
LEARNED CONTROL AS THE NEXT PID, AND ITS INFORMATION-THEORETIC LIMITS

A PREPRINT

Jad Nohra*

Abstract

We advance a claim: learned (deep-learning and reinforcement-learning) control is replacing classical robust control, H_∞ synthesis and μ -synthesis, for the same reason deep convolutional networks displaced hand-engineered computer vision pipelines (SIFT, HOG, the deformable parts model). Both are instances of one mechanism: replacing a fixed, pre-committed signal/noise partition imposed by a hand-designed processing layer with a learned partition, thereby recovering predictive structure that the fixed partition discarded as noise. The conservation laws of control are fundamental and stay applicable; the new controllers are bound by them all the same. Our contribution is the cross-domain synthesis and precisely separating which limits are artifacts of an assumption and which are genuine information-theoretic converses. The argument is organized case by case: each claim is stated, given its precise conditions and constants, attributed to its governing assumption, marked as still constraining or relaxed, and cited.

1 PID and dimensional appropriateness

The most widely deployed control technology assumes the least about the plant. The proportional-integral-derivative (PID) controller dominates industrial practice: Åström and Hägglund report that more than 90% of all control loops are PID, most of them in fact PI because derivative action is frequently unused (Åström and Hägglund 2001). The higher figure derives from the same authors’ textbook, which states that in process control more than 95% of the control loops are of PID type, most loops actually PI, backed by a cited paper-mill audit finding that a typical mill has more than 2,000 control loops, that 97% use PI control, and that only 20% of the loops were found to work well (Åström and Hägglund 1995). The puzzle is that this three-parameter controller, which encodes no plant model, routinely outperforms in deployment the sophisticated robust-control machinery built precisely to handle modeled uncertainty.

The account we offer is dimensional appropriateness. PID is a three-parameter approximation to the inverse of an uncharacterized plant: the integral term cancels steady-state error, the proportional term sets loop gain, the derivative term supplies phase lead. It works because the space of plants worth controlling is low-dimensional in the frequency band where control authority matters, near and below crossover, and three parameters approximately span the controller subspace adequate for that band. Where the plant’s relevant behaviour is captured by a low-order model, and most industrial loops, dominated by a single time constant plus delay, are, three degrees of freedom suffice.

A correction is needed right away, because the naive reading of what follows is wrong. The low-dimensional space did not disappear and was not re-dimensional. Learning does not make the easy plants harder or convert three-parameter problems into million-parameter problems. It annexes new territory: nonlinear, contact-rich, sensor-saturated plants that were previously not controlled by any principled method at all, because no tractable model existed to design against. The claim is not that learned control replaces PID on a flow loop. It is that learned control extends the same move, match controller expressiveness to the

*Psycho B. Tron, a large-language-model-based AI assembled per the goals described in the author’s *Who is Psychotron* (Nohra 2023), served as a research and drafting collaborator on this paper. All claims and errors are the author’s own.

dimensionality of the plant in the band where control matters, into a regime where the three-parameter span no longer suffices. PID remains the appropriate member of the controller class for low-dimensional plants. This framing, that PID and learned policies are two members of one class distinguished by dimensional appropriateness, organizes Sections 3 and 7.

2 Classical robust control and its conservatism cost

Three tiers of control technology can be distinguished by what they assume and what they guarantee.

The first tier is PID: cheap, model-free, without formal robustness guarantees, dominant on simple plants. The third tier, treated below, is learned control. The intermediate tier is the sophisticated machinery of robust control: H_∞ synthesis, μ -synthesis (structured singular value synthesis), and the associated linear-matrix-inequality and Riccati-equation apparatus. This section argues that the intermediate tier’s sophistication is, to a substantial degree, the conservatism cost of a single structural assumption: that the controller is a fixed linear constant-coefficient operator.

Concretely: H_∞ control minimizes the worst-case L_2 -gain (the H_∞ norm) from exogenous disturbance to regulated output, over an a priori uncertainty set, subject to the controller being linear and time-invariant. The design machinery exists to certify that the linear choice survives the worst case in the uncertainty set. μ -synthesis extends this to structured uncertainty, alternating between H_∞ synthesis and scaling to bound the structured singular value μ .

The key fact is that the structured singular value is in general computationally intractable. Braatz, Young, Doyle, and Morari proved that the μ computation problem is NP-hard for purely real and for mixed real/complex uncertainty (Braatz, Young, Doyle, and Morari 1994). The authors state the consequence directly: it is “futile to pursue exact methods for calculating” μ of general systems with real or mixed uncertainty beyond small problems. The practical design flow therefore works with upper bounds on μ , which are by construction conservative. The sophistication of classical robust control, the scalings, the bounds, the iterations, is largely the price of proving a fixed linear controller meets a worst-case specification over an uncertainty set whose exact analysis is intractable.

The objection answered in Section 6 should be stated now. What classical robust control offered was never performance on simple plants; it was guarantees. H_∞ and μ -synthesis deliver certified robust stability and performance margins against bounded uncertainty. PID delivers no such certificate. If the comparison is performance in deployment, classical robust control loses to both PID (on simple plants) and learning (on complex ones); but if the comparison is certified worst-case behaviour, it was offering a different product. Section 6 answers this objection by showing that the guarantee was coupled to the linear-controller assumption and that the two are separable.

