the three defined approaches, Eulerian methods are generally the least explored for high-dimensional nonlinear uncertainty propagation. This is most likely due to both the finite domain limitation for standard grid-based methods as well as the high computational intensity that accompanies the application of fluids-based finite volume methods to ($n > 3$)-dimensional probability density time-marching. However, Eulerian approaches do not require splitting procedures to maintain accuracy, do not succumb to particle degeneracy, and are extremely robust for chaotic dynamics models, underdetermined measurement models, and infrequent correction updates.

A novel Eulerian approach known as Grid-based Bayesian Estimation Exploiting Sparsity (GBEES) dynamically allocates grid cells in regions of non-negligible probability, unlocking propagation over all of phase space [16]. However, the algorithm’s $\mathcal{O}(n^2)$ time complexity poses computational challenges for high-dimensional systems. In this paper, we optimize the computational architecture of GBEES by:

1. Storing the dynamic grid in a hashtable
2. Time-marching with a CFL-minimized adaptive step size
3. Employing directional growing and pruning procedures
4. Implementing the algorithm in CUDA

These improvements result in GBEES-GPU, an efficient, high-dimensional parallel GPU algorithm for nonlinear uncertainty propagation. In Section 1 we outlined the landscape of nonlinear uncertainty propagation methods and defined our contributions towards the computational optimization of GBEES. In Section 2 the finite volume formulation underlying GBEES is extended to n -dimensions. Section 3 outlines the improvements made to the CPU implementation, while Section 4 describes its adaptation to the CUDA architecture. Section 5 presents the use cases employed for validation and testing, including the evaluation of the correct convergence and accuracy. Next, Section 6 compares the performance of the improved CPU and

GPU implementations against the legacy version. Conclusions are presented in Section 7, and the hardware specifications for the tests are provided in Appendix A.

2. Grid-based Bayesian Estimation Exploiting Sparsity

The equations of motion of a stochastic process $\mathbf{X}(t) \in \mathbb{R}^n$ governed by a combination of deterministic and random forces can be described by the following stochastic differential equation:

$$d\mathbf{X}(t) = \mathbf{f}(\mathbf{X}(t), t)dt + \mathbf{q}(\mathbf{X}(t), t)d\mathbf{W}(t), \quad (1)$$

where $d\mathbf{W}(t) = \boldsymbol{\xi}(t)dt$ is a Wiener process, meaning $\boldsymbol{\xi}(t)$ is zero-mean, uncorrelated white noise (i.e., $\mathbb{E}[\boldsymbol{\xi}(t)] = \mathbf{0}$ and $\mathbb{E}[\boldsymbol{\xi}(t + \tau)\boldsymbol{\xi}^T(t)] = \boldsymbol{\delta}(\tau)$). In continuous-time, the Fokker-Planck equation gives the evolution of the probability density function (PDF) $p(\mathbf{x}, t)$ of $\mathbf{X}(t)$ in Eq. (1) as follows:

$$\frac{\partial p(\mathbf{x}, t)}{\partial t} = - \sum_{j=1}^n \frac{\partial f_j(\mathbf{x}, t)p(\mathbf{x}, t)}{\partial x_j} + \frac{1}{2} \sum_{j=1}^n \sum_{\ell=1}^n \frac{\partial^2 Q_{j\ell}(\mathbf{x}, t)p(\mathbf{x}, t)}{\partial x_j \partial x_\ell} \quad (2)$$

where $\mathbf{x} = (x_1, \dots, x_n)$, $f_j(\mathbf{x}, t)$ is the j^{th} component of $\mathbf{f}(\mathbf{x}, t)$, $Q_{j\ell}(\mathbf{x}, t)$ is the $(j, \ell)^{\text{th}}$ component of $Q(\mathbf{x}, t) = \mathbf{q}(\mathbf{x}, t)\mathbf{q}^T(\mathbf{x}, t)$. If $Q(\mathbf{x}, t) > 0$, Eq. (2) is elliptic, but if $Q(\mathbf{x}, t)$ is relatively small compared to the deterministic forces, Eq. (2) is hyperbolic and satisfies the conservative form of the n -dimensional advection equation:

$$\frac{\partial p(\mathbf{x}, t)}{\partial t} + \sum_{j=1}^n \frac{\partial f'_j(p(\mathbf{x}, t))}{\partial x_j} = 0, \quad (3)$$

where $f'_j(p(\mathbf{x}, t)) = f_j(\mathbf{x}, t)p(\mathbf{x}, t)$. At discrete measurement intervals $t^{(k)}$ the PDF is updated via Bayes' theorem:

$$p(\mathbf{x}, t^{(k+)}) = \frac{p(\mathbf{y}^{(k)}|\mathbf{x})p(\mathbf{x}, t^{(k-)})}{C}, \quad (4)$$

where $p(\mathbf{x}, t^{(k+)})$ is the *a posteriori*, $p(\mathbf{y}^{(k)}|\mathbf{x})$ is the measurement likelihood, $p(\mathbf{x}, t^{(k-)})$ is the *a priori*, and C is a normalization constant. GBEES performs the accurate, mixed continuous/discrete time-marching of $p(\mathbf{x}, t)$ using numerical approximations of Eqs. (3) and (4). We first delve into the continuous-time prediction of $p(\mathbf{x}, t)$ via a fully discrete flux-differencing method.

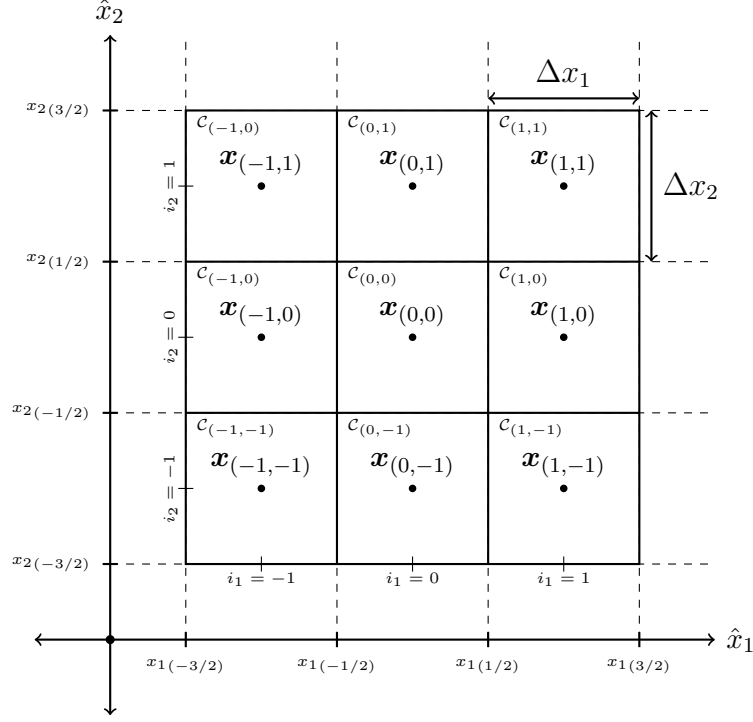


Figure 1: 2D schematic depicting the notation of GBEES.

2.1. Fully discrete flux-differencing methods

Consider a hyperrectangular grid cell of the form

$$\mathcal{C}_{(\mathbf{i})} = \prod_{j=1}^n [x_{j(i_j-1/2)}, x_{j(i_j+1/2)}] \quad (5)$$

where $\mathbf{i} = (i_1, \dots, i_n)$ is the grid cell index vector corresponding to the grid cell center coordinate vector $\mathbf{x}_{(\mathbf{i})}$. The grid width of $\mathcal{C}_{(\mathbf{i})}$ in the x_j -direction is then

$$\Delta x_{j(\mathbf{i})} = x_{j(i_j+1/2)} - x_{j(i_j-1/2)}; \quad (6)$$

for a uniform grid, $\Delta x_{j(\mathbf{i})}$ is equal for all $\mathbf{i}$ (since a uniform grid is utilized in this work, the grid width is referenced as Δx_j for the remainder of the paper). Additionally,

$$\mathbf{x}_{(\mathbf{i} \pm \hat{\mathbf{i}}_j/2)} = \mathbf{x}_{(\mathbf{i})} \pm \frac{\Delta x_j}{2} \hat{\mathbf{x}}_j \quad (7)$$

130 is the coordinate vector a half-grid width forward/backward relative to $\mathbf{x}_{(i)}$
 131 in the x_j -direction (Fig. 1 displays Eqs. (5)-(7) schematically). From Eq.
 132 (3), p is assumed to be conserved over $\mathcal{C}_{(i)}$, thus the integral of p varies only
 133 due to flux across the boundaries of $\mathcal{C}_{(i)}$:

$$\frac{d}{dt} \int_{\mathcal{C}_{(i)}} p(\mathbf{x}_{(i)}, t) d\mathbf{x} = \sum_{j=1}^n \left(\int_{\mathcal{C}_{(\mathbf{i} \sim i_j)}^\pm} \left[f'_j(p(\mathbf{x}_{(i \pm i_j/2)}, t)) - f'_j(p(\mathbf{x}_{(i - i_j/2)}, t)) \right] d\mathbf{x}_{\sim j} \right), \quad (8)$$

134 where

$$\begin{aligned} \mathbf{i}_{\sim j} &= (i_1, \dots, i_{j-1}, i_{j+1}, \dots, i_n), \\ d\mathbf{x} &= \prod_{j=1}^n dx_j, \quad d\mathbf{x}_{\sim j} = \prod_{\substack{\ell=1 \\ \ell \neq j}}^n dx_\ell, \\ \mathcal{C}_{(\mathbf{i} \sim i_j)}^\pm &= [x_{j(i_j \pm 1/2)}] \times \prod_{\substack{\ell=1 \\ \ell \neq j}}^n [x_{\ell(i_\ell - 1/2)}, x_{\ell(i_\ell + 1/2)}]; \end{aligned}$$

