

An Organoid World Model: Forecasting Self-Supervised Feature Dynamics for Early Chemosensitivity Prediction

Thomas Menard¹, Emilie Gontran², Jacques R.R. Mathieu², Jérôme Cartry², Paul-Henry Cournède¹, Fanny Jaulin², Stergios Christodoulidis¹, Véronique Le Chevalier¹, and Maria Vakalopoulou¹

¹ Université Paris-Saclay, CentraleSupélec, Gustave Roussy, INSERM, Cancer Data Science Unit, IHU PRISM National Precision Medicine Center in Oncology, Gif-sur-Yvette, France

thomas.menard@centralesupelec.fr

maria.vakalopoulou@centralesupelec.fr

² Université Paris-Saclay, Gustave Roussy, Inserm, Dynamique des cellules tumorales, 94805, Villejuif, France

Abstract. Predicting patient-specific drug sensitivity from organoid cultures remains a critical bottleneck in precision oncology, typically requiring days of observation and costly endpoint assays. In this paper, we present a framework that relies on a world model for patient-derived organoids (PDO): it forecasts their morphological evolution by predicting future self-supervised DINOv2 features from early time-lapse microscopy frames. In particular, masked autoencoders are used as a learned transition model to predict the future representations of the organoid’s state from previous dynamics. By operating in the representation space of a vision foundation model rather than in pixel space, the method learns the latent dynamics of the organoid and captures semantically meaningful morphological changes. Leveraging a private collection of colorectal and pancreatic cancer organoids, we evaluate the fidelity of these forecasted features through two downstream tasks: functional ATP-based drug response estimation, and a generative last-frame reconstruction that maps the predicted features back into pixel space. Predicted future features enable drug response estimation before the endpoint is physically reached – effectively rolling the world model forward to a virtual endpoint – while the reconstruction task confirms that structural and morphological integrity is preserved, opening a new avenue for early, non-destructive chemosensitivity screening.

Keywords: Medical World Model · Self-Supervised Feature Forecasting · Latent Dynamics · Precision Oncology Organoids · Virtual Endpoint.

1 Introduction

Precision oncology increasingly relies on patient-derived organoids (PDOs) to model individual drug responses and optimize personalized treatment regimens

[1–3]. PDOs are generated by isolating tissue cells, embedding them in an extracellular matrix, and culturing them under defined conditions to form self-organizing 3D structures. These organoids are then exposed to a panel of therapeutic agents to assess efficacy and toxicity, guiding personalized treatment selection. However, standard functional assays, such as ATP-based luminescence, which measures intracellular levels of adenosine triphosphate as a proxy for cell viability and metabolic activity, remain a major clinical bottleneck. These gold-standard measurements typically require multi-day observation periods and are inherently destructive, which makes longitudinal monitoring of the same biological sample impossible. A key challenge is therefore to infer treatment response before the experiment reaches its biochemical endpoint. One option is to regress the response directly from the early observed frames, ignoring the unobserved future; another is to forecast the organoid’s future appearance in pixel space with video-prediction models [4–6] and estimate the response from the predicted frames. However, generating images far ahead in time accumulates artifacts and does not explicitly capture the high-level biological state required for downstream clinical prediction [7, 8].

In this work, we instead formulate organoid response prediction as a world-modeling problem. Rather than predicting future images, we forecast the future evolution of self-supervised latent representations learned from multitemporal organoid microscopy. These latent states summarize biologically relevant morphology and its evolution through time while avoiding the instability of pixel-level forecasting. The predicted future representation is then used to estimate ATP response and can additionally be decoded into a future image for qualitative validation.

Contributions (i) We adapt the DINO-Foresight feature-forecasting framework to patient-derived organoids and use it as an organoid world model that forecasts the latent dynamics of self-supervised DINOv2 features to a virtual endpoint, enabling early drug-response prediction. (ii) we introduce a two-frame read-out that isolates the benefit of forecasting from that of simply observing longer, and evaluate it against a no-forecasting baseline and an oracle across Pearson r , MAE and C-index with multi-seed variability; and (iii) we show that this virtual endpoint estimates drug sensitivity about 39% earlier than the physical assay (96 h→59 h) while recovering most of its predictive information, offering a concrete, non-destructive tool for accelerated chemosensitivity screening.

