

Real-Time Event-Based Imaging Velocimetry for Closed-Loop Jet Control

Luca Franceschelli¹, Enrico Amico², Giuseppe Scirè C.², Christian E. Willert³, Marco Raiola¹, Gioacchino Cafiero² and Stefano Discetti¹

1: Department of Aerospace Engineering, Universidad Carlos III de Madrid, Avda. Universidad 30, Leganés, 28911, Madrid, Spain

2: Department of Mechanical and Aerospace Engineering, Politecnico di Torino, Corso Duca degli Abruzzi 24, Turin, 10129, Italy

3: DLR Institute of Propulsion Technology, German Aerospace Center, Linder Höhe, Köln, 51170, Germany

*Corresponding author: lfrances@ing.uc3m.es

Keywords: event-based imaging velocimetry, real-time velocimetry, flow control, jet flows, reinforcement learning.

ABSTRACT

This work presents a closed-loop flow control framework in which real-time event-based imaging velocimetry (RT-EBIV) serves as the primary sensing modality, operating directly within the feedback loop at an effective velocity-field update rate of approximately 200 Hz. Spatially resolved velocity fields are reconstructed online from pulsed-illumination event data, providing simultaneous planar access to an extended region of the flow. This single spatially resolved acquisition is exploited to define, within distinct sub-regions of the same instantaneous field, both the state representation used for controller inference and the reward signal used for policy optimisation. This dual role would not be accessible to point-wise sensing approaches. The approach is demonstrated on the mixing enhancement of a submerged round water jet at a Reynolds number based on the bulk velocity and jet diameter ≈ 8000 . The jet is actuated by pulsed radial jets injecting momentum into the shear layer. A compact state representation is constructed by projecting instantaneous RT-EBIV velocity fields onto a Proper Orthogonal Decomposition basis computed from unforced data, enabling high-frequency inference while suppressing measurement noise. Closed-loop control is implemented using the soft actor-critic algorithm, trained directly on the experimental velocity field to maximize the turbulent kinetic energy (TKE). The agent progressively learns actuation strategies that increase fluctuation-energy levels in the reward region, as confirmed independently by TKE distributions computed from time-averaged velocity fields. These results demonstrate the viability of RT-EBIV as a spatially resolved, real-time sensing modality for learning-driven closed-loop flow control, and validate the end-to-end sensing–decision–actuation pipeline under experimental and fully turbulent conditions.

1. Introduction

Active Flow Control (AFC) plays a central role in improving the performance of turbulent flows in engineering applications, as actuation can be used to suppress undesired dynamics, enhance

robustness, and improve efficiency (Gad-el Hak, 1989). In particular, closed-loop strategies are attractive because they adapt the actuation to the instantaneous flow state, enabling performance gains under changing operating conditions and unmodelled disturbances (Brunton and Noack, 2015; Brunton et al., 2020).

A central bottleneck for practical closed-loop implementations remains sensing. Conventional feedback strategies typically rely on high-rate point probes (e.g. hot wires, microphones, or pressure taps) because their signals can be acquired and processed with limited latency (Amico et al., 2022; Zong et al., 2025). However, such measurements offer only a sparse view of the flow, and the resulting observability of the underlying, high-dimensional dynamics is often limited, especially when coherent structures evolve in space and time (Brunton and Noack, 2015).

Imaging-based measurements are therefore attractive, as they can provide spatially resolved information, allowing for better inference of the flow physics. In practice, however, their use in real-time closed-loop control has been constrained by data rate, processing requirements, experimental complexity, and the need for powerful illumination and optically demanding setups for most of the techniques. As a result, while a number of efforts have explored imaging-enabled feedback (Willert et al., 2010; McCormick et al., 2024; Pimienta and Aider, 2025), the integration of flow-field measurements into robust, high-speed, real-time control loops remains an active and evolving research direction.

Event-Based Vision (EBV) offers a distinct sensing paradigm compared with conventional frame-based cameras: rather than reading out full pixel arrays at a fixed frame rate, each pixel operates independently and asynchronously, generating an event whenever the local log-intensity variation exceeds a prescribed threshold (Gallego et al., 2020). In flow velocimetry, this principle has been exploited through pulsed event-based imaging velocimetry (EBIV), where tracer particles are illuminated by short laser pulses and the corresponding events are accumulated into pulse-synchronised particle images suitable for correlation-based velocity estimation (Willert, 2023). This formulation is particularly attractive for real-time measurements because the velocity-field update rate is primarily imposed by the laser pulse repetition frequency, rather than by full-frame camera readout. At the same time, the data stream remains sparse, since events are generated only where intensity changes occur. This combination of high temporal resolution and activity-driven data throughput is difficult to achieve with conventional frame-based real-time particle image velocimetry (RT-PIV) systems, for which increasing the acquisition rate typically implies reading, transferring, and processing complete image frames (Willert et al., 2010; McCormick et al., 2024).

Building on pulsed EBIV, subsequent work has progressively extended the technique toward real-time closed-loop applications. Franceschelli et al. (2025b) demonstrated that EBIV measurements can accurately recover low-dimensional flow coordinates, benchmarking the technique against conventional particle image velocimetry (PIV) and establishing its suitability as a reduced-order sensing modality in turbulent flows. Franceschelli et al. (2025a) explored the application of EBIV to open-loop jet flow control, providing an experimental characterisation of the technique in an actively forced configuration. The real-time implementation of EBIV was then introduced by Willert

et al. (2026), who demonstrated stable online velocity-field reconstruction on standard hardware and established the computational basis for closed-loop integration. Building on this development, Franceschelli et al. (2026) proposed a data-driven framework for real-time inference of high-resolution flow representations and reduced-order coordinates from real-time event-based imaging velocimetry (RT-EBIV) data, providing compact and physically interpretable state variables compatible with high-frequency closed-loop operation.