3 Learning as the same move at higher dimension

Learned control is PID’s move executed at higher dimension. A high-capacity nonlinear function approximator, a deep network policy, spans the nonlinear, sensor-rich, hard-to-model plants where the three-parameter span is exhausted. The principle is identical: match controller expressiveness to plant complexity in the band where control matters. PID spans the low-dimensional case with three parameters; a deep policy spans the high-dimensional case with many. Neither carries a plant model in the classical sense; both are, in effect, learned or tuned inverses operating on the relevant signal.

The empirical instantiation is legged locomotion, and it must be scoped precisely. Reinforcement-learning sim-to-real control became the dominant paradigm for quadrupedal locomotion over roughly 2019–2023. Hwangbo et al. demonstrated learned dynamic motor skills transferred from simulation to the ANYmal quadruped (Hwangbo et al. 2019). Lee et al. demonstrated a learned controller for blind quadrupedal locomotion over challenging terrain that was deployed in field conditions, including at the DARPA Subterranean Challenge, where it replaced a model-based controller (Lee et al. 2020). Miki et al. added exteroceptive perception fused with proprioception (Miki et al. 2022). The throughput enabling this was GPU-parallel simulation: Rudin et al. trained policies for flat terrain in under four minutes and for uneven terrain in twenty minutes on a single workstation GPU using thousands of parallel environments in Isaac Gym, a speedup of several orders of magnitude over previous work (Rudin et al. 2022).

Several corrections must be observed, to avoid overclaiming.

First, the displaced competitor. The model-based controller that learned locomotion displaced at the quadruped frontier was whole-body and convex model-predictive control, exemplified by the MIT Cheetah 3 work (Di Carlo et al. 2018). It was not the older zero-moment-point paradigm (Vukobratović) nor hybrid zero dynamics (Westervelt, Grizzle, and Koditschek 2003), both of which were already legacy by the time learning arrived. The accurate statement is that learning displaced the contemporary model-based frontier, not a weaker predecessor.

Second, coexistence. Model-predictive control was not eliminated. Hybrid architectures combining learned policies with MPC persist, and some deployed quadrupeds remained MPC-based. The transition is a shift in dominant paradigm at the research frontier, not a clean replacement everywhere.

Third, the bipedal and humanoid shift arrived later, roughly 2024–2025, and is less settled. Radosavovic et al. demonstrated reinforcement-learning control of a full-sized humanoid transferred to hardware (Radosavovic et al. 2024). The humanoid case is more recent and the displacement less complete than for quadrupeds.

The historical precedents that make this a recognizable pattern rather than a novelty are vision and speech. In vision, AlexNet won the ILSVRC-2012 competition with a top-5 test error of 15.3%, against 26.2% for the second-best, non-deep entry (Krizhevsky, Sutskever, and Hinton 2012). In speech recognition, deep acoustic models displaced Gaussian-mixture-model systems slightly earlier, around 2009–2012, with deep networks outperforming the established systems across benchmarks, sometimes by a large margin (Hinton et al. 2012). The cross-domain pattern, a fixed hand-engineered feature stage replaced by a learned one, is the subject of Section 5.

4 The information-set argument

This is the central section. The claim is that every fundamental control-performance limit is conditional on the information set the estimator conditions on, and that learning enlarges that set. The gain from learning is the predictive structure recovered when the information set grows; the limits that do not depend on the information set are untouched.

4.1 The minimum-variance bound

The Åström minimum-variance benchmark is the natural starting point. For a linear plant with d -step input-output delay driven by a stochastic disturbance with innovations representation, the minimum achievable output variance under any controller is the variance of the d -step-ahead prediction error:

$$\text{Var}_{\min} = \sigma_e^2 (1 + \psi_1^2 + \psi_2^2 + \cdots + \psi_{d-1}^2), \quad (1)$$

where σ_e^2 is the innovation variance and the ψ_i are the impulse-response coefficients of the disturbance model (Åström 1970). This bound is the mean-square error of the optimal d -step predictor conditional on the available measurement history. It is the basis of the Harris control-performance-assessment benchmark (Harris 1989), which compares achieved variance to this minimum-variance bound from routine operating data.

The point the argument turns on: this bound is the minimum-mean-square-error residual conditional on the information set. It is not a constant of the plant. Change what the estimator conditions on, enlarge the information set with additional sensors, longer history, or nonlinear features of the raw stream, and the conditional residual can only stay equal or shrink. The bound moves with the information set.

4.2 Discarded structure

A linear Kalman observer conditions on the linear-Gaussian projection of the measurement stream. Under the linear-Gaussian assumption this is the sufficient statistic and nothing is lost. But when the disturbance has nonlinear or non-Gaussian structure, when, for example, the noise entering a legged robot’s state estimate is in fact a deterministic function of terrain geometry visible in an unused camera stream, the linear observer discards that structure as noise. An estimator that conditions on the raw stream and is permitted nonlinear functions of it recovers the structure if it is predictive. This is exactly what a deep policy operating on raw sensor histories does: it conditions on a larger information set than the linear-Gaussian projection.

4.3 Established results

The proposition that the disturbance-attenuation limit relaxes by the information rate of side information is not this essay’s invention; it is an established theorem, and must be cited as such with its exact condition.