2 Related work

World Models and Latent Dynamics. World models learn a compact latent representation of an environment together with a transition model that predicts its temporal evolution, enabling forecasting directly in representation space rather than in raw observations [9, 10]. Recent frameworks show that forecasting semantic representations is often more effective than reconstructing pixels: JEPA [11] advocates prediction in latent space, V-JEPA [12] extends this to video, and

DINO-Foresight [13] forecasts self-supervised DINOv2 embeddings to model natural video dynamics. The medical imaging community has begun adopting these ideas – EchoWorld [14] for echocardiographic navigation, SurgFUTR [15] for future surgical videos – but these target short-term or large-magnitude changes, whereas organoid development involves subtle morphological evolution over several days.

Foundation Models and Organoid Phenotyping. Foundation models such as DINOv2 [16] provide transferable self-supervised features that are effective across medical imaging tasks [17–20], but are applied to static images; we instead extract them from every frame of a time-lapse sequence and model their evolution. Medical video generation has been tackled with GANs for cross-modality translation [21–24], diffusion models [25–27], and deep temporal models for tumour-growth prediction [28]; these operate at the image level and are restricted to already-observed or short-horizon settings. In organoid analysis, Fillioux et al. [29] classified drug efficacy from complete microscopy sequences. Unlike these approaches, we do not analyse an already observed trajectory but forecast its evolution: from an abbreviated observation window we estimate ATP response before the biological endpoint is reached.

3 Methods

An overview of our method is presented in Fig. 1. The core of our approach builds on the DINO-Foresight architecture, which we cast explicitly as a world model with four components: an observation model (frozen DINOv2 encoder + PCA), a latent transition model (masked transformer), a functional outcome head (ATP regressor), and an observation decoder (image reconstructor).

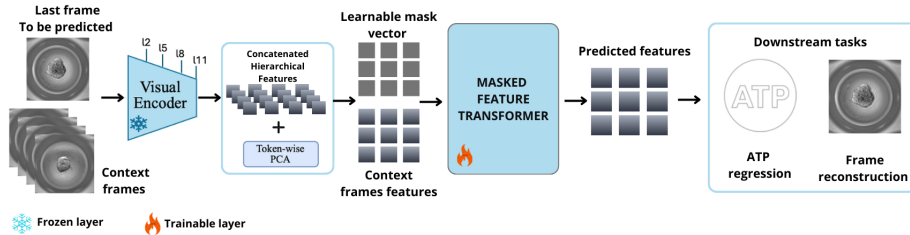


Fig. 1: The organoid world model. A sequence of $T - 1$ context frames is encoded by a frozen DINOv2 backbone and PCA-compressed into latent states (observation model). A Masked Feature Transformer (transition model) forecasts the terminal-frame state, which is read out by an image decoder for visual reconstruction and by the two-frame ATP head of Sec. 3.2 for ATP-based functional prediction.

3.1 Formulation of the world model

Overview of the proposed world-model. Let $o_1, \dots, o_T$ be T frames uniformly sampled from a 50-frame time-lapse video. The observation model ϕ maps each frame to a compact latent state $s_t = \phi(o_t) \in \mathbb{R}^{N \times K}$, where N is the number of spatial patches and K the compressed feature dimension. The transition model f_θ predicts the terminal latent state from the observed context,

$$\hat{s}_T = f_\theta(s_1, \dots, s_{T-1}), \quad (1)$$

and two read-out functions decode the model’s belief about the future: a functional head $h_\psi(s_1, \hat{s}_T)$ estimating ATP viability from the initial and terminal states, and a generative decoder $g_\xi(\hat{s}_T)$ mapping the predicted terminal state back to pixels for visual validation.

Observation model: DINOv2 features and PCA state. Each frame is passed through a frozen DINOv2 ViT-B/14 encoder (with registers) [30]. To capture multi-scale semantics we extract patch tokens from intermediate layers and concatenate them. Because this concatenated descriptor is high-dimensional, forecasting it directly would make the transition model costly to train and prone to overfitting, so we first compress it. A frozen, offline-fitted PCA then projects each descriptor onto its top K principal components. This yields a compact state that keeps the transition model tractable while retaining the biological detail relevant to drug response.