135 $\mathcal{C}_{(\mathbf{i} \sim i_j)}^\pm$ is the $(n-1)$ -dimensional grid cell interface a half-step forward/backward
 136 in the x_j -direction at $x_j = x_{j(i_j \pm 1/2)}$ with normal vector pointing paral-
 137 lel/antiparallel to $\hat{x}_j$. Integrating Eq. (8) from $t^{(k)}$ to $t^{(k+1)}$ and dividing
 138 by the grid cell area leads to the fully discrete flux-differencing method

$$P_{(i)}^{(k+1)} = P_{(i)}^{(k)} - \sum_{j=1}^n \frac{\Delta t}{\Delta x_j} \left[F_{j(\mathbf{i} + i_j/2)}^{(k)} - F_{j(\mathbf{i} - i_j/2)}^{(k)} \right] \quad (9)$$

139 where $P_{(i)}^{(k+1)}$ is the discrete, updated probability cell average at grid cell $\mathcal{C}_{(i)}$,
 140 with

$$F_{j(\mathbf{i} \pm i_j/2)}^{(k)} \approx \frac{1}{\Delta t \prod_{\substack{\ell=1 \\ \ell \neq j}}^n \Delta x_\ell} \int_{t^{(k)}}^{t^{(k+1)}} \int_{\mathcal{C}_{(\mathbf{i} \sim i_j)}^\pm} f'_j(p(\mathbf{x}_{(i \pm i_j/2)}, t)) d\mathbf{x}_{\sim j} dt.$$

141 The numerical fluxes $F_j^{(k)}$ parallel to the $(n-1)$ -dimensional cell faces are
 142 approximated via a Godunov-type finite volume method known as Corner
 143 Transport Upwinding (CTU).

144 *2.2. The Corner Transport Upwind method*

145 First, the Donor Cell Upwind (DCU) method is used to calculate the
 146 first-order accurate numerical fluxes. For $\mathbf{i} \in \mathcal{I}$, where $\mathcal{I}$ represents the set
 147 of grid cell index vectors in the grid, in all directions $j = 1, \dots, n$, the upwind
 148 flux in each direction x_j at time step k is calculated:

$$F_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)} = f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^+ P_{(\mathbf{i}-\hat{\mathbf{i}}_j)}^{(k)} + f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^- P_{(\mathbf{i})}^{(k)}, \quad (10)$$

149 where

$$\begin{aligned} f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^+ &= \max(f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)}, 0), \\ f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^- &= \min(f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)}, 0), \\ f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)} &= f_j(\mathbf{x}_{(\mathbf{i}-\hat{\mathbf{i}}_j/2)}, t^{(k)}). \end{aligned}$$

150 Eq. (10) only considers probability flowing normal to the grid cell interface,
 151 but in general, probability may flow oblique to the interfaces of $\mathcal{C}_{(\mathbf{i})}$. To
 152 obtain second-order accuracy, we must account for this with flux corrections.

153 We consider $\mathcal{C}_{(\mathbf{i}-\hat{\mathbf{i}}_1)}^-$, the $(n-1)$ -dimensional grid cell interface a half-step
 154 backward in the x_1 -direction at $x_1 = x_{1(\mathbf{i}-\hat{\mathbf{i}}_1/2)}$. Generally, the direction of
 155 advection at this interface may not be perpendicular to this interface; thus,
 156 probability may propagate in any of the four ordinal directions depending
 157 on the signs of the components of $\mathbf{f}_{(\mathbf{i}-\hat{\mathbf{i}}_1/2)}^{(k)}$, as shown schematically in Fig.
 158 2. For second-order accuracy, this corresponds to $4(n-1)$ possible updates
 159 to neighboring fluxes in the $(n-1)$ -directions, excluding the x_1 -direction.
 160 Because the n -dimensional CTU method is notationally complex, we provide
 161 it in full in Algorithm 1.

162 The CTU method is still not second-order accurate, as it is missing high-
 163 resolution correction terms. For $\mathbf{i} \in \mathcal{I}$, in all directions x_j , the high-resolution
 164 correction terms are added to the (DCU + CTU)-calculated fluxes:

$$F_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)} \leftarrow F_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)} + \frac{1}{2} \left| f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)} \right| \left(1 - \frac{\Delta t}{\Delta x_j} \left| f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)} \right| \right) \frac{P_{(\mathbf{i})}^{(k)} - P_{(\mathbf{i}-\hat{\mathbf{i}}_j)}^{(k)}}{\Delta x_j} \phi(\theta_{(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)}) \quad (11)$$

165 where

$$\begin{aligned} \theta_{(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)} &= \begin{cases} (P_{(\mathbf{i}-\hat{\mathbf{i}}_j)}^{(k)} - P_{(\mathbf{i}-2\hat{\mathbf{i}}_j)}^{(k)}) / (P_{(\mathbf{i})}^{(k)} - P_{(\mathbf{i}-\hat{\mathbf{i}}_j)}^{(k)}) & \text{if } f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)} \geq 0, \\ (P_{(\mathbf{i}+\hat{\mathbf{i}}_j)}^{(k)} - P_{(\mathbf{i})}^{(k)}) / (P_{(\mathbf{i})}^{(k)} - P_{(\mathbf{i}-\hat{\mathbf{i}}_j)}^{(k)}) & \text{if } f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)} < 0, \end{cases} \\ \phi(\theta) &= \max\left(0, \min\left[(1+\theta)/2, 2, 2\theta\right]\right); \end{aligned}$$

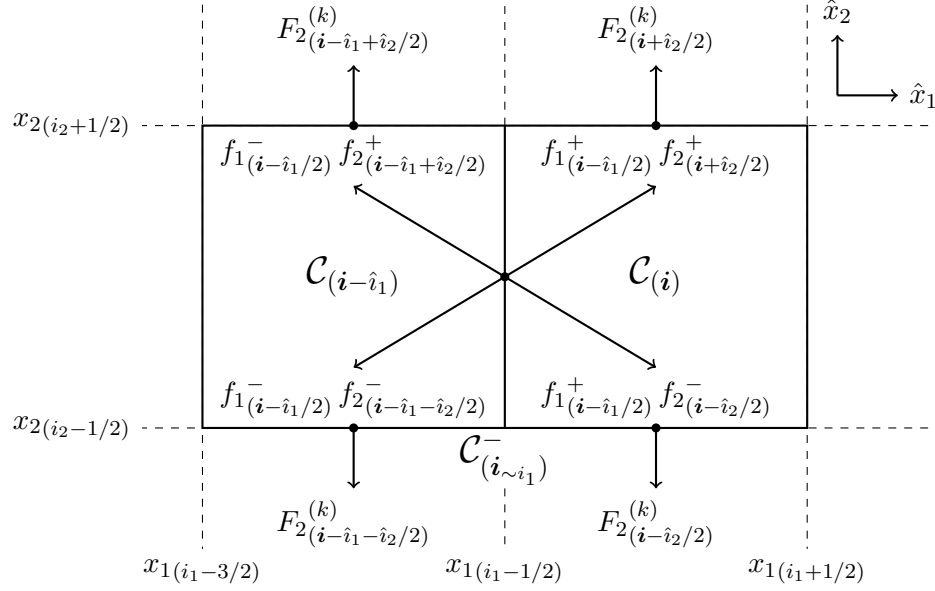


Figure 2: (x_1, x_2) -plane of an n -dimensional grid depicting the notation of the CTU method.

Algorithm 1 Corner Transport Upwind

Require: Perform DCU for $\mathbf{i} \in \mathcal{I}$

```

1: for  $\mathbf{i} \in \mathcal{I}$  do
2:   for  $j = 0$  to  $n$  do
3:      $F^* \leftarrow \frac{\Delta t}{2\Delta x_j} (P_{(\mathbf{i})}^{(k)} - P_{(\mathbf{i}-\hat{\mathbf{i}}_j)}^{(k)})$ 
4:     for  $\ell = 0$  to  $n$ ,  $\ell \neq j$  do
5:        $F_{\ell(\mathbf{i}+\hat{\mathbf{i}}_\ell/2)}^{(k)} \leftarrow F_{\ell(\mathbf{i}+\hat{\mathbf{i}}_\ell/2)}^{(k)} - f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^+ f_{\ell(\mathbf{i}+\hat{\mathbf{i}}_\ell/2)}^+ F^*$ 
6:        $F_{\ell(\mathbf{i}-\hat{\mathbf{i}}_\ell/2)}^{(k)} \leftarrow F_{\ell(\mathbf{i}-\hat{\mathbf{i}}_\ell/2)}^{(k)} - f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^+ f_{\ell(\mathbf{i}-\hat{\mathbf{i}}_\ell/2)}^- F^*$ 
7:        $F_{\ell(\mathbf{i}-\hat{\mathbf{i}}_j+\hat{\mathbf{i}}_\ell/2)}^{(k)} \leftarrow F_{\ell(\mathbf{i}-\hat{\mathbf{i}}_j+\hat{\mathbf{i}}_\ell/2)}^{(k)} - f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^- f_{\ell(\mathbf{i}-\hat{\mathbf{i}}_j+\hat{\mathbf{i}}_\ell/2)}^+ F^*$ 
8:        $F_{\ell(\mathbf{i}-\hat{\mathbf{i}}_j-\hat{\mathbf{i}}_\ell/2)}^{(k)} \leftarrow F_{\ell(\mathbf{i}-\hat{\mathbf{i}}_j-\hat{\mathbf{i}}_\ell/2)}^{(k)} - f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^- f_{\ell(\mathbf{i}-\hat{\mathbf{i}}_j-\hat{\mathbf{i}}_\ell/2)}^- F^*$ 
9:     end for
10:   end for
11: end for

```

166 $\phi(\theta)$ is a monotonized central difference limiter used for representing proba-
167 bility discontinuities with sharper resolution [17]. After the fluxes have been

168 calculated and plugged into Eq. (9), the discretized PDF is normalized

$$P_{(i)}^{(k+1)} = \left(\sum_{i \in \mathcal{I}} P_{(i)}^{(k+1)} \right)^{-1} P_{(i)}^{(k+1)}. \quad (12)$$

169 Together, Eq. (10), Algorithm 1, and Eqs. (11) and (12) form a second-
 170 order accurate finite volume method employed by GBEES for numerically
 171 time-marching a discretized, n -dimensional PDF.