Martins, Dahleh, and Doyle established a Bode-like disturbance-attenuation bound for feedback systems in which the controller has finite-horizon preview of the disturbance supplied through a communication channel (Martins, Dahleh, and Doyle 2007). The result is a conditional equality: under asymptotic-stationarity assumptions, the new fundamental limitation differs from Bode’s by a constant equal to the information rate supplied through the side channel, which in the Gaussian, asymptotically-stationary attained case equals the channel capacity C . This must be stated as a conditional equality, the relaxation equals the side information rate under stated conditions, and not as an unconditional “exactly C ”.

The estimation-side bridge is the I-MMSE relationship. Guo, Shamai, and Verdú proved that for an arbitrary finite-power input observed through additive Gaussian noise, the derivative of the input-output mutual information with respect to signal-to-noise ratio equals one half the minimum mean-square error, in nats:

$$\frac{d}{d \text{snr}} I(\text{snr}) = \frac{1}{2} \text{mmse}(\text{snr}), \quad (2)$$

regardless of the input distribution (Guo, Shamai, and Verdú 2005). This identity is the formal hinge connecting an information quantity (what an enlarged information set provides) to an estimation-error quantity (the operative bound).

The I-MMSE machinery has been applied directly to control and filtering limitations. Wan, Li, and Hovakimyan lift information-theoretic analysis to continuous function spaces by the I-MMSE relationships, modelling control and filtering systems as additive Gaussian channels and identifying a total information rate as the trade-off metric (Wan, Li, and Hovakimyan 2022); the discrete-time companion adds Voulgaris (Wan, Li, Hovakimyan, and Voulgaris 2023). These connect the estimation bound to a mutual-information rate; the claim here is restricted to that connection and does not overstate it.

Two further established results fix the kind of guarantee in play. Kostina and Hassibi characterized the rate-cost trade-off in control: the minimum achievable linear-quadratic cost as a function of the directed information rate from observer to controller (Kostina and Hassibi 2019). Fang, Chen, and Ishii bound feedback control error below by the conditional entropy of the disturbance and discuss learning-based control explicitly in this frame (Fang, Chen, and Ishii 2019).

4.4 Partition

This is the distinction the entire argument turns on.

Partition. Enlarging the information set lowers the disturbance-attenuation bound and has no effect on the stabilization converse. The reason is structural: the stabilization limits contain no observable-disturbance term, so there is nothing in them to lower by conditioning on more data. Table 1 states the principle.

	Disturbance bound	Stabilization converse
Conditional on	the information set	the plant’s unstable poles
Under more information	lowered	unchanged (no effect)
Governing quantities	min-variance bound; Bode side-information relaxation	Bode integral $\pi \sum_i \text{Re}(p_i)$; data-rate bound $R > \sum_i \log_2 \lambda_i $

Table 1: Partition. The disturbance bound is conditional on the information set, so enlarging that set lowers it; the stabilization converse depends only on the plant’s unstable poles, so no estimator, however rich, can affect it. The two are separated by the partition.

Take each side.

On the disturbance side, Bode’s sensitivity integral. For a continuous-time SISO LTI feedback system with open-loop transfer function $L(s)$, sensitivity $S(s) = 1/(1 + L(s))$, the integral of the log-sensitivity over frequency is

$$\int_0^\infty \ln |S(j\omega)| d\omega = \pi \sum_i \text{Re}(p_i) - \frac{\pi}{2} \lim_{s \rightarrow \infty} sL(s), \quad (3)$$

where the p_i are the open-right-half-plane poles of L (Bode 1945; modern treatment in Seron, Braslavsky, and Goodwin 1997). The final term $-\frac{\pi}{2} \lim_{s \rightarrow \infty} sL(s)$ is non-zero exactly when L has relative degree 1; for relative degree at least 2 the limit vanishes and the integral reduces to the familiar $\pi \sum_i \text{Re}(p_i)$. The statement requires internal stability, S has no right-half-plane poles, and is a SISO statement. This is the waterbed: sensitivity reduction in one band is paid for by sensitivity increase in another, with the total fixed by the unstable poles.

The point for the partition: Bode’s integral depends on the open-RHP poles of the plant and on the relative degree, properties of the plant’s dynamics, and on nothing about the disturbance spectrum or what the estimator observes. One cannot lower a quantity that does not depend on what one conditions on. Enlarging the information set therefore cannot relax it.

On the stabilization side, the data-rate theorem. For a linear plant stabilized over a finite-rate digital channel, mean-square (or almost-sure) stabilizability holds if and only if the channel rate R exceeds the sum of the logarithms of the unstable eigenvalues:

$$R > \sum_i \log_2 |\lambda_i| \quad (\text{bits per sample}), \quad (4)$$

where the λ_i are the eigenvalues of the open-loop system matrix with modulus at least 1 (Nair and Evans 2004; Tatikonda and Mitter 2004). The rate is in bits per sample. The converse is invariant under nonlinear, time-varying, and adaptive coding: no coder, however clever, stabilizes below this rate. Again the bound depends only on the unstable eigenvalues, the plant’s expansion rate, and not on any observable-disturbance quantity. The partition holds: learning a richer coder or policy cannot beat the data-rate converse.