Latent transition model: masked transformer, mask tokens and 3D position. The observed states are processed by a masked transformer that serves as the world model’s transition function. Forecasting is framed as masked prediction: every patch token of the terminal frame is replaced by a shared, learnable mask token $m \in \mathbb{R}^K$, while the observed context tokens pass through unchanged. Since this token carries no location of its own, we add factorized 3D positional embeddings – separate spatial and temporal components – to every token, so each masked query knows where and when its missing patch belongs. Self-attention then lets the masked queries attend over the full observed spatio-temporal context to infer the terminal state $\hat{s}_T$, trained to match the true s_T in the compressed feature space. At inference, only the $T - 1$ observed frames are available, and the model forecasts $\hat{s}_T$.

3.2 Downstream Tasks

Functional outcome head: ATP regression. The ATP head estimates terminal cell viability and, by design, always ingests exactly two latent frame states: the initial frame s_1 and a terminal state $s^\star$. We intentionally fix the two-frame interface so that the head architecture remains identical across settings, isolating the impact of replacing the observed terminal state with its forecast. Following a Multiple Instance Learning (MIL) [31], each frame is summarized into a

compact descriptor by a convolution and global average pooling over its patch tokens; the two descriptors are then concatenated and passed to a 4-layer MLP that regresses the scalar ATP value:

$$\hat{y} = h_{\psi}(s_1, s^*). \quad (2)$$

The choice of terminal state s^* defines the two settings that we compare throughout Sec. 4. Without forecasting, s^* is the last observed frame s_{T-1} ; the head uses only real data, so a prediction can only be made once that frame has been physically acquired.

$$\hat{y}_{\text{nf}} = h_{\psi}(s_1, s_{T-1}). \quad (3)$$

With forecasting, s^* is the transition model’s predicted terminal state $\hat{s}_T$, the virtual endpoint, so the same head reads a state that has not yet occurred,

$$\hat{y}_{\text{fc}} = h_{\psi}(s_1, \hat{s}_T) \text{ with } \hat{s}_T = f_{\theta}(s_1, \dots, s_{T-1}). \quad (4)$$

Because both regimes share the identical first/last-frame interface and weights, any performance gap is attributable solely to forecasting the endpoint rather than waiting to observe it.

Observation decoder: image reconstruction. To visually validate the world model’s foresight, a convolutional decoder maps the PCA-compressed feature grid back to pixels. Following a U-Net-inspired design [32], it progressively up-samples the latent grid via bilinear interpolation and convolutions to a grayscale output, terminating in a Tanh layer mapping to $[-1, 1]$. It is trained with a composite loss

$$\mathcal{L}_{\text{total}} = \lambda_{L1} \mathcal{L}_{L1} + \lambda_{\text{LPIPS}} \mathcal{L}_{\text{LPIPS}}, \quad (5)$$

where $\mathcal{L}_{L1}$ ensures global intensity fidelity and $\mathcal{L}_{\text{LPIPS}}$ [33] penalizes structural discrepancies via deep perceptual features, preserving high-frequency morphological cues important for downstream regression.

3.3 Implementation details

Features were extracted from layers 2, 5, 8, and 11 of a frozen DINOv2 ViT-B/14 (with registers) backbone. An offline PCA reduced concatenated descriptors from 3,072 to 768 dimensions, preserving 91% of the total variance, ensuring that essential morphological characteristics remained intact. For the forecasting task, we utilize a Masked Transformer with 12 layers, 8 attention heads, and a hidden dimension of 768. The model was optimized with AdamW (lr = 1e−4, weight decay = 0.05) over 20 epochs. Masking is deterministic: all N terminal-frame tokens are masked and all context tokens kept visible.