The present contribution builds on these developments to deploy a complete sensing–decision–actuation loop in which RT-EBIV serves simultaneously as the state sensor and the reward evaluator. The approach is demonstrated on a submerged round water jet at $Re_D = U_b D / \nu \approx 8000$, actuated by pulsed radial jets injecting momentum into the jet shear layer. RT-EBIV velocity fields are reconstructed online at a frequency of $\mathcal{O}(10^2 \text{ Hz})$ and projected onto a Proper Orthogonal Decomposition (POD) basis computed from unforced data, yielding a low-dimensional state vector that drives a Soft Actor-Critic (SAC) agent trained directly on the experimental flow. The control objective is the enhancement of the jet turbulent kinetic energy (TKE), downstream the actuation. Here the TKE is used as a spatially resolved proxy for increased jet mixing: higher fluctuation energy in the reward region implies stronger entrainment and broader momentum spreading, consistent with mixing enhancement. The reward is evaluated instantaneously using spatial non-uniformity of the fluctuations as a fast-evaluated proxy for TKE, without requiring any time-averaged reference. Results demonstrate that the agent progressively learns actuation strategies that increase the reward during training, and that the corresponding modification of the downstream jet dynamics is confirmed by an independent TKE assessment based on time-averaged velocity fields.

2. Methodology

2.1. Experimental Setup

The experiments are conducted on a submerged round water jet issuing into a water tank (dimensions $900 \times 500 \times 500 \text{ mm}^3$), as shown in Figure 1(b). The jet nozzle has a diameter $D = 30 \text{ mm}$ and the flow operates at a Reynolds number $Re_D \approx 8000$, defined as $Re_D = U_0 D / \nu$, where U_0 is the bulk velocity at the nozzle exit and ν is the kinematic viscosity of water. The tank is equipped with overflow holes positioned at 450 mm from the bottom, which discharge excess water into a secondary side tank, ensuring a constant free-surface level throughout the experiment. From the side tank, two pumps operate in parallel: the first feeds the main jet, supplying the bulk flow through the nozzle; the second fills a fixed elevated reservoir mounted on top of the actuation structure, which in turn feeds eight individual small reservoirs, one per valve. Since the actuator outlets are arranged circumferentially around the jet axis, they lie at different depths within the tank, which would otherwise introduce unequal supply pressures across the actuation circle. To compensate, each individual reservoir is mounted on an independent vertical rail and its height is adjusted manually prior to the experiment, so that the hydrostatic head between each small reservoir and its corresponding valve outlet remains constant across all actuators.

The coordinate system is defined such that x denotes the streamwise direction aligned with the jet axis, while y and z are the transverse directions. The measurement plane coincides with the (x, y) plane and intersects the jet centerline. Figure 1 presents a schematic overview of the experimental setup and control loop.

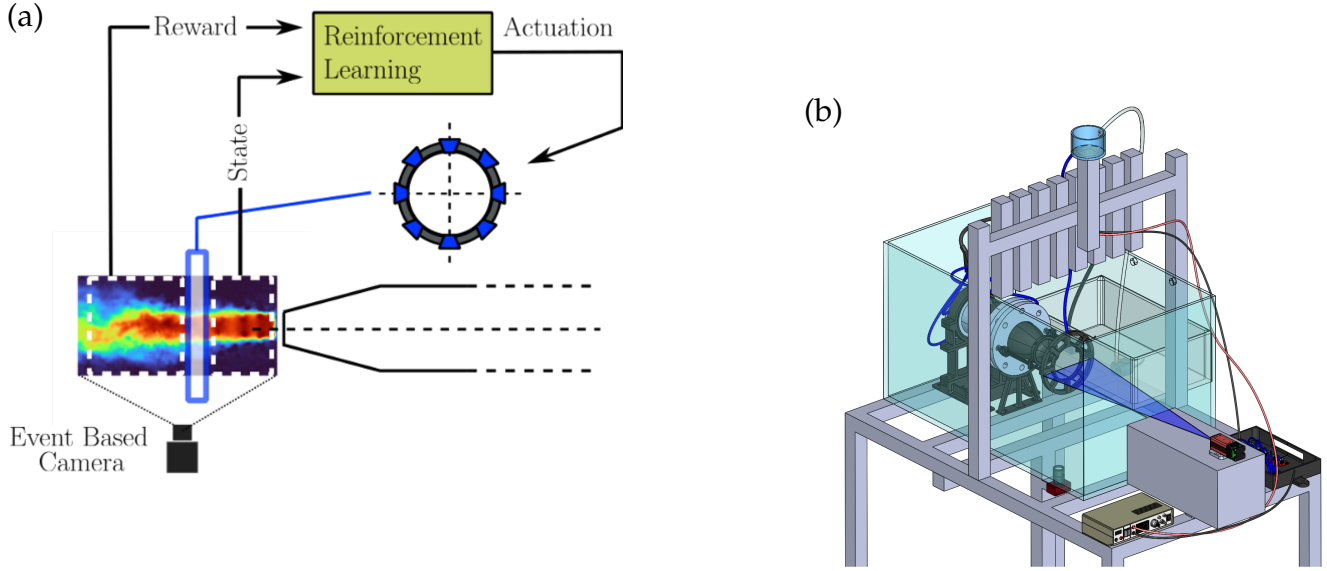


Figure 1. (a) Schematic representation of the closed-loop control architecture, illustrating the sensing–decision–actuation pipeline. (b) Sketch of the experimental setup.