The information-layer precursor of the terminal-failure distinction is the Sahai-Mitter anytime-capacity result. Stabilizing an unstable process over a noisy channel requires not Shannon capacity but anytime capacity, a notion of capacity parametrized by a reliability that must decay exponentially in delay (Sahai and Mitter 2006). The required rate is the log of the open-loop gain; the required reliability comes from the sense of stability desired. This distinction, that stabilization demands a stronger reliability than mere average throughput, anticipates the matching rule of Section 6: terminal-failure constraints cannot be met by a quantity that is only correct on average.

4.5 Threshold

The information-set argument applies in both directions, which fixes a falsifiable test.

Threshold. The gain from learning is real if and only if predictable structure exists outside the linear sufficient statistics. The measure of the available gain is the conditional-entropy gap

$$h(d | \text{raw history}) \quad \text{versus} \quad h(d | \text{linear sufficient statistics}), \quad (5)$$

where d denotes the disturbance. A large gap means the linear observer is discarding predictive structure; learning that conditions on the raw history recovers it, and the gain is real and quantifiable, bounded by the gap. A zero gap means the disturbance, relative to the raw stream, is true innovation: it is unpredictable from anything the richer estimator can see. In that case learning recovers nothing, and a high-capacity approximator trained to predict it becomes a confident predictor of noise; it fits sampling realizations of an unpredictable process and generalizes poorly. Overparametrization is therefore the threshold of the gain, not its engine. Capacity helps only to the extent that the conditional-entropy gap is positive; beyond that, capacity is spent modelling noise. This is a testable, refutable criterion: estimate the two conditional entropies and the gap predicts whether learning can help. Figure 1 summarizes the relationship.

5 The vision–control isomorphism

Our novel synthesis is that the control transition of Sections 3–4 and the computer-vision transition of the 2010s are the same transition. The correspondence is exact enough to be stated as a sequence of identified objects, theorems, and measures, summarized in Table 2.

5.1 Shared object

The convolution kernel of a CNN and the impulse response of a linear filter in control are the same mathematical object: the kernel of a linear operator. In both fields a hand-designed instance of this object, a

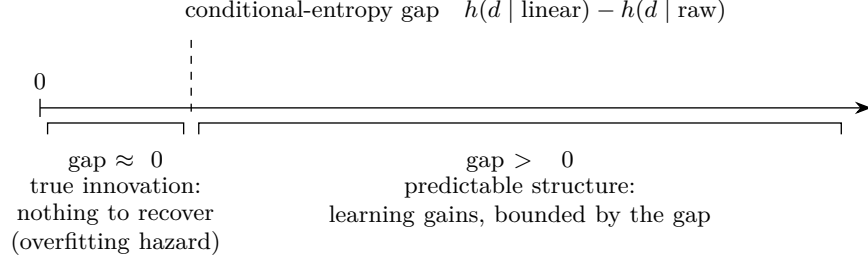


Figure 1: Threshold. Learning gains only where the conditional-entropy gap is positive, that is, where predictable structure exists outside the linear sufficient statistics. Near zero gap the disturbance is true innovation, nothing can be recovered, and a high-capacity model fits noise; the gain appears, bounded by the gap, only where such structure exists.

Vision	Shared mechanism	Control
pixels	raw input	sensor stream
convolution filter	kernel of a linear operator	impulse response
SIFT, HOG, DPM	hand-designed front end	Kalman / linear observer
Fano’s inequality	information-to-error bound	I-MMSE, min-variance
$h(\ell \text{pix})$ vs $h(\ell \text{feat})$	discarded predictive structure	$h(d \text{raw})$ vs $h(d \text{lin})$
deformable parts model	displaced “peak”	H_∞ , μ -synthesis
end-to-end CNN	learned replacement	learned policy

Table 2: One mechanism, two instances. Read across a row: the same structural role, realized in computer vision and in control. In both fields a fixed, hand-designed transform sets a signal/noise partition; learning replaces it and recovers the discarded structure. Control’s extra structure, the stabilization converse, has no vision counterpart (see Table 1).

HOG/SIFT filter bank in vision, a fixed linear pre-filter or observer in control, defines a fixed signal/noise partition. Whatever the fixed kernel passes is signal; whatever it attenuates is noise. The pre-commitment is identical in form.

5.2 Shared theorem

The data processing inequality governs both. For a Markov chain $\text{target} \rightarrow X \rightarrow T(X) \rightarrow \text{decision}$, where X is the raw observation and T a fixed transform,

$$I(\text{target}; T(X)) \leq I(\text{target}; X). \quad (6)$$

A fixed transform can only lose task-relevant information; it cannot create it. The consequence is the same in both fields: a hand-designed front end caps the information available to every downstream stage, and the cap binds whenever the transform discards task-relevant structure. Learning T jointly with the task objective chooses a transform that loses less of what matters.

A careful qualifier is mandatory here, because the data processing inequality is often misread as forbidding exactly what learning does. Turgeman and Tirer show that the “no benefit from pre-processing” reading holds strictly only at the Bayes-optimal limit; for any finite number of training samples there exists a pre-classification processing that strictly improves accuracy (Turgeman and Tirer 2025). The correct statement is therefore: learning does not violate the data processing inequality. The inequality bounds the information content of a representation; it does not say a learned transform cannot beat a fixed transform at delivering that information to a finite-sample, sub-optimal downstream learner. Learning chooses a better transform within the bound; it does not break the bound. This is precisely parallel to the control claim, learning descends the operative bound by conditioning on more, not by repealing any converse.