The ATP head follows the two-frame interface of Sec. 3.2: the tokens of s_1 and s^* are each reduced to a descriptor by a convolution and global average pooling, concatenated, and processed by a 4-layer MLP. Each hidden layer of the MLP incorporates Batch Normalization, ReLU activation, and a Dropout layer

($p=0.3$). This head was trained for 30 epochs using an L1 loss, optimized via AdamW ($\text{lr} = 1\text{e-}3$, $\text{weight decay}=1\text{e-}3$) with a batch size of 16, on observed terminal states s_T , and applied unchanged with $s^* = s_{T-1}$ or $s^* = \hat{s}_T$. For visual reconstruction, we used a composite loss including a Learned Perceptual Image Patch Similarity (LPIPS) component based on a pre-trained AlexNet backbone. Both the DINOv2 encoder and the AlexNet weights remained frozen during training to ensure the model optimized solely for temporal forecasting and decoding. We empirically set the loss weights to $\lambda_{L1} = \lambda_{\text{LPIPS}} = 0.5$ to balance low-level intensity accuracy with high-level feature fidelity.

All the above models were implemented in PyTorch [34] and trained on a single NVIDIA H100 GPU equipped with 80 GB of VRAM.

4 Results and discussion

4.1 Dataset and preprocessing

Our cohort of 53 patients (26 colorectal, 27 pancreatic), exposed to a panel of 13 drugs, is split at the patient level into 37 patients for training, 10 for validation, and 6 for independent testing, corresponding to 64,480, 20,571, and 13,045 organoid videos, respectively. The forecasting model is trained on these individual organoid videos, treating each as an independent time-lapse sequence. Each raw sequence contains 50 frames acquired over 96 hours. To efficiently capture long-range morphological dynamics while reducing computational cost, we uniformly downsample each sequence to $T = 8$ representative frames. This choice is motivated by the ablation in Table 1 : beyond 7 context frames, cosine similarity improves by less than 0.002 while epoch time grows $2.6\times$. At the end of the experiment (96 h), drug response is quantified by ATP luminescence. Each organoid video is associated with a single ATP measurement, providing a biological endpoint for every forecasted sequence.

Table 1: Ablation on the number of context frames. The 49th (last) frame is forecast from 1 to 13 context frames, uniformly sampled from the 50-frame videos. Bold marks the retained configuration, not the per-column best: beyond $T - 1 = 7$, epoch time grows $2.6\times$ for a 1.2% L1 gain.

Number of context frames ($T - 1$)	L1 ↓	RMSE ↓	Cosine Similarity ↑	Time for one epoch (min) ↓
1	0.404	0.559	0.955	15
4	0.346	0.379	0.969	35
7	0.324	0.319	0.974	62
10	0.324	0.310	0.975	94
13	0.320	0.297	0.976	163

4.2 Forecasting vs. no forecasting

The central question is whether forecasting the terminal state ($\hat{y}_{fc}$) is better than simply observing for longer and predicting ATP from the last available frame $\hat{y}_{nf}$. Using the two-frame ATP head defined above, we vary the observation time before prediction and compare both regimes against the ceiling given by the real terminal frame, $\hat{y}_{gt}$.

Figure 2 shows that forecasting matches or outperforms no forecasting across all three metrics and essentially all horizons. At the very earliest observation times (around 12–23h) the two settings stay close: from only a handful of initial frames, forecasting the terminal state is too challenging to provide a clear advantage. The gap then opens at intermediate horizons, where anticipation matters most – e.g. at ~ 35 h forecasting lifts Pearson r from 0.32 to 0.49 and reduces MAE from 0.30 to 0.25, with non-overlapping error bars. As observation time grows the curves converge again, since predicting an already-imminent endpoint offers diminishing returns. Around 59h the forecasting model reaches $r \approx 0.60$, $MAE \approx 0.23$ and $C\text{-index} \approx 0.70$ while using roughly 39% less observation time than the 96h assay. This 59h virtual endpoint retains about 95% of the ceiling Pearson correlation and comes within 1.4% of the oracle MAE (0.225 vs. 0.222), advancing the read-out by roughly one working day while leaving the culture intact. Forecasting thus recovers most of the information of the unobserved endpoint: the transition model has learned latent dynamics predictive enough to substitute a simulated endpoint for a physical one. Per-strategy numbers for all sampling configurations are reported in the supplementary material (Table S1).