Flow actuation is provided by eight radially oriented pulsed water jets arranged symmetrically around the nozzle. The actuators are located approximately $4D$ downstream the jet exit and inject momentum into the jet shear layer. Each actuator is independently driven, enabling spatially distributed forcing of the flow. In the preliminary validation experiments presented here, actuation is restricted to the two pulsed jets aligned with the EBIV measurement plane in order to reduce the dimensionality of the control problem and facilitate validation of the closed-loop framework.

The flow is seeded with $56\text{ }\mu\text{m}$ -diameter tracer particles. Illumination is provided by a 20 W continuous-wave blue diode laser module (LaserTree LT-4LDS-V2, $\lambda = 450\text{ nm}$), whose beam is collimated and expanded into a light sheet of approximately 1 mm thickness in the jet longitudinal plane using a combination of cylindrical and spherical lenses. The laser is operated in pulsed mode at a frequency $f_{\text{EBIV}} = 200\text{ Hz}$ with a duty cycle of 10%, driven by an external pulse width modulation (PWM) signal generated by a Digilent Analog Discovery 3, which also provides the hardware synchronisation signal to the Prophesee EVK4 event-based camera, ensuring consistent event-to-pulse timing across all acquisition runs.

All stages of acquisition and processing are executed online, using an in-house developed code. Event accumulation, velocity reconstruction, feature extraction, reward evaluation, and controller inference are performed in real time, resulting in an effective velocity-field update rate of approximately 200 Hz.

2.2. Real-Time Event-Based Imaging Velocimetry

An in-house implementation of RT-EBIV provides spatially resolved velocity measurements with low latency. Events generated by the camera are accumulated over a short temporal window centred on each laser pulse and mapped onto an image-like representation. These reconstructed event frames contain particle signatures suitable for correlation-based velocimetry. The pipeline is implemented in software and executed on a standard CPU, sustaining stable operation up to 500 Hz. For the present experiments, the acquisition frequency is set to $f_{\text{EBIV}} = 200$ Hz, which is sufficient to resolve the dominant convective time scales of the jet while remaining consistent with the controller update rate.

Velocity fields are computed using a single-pass PIV algorithm applied to pairs of consecutive event frames separated by a pulse-to-pulse time interval δt . The local displacement vector $\Delta \mathbf{p}(x, y)$ is obtained from the correlation peak and converted into a velocity estimate as

$$\mathbf{u}(x, y, t) = \frac{\Delta \mathbf{p}(x, y)}{\delta t} s, \quad (1)$$

where s is the spatial calibration factor. Basic vector validation procedures (Westerweel and Scarano, 2005) are applied online to ensure robustness while maintaining low latency. Velocity fields are computed on a grid with 16 px pitch, using interrogation windows of 32×32 px² and 50% overlap.

The EBIV measurement domain is spatially divided into two regions with different roles in the control loop. As shown in Figure 1, the instantaneous flow state is evaluated considering a sensing region Ω_s , located within $0.2 < x/D < 3.8$. A separate reward region Ω_r , located downstream and beyond the actuation zone, in the area $4.5 < x/D < 10$, is used to quantify control performance. This separation prevents trivial feedback strategies and ensures that the reward is based on the downstream impact of actuation rather than local forcing effects.

From the velocity field in Ω_s , a compact state representation is constructed and supplied to the controller. The state consists of spatially resolved features derived from the RT-EBIV measurements, chosen to balance physical relevance, robustness to noise, and computational efficiency. The resulting state proxy retains essential information about the jet dynamics while remaining compatible with real-time closed-loop operation.

2.3. Closed-Loop Control via Reinforcement Learning (RL)

Closed-loop flow control is implemented using a RL approach (Kaelbling et al., 1996). The control problem is formulated as a Markov decision process in which the agent observes the flow state inferred from RT-EBIV measurements and selects actuation commands to maximize a cumulative reward.

At each control step t , the velocity field provided by RT-EBIV is processed to extract a state vector s_t . Based on this observation of the state, the agent samples an action a_t , which specifies the

actuation amplitudes applied to the pulsed jets during the next control interval. After actuation, a scalar reward r_t is computed from the velocity field measured in the downstream reward region Ω_r . The objective of the agent is to maximize the expected discounted return

$$J = \mathbb{E} \left[\sum_{t=0}^T \gamma^t r_t \right], \quad (2)$$

where γ is the discount factor and T is the episode length.

The controller is implemented using the SAC algorithm (Haarnoja et al., 2018), an off-policy, model-free RL method based on maximum-entropy optimization. SAC is well suited for experimental flow-control applications due to its robustness to measurement noise, ability to reuse past data efficiently, and stable learning behavior. The algorithm learns a stochastic policy that maps the observed state to continuous-valued actuator commands while simultaneously estimating the associated action-value functions.

To ensure stable and safe operation during learning, actuator commands are constrained within physical limits and the control loop is designed to account explicitly for sensing and actuation latency. Training is organized in episodes of fixed duration, allowing repeated interaction with the flow under controlled initial conditions.