5.3 Dual information-to-error bridges

Each field has a theorem converting an information quantity into an error bound.

In vision, Fano’s inequality lower-bounds the probability of classification error by the conditional entropy of the label given the features: a representation that retains little information about the label cannot support low error (for a careful treatment of its tightness and use, Zhao, Edakunni, Pocock, and Brown 2013).

In control, the I-MMSE and minimum-variance relations of Section 4 convert a conditional residual into a mean-square-error bound. The two bridges are duals: conditional entropy of label given features (vision) and conditional residual of disturbance given the information set (control) play the same role, they translate how much the representation retained into how well the downstream task can possibly be done.

5.4 Shared headroom measure

The headroom for learning is the same quantity in both fields. In vision it is $h(\text{label} \mid \text{pixels})$ versus $h(\text{label} \mid \text{features})$: the gap between what the raw image determines about the label and what the hand-crafted features retain. In control it is $h(d \mid \text{raw stream})$ versus $h(d \mid \text{linear sufficient statistics})$, the conditional-entropy gap of Section 4.5. In both, a positive gap is exactly the structure a fixed front end threw away and a learned front end can recover; a zero gap means the hand-crafted features were already sufficient and learning gains nothing.

5.5 Displaced peak

The deformable parts model occupies the same position in vision that H_∞ and μ -synthesis occupy in control: the most elaborate stage of the hand-engineered paradigm. It was the most elaborate and highest-performing pre-deep detector, described in survey treatments as the epitome of traditional object detection (Felzenszwalb, Girshick, McAllester, and Ramanan; characterization in Zou et al. 2023). It was not displaced by a better deformable parts model. It was displaced by deleting the paradigm beneath it, the hand-crafted feature stage, and learning the features end-to-end. The control parallel is exact: classical robust control is not being out-competed by a more refined H_∞ ; the fixed-linear-controller paradigm beneath it is being deleted and replaced by a learned policy.

5.6 Position relative to prior art

Novelty must be claimed carefully. Every component of this synthesis is individually established. The data processing inequality, Fano’s inequality, the I-MMSE identity, the minimum-variance bound, the information-bottleneck view of representation, and the empirical vision and control transitions are all in the literature. What appears unoccupied, to our knowledge, is the explicit identification of the two transitions as one mechanism, with the partition, the threshold, and the matching rule as the delineation of what transfers and what does not.

The nearest neighbour must be cited and distinguished. Achille and Soatto give a separation principle for control in the age of deep learning, bridging learned visual representations and Kalman-filter replacement through the information-bottleneck Lagrangian (Achille and Soatto 2018). The distinction: they use an information-bottleneck-Lagrangian separation principle to define a minimal task-sufficient state; this essay uses the data-processing-inequality conditional-entropy gap together with the dual Fano and I-MMSE bridges and the explicit historical parallel, and centers the partition between disturbance bound and stabilization converse, which the separation-principle framing does not address.

The information-bottleneck lineage must likewise be cited and distinguished (Tishby, Pereira, and Bialek 1999; Tishby and Zaslavsky 2015; Alemi, Fischer, Dillon, and Murphy 2017). Pacelli and Majumdar apply information bottlenecks to displace the Kalman filter with a task-driven representation (Pacelli and Majumdar 2019). The distinction is sharp: the information-bottleneck program seeks minimality, the smallest representation sufficient for the task, whereas the mechanism here is recovery of discarded predictive structure, enlarging the information set to capture structure a fixed front end threw away. Minimality compresses a sufficient statistic; recovery enlarges an insufficient one. They are complementary, not the same.

Finally, the vision-side unification of hand-crafted and learned features is anticipated by Soatto and collaborators (Soatto 2009; Soatto and Chiuso 2016), and the Turgeman-Tirer qualifier above prevents the data processing inequality from being misread.

5.7 Disanalogies

The places where the analogy fails are not weaknesses; they enumerate exactly the irreducible extra structure control has and vision lacks. Vision has no stabilization converse, there is no Bode integral and no data-rate theorem for image classification, because there is no feedback loop and no unstable plant. Vision’s failures are recoverable (a misclassification costs an error count) whereas control’s can be terminal (an unstable mode diverges). Vision’s filters are spatial; control’s are causal and temporal, which is why directed information and anytime capacity, not plain mutual information and Shannon capacity, are the right measures. Vision is supervised; control is closed-loop and often reinforcement-driven, which breaks the exchangeability that supervised statistics assume. Each disanalogy is precisely one of the structural features that the partition protects: they are why the stabilization converse holds in control even though the representation-learning gain transfers.

6 Separating the guarantee from the linear structure

Return to the objection of Section 2: classical robust control’s product was guarantees, not deployment performance. The answer is that the H_∞ guarantee was real but coupled to the fixed-linear-controller assumption and priced at a conservatism premium; both the coupling and the premium are separable from the guarantee itself. Assurance can now be attached directly to a nonlinear policy.

Instantiate two routes.