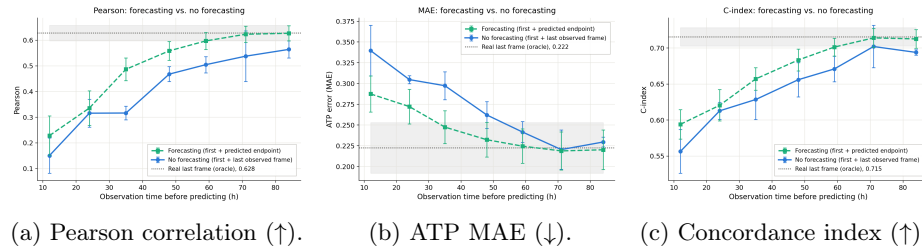


Fig. 2: Forecasting vs. no forecasting as a function of observation time before prediction. Green: the ATP head reads the first frame and the predicted endpoint; blue: the head reads the first and the last observed frame (no forecasting). The dotted line is the ceiling using the real 96h frame. Error bars show ± 1 standard deviation across random seeds.

4.3 Decoder results

Decoding the predicted terminal state into pixels (Fig. 3), distant horizons yield a diffuse, blurry stain, matching the high MSE and the failure of ATP regres-

sion there, whereas our chosen sampling (≈ 59 h) gives anatomically grounded reconstructions closely mirroring the ground truth (GT).

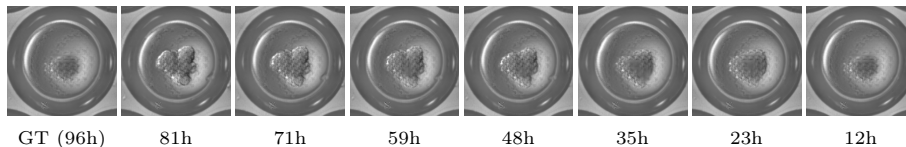


Fig. 3: Qualitative validation of terminal-frame forecasting. Values are the observation time $t_i = i \Delta t$ of the last context frame (indices 42, 37, 31, 25, 18, 12, 6). The proposed sampling maintains structural integrity, whereas forecasting from very early frames yields a blurry semantic stain.

5 Conclusion

We reframed self-supervised feature forecasting as a compact organoid world model, where a DINOv2 observation model and a masked-transformer transition model turn a predicted latent endpoint into an actionable drug-response estimate. Forecasting beats a no-forecasting baseline across Pearson r , MAE and C-index, and recovers most of the endpoint information while cutting observation time by roughly 39% (96 h $\rightarrow$ 59 h) – enabling faster, non-destructive screening and offering a novel organoid world model. Our evaluation nonetheless uses a single private cohort with $n = 6$ test patients; the reported variability across random seeds reflects training stability rather than patient-level generalization, so external multi-centre validation on more patients remains necessary, and generalization to other types of organoid time-lapse modalities is untested. All drugs of the panel are moreover seen during training, so generalization to unseen compounds, comparison against hand-crafted morphological descriptors, and external benchmarking remain open.

Data and code availability. Code will be released upon acceptance; the cohort is governed by patient-consent and institutional agreements and is available on reasonable request under a data-use agreement.

Acknowledgments. This work has benefited from a government grant managed by the National Research Agency (ANR), integrated into France 2030 with the reference ANR-21-RHUS-0003. This work was performed using HPC resources from GENCI–IDRIS.

Disclosure of Interests. The authors have no competing interests to declare that are relevant to the content of this article.