2.4. Reduced order state representation

The most direct way to define the flow state would be to use the velocity vectors obtained from the PIV reconstruction over the full region downstream of the actuation. However, this choice would result in a dense, matrix-valued state representation, whose dimensionality is not well-suited for a high-frequency Deep Reinforcement Learning (DRL) framework. In particular, directly feeding the full velocity field to the agent would increase the computational cost of inference and learning, while also making the controller more sensitive to small-scale measurement noise. To overcome this limitation, a reduced-order state representation is constructed by projecting the instantaneous velocity field onto a POD basis computed from unforced data. This choice is motivated by the observation that the leading POD modes capture the dominant coherent structures of the unforced jet, which are the primary targets for mixing enhancement: manipulating such structures is indeed widely recognised as an effective mechanism for increasing downstream TKE (Benard et al., 2008; Oberleithner et al., 2011). It should be acknowledged that the unforced POD basis may not optimally represent structures introduced by actuation; however, since the control objective is the enhancement of large-scale mixing, the leading modes of the unforced flow provide a physically grounded and computationally efficient state representation. A long sequence of velocity fields was recorded using the same RT-EBIV framework adopted during closed-loop operation, including the same spatial resolution, PIV grid, and acquisition frequency. Let the instantaneous velocity

field be written as

$$\mathbf{u}(\mathbf{x}, t) = \begin{bmatrix} u(\mathbf{x}, t) \\ v(\mathbf{x}, t) \end{bmatrix}, \quad (3)$$

where $\mathbf{x}$ denotes the spatial position in the sensing region. From the unforced dataset, the mean velocity field is first computed as

$$\bar{\mathbf{u}}(\mathbf{x}) = \frac{1}{N} \sum_{n=1}^N \mathbf{u}(\mathbf{x}, t_n), \quad (4)$$

and the corresponding fluctuation fields are defined as

$$\mathbf{u}'(\mathbf{x}, t_n) = \mathbf{u}(\mathbf{x}, t_n) - \bar{\mathbf{u}}(\mathbf{x}). \quad (5)$$

A POD decomposition was then applied to the ensemble of fluctuation fields, leading to the expansion

$$\mathbf{u}'(\mathbf{x}, t) \simeq \sum_{i=1}^{N_m} a_i(t) \phi_i(\mathbf{x}), \quad (6)$$

where $\phi_i(\mathbf{x})$ are the POD spatial modes, $a_i(t)$ are the corresponding modal coefficients, and N_m is the number of retained modes. The POD modes provide an energetically optimal basis for representing the dominant flow structures in the measurement domain. Since the control objective is primarily associated with the manipulation of large-scale coherent structures, only the leading POD modes were retained. The number of modes was selected such that the cumulative reconstructed kinetic energy exceeded 50% of the total fluctuation energy contained in the unforced dataset, namely

$$\frac{\sum_{i=1}^{N_m} \lambda_i}{\sum_{i=1}^{N_{\text{POD}}} \lambda_i} > 0.5, \quad (7)$$

where λ_i is the energy associated with the i -th POD mode and N_{POD} is the total number of available modes. In this case, $N_m = 30$ modes are retained. At each control step, the instantaneous velocity field measured by RT-EBIV is projected onto the reduced POD basis. The modal coefficients are obtained as

$$a_i(t) = \langle \mathbf{u}'(\mathbf{x}, t), \phi_i(\mathbf{x}) \rangle, \quad i = 1, \dots, N_m, \quad (8)$$

where $\langle \cdot, \cdot \rangle$ denotes the spatial inner product over the sensing region. The state supplied to the RL agent is then defined as

$$\mathbf{s}_t = \begin{bmatrix} a_1(t) & a_2(t) & \dots & a_{N_m}(t) \end{bmatrix}^\top. \quad (9)$$

In this way, the original high-dimensional velocity field is replaced by a compact state representation whose dimension is equal to the number of retained POD modes, while preserving the dominant large-scale information relevant to the control task. The construction procedure is schematically illustrated in Figure 2.

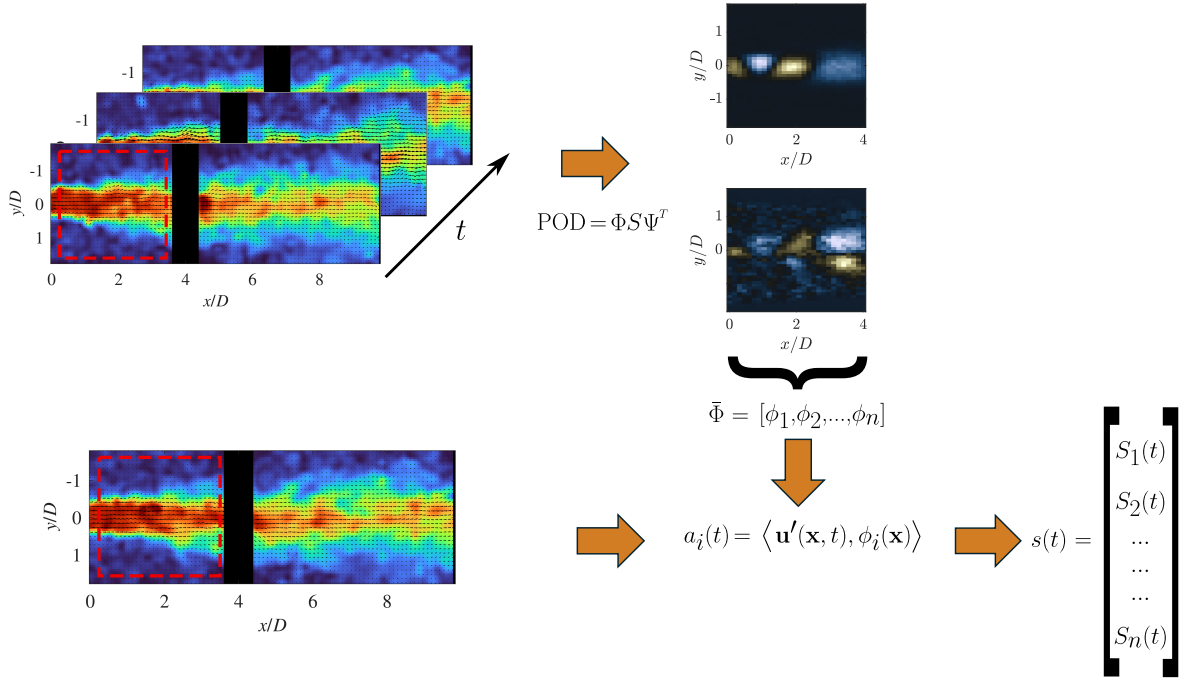


Figure 2. Construction of the reduced-order state representation. A sequence of RT-EBIV velocity fields acquired in the unforced condition is decomposed via POD, yielding a set of energy-ranked spatial modes $\bar{\Phi} = [\varphi_1, \varphi_2, \dots, \varphi_{N_m}]$.