172 3. Advancements made to the CPU implementation

173 GBEES excels where other finite volume methods fail by dynamically
 174 allocating grid cells where probability is above some threshold. However, the
 175 legacy algorithm architecture contains structures and subprocesses that are
 176 ripe for optimization. Prior to detailing the GPU implementation, we discuss
 177 the efficiency improvements made to the CPU implementation.

178 3.1. *Dynamic grid stored in hashtable*

179 The legacy implementation of GBEES stores the dynamic grid in a nested
 180 list data structure. Many functions within GBEES require a searching pro-
 181 cedure to check if a given $\mathcal{C}_{(i)}$ exists in the grid. The time complexity of
 182 searching a nested list is $\mathcal{O}(n^2)$, which will result in computational bottle-
 183 necks for high-dimensional systems. The first attempt to address this issue
 184 employed a binary search tree [18], but overhead of the conversion from grid
 185 cell index vector $\mathbf{i}$ to unique, positive key value proved too large. Instead,
 186 a hashtable was utilized, as the structure allows for collisions between map-
 187 pings, thus removing the overhead from ensuring bijectivity. Additionally,
 188 the time complexity of search for a hashtable is $\mathcal{O}(1)$. Hashtables are dis-
 189 cussed further in Section 4.2.

190 3.2. *CFL-minimized adaptive time-marching*

191 For finite volume methods, the Courant–Friedrichs–Lewy (CFL) condi-
 192 tion [19]

$$\Delta t \left(\sum_{j=1}^n \frac{f_j}{\Delta x_j} \right) \leq C_{\max}$$

193 must be satisfied in order for the method to be stable, where $C_{\max} = 1$ for
 194 explicit methods (as GBEES is an explicit finite volume method, we assume

195 $C_{\max} = 1$ for the remainder of the paper). In general, because the advection
 196 is variable across spatial coordinates, Δt must be small enough such that this
 197 condition holds for $\mathbf{i} \in \mathcal{I}$:

$$\Delta t \leq \left(\sum_{j=1}^n \frac{f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}}{\Delta x_j} \right)^{-1}. \quad (13)$$

198 The legacy implementation of GBEES employed a static, over-conservative
 199 time step to ensure stability for the entire propagation period. The new
 200 implementation of GBEES finds the minimum allowable Δt such that Eq.
 201 (13) is true for $\mathbf{i} \in \mathcal{I}$ at time step k :

$$\Delta t^{(k)} = \min_{\mathbf{i} \in \mathcal{I}} \left[\left(\sum_{j=1}^n \frac{f_{j(\mathbf{i}-\hat{\mathbf{i}}_j/2)}^{(k)}}{\Delta x_j} \right)^{-1} \right]. \quad (14)$$

202 Because of the dynamic nature of the grid, $\Delta t^{(k)}$ may change depending
 203 on where in phase space probability is focused. Implementing Eq. (14)
 204 maximizes the time step size while ensuring the stability of the explicit finite
 205 volume method.

206 3.3. Directional growing and pruning

207 To exploit the sparsity of an n -dimensional PDF over phase space, grid
 208 cells are tracked where probability is above some threshold p^* . To ensure
 209 probability is not lost during time-marching, grid cells neighboring those
 210 above threshold are also tracked. In the legacy implementation of GBEES,
 211 during the growing procedure, the algorithm loops through all existing grid
 212 cells and checks if any of the $3^n - 1$ neighbors that do not exist must be
 213 inserted, regardless if probability is likely to flow into the new grid cell in the
 214 following step. This can create irrelevant grid cells that are deleted in future
 215 steps without ever increasing in probability. In the new GBEES implemen-
 216 tation, the direction of the advection is used to determine if a neighboring
 217 grid cell is required for the next time step. As is demonstrated in Fig. 3,
 218 only the *downwind* grid cells are created; this process saves on the number
 219 of cells that are inserted in each growth step.

220 Similarly, during the pruning procedure, the algorithm loops through all
 221 existing grid cells, looking for those that are below threshold p^* . In the
 222 legacy GBEES implementation, before deleting the negligible cell, the algo-
 223 rithm checks each of the $3^n - 1$ neighbors to see if any are above p^* . Again,

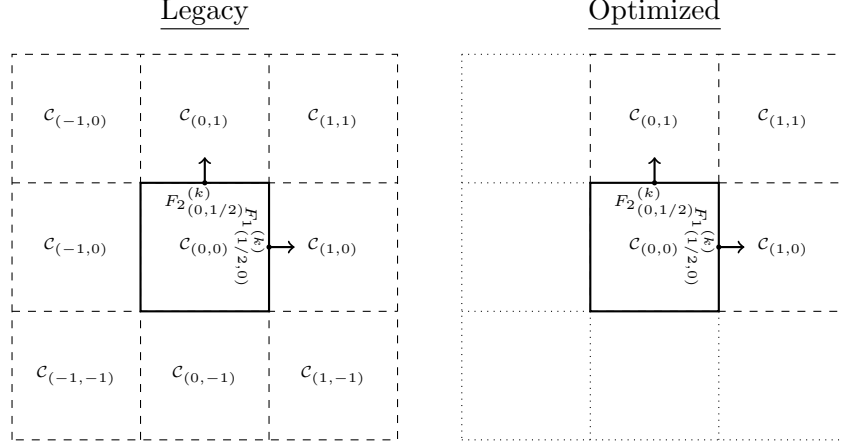


Figure 3: 2D schematic demonstrating the differences in the growing procedure for the legacy and optimized implementations of GBEES. Solid border cells are those with probability above threshold, dashed border cells are those set to be created during the growing procedure, and dotted border cells are neglected.

224 this results in redundant cells being saved, as even if a neighboring cell is
 225 above threshold, it does not necessarily mean that in the following time steps,
 226 it will flow probability into the considered cell. Instead, the new GBEES im-
 227 plementation takes a directional-approach to the growth procedure, wherein
 228 only the neighboring grid cells that are *upwind* are checked for probability
 229 above p^* . Fig. 4 shows that this requires the algorithm to check a fraction
 230 of the total neighbor cells, while ensuring that negligible cells are not saved,
 231 again contribution to the efficiency of the new algorithm.

232 4. CUDA Implementation

233 The CUDA implementation builds upon the enhancements made to the
 234 CPU version described in the previous section. This section describes the
 235 specifics of the CUDA architecture [20] and the optimization strategies em-
 236 ployed.

237 The CUDA architecture is designed for parallel processing, with GPUs
 238 containing multiple Streaming Multiprocessors (SMs). Each SM can execute
 239 groups of threads organized into blocks, which are further divided into smaller
 240 units of 32 threads called warps. Warps execute instructions in a single-
 241 instruction, multiple-thread (SIMT) fashion, meaning all threads in a warp
 242 perform the same instruction simultaneously on different data.

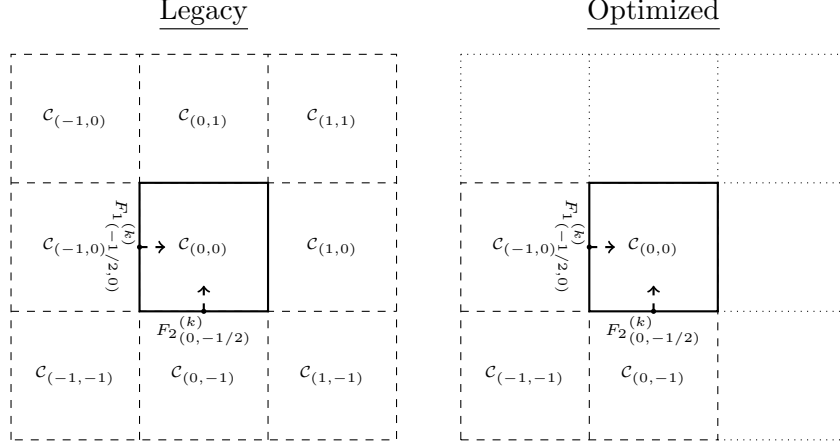


Figure 4: 2D schematic demonstrating the differences in the pruning procedure for the legacy and optimized implementations of GBEES. Solid border cells are those with probability below threshold, dashed border cells are those checked during the pruning procedure, and dotted border cells are neglected.

243 CUDA GPUs also feature a layered memory hierarchy. Global (or de-
 244 vice) memory provides large storage capacity but has higher latency. Each
 245 SM has faster, limited shared memory that is accessible only to threads
 246 within a block, enabling efficient data sharing. Additionally, each thread has
 247 its own private registers, the fastest type of memory, used for intermediate
 248 calculations.

249 As described in Section 2, the GBEES method requires a dynamic grid
 250 in which cells are added or removed throughout the integration steps. This
 251 dynamic nature prevents establishing a fixed mapping between execution
 252 threads and cells. In traditional finite volume software using a static grid,
 253 the grid is partitioned into subdomains —each one including also a bound-
 254 ary with additional halo cells— and these subdomains are then assigned to
 255 thread blocks on the GPU [21]. However, with a dynamic grid, a flexible
 256 assignment of cells to threads is required, along with additional synchroniza-
 257 tion, as thread blocks can no longer operate independently within isolated
 258 subdomains. This dynamic structure impacts all aspects of the GPU imple-
 259 mentation, necessitating more complex synchronization mechanisms and a
 260 highly efficient memory layout for storing the grid.