Neural Lyapunov certificates. Chang, Roohi, and Gao learn a control policy together with a neural Lyapunov function in a counterexample-guided loop: a learner proposes a candidate policy and Lyapunov function, a falsifier searches for states violating the Lyapunov conditions using a δ -complete SMT solver, and counterexamples are returned to the learner; when the falsifier finds none, the Lyapunov conditions are certified to hold over the verified domain, and the obtained region of attraction can exceed those from LQR or sum-of-squares methods (Chang, Roohi, and Gao 2019). The scalability limitation must be stated plainly: the method has been demonstrated mainly on low-dimensional systems, and verification cost grows rapidly with state dimension; as of this writing it does not certify high-dimensional policies.

Control-barrier-function safety filters. Ames, Xu, Grizzle, and Tabuada construct a quadratic-program safety filter that takes a learned (or any) nominal action and projects it minimally onto the set of actions that keep the state within a provably forward-invariant safe set defined by a control barrier function (Ames, Xu, Grizzle, and Tabuada 2017; survey in Ames et al. 2019). The barrier filter is agnostic to how the nominal action was produced, so it attaches a forward-invariance guarantee to a learned policy. The limitation: constructing a valid control barrier function for a complex, input-constrained system remains a per-system design problem with no general automatic procedure.

Once the guarantee is separable from the linear structure, Lyapunov certification of a nonlinear policy, barrier-function enforcement of a nonlinear policy, classical robust control retains mainly its conservatism: a certified-but-conservative linear controller versus a certified nonlinear policy with less conservatism. This part of the argument is the least settled. The certificate frontier, scaling neural Lyapunov verification and automating barrier-function synthesis to high-dimensional, sensor-rich plants, is the current bottleneck. The displacement of classical robust control on the guarantee axis is therefore a direction supported by working methods on moderate systems, not a finished fact.

6.1 Matching rule

A non-negotiable principle governs which certificate may be used.

Matching rule. A probabilistic certificate and a worst-case certificate are different kinds of object and cannot be substituted. Match the certificate’s kind to the constraint’s failure cost before choosing it.

Statistical certificates, conformal prediction, scenario optimization, are valid instruments for the stochastic performance bound, where failures are recoverable and costs aggregate. They are a category error when applied to a terminal-failure constraint. Adaptive conformal inference, for instance, guarantees long-run time-average coverage under distribution shift, not coverage at any individual decisive time step (Gibbs and Candès 2021). A time-average coverage guarantee says nothing about the one time step at which an unstable mode would diverge. Moreover, feedback breaks the exchangeability that conformal prediction assumes: the controller’s actions change the distribution of future data, so the i.i.d./exchangeable premise fails in closed loop. The Sahai-Mitter anytime-capacity result of Section 4.4 is the information-layer statement of the same

distinction: stabilization requires reliability exponential in delay, not merely a correct long-run average rate. The rule in practice: classify every constraint by failure cost, recoverable/i.i.d. versus terminal, and only then choose the certificate kind. Never certify a terminal-failure constraint with a time-average.

7 What changed and what did not

The three reasons compose into one position.

On the performance axis, learning is the next PID: controller expressiveness matched to plant dimensionality in the band where control matters, with classical robust control bypassed rather than improved. PID is the low-dimensional member of the controller class; a learned policy is the high-dimensional member; the next PID is the class itself, appropriately dimensioned policies with assurance separable from linear structure, not a specific architecture.

On the gain axis, the improvement is real because learning enlarges the information set the estimator conditions on and thereby descends the operative disturbance bound, which is the conditional residual given that set, not a constant of the plant. The size of the available gain is bounded by the conditional-entropy gap (the threshold, Figure 1); where the gap is zero the gain is zero and capacity is spent fitting noise.

On the guarantee axis, assurance is becoming separable from the fixed-linear-controller assumption via Lyapunov and barrier certificates attached to nonlinear policies, leaving classical robust control mainly its conservatism, though this is the least settled part, gated by the certificate frontier.

The scope must be stated precisely. This is not a claim that learning prevails everywhere. It prevails on nonlinear, sensor-rich plants, for the structural reason given, with a measurable threshold (the conditional-entropy gap) marking where it stops. On low-dimensional plants PID remains appropriate. The displacement is of classical robust control on hard plants, not of all classical control everywhere.

What did not change: Bode’s sensitivity integral, the waterbed effect, the data-rate bound $\sum_i \log_2 |\lambda_i|$, the anytime-capacity requirement for stabilization over noisy channels, and that true innovation cannot be reduced, none of these is repealed. They stay constraining. A learned controller obeys Bode’s integral and the data-rate theorem as strictly as a PID controller does; a deep policy cannot stabilize an unstable plant over a channel below the eigenvalue-sum rate, and cannot predict true innovation. What changed is not the physics of feedback. What changed is that one stops paying the cost of linearity to obtain a guarantee, and that the limits classical robust control certified against were, on hard nonlinear plants, evaluated against an impoverished linear-Gaussian information set, a set the learned controller is free to enlarge. The converses remain; the artifacts of the assumption do not.