At each control step, the instantaneous fluctuation field is projected onto the retained modes to obtain the modal coefficient vector $\mathbf{s}(t) = [a_1(t), a_2(t), \dots, a_{N_m}(t)]^T$, which constitutes the state supplied to the RL agent.

2.5. Reward estimation

The control objective considered in this work is the enhancement of TKE in the region downstream of the actuation. In its classical definition, TKE is computed from velocity fluctuations with respect to a mean velocity field. However, in a real-time closed-loop experiment, the mean flow is not available a priori and cannot be reliably estimated online over the time scales required by the controller. For this reason, an instantaneous TKE proxy $\tilde{k}$ is introduced and used to define the reward signal, as suggested in Chauhan et al. (2014) for a turbulent boundary layer and a planar jet. A fixed region of interest Ω_r , located downstream of the actuation region, is selected for the evaluation of the reward, as illustrated in Figure 3. This domain is divided into m non-overlapping subwindows,

$$\Omega_r = \bigcup_{j=1}^m \Omega_j, \quad (10)$$

where each subwindow Ω_j contains k PIV grid points. At each control step, a local reference velocity is computed in each subwindow as

$$\bar{\mathbf{u}}_j(t) = \frac{1}{k} \sum_{\mathbf{x}_i \in \Omega_j} \mathbf{u}(\mathbf{x}_i, t). \quad (11)$$

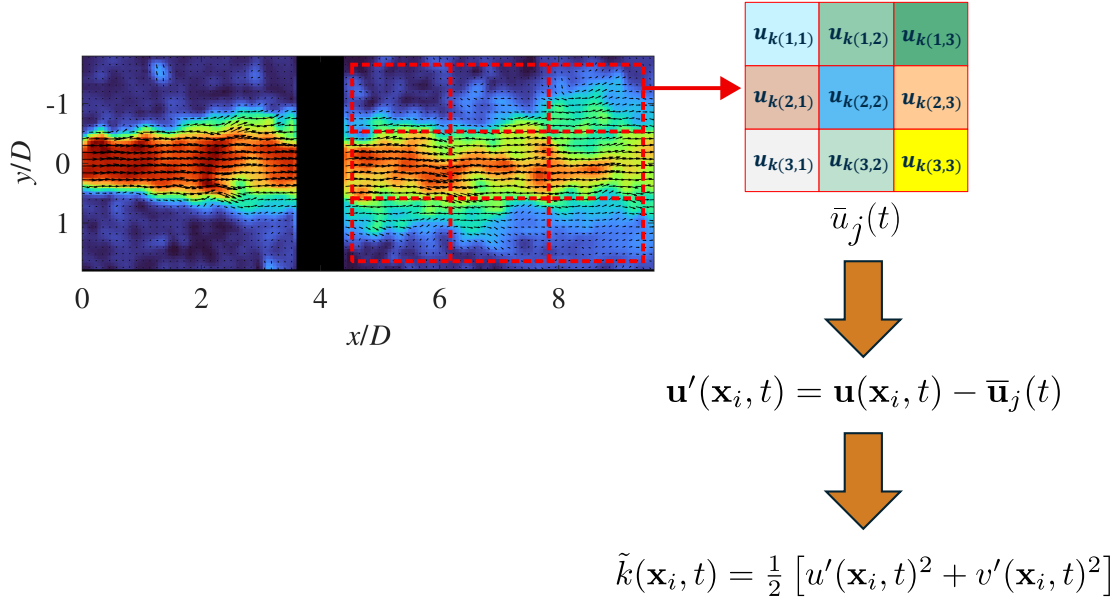


Figure 3. Computation of the instantaneous TKE proxy $\tilde{k}$ in the reward region Ω_r . The reward domain is partitioned into non-overlapping subwindows, each containing k PIV grid points. Within each subwindow Ω_j , a local spatial mean $\bar{\mathbf{u}}_j(t)$ is subtracted from the instantaneous velocity to obtain local fluctuations $\mathbf{u}'(\mathbf{x}_i, t)$. The TKE proxy $\tilde{k}$ is then evaluated pointwise and summed over Ω_r to yield the scalar reward r_t .

The local velocity fluctuations are then defined by subtracting this subwindow-wise spatial mean,

$$\mathbf{u}'(\mathbf{x}_i, t) = \mathbf{u}(\mathbf{x}_i, t) - \bar{\mathbf{u}}_j(t), \quad \mathbf{x}_i \in \Omega_j. \quad (12)$$

This operation removes the local mean contribution without requiring a time-averaged reference field. The TKE proxy is then computed from the in-plane velocity fluctuations as

$$\tilde{k}(\mathbf{x}_i, t) = \frac{1}{2} [u'(\mathbf{x}_i, t)^2 + v'(\mathbf{x}_i, t)^2], \quad (13)$$

and the instantaneous reward is defined as the total fluctuation energy within the reward region,

$$r_t = \sum_{j=1}^m \sum_{\mathbf{x}_i \in \Omega_j} \tilde{k}(\mathbf{x}_i, t). \quad (14)$$

This formulation provides an instantaneous measure of local velocity fluctuations that can be evaluated directly from each RT-EBIV velocity field. The use of subwindow-wise spatial averages preserves the energetic content associated with spatial velocity variations, while remaining compatible with high-frequency closed-loop operation. This TKE proxy does not strictly recover the classical TKE, which requires temporal averages unavailable during real-time operation; nonetheless, the results presented in Section 3 support its practical viability as a reward signal for the closed-loop framework.