261 This extra global-level synchronization required by the dynamic grid is
 262 provided by the Cooperative Kernel abstraction in CUDA [22]. A Coop-

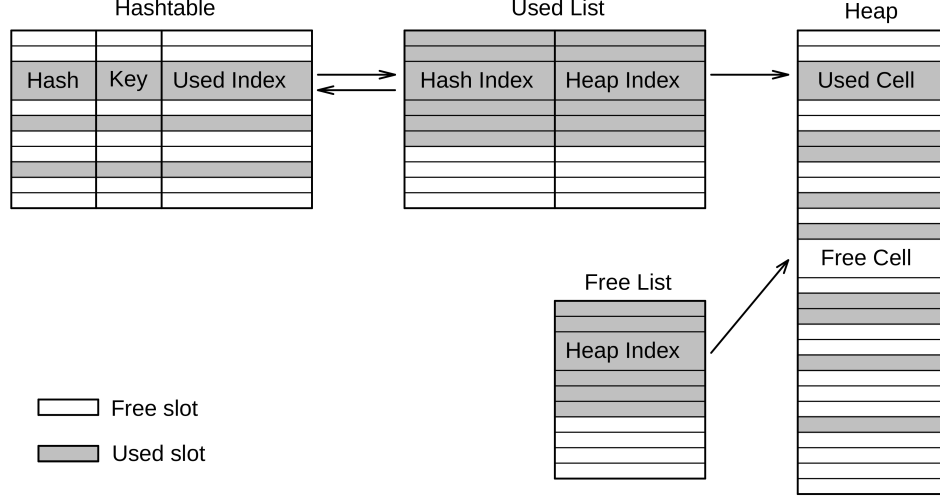


Figure 5: Grid data structure. Notice that the Used List and the Free List are maintained compact, with all the used slots at the beginning and all the free slots at the end.

erative Kernel requires all threads to be active concurrently, enabling the establishment of global synchronization barriers. This means the maximum number of threads equals the GPU device’s maximum simultaneous threads. If the grid contains more cells than this limit, each thread must process multiple cells sequentially.

4.1. Grid data structure

Given the synchronization requirements, the dynamic thread-cell assignments, and the need for each thread to process a variable number of cells, we propose the memory structure for the grid depicted in Fig. 5.

The data structure consists of a hashtable, a list to track used cells, another list for unused cells, and a heap table for cell storage.

The hashtable provides fast access to a cell by its key —the state coordinates $\mathbf{i} = (i_1, \dots, i_n)$. This random access is crucial when accessing the neighbor cells in each dimension.

The list of used cells in the grid serves two main purposes. First, it enables quick access to used cells for functions that process individual cells. More importantly, it allows an even distribution of workload among active threads.

The Free List maintains a record of not used cells in the heap, reducing the time needed to locate a free slot.

283 Finally, the heap stores the cells themselves. In CUDA, it is not possible
 284 to allocate memory directly within a device function in a kernel. Therefore,
 285 and for performance reasons, all memory structures are of fixed size, which,
 286 for a given configuration, sets the maximum number of cells in the grid.

287 The relationships among these memory structures are depicted in Fig. 5.
 288 In addition to the depicted relationships, each cell contains pointers to its
 289 neighboring cells in each dimension by directly storing their indexes of the
 290 Used List.

291 *4.2. Hashtable*

292 Because the maximum number of grid cells is fixed, we can set the
 293 hashtable size to ensure a maximum occupancy level. The hashtable size
 294 is configurable in the software as a multiple of the grid’s maximum size,
 295 with a default setting of twice that size. This default configuration ensures
 296 a maximum occupancy factor $\alpha = 0.5$. This bounded occupancy allows us
 297 to use a simple open-addressing scheme with linear probing [23]. With a
 298 well-randomized hash function, the expected number of probes during an
 299 unsuccessful search is [24]

$$(1 + 1/(1 - \alpha)^2)/2$$

300 and for a successful search

$$(1 + 1/(1 - \alpha))/2.$$

301 The open addressing scheme requires marking elements as deleted [23]. In
 302 the GBEEES method, all cell deletions occur during the prune grid operation.
 303 Since the prune operation is executed only for a subset of integration steps,
 304 rehashing the hashtable after each grid prune has minimal impact on per-
 305 formance while ensuring that the maximum occupancy level is maintained.
 306 Moreover, in the CUDA implementation, this rehashing is fully parallelized,
 307 with all execution threads sharing the workload to rehash the hashtable en-
 308 tries concurrently.

309 Finally, achieving the expected efficiency requires a well-randomized hash
 310 function. Fig. 6 shows collision graphs for different hash functions, with each
 311 plot representing the number of collisions per entry in a 2D projection for the
 312 3D Lorenz ’63 model—one of the test cases included in this study. Based
 313 on these patterns, the BuzHash algorithm, also known as hashing by cyclic
 314 polynomial, was selected and implemented as described in [25].

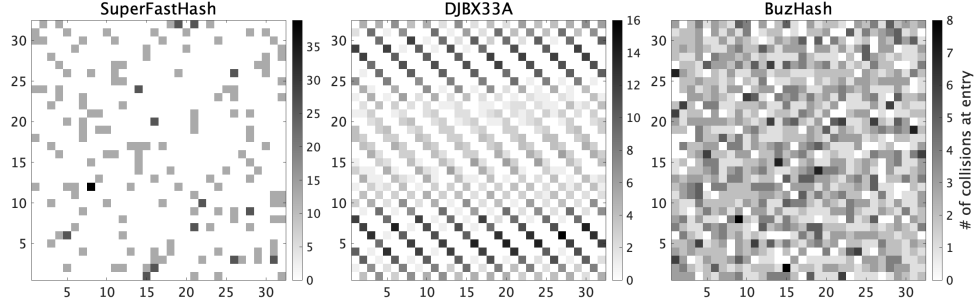


Figure 6: Number of collisions for different hash functions represented as 2D projections, where the capacity of the hashtable is set to $1024 = 32 \times 32$. For these examples, the hashtable stores the initial grid from the Lorenz '63 example, explained further in Section 5.1.1.

315 4.3. Main code blocks

316 From an implementation perspective, we can divide the main code blocks
 317 into operations that act on individual cells and those that act on the grid
 318 as a whole. The first group is straightforward to implement, as each thread
 319 modifies only its assigned cells without significant synchronization issues,
 320 requiring only quick access to the cell and its neighbors. This rapid access
 321 is achieved through the Used List. This category includes operations such
 322 as cell initialization, updating references to the neighbor cells, updating the
 323 time step based on the CFL condition —Eq. (14)—, computing the DCU
 324 and CTU —Eq. (10), Algorithm (1), and Eq. (11)—, probability distribution
 325 normalization —Eq. (12)—, and applying new measurements —Eq. (4)—.

326 The second category involves grid-wide operations, specifically the grid
 327 growth and grid pruning. These operations modifies a shared global re-
 328 source —the grid— and therefore require careful synchronization. To opti-
 329 mize CUDA performance, all synchronization is managed using atomic op-
 330 erations and synchronization barriers, either at the block or device level.

331 The following sections detail the key synchronization aspects and parallel
 332 techniques applied in these code blocks.

333 4.4. Synchronization aspects

334 4.4.1. Grow grid operation

335 To maximize efficiency in the grid growth operation, a concurrent cell
 336 insertion method is required to avoid blocking threads during simultaneous
 337 cell creations. Additionally, when exploring different dimensions in the phase

338 space, the algorithm frequently attempts to create cells that already exist.
 339 Algorithm 2 presents the chosen compromise solution, which ensures correct
 340 synchronization without thread blocking by utilizing atomic operations.

341 This implementation delays the complete initialization of the cell —using
 342 a callback function— until it is confirmed that the cell does not already exist
 343 in the grid, thereby improving performance.

344 However, the selected approach has a trade-off: it can only check the
 345 existence of a new cell against the previous state of the grid and cannot
 346 guarantee successful checking with the other concurrent insertions. To ad-
 347 dress this limitation, the grid growth operation adopts a staged, directional
 348 cell growth strategy. Specifically:

- 349 • Growth is first performed along the forward axis of all dimensions. A
 350 global synchronization barrier is then executed.
- 351 • Growth is subsequently performed along the backward axis of all di-
 352 mensions, followed by another global synchronization step.
- 353 • Finally, edge growth is carried out in a similar staged manner in the four
 354 diagonal directions —forward-forward, forward-backward, backward-
 355 forward, and backward-backward.

356 This staged approach ensures that no concurrent thread attempts to insert
 357 the same cell at the same time.

358 4.4.2. *Prune grid operation*

359 The prune grid operation, outlined in Algorithm 3, is less performance-
 360 critical as it is executed only once every several integration steps. However,
 361 it requires specialized techniques to be performed in parallel by all threads.

362 The operation begins by marking cells whose probability values fall below
 363 a specified threshold, identifying them as candidates for pruning. The next
 364 step involves performing a parallel prefix sum operation [26, 27] to compact
 365 the Used List. The prefix sum, also known as scan, computes cumulative
 366 sums over a list to facilitate parallel data compaction. This scan is carried
 367 out by the active threads, as detailed in the following section on specific
 368 parallel techniques.

369 After the scan, the Used List is compacted using a double-buffer scheme,
 370 and the freed slots are added to the Free List via atomic operations.

371 Finally, the Hashtable is rehashed, also employing a double-buffer scheme
 372 and distributing the rehashing workload across all active threads.

Algorithm 2 Concurrent Cell Creation

Require: The Used List and the Free List are compact.