References

1. Vlachogiannis, G., et al.: Patient-derived organoids model treatment response of metastatic gastrointestinal cancers. *Science* 359(6378), 920–926 (2018)
2. Polak, R., et al.: Cancer organoids 2.0: modelling the complexity of the tumour immune microenvironment. *Nat. Rev. Cancer* 24(9), 617–636 (2024)
3. Cartry, J., et al.: Leveraging large-scale PDO-based assays to optimize antibody-drug conjugate efficacy in digestive cancer. In: *Proc. AACR* (2024)
4. Shi, X., et al.: Convolutional LSTM Network: A Machine Learning Approach for Precipitation Nowcasting. In: *Proc. NeurIPS* (2015)
5. Lotter, W., et al.: Deep Predictive Coding Networks for Video Prediction and Unsupervised Learning. In: *Proc. ICLR* (2017)
6. Ho, J., et al.: Video Diffusion Models. *arXiv preprint arXiv:2204.03458* (2022)
7. Cohen, J.P., et al.: Distribution Matching Losses Can Hallucinate Features in Medical Image Translation. In: *Proc. MICCAI*, pp. 529–536 (2018)
8. Moyaert, P., et al.: Risks of Hallucination in Diffusion-Based Medical Image Synthesis. *Medical Image Analysis* (2024)
9. Ha, D., Schmidhuber, J.: Recurrent World Models Facilitate Policy Evolution. In: *Proc. NeurIPS* (2018)
10. Hafner, D., Pasukonis, J., Ba, J., Lillicrap, T.: Mastering Diverse Domains through World Models. *arXiv preprint arXiv:2301.04104* (2023)
11. LeCun, Y.: A Path Towards Autonomous Machine Intelligence, Version 0.9.2. *OpenReview* (2022)
12. Bardes, A., Ponce, J., LeCun, Y.: Revisiting Feature Prediction for Learning Visual Representations from Video. *arXiv preprint arXiv:2404.08471* (2024)
13. Karypidis, E., Kakogeorgiou, I., Gidaris, S., Komodakis, N.: DINO-Foresight: Looking into the Future with DINO. In: *Proc. NeurIPS* (2025)
14. Yue, Y., Wang, Y., Jiang, H., Liu, P., Song, S., Huang, G.: EchoWorld: Learning Motion-Aware World Models for Echocardiography Probe Guidance. In: *Proc. CVPR* (2025)
15. Sharma, S., et al.: State-Change Learning for Prediction of Future Events in Endoscopic Videos. *arXiv preprint arXiv:2510.12904* (2025)
16. Oquab, M., et al.: DINOv2: Learning Robust Visual Features without Supervision. *Trans. Mach. Learn. Res.* (2023)
17. Chen, R.J., et al.: Towards a General-Purpose Foundation Model for Computational Pathology. *Nature Medicine* **30**, 312–325 (2024)
18. Moutakanni, T., Bojanowski, P., et al.: Advancing Human-Centric AI for Robust X-ray Analysis Through Holistic Self-supervised Learning. In: *Proc. MICCAI*, pp. 25–35. Springer (2024)
19. Huang, Y., et al.: Comparative Analysis of ImageNet Pre-Trained Deep Learning Models and DINOv2 in Medical Imaging Classification. *arXiv preprint arXiv:2402.07595* (2024)
20. Baharoon, M., et al.: Towards General Purpose Vision Foundation Models for Medical Image Analysis: An Experimental Study of DINOv2 on Radiology Benchmarks. In: *Proc. MICCAI* (2024)
21. Goodfellow, I., et al.: Generative Adversarial Nets. In: *Proc. NeurIPS* (2014)
22. Isola, P., et al.: Image-to-Image Translation with Conditional Adversarial Networks. In: *Proc. CVPR*, pp. 1125–1134 (2017)
23. Wolterink, J.M., et al.: Deep MR to CT Synthesis Using Adversarial Training. *IEEE Trans. Med. Imaging* **36**(12), 2578–2589 (2017)

24. Nie, D., et al.: Medical Image Synthesis with Context-Aware Generative Adversarial Networks. *Medical Image Analysis* 48, 178–192 (2018)
25. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-Resolution Image Synthesis with Latent Diffusion Models. In: *Proc. CVPR* (2022)
26. Ho, J., Jain, A., Abbeel, P.: Denoising Diffusion Probabilistic Models. In: *Proc. NeurIPS* (2020)
27. Pinaya, W.H., et al.: Brain Imaging Generation with Latent Diffusion Models. In: *MICCAI Workshop on Deep Generative Models* (2022)
28. Zhao, Y., et al.: Deep Learning for Predicting Tumor Growth from Medical Images. In: *Proc. MICCAI* (2017)
29. Fillioux, L., Gontran, E., Cartry, J., Mathieu, J.R., Bedja, S., Boilève, A., Cournède, P.H., Jaulin, F., Christodoulidis, S., Vakalopoulou, M.: Spatio-temporal Analysis of Patient-Derived Organoid Videos Using Deep Learning for the Prediction of Drug Efficacy. In: *Proc. ICCV Workshops*, pp. 3930–3939 (2023)
30. Darcet, T., Oquab, M., Mairal, J., Bojanowski, P.: Vision Transformers Need Registers. *arXiv preprint arXiv:2309.16588* (2023)
31. Ilse, M., Tomczak, J.M., Welling, M.: Attention-based Deep Multiple Instance Learning. In: *Proc. ICML*, pp. 2127–2136 (2018)
32. Ronneberger, O., Fischer, P., Brox, T.: U-Net: Convolutional Networks for Biomedical Image Segmentation. In: *Proc. MICCAI*, pp. 234–241 (2015)
33. Zhang, R., Isola, P., Efros, A.A., Shechtman, E., Wang, O.: The Unreasonable Effectiveness of Deep Features as a Perceptual Metric. In: *Proc. CVPR*, pp. 586–595 (2018)
34. Paszke, A., et al.: PyTorch: An Imperative Style, High-Performance Deep Learning Library. In: *Proc. NeurIPS* (2019)