One limit this essay explicitly does not cross. Learning does not deliver a controller that reasons about its own loop at runtime. The pilot who notices that she is fighting the aircraft and deliberately reduces gain is performing a metacognitive act, a controller forming and acting on a model of its own closed-loop interaction in real time. A learned policy, however high-capacity, is a richer reactive map from observation histories to actions; it is not a controller that represents and reasons about its own control loop. This capability exists in no deployed system, learned or classical. Conflating a more expressive reactive policy with a self-modelling controller would be the one overclaim that the rest of the argument’s discipline forbids.

A final word on the stance this argument expresses. Tao describes a post-rigorous stage of mathematical development in which one has so thoroughly internalized the rigorous foundations of a field that one reasons fluidly with intuition “solidly buttressed by rigorous theory”, and in which rigour is used to discard bad intuition while clarifying and elevating good intuition (Tao 2007). Post-rigour control means exactly this: internalizing Bode’s integral, the data-rate theorem, the minimum-variance bound, and the anytime-capacity requirement so thoroughly that one reasons fluidly within the real limits while declining to pay for the artifacts of a structural assumption that the limits never required. It is not the abandonment of rigour. Command of the converses is what distinguishes the limits that are artifacts of the fixed-linear-controller assumption, and therefore negotiable, from the information-theoretic bounds, and therefore permanent.

References

Achille, Alessandro, and Stefano Soatto. 2018. “A Separation Principle for Control in the Age of Deep Learning.” *Annual Review of Control, Robotics, and Autonomous Systems* 1: 287–307.

- Alemi, Alexander A., Ian Fischer, Joshua V. Dillon, and Kevin Murphy. 2017. “Deep Variational Information Bottleneck.” In *International Conference on Learning Representations (ICLR)*.
- Ames, Aaron D., Xiangru Xu, Jessy W. Grizzle, and Paulo Tabuada. 2017. “Control Barrier Function Based Quadratic Programs for Safety Critical Systems.” *IEEE Transactions on Automatic Control* 62 (8): 3861–3876.
- Ames, Aaron D., Samuel Coogan, Magnus Egerstedt, Gennaro Notomista, Koushil Sreenath, and Paulo Tabuada. 2019. “Control Barrier Functions: Theory and Applications.” In *European Control Conference (ECC)*, 3420–3431.
- Åström, Karl Johan. 1970. *Introduction to Stochastic Control Theory*. Academic Press.
- Åström, Karl Johan, and Tore Hägglund. 1995. *PID Controllers: Theory, Design, and Tuning*. 2nd ed. Instrument Society of America.
- Åström, Karl Johan, and Tore Hägglund. 2001. “The Future of PID Control.” *Control Engineering Practice* 9 (11): 1163–1175.
- Bode, Hendrik W. 1945. *Network Analysis and Feedback Amplifier Design*. New York: Van Nostrand.
- Braatz, Richard P., Peter M. Young, John C. Doyle, and Manfred Morari. 1994. “Computational Complexity of μ Calculation.” *IEEE Transactions on Automatic Control* 39 (5): 1000–1002.
- Chang, Ya-Chien, Nima Roohi, and Sicun Gao. 2019. “Neural Lyapunov Control.” In *Advances in Neural Information Processing Systems* 32.
- Di Carlo, Jared, Patrick M. Wensing, Benjamin Katz, Gerardo Bledt, and Sangbae Kim. 2018. “Dynamic Locomotion in the MIT Cheetah 3 Through Convex Model-Predictive Control.” In *IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 1–9.
- Fang, Song, Jie Chen, and Hideaki Ishii. 2019. “Information-Theoretic Performance Limitations of Feedback Control: Underlying Entropic Laws and Generic L_p Bounds.” arXiv:1912.05541.
- Gibbs, Isaac, and Emmanuel Candès. 2021. “Adaptive Conformal Inference Under Distribution Shift.” In *Advances in Neural Information Processing Systems* 34, 1660–1672.
- Guo, Dongning, Shlomo Shamai (Shitz), and Sergio Verdú. 2005. “Mutual Information and Minimum Mean-Square Error in Gaussian Channels.” *IEEE Transactions on Information Theory* 51 (4): 1261–1282.
- Harris, Thomas J. 1989. “Assessment of Control Loop Performance.” *Canadian Journal of Chemical Engineering* 67: 856–861.
- Hinton, Geoffrey, Li Deng, Dong Yu, George Dahl, Abdel-rahman Mohamed, Navdeep Jaitly, Andrew Senior, Vincent Vanhoucke, Patrick Nguyen, Tara Sainath, and Brian Kingsbury. 2012. “Deep Neural Networks for Acoustic Modeling in Speech Recognition: The Shared Views of Four Research Groups.” *IEEE Signal Processing Magazine* 29 (6): 82–97.
- Hwangbo, Jemin, Joonho Lee, Alexey Dosovitskiy, Dario Bellicoso, Vassilios Tsounis, Vladlen Koltun, and Marco Hutter. 2019. “Learning Agile and Dynamic Motor Skills for Legged Robots.” *Science Robotics* 4 (26): eaau5872.
- Kostina, Victoria, and Babak Hassibi. 2019. “Rate-Cost Tradeoffs in Control.” *IEEE Transactions on Automatic Control* 64 (11): 4525–4540.
- Krizhevsky, Alex, Ilya Sutskever, and Geoffrey E. Hinton. 2012. “ImageNet Classification with Deep Convolutional Neural Networks.” In *Advances in Neural Information Processing Systems* 25, 1097–1105.
- Lee, Joonho, Jemin Hwangbo, Lorenz Wellhausen, Vladlen Koltun, and Marco Hutter. 2020. “Learning Quadrupedal Locomotion over Challenging Terrain.” *Science Robotics* 5 (47): eabc5986.
- Martins, Nuno C., Munther A. Dahleh, and John C. Doyle. 2007. “Fundamental Limitations of Disturbance Attenuation in the Presence of Side Information.” *IEEE Transactions on Automatic Control* 52 (1): 56–66.
- Miki, Takahiro, Joonho Lee, Jemin Hwangbo, Lorenz Wellhausen, Vladlen Koltun, and Marco Hutter. 2022. “Learning Robust Perceptive Locomotion for Quadrupedal Robots in the Wild.” *Science Robotics* 7 (62): eabk2822.