- 1: Obtain the hashtable slot by applying the hash function to the cell key.
 - 2: **loop**
 - 3: Check if the hashtable slot is free using an atomicCAS() operation. If it is free, reserve it with the RESERVED flag; if not, obtain the current Used Index.
 - 4: **if** The existing cell is the same **then**
 - 5: break
 - 6: **end if**
 - 7: **if** An empty slot is reached **then**
 - 8: Reserve a Used List slot using an atomicAdd() operation on the list size.
 - 9: Obtain a free heap slot from the Free List using an atomicDec() operation on the list size.
 - 10: Update the Hashtable and Used List contents.
 - 11: Update the heap content and complete cell initialization with a callback function.
 - 12: **end if**
 - 13: Move to the next hashtable slot.
 - 14: **end loop**
-

Algorithm 3 Grid Prune Operation

- 1: Mark the negligible cells in the heap.
- 2: Perform a prefix sum process of the Used List in shared memory.
- 3: Complete the prefix sum of the Used List in global memory.
- 4: Compact the Used List and update the Free List.
- 5: Rehash the Hashtable.

Ensure: To perform a global synchronization at the end of each step.

373 *4.5. Specific parallel techniques*

374 The GBEEES-GPU implementation employs two high-level parallel tech-
375 niques: parallel reduction and parallel scan. Parallel reduction is utilized
376 to compute the sum of grid cell probabilities for normalizing the distribu-
377 tion following Eq. (12). Parallel scan is applied during the prune operation
378 to compact the Used List. These techniques are widely recognized as stan-
379 dard methods [27], and only a brief description is provided here, focusing

380 on their adaptation to the GBEES kernel’s context, which involves multiple
381 concurrent blocks and threads processing several cells each.

382 In the case of parallel reduction, each thread begins by summing the
383 probability value of all the cells assigned to it. Next, a parallel reduction
384 is performed within each thread block, utilizing shared memory. This intra-
385 block reduction employs a sequential addressing scheme to obtain an opti-
386 mal shared memory access. Once the reduction within shared memory is
387 complete, a global reduction is performed, involving the first thread of each
388 block. Unlike the intra-block reduction, which uses thread synchronization,
389 the outer reduction relies on global barriers. The final result of the reduction
390 is the sum of the probabilities of all cells.

391 For the scan operation required to compact the Used List, the process
392 begins with a per-block scan using a double buffer in shared memory. Specif-
393 ically, an inclusive scan with sequential addressing is employed. Unlike par-
394 allel reduction, it is not possible to pre-accumulate the values of all cells
395 processed by each thread. Instead, multiple intra-block scans are performed
396 within each block, with the sums orderly accumulated into a global array.
397 Following this, a second outer scan is conducted at the global level by the
398 first thread of each block. Once the corresponding prefix sums are obtained,
399 each thread populates the compacted Used List in parallel and updates the
400 Free List to account for unused or deleted cells.

401 5. Validation

402 5.1. Use cases

403 In order to validate the GBEES implementation we consider the following
404 use cases

- 405 • The Lorenz ’63 model (three-dimensional)
- 406 • The Lorenz ’96 model (six-dimensional)

407 5.1.1. Lorenz ’63

408 The Lorenz ’63 model, colloquially referred to as the Butterfly Effect, is
409 often employed to validate the accuracy of uncertainty propagation meth-
410 ods because of the highly non-Gaussian behavior exhibited [28]. The three-

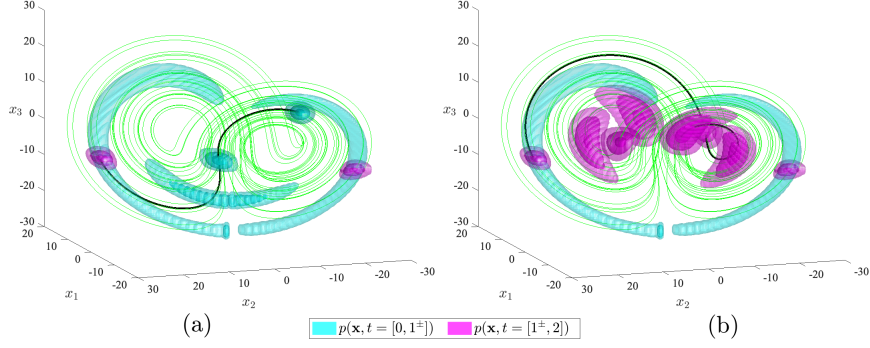


Figure 7: PDF isosurfaces governed by the Lorenz '63 model in (x_1, x_2, x_3) -space at $p = 0.607$, $p = 0.135$, and $p = 0.011$ with $\Delta x_j = 0.5$ for $j = 1, 2, 3$. On the left (a), the isosurfaces are at $t = 0$, $t = 1/3$, $t = 2/3$, and $t = 1^\pm$ and on the right (b), the isosurfaces are at $t = 1^\pm$, $t = 4/3$, $t = 5/3$ and $t = 2$.

dimensional state and equations of motion are defined as

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}, \quad \frac{d\mathbf{x}}{dt} = \mathbf{f}(\mathbf{x}) = \begin{bmatrix} \sigma(x_2 - x_1) \\ -x_2 - x_1x_3 \\ -bx_3 + x_1x_2 - br \end{bmatrix},$$

where $(\sigma, b, r) = (4, 1, 48)$ results in the system being chaotic. In Fig. 7, a 3D Gaussian PDF is initialized at $\mathbf{x}^{(0)} = (-11.5, -10, 9.5)$ with standard deviation $\sigma_{x_j} = 1$ for $j = 1, 2, 3$. The uncertainty is then continuous-time propagated using GBEES until $t^{(1)} = 1$, where a discrete measurement update is performed with measurement $y^{(1)} = -8$, where the measurement model is

$$y = h(\mathbf{x}) = x_3,$$

and the measurement uncertainty $\sigma_y = 1$. The uncertainty is then continuous-time propagated till $t = 2$, when the simulation ends. Fig. 7 illustrates the rapid evolution of the PDF from Gaussian to highly non-Gaussian, with the PDF naturally bifurcating at $t = 1$. Since this model was also used to validate the legacy GBEES implementation, it serves as a valuable basis for performance comparison, as discussed further in Section 6.

5.1.2. Lorenz '96

As an analog to the Lorenz '63 validation first performed in [16], the CPU-optimized and GPU implementations of GBEES are validated on the Lorenz

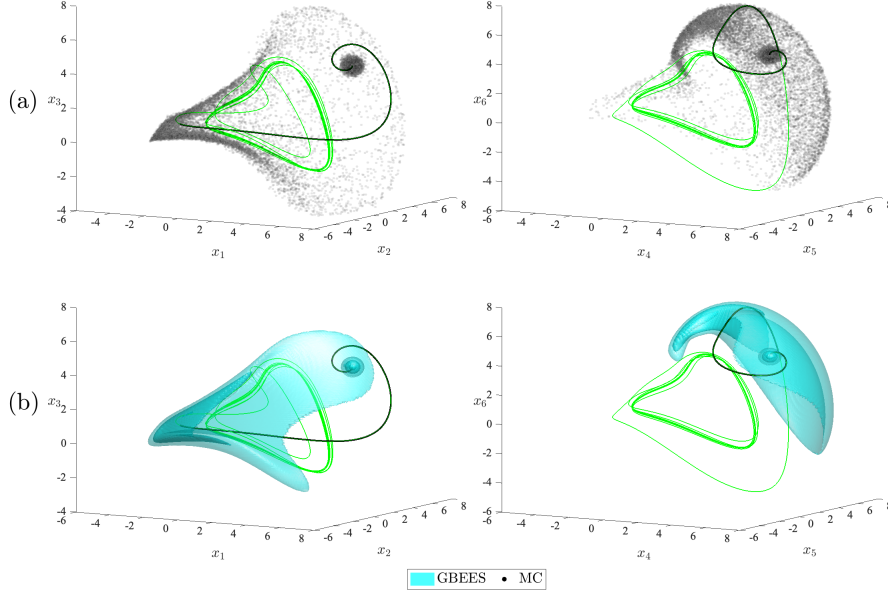


Figure 8: PDF representations in (x_1, x_2, x_3) - and (x_4, x_5, x_6) -spaces governed by the Lorenz '96 model with $\Delta x_j = 0.01$ for $j = 1, \dots, 6$ and $F = 4$. On top (a), the MC point-mass distributions with 10,000 particles at $t = 0$ and $t = 1.3$ and on bottom (b), the GBEES isosurfaces at $p = 0.607$, $p = 0.135$, and $p = 0.011$ for $t = 0$ and $t = 1.3$.

427 '96 model, a generalized dynamical system that exhibits chaotic behavior
 428 [29]. The n -dimensional state and equations of motion are defined as

$$\mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_j \\ \vdots \\ x_n \end{bmatrix}, \quad \frac{d\mathbf{x}}{dt} = \mathbf{f}(\mathbf{x}) = \begin{bmatrix} (x_2 - x_{n-1})x_n - x_1 + F \\ \vdots \\ (x_{j+1} - x_{j-2})x_{j-1} - x_j + F \\ \vdots \\ (x_1 - x_{n-2})x_{n-1} - x_n + F \end{bmatrix},$$

429 where $(F, \dots, F)$ is an unstable equilibrium, with F being a forcing constant.
 430 A 6D Gaussian PDF is initialized at $\mathbf{x}^{(0)} = (F + 0.5, F, \dots, F)$ where $F = 4$,
 431 with standard deviation $\sigma_{x_j} = 0.02$ for $j = 1, \dots, 6$. The uncertainty is then
 432 propagated using GBEES until $t = 1.3$. No measurement update is per-
 433 formed in this simulation. To validate the accuracy of GBEES qualitatively,
 434 an MC simulation with identical initial conditions is plotted, depicted in Fig.