Supplementary Material

Table S1: Comprehensive evaluation of the pipeline under different context-frame sampling strategies. Foresight fidelity: cosine similarity (CS), MSE, L1 between predicted and ground-truth last-frame embeddings. ATP prediction is reported with MAE, Pearson r and concordance index (C) for the no-forecasting baseline ($\text{MAE}_{\text{nf}}, r_{\text{nf}}, C_{\text{nf}}$; ATP from the observed frames) and for forecasting ($\text{MAE}_{\text{pred}}, r_{\text{pred}}, C_{\text{pred}}$; ATP from the predicted endpoint). The window-independent oracle (first + real last frame) attains MAE 0.222 ± 0.031 , r 0.628 ± 0.030 and C 0.715 ± 0.013 . Visual fidelity: PSNR and SSIM. All ATP entries are mean $\pm$ standard deviation over 3 seeds; bold marks the chosen 59 h operating point. This is the numeric counterpart of Fig. 2.

Context Frame Indices	Foresight Fidelity			ATP Prediction						Pixel-wise Prediction	
	CS $\uparrow$	MSE $\downarrow$	L1 $\downarrow$	$\text{MAE}_{\text{nf}} \downarrow$	$r_{\text{nf}} \uparrow$	$C_{\text{nf}} \uparrow$	$\text{MAE}_{\text{pred}} \downarrow$	$r_{\text{pred}} \uparrow$	$C_{\text{pred}} \uparrow$	PSNR $\uparrow$	SSIM $\uparrow$
[0, 7, 14, 21, 28, 35, 42]	0.974	0.319	0.324	0.230 ± 0.006	0.564 ± 0.033	0.694 ± 0.004	0.220 ± 0.024	0.627 ± 0.029	0.713 ± 0.013	23.04	0.802
[0, 6, 12, 18, 24, 30, 37]	0.971	0.363	0.339	0.220 ± 0.024	0.538 ± 0.099	0.702 ± 0.029	0.219 ± 0.023	0.623 ± 0.032	0.714 ± 0.013	21.95	0.763
[0, 5, 10, 15, 20, 25, 31]	0.967	0.404	0.353	0.242 ± 0.013	0.505 ± 0.032	0.671 ± 0.018	0.225 ± 0.021	0.597 ± 0.034	0.701 ± 0.012	21.27	0.742
[0, 4, 8, 12, 16, 20, 25]	0.966	0.421	0.355	0.262 ± 0.016	0.468 ± 0.029	0.656 ± 0.024	0.232 ± 0.021	0.558 ± 0.037	0.683 ± 0.015	20.34	0.708
[0, 3, 6, 9, 12, 15, 18]	0.963	0.461	0.371	0.297 ± 0.017	0.316 ± 0.026	0.629 ± 0.028	0.247 ± 0.020	0.487 ± 0.044	0.657 ± 0.016	19.53	0.655
[0, 2, 4, 6, 8, 10, 12]	0.961	0.482	0.378	0.305 ± 0.005	0.316 ± 0.054	0.613 ± 0.012	0.272 ± 0.021	0.336 ± 0.067	0.621 ± 0.022	19.78	0.662
[0, 1, 2, 3, 4, 5, 6]	0.960	0.489	0.378	0.340 ± 0.030	0.151 ± 0.069	0.556 ± 0.030	0.287 ± 0.022	0.228 ± 0.077	0.594 ± 0.020	19.23	0.647