3. Results and Discussion

Before each training episode, a short sanity-check phase was performed to verify that both the seeding conditions and the RT-EBIV reconstruction quality were suitable for closed-loop operation. This preliminary check was introduced to prevent the agent from being trained on corrupted or poorly resolved velocity fields, which would otherwise result in unreliable state estimates and noisy reward signals.

To provide a reference for the uncontrolled flow, Figure 4(a) reports the TKE distribution computed from velocity fluctuations with respect to the time-averaged velocity field. This field represents the baseline fluctuation-energy distribution before learning and is used as a reference to assess the effect of the learned actuation strategy.

The RL training was performed over 95 episodes, each with a length of $T = 1000$ timesteps. Figure 4(b) reports the evolution of the episode reward during training, together with its moving-average trend. The reward exhibits a clear increasing tendency throughout the experiment, indicating that the SAC agent progressively learns actuation strategies that enhance the velocity fluctuations in the downstream reward region. Although significant episode-to-episode variability is observed, as expected in an experimental turbulent flow-control problem, the overall trend shows that the controller is able to improve its performance through direct interaction with the flow.

To further assess whether the increase in reward corresponds to a physical increase in fluctuation energy, Figure 5 reports the percentage increase of both the online proxy $\tilde{k}$ and the TKE evaluated a posteriori, with respect to the uncontrolled case. Both quantities are averaged over each episode and over the reward domain Ω_r . The classical TKE estimate follows the same overall trend as the online proxy $\tilde{k}$, supporting the use of the proxy as a real-time reward signal. This comparison is important because the controller is trained using only $\tilde{k}$, whereas the TKE is evaluated independently after the experiment from velocity fluctuations with respect to the episode-averaged velocity field. The observed increase in TKE therefore confirms that the learned strategy produces a measurable amplification of the downstream fluctuation energy. In less than 100 episodes, the averaged TKE in the reward region more than doubles with respect to the uncontrolled case.

The corresponding spatial effect on the flow field is illustrated in Figure 6, which compares the TKE distribution in the reward region between an early training episode and episode 92. In contrast to the online reward, which is based on the instantaneous proxy $\tilde{k}$, here the TKE is evaluated from velocity fluctuations with respect to the episode-averaged velocity field $\bar{\mathbf{u}}(\mathbf{x}) = \frac{1}{T} \sum_t \mathbf{u}(\mathbf{x}, t)$, taken over the T timesteps of the episode and without any subwindow-wise spatial averaging. A comparison between the early episode (Figure 6(a)) and episode 92 (Figure 6(b)) reveals a visible increase in TKE over a large portion of the reward region. This indicates that the growth in episode reward and in the averaged TKE is associated with a physical modification of the downstream jet dynamics, rather than only reflecting the numerical definition of the proxy used for online reward evaluation.

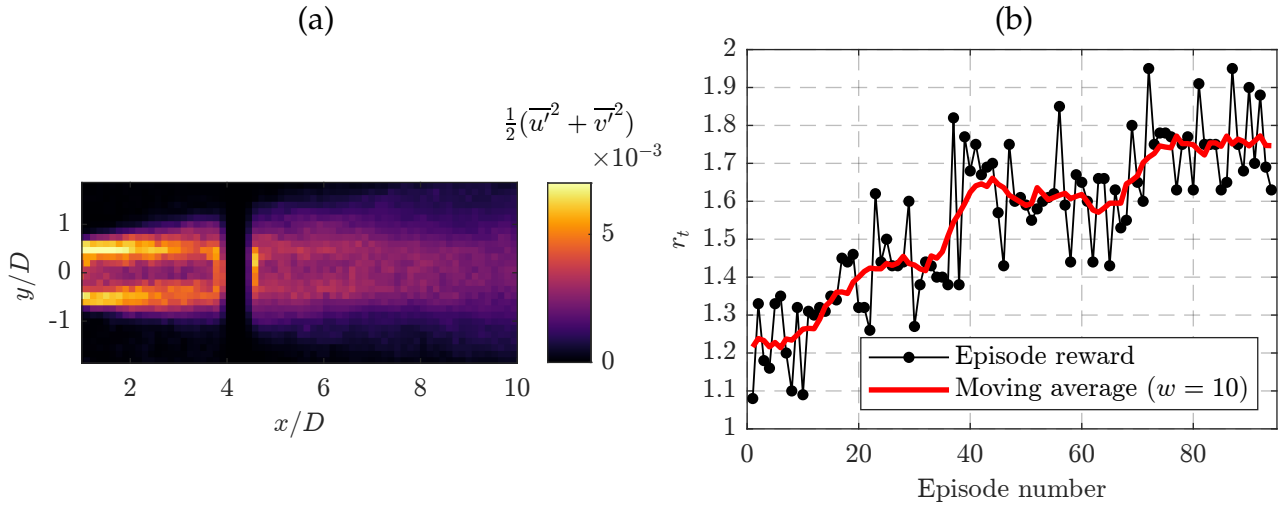


Figure 4. (a) TKE distribution computed from velocity fluctuations with respect to the time-averaged velocity field in the uncontrolled case. (b) Evolution of the episode reward during DRL training, together with its moving-average trend.