435 8(a). The 3D PDFs in the (x_1, x_2, x_3) - and (x_4, x_5, x_6) -spaces, shown in Fig.
 436 8(b), are calculated from the discretized 6D PDF by numerically integrating
 437 over the (x_4, x_5, x_6) - and (x_1, x_2, x_3) -spaces, respectively:

$$p(x_1, x_2, x_3, t) = \int_{\min(x_6)}^{\max(x_6)} \int_{\min(x_5)}^{\max(x_5)} \int_{\min(x_4)}^{\max(x_4)} p(\mathbf{x}, t) dx_4 dx_5 dx_6,$$

$$p(x_4, x_5, x_6, t) = \int_{\min(x_3)}^{\max(x_3)} \int_{\min(x_2)}^{\max(x_2)} \int_{\min(x_1)}^{\max(x_1)} p(\mathbf{x}, t) dx_1 dx_2 dx_3.$$

438 5.2. Convergence and accuracy

439 In addition to the qualitative validation against the MC simulation shown
 440 in Figs. 7 and 8, a quantitative validation is conducted by comparing the
 441 results to a propagation with a very small step size, which serves as the
 442 reference “true” solution for this test.

443 This test has two primary objectives. First, to ensure the GBEES imple-
 444 mentation converge correctly, as demonstrated by decreasing errors relative
 445 to the reference propagation when step sizes are reduced. Second, to assess
 446 the impact of the non-deterministic execution order of operations in CUDA.

447 In the CUDA architecture, the execution order of different threads is non-
 448 deterministic [22]. Due to the non-associativity of floating-point arithmetic
 449 caused by the rounding errors, small variations in the computed probability
 450 distribution are expected.

451 The errors in the probability distribution are assessed using two metrics.
 452 The first metric measures the sum of the absolute differences between the
 453 probability values in each cell. Since the mesh is equispaced, there is no need
 454 to assign different weights to the cells. Additionally, because the total sum
 455 of the probabilities in the cells is 1, this metric is already normalized. The
 456 error for this metric is calculated as follows:

$$E(t) = \sum_{i \in \mathcal{I}} |p(\mathbf{x}_{(i)}, t) - p^0(\mathbf{x}_{(i)}, t)|,$$

457 where $p^0(\mathbf{x}_{(i)}, t)$ represents the discrete probability distribution obtained
 458 from the reference propagation.

459 The second metric is the Bhattacharyya Coefficient [30] defined as

$$BC(t) = \sum_{i \in \mathcal{I}} \sqrt{p(\mathbf{x}_{(i)}, t) p^0(\mathbf{x}_{(i)}, t)},$$

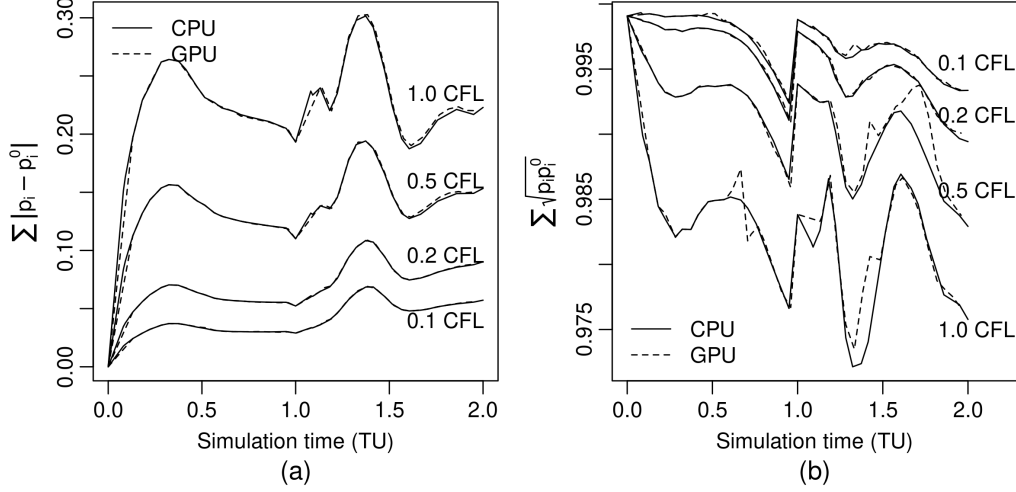


Figure 9: Comparison of the discrete probability distributions over simulated time with a reference propagation of step size $0.01 \times \text{CFL}$. On the left (a), the comparison is based on the absolute probability differences between cells $E(t)$. On the right (b), the comparison utilizes the Bhattacharyya Coefficient $BC(t)$.

that measures the similarity between two probability distributions with a value of 1 indicating that the distributions coincide and 0 when the distributions are entirely distinct.

Using these metrics, the reference propagation of a variable step size of $0.01 \times \text{CFL}$ was compared to CPU and GPU propagations with step sizes of 1.0, 0.5, 0.2, and 0.1 times CFL.

The results of these comparisons are shown in Fig. 9. The curves demonstrate proper convergence, with progressively more accurate values as the step size is reduced. Additionally, the differences caused by the non-deterministic execution order in CUDA are minimal compared to the influence of other error sources in the simulation, such as step size variation.

6. Performance

Performance improvements for both the new CPU and CUDA versions are evaluated using the same validation cases described in Section 5. Their computational effort is summarized in Table 1. The significant difference in computational load between the two cases is primarily due to the dimensionality of the models: the Lorenz '63 is formulated over a three-dimensional

	Lorenz '63	Lorenz '96
Maximum grid size	$\approx 24k$	$\approx 50M$
Integration steps	≈ 960	≈ 1050
Total cell computations	$\approx 6.5M$	$\approx 30G$

Table 1: Computational burden for the Lorenz '63 and Lorenz '96 models with a variable step size corresponding to a $1.0 \times CFL$ and a cell size equal to half the standard deviation in each dimension of the first measurement.

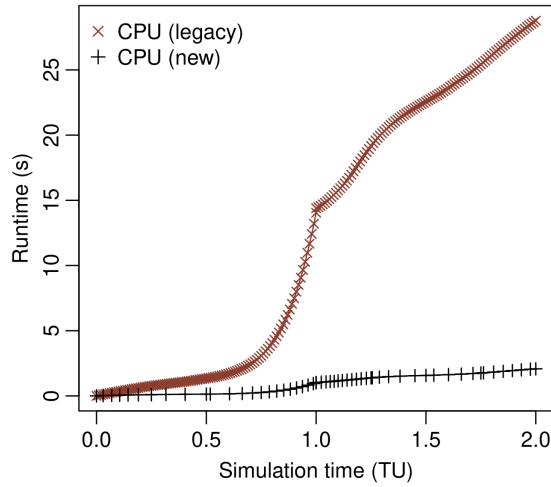


Figure 10: Runtime comparison between the legacy and the new CPU version for the Lorenz '63 model.

477 phase space, while the Lorenz '96 uses a six-dimensional one.

478 Fig. 10 shows the runtime comparison between the legacy and the new
479 CPU versions for the Lorenz '63 model. The legacy version performs a fixed-
480 step integration, while the new CPU version uses a variable-step scheme.
481 To ensure a fair comparison, the fixed step size of the legacy CPU version
482 was adjusted so that, during the integration, the Bhattacharyya coefficients,
483 computed using the same procedure as in Section 5, reach a similar maximum
484 value.

485 The runtime results of these executions are also included in Table 2. The
486 relative speed-up of the new CPU version relative to the legacy version is
487 approximately 13.85 times faster (calculated as $1/0.072$).

488 To assess performance in the CUDA version, it is essential first to outline
489 the launch configuration parameters and explain how these settings influence

490 overall performance.

491 The GPU launch configuration is determined by three key parameters:
 492 the number of blocks, the number of threads per block, and the number
 493 of cells each thread processes. The objective of the launch configuration
 494 is to maximize the GPU occupancy. Since we have a Cooperative Kernel,
 495 this is achieved by launching a total number of threads equal to the GPU’s
 496 maximum simultaneous thread capacity.

Device	Runtime (ms)	Cells/s	Speed-up
CPU-legacy: Apple M2 MAX	28777	$\approx 0.54\text{M/s}$	0.072
CPU-optimized: Apple M2 MAX	2077	$\approx 3.13\text{M/s}$	1
GPU 1: Tesla V100	244	$\approx 26.6\text{M/s}$	8.5
GPU 2: NVIDIA A100	258	$\approx 25.2\text{M/s}$	8.1
GPU 3: NVIDIA H100	226	$\approx 28.8\text{M/s}$	9.2
GPU 4: NVIDIA H200	230	$\approx 28.3\text{M/s}$	9.0

Table 2: Total runtime, number of cells processed per second, and relative speed-up compared to a single-core CPU running the optimized version of GBEES for the Lorenz ’63 model.

497 If the maximum grid size exceeds this capacity, each thread must process
 498 multiple cells, requiring the parameter for cells processed per thread to be set
 499 to a value greater than one. Conversely, if the maximum grid size is smaller
 500 than this capacity, the configuration should launch only as many threads as
 501 the grid requires, with each thread processing only one cell. In this last case,
 502 the achieved occupancy will be less than the theoretical maximum because
 503 the model does not expose sufficient parallelism to fully utilize the GPU.

504 Therefore, the launch configuration strategy is to keep the number of
 505 cells processed by each thread as low as possible. If the model is sufficiently
 506 large, the product of the number of blocks and the number of threads per
 507 block should equal the GPU’s maximum thread capacity. This product can
 508 be achieved through various combinations of blocks and threads per block.

509 This balance between blocks and threads per block is very subtle. Setting
 510 the number of threads per block to the maximum (1024 in the current CUDA
 511 architectures) benefits the parallel reduction and scan processes described
 512 in Section 4.5, as more computation is performed at the block level using
 513 shared memory. However, the performance differences are minimal, and in
 514 some tests, using fewer threads per block than the maximum has resulted in
 515 slightly better runtimes.