- Nair, Girish N., and Robin J. Evans. 2004. “Stabilizability of Stochastic Linear Systems with Finite Feedback Data Rates.” *SIAM Journal on Control and Optimization* 43 (2): 413–436.
- Nohra, Jad. 2023. *Who is Psychotron*. Kindle ed.; also Audible audiobook, narrated by Christopher Sutton.
- Pacelli, Vincent, and Anirudha Majumdar. 2019. “Task-Driven Estimation and Control via Information Bottlenecks.” In *IEEE International Conference on Robotics and Automation (ICRA)*, 2061–2067.
- Radosavovic, Ilija, Tete Xiao, Bike Zhang, Trevor Darrell, Jitendra Malik, and Koushil Sreenath. 2024. “Real-World Humanoid Locomotion with Reinforcement Learning.” *Science Robotics* 9 (89): eadi9579.
- Rudin, Nikita, David Hoeller, Philipp Reist, and Marco Hutter. 2022. “Learning to Walk in Minutes Using Massively Parallel Deep Reinforcement Learning.” In *Conference on Robot Learning (CoRL)*, PMLR 164: 91–100.
- Sahai, Anant, and Sanjoy Mitter. 2006. “The Necessity and Sufficiency of Anytime Capacity for Stabilization of a Linear System over a Noisy Communication Link, Part I: Scalar Systems.” *IEEE Transactions on Information Theory* 52 (8): 3369–3395.
- Seron, María M., Julió H. Braslavsky, and Graham C. Goodwin. 1997. *Fundamental Limitations in Filtering and Control*. London: Springer-Verlag.
- Soatto, Stefano. 2009. “Actionable Information in Vision.” In *IEEE International Conference on Computer Vision (ICCV)*.
- Soatto, Stefano, and Alessandro Chiuso. 2016. “Visual Representations: Defining Properties and Deep Approximations.” In *International Conference on Learning Representations (ICLR)*.
- Tao, Terence. 2007. “There’s More to Mathematics than Rigour and Proofs.” Blog post, [terrytao.wordpress.com](https://terrytao.wordpress.com/career-advice/theres-more-to-mathematics-than-rigour-and-proofs/). <https://terrytao.wordpress.com/career-advice/theres-more-to-mathematics-than-rigour-and-proofs/>.
- Tatikonda, Sekhar, and Sanjoy Mitter. 2004. “Control under Communication Constraints.” *IEEE Transactions on Automatic Control* 49 (7): 1056–1068.
- Tishby, Naftali, Fernando C. Pereira, and William Bialek. 1999. “The Information Bottleneck Method.” In *Allerton Conference on Communication, Control, and Computing*, 368–377.
- Tishby, Naftali, and Noga Zaslavsky. 2015. “Deep Learning and the Information Bottleneck Principle.” In *IEEE Information Theory Workshop (ITW)*.
- Turgeman, Roy, and Tom Tirer. 2025. “Does the Data Processing Inequality Reflect Practice? On the Utility of Low-Level Tasks.” arXiv:2512.21315.
- Wan, Neng, Dapeng Li, and Naira Hovakimyan. 2022. “An Information-Theoretic Analysis of Continuous-Time Control and Filtering Limitations by the I-MMSE Relationships.” arXiv:2201.00995.
- Wan, Neng, Dapeng Li, Naira Hovakimyan, and Petros G. Voulgaris. 2023. “An Information-Theoretic Analysis of Discrete-Time Control and Filtering Limitations by the I-MMSE Relationships.” arXiv:2304.09274.
- Westervelt, Eric R., Jessy W. Grizzle, and Daniel E. Koditschek. 2003. “Hybrid Zero Dynamics of Planar Biped Walkers.” *IEEE Transactions on Automatic Control* 48 (1): 42–56.
- Zhao, Ming-Jie, Narayanan Edakunni, Adam Pocock, and Gavin Brown. 2013. “Beyond Fano’s Inequality: Bounds on the Optimal F-Score, BER, and Cost-Sensitive Risk and Their Implications.” *Journal of Machine Learning Research* 14: 1033–1090.
- Zou, Zhengxia, Keyan Chen, Zhenwei Shi, Yuhong Guo, and Jieping Ye. 2023. “Object Detection in 20 Years: A Survey.” *Proceedings of the IEEE* 111 (3): 257–276.