4. Conclusions

This work presented a closed-loop flow-control framework in which RT-EBIV is directly integrated into an experimental feedback loop. The proposed approach combines spatially resolved planar velocity measurements with a RL controller to manipulate the dynamics of a submerged turbulent jet. A key feature of the framework is that the single RT-EBIV acquisition simultaneously serves two distinct roles: the instantaneous velocity field is partitioned into spatially separated sub-regions, one used for state estimation and one for reward evaluation, within the same snapshot. Beyond this dual exploitation, the spatially resolved nature of the measurement enables richer and more physically grounded definitions of both sensing and reward variables (such as modal projections, spatial energy distributions, or structure-aware quantities) that are inherently inaccessible to point-wise sensing approaches, which would additionally require dedicated and spatially separated probes to achieve equivalent functional separation.

Preliminary experiments demonstrated that RT-EBIV can provide sufficiently robust and low-latency velocity-field estimates to be used for both state estimation and online reward evaluation. Using a SAC algorithm, the controller was able to learn actuation strategies that increase the fluctuation-energy levels in the downstream reward region. This behaviour is evidenced by the progressive increase of the episode reward during training and by the comparison of the actual TKE distributions computed for an early and a late episode.

An additional assessment based on the classical TKE, computed from velocity fluctuations with respect to the time-averaged flow field, further confirms that the learned actuation produces a measurable modification of the downstream jet dynamics. These results validate the end-to-end functionality of the sensing–decision–actuation pipeline and demonstrate the feasibility of embedding EBIV measurements within a learning-driven control loop.

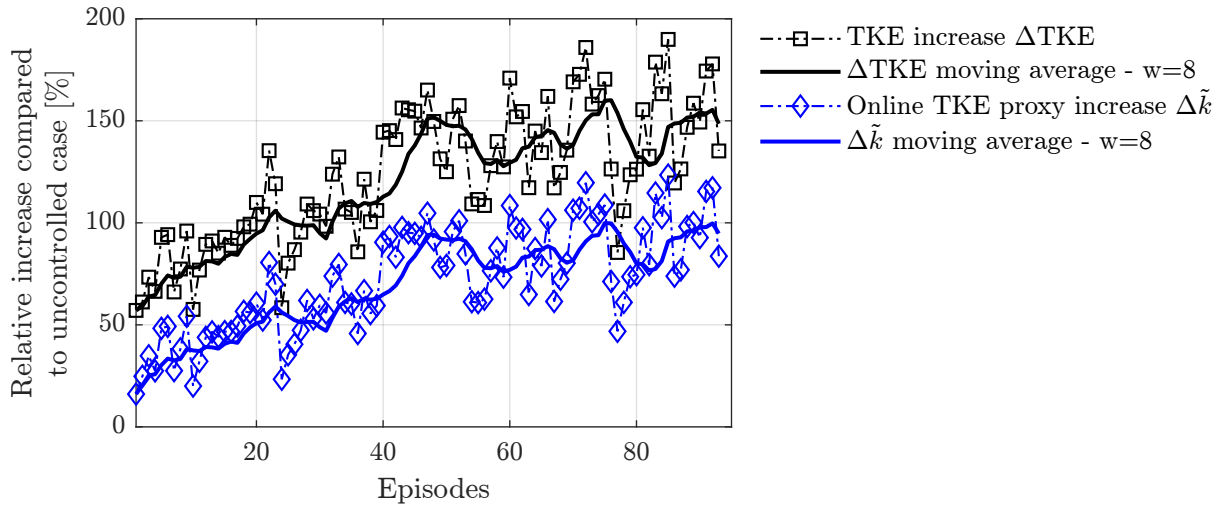


Figure 5. Percentage increase of the TKE and of the online proxy $\tilde{k}$ with respect to the uncontrolled case. Both quantities are averaged over each episode and over the reward region Ω_r . The corresponding moving-average trends are also shown.

The present results should be interpreted as a proof of feasibility rather than as an assessment of optimal control performance. The control objective was deliberately chosen to be simple, allowing the main challenges associated with real-time sensing, latency, reward evaluation, and experimental learning stability to be isolated.

Overall, the results indicate that EBIV can serve as a viable sensing modality for closed-loop flow-control applications when combined with RL. Future work will focus on more challenging jet-control objectives, including the regulation of mixing and integral performance metrics, as well as on richer state representations and extended actuation strategies.

Acknowledgment

Funded by the European Union, grant agreement n.101292611. Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Council Executive Agency (ERCEA). Neither the European Union nor the granting authority can be held responsible for them.

References

- Amico, E., Cafiero, G., and Iuso, G. (2022). Deep reinforcement learning for active control of a three-dimensional bluff body wake. *Physics of Fluids*, 34(10).
- Benard, N., Balcon, N., Touchard, G., and Moreau, E. (2008). Control of diffuser jet flow: turbulent kinetic energy and jet spreading enhancements assisted by a non-thermal plasma discharge. *Experiments in fluids*, 45(2):333–355.