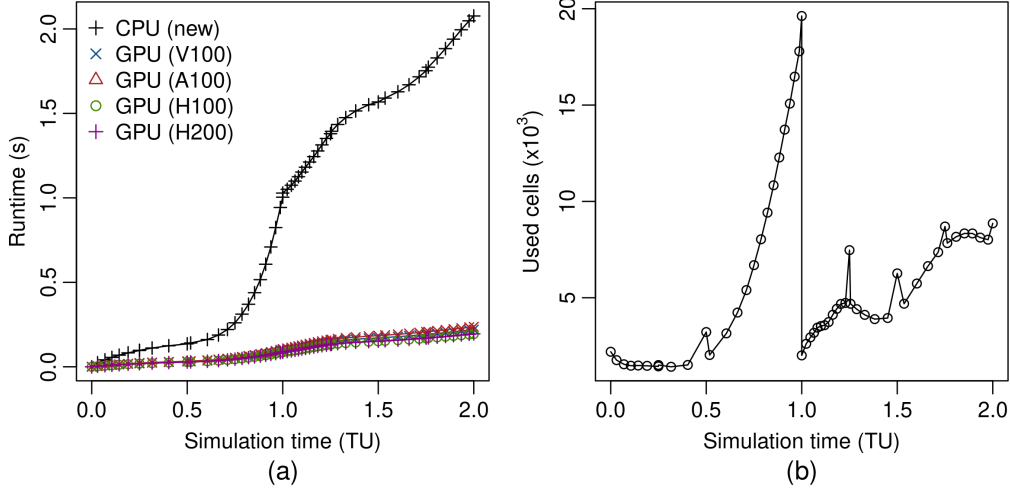


Figure 11: Runtime comparison between the new CPU version and GBEES-GPU on various GPU devices (a) and number of used cells (b) for the Lorenz '63 model.

Fig. 11 represents, for the Lorenz '63 model, the runtime comparison between the new CPU version and the CUDA implementation running on different GPU devices. Fig. 11(a) represents the program runtime as a function of the simulated time, where steeper regions correspond to moments when the grid contains more cells. The number of cells during the simulation is plotted in Fig. 11(b). The drop in the number of cells at $t = 1.0$ TU corresponds to a discrete measurement update.

The graph in Fig. 11(a) shows a significant performance boost from parallelizing and executing the algorithm on the GPU. Table 2 summarizes the total runtime, the number of processed cells per second, and the speed-up values. The Lorenz '63 model represents a case where the grid size is not large enough to achieve maximum occupancy on the tested GPUs. This limitation causes the performance to be similar across all devices. Despite this, the speed-up achieved by using the CUDA version ranges from 8.5 to 9.0, depending on the specific GPU tested.

The Lorenz '96 model requires a high computational burden and exposes enough parallelism to fully utilize the performance of the tested GPUs. This results in a significant performance difference between the CPU and CUDA versions. To highlight this difference and facilitate representation, Fig. 12 first presents the runtime comparison between the new CPU version and the

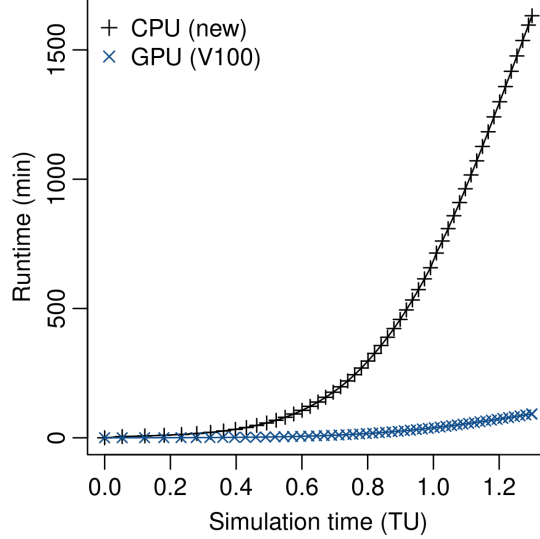


Figure 12: Runtime comparison between the new CPU version and the GBEES-GPU execution on a V100 GPU device for the Lorenz '96 model.

536 CUDA implementation running on a V100 GPU device, which is the slowest
 537 among the tested GPUs.

538 The runtime data for this comparison, along with comparisons to other
 539 GPU devices, are included in Table 3. The execution of the Lorenz '96 model
 540 is 17.8 times faster in the CUDA version on the V100 device compared to
 541 the new CPU version.

542 For the other devices, Fig. 13 shows that, as the model fully utilizes the
 543 GPUs, there is a progressive reduction in execution times corresponding to
 544 the increasing computing power of the different test devices. The speed-ups
 545 achieved are 56.4 times for the A100 device, 106.6 times for the H100, and
 546 132.5 times faster for the H200 device.

547 The observed execution time improvements surpass one order of magni-
 548 tude between the legacy and new CPU versions —as demonstrated in the
 549 Lorenz '63 use case— and two orders of magnitude between the new CPU
 550 version and the GPU implementation. Together, these results indicate that
 551 the enhanced GBEES algorithm achieves a total performance improvement
 552 of more than three orders of magnitude.

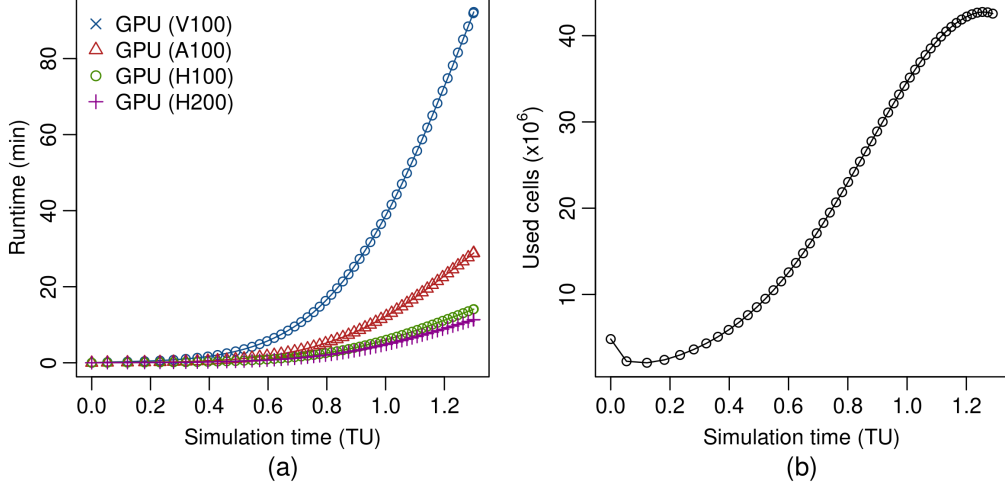


Figure 13: Runtime comparison between the new CPU version and GBEES-GPU on various GPU devices (a) and number of used cells (b) for the Lorenz '96 model.

Device	Runtime (s)	Cells/s	Speed-up
CPU-optimized: Apple M2 MAX	97927	$\approx 0.3\text{M/s}$	1
GPU 1: Tesla V100	5513	$\approx 5.4\text{M/s}$	17.8
GPU 2: NVIDIA A100	1736	$\approx 17.3\text{M/s}$	56.4
GPU 3: NVIDIA H100	919	$\approx 32.6\text{M/s}$	106.6
GPU 4: NVIDIA H200	739	$\approx 40.6\text{M/s}$	132.5

Table 3: Total runtime, number of cells processed per second, and relative speed-up compared to a single-core CPU running the new version of GBEES for the Lorenz '96 model.

553 7. Conclusions

554 This paper presents a CPU-optimized implementation and a GPU imple-
555 mentation of GBEES, a second-order accurate, Eulerian algorithm for robust,
556 nonlinear uncertainty propagation. To address the computational limitations
557 associated with the CPU implementation of GBEES, the main data struc-
558 ture was changed from a linked list to a hashtable, a CFL-minimized adaptive
559 time step was used, and the grid growing and pruning procedures were ad-
560 justed to consider the advection direction. Once the CPU implementation
561 was optimized, the algorithm was translated to CUDA for GPU execution.

562 The CUDA implementation is heavily influenced by the dynamic nature of
563 the grid required by the GBEES method. This dynamic grid demands more

sophisticated synchronization mechanisms and an efficient memory layout for its storage. To address these challenges, the CUDA implementation employs the Cooperative Kernel abstraction, an ad-hoc data structure to store the grid, non-blocking algorithms to modify it, and parallel-specific optimization techniques.

For validation, we test the two novel implementations on the Lorenz ‘63 model, a three-dimensional chaotic system utilized in the original GBEES paper. Two quantitative metrics are analyzed: the sum of the absolute differences and the Bhattacharyya Coefficient. For a truth model, we time-march an initially Gaussian PDF at $1/100^{\text{th}}$ the expected stable time step size subject to the chaotic dynamics. Both metrics indicate that the propagated probability distribution converges to the truth as the time step decreases for both the CPU-optimized and GPU implementations. For this low-dimensional example, the performance increase relative to the CPU implementation is one order of magnitude for the CPU-optimized implementation and two orders of magnitude for the GPU implementation.

Finally, we demonstrate the full capability of the new implementations on a six-dimensional variation of the Lorenz ‘96 model, an n -dimensional chaotic system. The application of finite volume methods to this dimensionality is uncommon, but necessary for the uncertainty propagation of second-order dynamical systems. For this example, we validate with a densely-populated MC simulation, qualitatively confirming that the point mass distributions and probability isosurfaces match at the start and end epochs. The high dimensionality of the system highlights the efficacy of the parallelized algorithm; for the Lorenz ‘96 system, the GPU implementation has a performance increase of two orders of magnitude relative to the CPU-optimized implementation, implying a 1000-fold increase in performance relative to the original CPU implementation.

CRediT authorship contribution statement

Benjamin L. Hanson: Writing – original draft, Methodology, Software, Validation, Visualization. **Carlos Rubio:** Writing – original draft, Methodology, Software, Validation. **Adrián García-Gutiérrez:** Validation. **Thomas Bewley:** Writing – review & editing, Conceptualization, Methodology, Supervision, Project administration.

598 Declaration of competing interest

599 The authors declare that they have no known competing financial inter-
600 ests or personal relationships that could have appeared to influence the work
601 reported in this paper.

602 Code availability

603 In the interest of facilitating further research, promoting its use, and
604 allowing the reproduction of the experiments, the complete source code is
605 available in the following repositories:

- 606 • The GBEES CPU-optimized code is available at <https://github.com/bhanson10/gbees-hash> under the BSD 3-Clause License.
- 608 • The GBEES-GPU code is available at <https://github.com/Cx-Rubio/gbees-cuda> under the BSD 3-Clause License.

610 Acknowledgements

611 The authors thank the San Diego Supercomputer Center for a computing
612 time allocation on the Triton Shared Computing Cluster.

613 Appendix A. Benchmark hardware specifications

614 The CPU execution times were measured on a system with the following
615 specifications:

- 616 • CPU: Apple M2 Max, 12-core CPU. Clock frequency 8 cores $\times$ 3.7GHz,
617 4 cores $\times$ 3.4 GHz. L2 cache size 36 MB.

618 The GPU performance tests were executed in the next GPU devices:

- 619 • GPU 1: Tesla V100-SXM2-32GB, CUDA architecture Volta. Stream
620 multiprocessors (SMs) 80. Maximum threads per SM 2048. SM clock
621 frequency 1.530 GHz. Memory clock frequency 0.877 GHz. Memory 32
622 GB HBM2.
- 623 • GPU 2: NVIDIA A100-SXM4-40GB, CUDA architecture Ampere. Stream
624 multiprocessors (SMs) 108. Maximum threads per SM 2048. SM clock
625 frequency 1.410 GHz. Memory clock frequency 1.215 GHz. Memory 40
626 GB HBM2e.

- 627 • GPU 3: NVIDIA H100-80GB, CUDA architecture Hopper. Stream
628 multiprocessors (SMs) 132. Maximum threads per SM 2048. SM clock
629 frequency 1.980 GHz. Memory clock frequency 2.619 GHz. Memory 80
630 GB HBM3.
- 631 • GPU 4: NVIDIA H200-141GB, CUDA architecture Hopper. Stream
632 multiprocessors (SMs) 132. Maximum threads per SM 2048. SM clock
633 frequency 1.980 GHz. Memory clock frequency 3.201 GHz. Memory
634 141 GB HBM3e.

635 References

- 636 [1] R. E. Kalman, A New Approach to Linear Filtering and Prediction
637 Problems, *Journal of Basic Engineering* 82 (1) (1960) 35–45. doi:10.
638 1115/1.3662552.
- 639 [2] S. J. Julier, et al., Unscented filtering and nonlinear estimation, *Pro-*
640 *ceedings of the IEEE* 92 (3) (2004) 401–422. doi:10.1109/JPROC.2003.
641 823141.
- 642 [3] G. Evensen, The ensemble kalman filter: Theoretical formulation and
643 practical implementation, *Ocean dynamics* 53 (2003) 343–367. doi:
644 10.1007/s10236-003-0036-9.
- 645 [4] H. W. Sorenson, D. L. Alspach, Recursive bayesian estimation us-
646 ing gaussian sums, *Automatica* 7 (4) (1971) 465–479. doi:10.1016/
647 0005-1098(71)90097-5.
- 648 [5] K. J. DeMars, et al., Entropy-based approach for uncertainty propaga-
649 tion of nonlinear dynamical systems, *Journal of Guidance, Control, and*
650 *Dynamics* 36 (4) (2013) 1047–1057. doi:10.2514/1.58987.
- 651 [6] J. M. Hammersley, K. W. Morton, Poor man’s monte carlo, *Journal of*
652 *the Royal Statistical Society Series B: Statistical Methodology* 16 (1)
653 (1954) 23–38. doi:10.1111/j.2517-6161.1954.tb00145.x.
- 654 [7] N. J. Gordon, D. J. Salmond, A. F. Smith, Novel approach to
655 nonlinear/non-gaussian bayesian state estimation, in: *IEE proceedings*
656 *F (radar and signal processing)*, Vol. 140, IET, 1993, pp. 107–113.
657 doi:10.1049/ip-f-2.1993.0015.

- 658 [8] M. K. Pitt, et al., Filtering via simulation: Auxiliary particle filters,
659 Journal of the American statistical association 94 (446) (1999) 590–599.
660 doi:10.1080/01621459.1999.10474153.
- 661 [9] J. V. Candy, Bootstrap particle filtering, IEEE Signal Processing Mag-
662 azine 24 (4) (2007) 73–85. doi:10.1109/MSP.2007.4286566.
- 663 [10] J. L. Anderson, et al., A monte carlo implementation of the nonlin-
664 ear filtering problem to produce ensemble assimilations and forecasts,
665 Monthly weather review 127 (12) (1999) 2741–2758. doi:10.1175/
666 1520-0493(1999)127<2741:AMCIOT>2.0.CO;2.
- 667 [11] R. J. Leveque, Finite difference methods for ordinary and partial dif-
668 ferential equations: steady-state and time-dependent problems, SIAM,
669 2007.
- 670 [12] B. Cockburn, C.-W. Shu, C. Johnson, E. Tadmor, C.-W. Shu, Essen-
671 tially non-oscillatory and weighted essentially non-oscillatory schemes
672 for hyperbolic conservation laws, Springer, 1998.
- 673 [13] P. Lax, B. Wendroff, Systems of conservation laws, in: Selected Papers
674 Volume I, Springer, 2005, pp. 263–283.
- 675 [14] S. K. Godunov, A difference scheme for numerical solution of discon-
676 tinuous solution of hydrodynamic equations, Math. Sbornik 47 (1959)
677 271–306.
678 URL <https://cir.nii.ac.jp/crid/1573387449783562752>
- 679 [15] P. K. Sweby, High resolution schemes using flux limiters for hyperbolic
680 conservation laws, SIAM journal on numerical analysis 21 (5) (1984)
681 995–1011. doi:10.1137/0721062.
- 682 [16] T. R. Bewley, et al., Efficient grid-based bayesian estimation of nonlinear
683 low-dimensional systems with sparse. non-gaussian pdfs, Automatica
684 48 (7) (2012) 1286–1290. doi:10.1016/j.automatica.2012.02.039.
- 685 [17] B. Van Leer, Towards the ultimate conservative difference scheme. iv. a
686 new approach to numerical convection, Journal of computational physics
687 23 (3) (1977) 276–299. doi:10.1016/0021-9991(77)90095-X.

- [18] B. L. Hanson, A. J. Rosengren, T. R. Bewley, State estimation of chaotic trajectories: A higher-dimensional, grid-based, bayesian approach to uncertainty propagation, in: AIAA SCITECH 2024 Forum, 2024, p. 0426. doi:10.2514/6.2024-0426.
- [19] R. Courant, K. Friedrichs, H. Lewy, On the partial difference equations of mathematical physics, IBM journal of Research and Development 11 (2) (1967) 215–234. doi:10.1147/rd.112.0215.
- [20] NVIDIA-Corporation, About CUDA, accessed on 19 November 2024 (2024).
URL <https://developer.nvidia.com/about-cuda>
- [21] K. Karzhaubayev, L.-P. Wang, D. Zhakebayev, DUGKS-GPU: An efficient parallel GPU code for 3D turbulent flow simulations using Discrete Unified Gas Kinetic Scheme, Computer Physics Communications 301 (2024) 109216. doi:10.1016/j.cpc.2024.109216.
- [22] NVIDIA-Corporation, CUDA C++ Programming Guide, accessed on 19 November 2024 (2024).
URL <https://docs.nvidia.com/cuda/cuda-c-programming-guide/index.html>
- [23] D. P. Mehta (Ed.), Handbook of data structures and applications, no. 3 in Chapman & Hall/CRC computer and information science series, Chapman & Hall/CRC, Boca Raton, Fla, 2005.
- [24] D. E. Knuth, The art of computer programming, second ed. Edition, Vol. 3, Addison-Wesley publ, Reading (Mass.) London Manila [etc.], 1973.
- [25] J. D. Cohen, Recursive hashing functions for n-grams, ACM Transactions on Information Systems (TOIS) 15 (3) (1997) 291–320. doi:10.1145/256163.256168.
- [26] R. E. Ladner, M. J. Fischer, Parallel Prefix Computation, J. ACM 27 (4) (1980) 831–838. doi:10.1145/322217.322232.
- [27] W. Hwu, D. Kirk, I. El Hajj, Programming massively parallel processors: a hands-on approach, fourth ed. Edition, Elsevier : Morgan Kauffmann, Cambridge, MA, 2023.

- 720 [28] E. N. Lorenz, Deterministic nonperiodic flow, *Journal of atmospheric*
 721 *sciences* 20 (2) (1963) 130–141. doi:10.1175/1520-0469(1963)
 722 020<0130:DNF>2.0.CO;2.
- 723 [29] E. N. Lorenz, Predictability: A problem partly solved, in: *Proc. Semi-*
 724 *nar on predictability*, Vol. 1, Reading, 1996, pp. 40–58. doi:10.1017/
 725 CB09780511617652.004.
- 726 [30] A. Bhattacharyya, On a Measure of Divergence between Two Multino-
 727 *mial Populations*, *Sankhyā: The Indian Journal of Statistics* 7 (4) (1960)
 728 401–406.
 729 URL <http://www.jstor.org/stable/25047882>