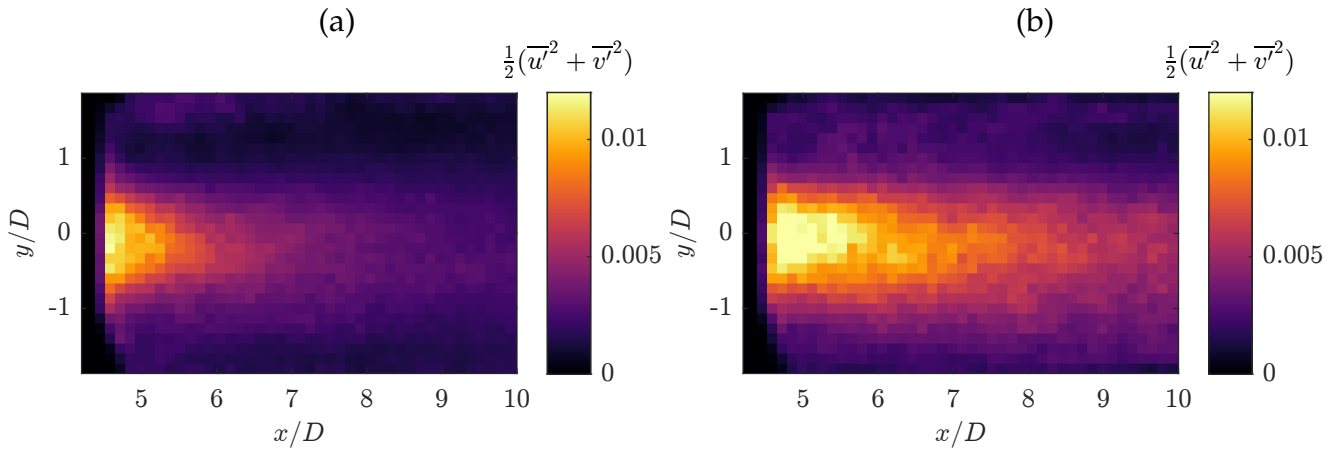


Figure 6. Spatial distribution of the TKE in the reward region Ω_r , computed from velocity fluctuations with respect to the episode-averaged velocity field $\bar{\mathbf{u}}(\mathbf{x})$ (the temporal mean over the T timesteps of the episode), without subwindow-wise spatial averaging: (a) early training episode, representative of the unlearned actuation policy; (b) episode 92, representative of the learned actuation strategy.

Brunton, S. L. and Noack, B. R. (2015). Closed-loop turbulence control: Progress and challenges. *Applied Mechanics Reviews*, 67(5):050801.

Brunton, S. L., Noack, B. R., and Koumoutsakos, P. (2020). Machine learning for fluid mechanics. *Annual review of fluid mechanics*, 52(1):477–508.

Chauhan, K., Philip, J., De Silva, C. M., Hutchins, N., and Marusic, I. (2014). The turbulent/non-turbulent interface and entrainment in a boundary layer. *Journal of Fluid Mechanics*, 742:119–151.

Franceschelli, L., Amico, E., Raiola, M., Willert, C., Serpieri, J., Cafiero, G., and Discetti, S. (2025a). Event-based imaging velocimetry for jet flow control. In *Proceedings of the 16th International Symposium on Particle Image Velocimetry (ISPIV 2025)*, Tokyo, Japan.

Franceschelli, L., Amico, E., Willert, C. E., Raiola, M., Cafiero, G., and Discetti, S. (2026). Real-time estimation of high-resolution flow fields and reduced-order coordinates from event-based imaging velocimetry. *arXiv preprint arXiv:2605.04186*.

Franceschelli, L., Willert, C. E., Raiola, M., and Discetti, S. (2025b). An assessment of event-based imaging velocimetry for efficient estimation of low-dimensional coordinates in turbulent flows. *Experimental Thermal and Fluid Science*, 164:111425.

Gad-el Hak, M. (1989). Flow control. *Applied Mechanics Reviews*, 42(10):261–293.

Gallego, G., Delbrück, T., Orchard, G., Bartolozzi, C., Taba, B., Censi, A., Leutenegger, S., Davison, A. J., Conradt, J., Daniilidis, K., et al. (2020). Event-based vision: A survey. *IEEE transactions on pattern analysis and machine intelligence*, 44(1):154–180.

- Haarnoja, T., Zhou, A., Abbeel, P., and Levine, S. (2018). Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor. In *International Conference on Machine Learning*, pages 1861–1870. PMLR.
- Kaelbling, L. P., Littman, M. L., and Moore, A. W. (1996). Reinforcement learning: A survey. *Journal of Artificial Intelligence Research*, 4:237–285.
- McCormick, F., Gibeau, B., and Ghaemi, S. (2024). Reactive control of velocity fluctuations using an active deformable surface and real-time piv. *Journal of Fluid Mechanics*, 985:A9.
- Oberleithner, K., Sieber, M., Nayeri, C., and Paschereit, C. (2011). On the control of global modes in swirling jet experiments. In *Journal of Physics: Conference Series*, volume 318, page 032050.
- Pimienta, J. and Aider, J.-L. (2025). High resolution and high-speed live optical flow velocimetry. *arXiv preprint arXiv:2509.25924*.
- Westerweel, J. and Scarano, F. (2005). Universal outlier detection for piv data. *Experiments in fluids*, 39(6):1096–1100.
- Willert, C. (2023). Event-based imaging velocimetry using pulsed illumination. *Experiments in Fluids*, 64(5):98.
- Willert, C., Franceschelli, L., Amico, E., Raiola, M., Cafiero, G., and Discetti, S. (2026). Real-time event-based particle image velocimetry for active flow control. In *Journal of Physics: Conference Series*, volume 3173, page 012001. IOP Publishing.
- Willert, C., Munson, M., and Gharib, M. (2010). Real-time particle image velocimetry for closed-loop flow control applications. In *Proceedings of the 15th Int Symp on Applications of Laser Techniques to Fluid Mechanics*, Lisbon, Portugal.
- Zong, H., Wu, Y., Li, J., Su, Z., and Liang, H. (2025). Closed-loop supersonic flow control with a high-speed experimental deep reinforcement learning framework. *Journal of Fluid Mechanics*, 1009:A3.