

# Seldon Manifolds — Psychological Grounding

## Forward Models, Acceptable Regions, and the Aperture of One Embodiment

Matthew Parslow  
Independent Researcher

**Draft — draft-2026-09-06.1**  
6 September 2026 • document revision 14

**DRAFT — draft-2026-09-06.1r14** (6 September 2026). Circulated for comment, not as a finished result. It is the grounding companion to *Seldon Manifolds — Theory* (referred to below as *Seldon*) and, behind that, to *The Law of Graphic Equalisation — Theory* (referred to below as *Graphic Equalisation*), and inherits their definitions, epistemic boundary, and level-of-claims. Those are the only forms used below. *Seldon* sets out what a Seldon manifold is and how it works; this paper asks only whether one particular substrate implements anything of that shape. Sections marked as falsifiability commitments are the intended points of attack.

### Abstract

The Seldon manifold is a construct built to do forecasting, not a model of human cognition, so nothing here bears on whether it works. What this companion asks is narrower and worth asking on its own account: does a biological nervous system implement anything of that shape, and if so with what horizon, what resolution, and over which dimensions. The answer is that several components are implemented and individually measured — a forward model, an acceptable *region* rather than a target with correction only on deviations that would leave it, resolution degrading with temporal distance in a known functional form, and boundaries indexed to the body rather than to the world — while the composite is not, and while the evidence is uneven across the components, some carried by robust replicated literatures and some by only one or two studies each. The measured motor horizon is on the order of a hundred and fifty milliseconds and exists to cancel sensory delay; deliberative search runs about three steps deep and is truncated by a reflex rather than by a cost calculation; and within-frame forecasting is good where a domain has recurrent structure to hold and worth nothing where it does not. The honest summary is that the pieces are real, the human instantiation is narrow, and its quality is frame-relative rather than general. One line of evidence the construct would most naturally have leaned on — that representations coarsen with psychological distance — did not survive its best-powered test and is in dispute, and is not relied upon here.

## 1 Scope

This paper is a grounding and not a validation, and the distinction is load-bearing. The construct it grounds was built as engineering, to forecast; it makes no claim about how people think, so evidence that people forecast badly would not count against it. What such evidence does is characterise one implementation. Where the psychology is unflattering it is therefore reported as a measurement of that implementation's limits rather than as an objection, and where it is flattering it is reported without the suggestion that the construct needed it.

Two consequences follow for how this paper should be read. Findings are organised by the component they bear on — horizon, resolution, represented dimensions, fields, the acceptable region, and the capacity to reshape rather than navigate — rather than by the literature they

come from, so several mature research traditions appear here in pieces. And a component may be well evidenced while the composite is not; nothing below argues that any nervous system holds a Seldon manifold, only that each part of one has an implemented counterpart somewhere, mostly in different systems and at wildly different timescales. The studies below therefore support component-specific horizons, regions, and resolution effects, not a universal cognitive threshold or an integrated human manifold. Where a study reports a numerical horizon or cutoff, it is retained as that study’s measured operating range; this paper does not promote it to a cross-task threshold without direct support.

## 2 Forward models, and a horizon of milliseconds

The least controversial forward model in biology is the motor one. A forward model in this sense takes the current state and the outgoing motor command and predicts the next state (Miall and Wolpert, 1996), and the evidence that the nervous system runs one is behavioural rather than inferential: sensory consequences of self-produced movement are attenuated in proportion to how well they are predicted, so that perturbing the timing or trajectory of a self-administered tactile stimulus restores the sensation it would otherwise have suppressed (Blakemore et al., 1999), with a delay of roughly a hundred milliseconds sufficient to break the attenuation (Kilteni et al., 2019). Prediction is also learned before control is (Flanagan et al., 2003), which separates the forward model from the controller that uses it.

The horizon of that model has been measured, and it is short. In one study (Miall et al., 2007), disrupting the lateral cerebellum by stimulation made reaches behave as though planned from a hand position around 138ms out of date, a figure that study reports as replicated at 134ms with its own upper estimate at 150–160ms. That figure is not an incidental limitation but the model’s reason for existing: the feedback loops it compensates run at twenty to sixty milliseconds spinally and around a hundred visually, and a predictor that cancels them needs exactly that much lead. The motor forward model is a delay compensator rather than a simulator of futures, and the gap between it and a deliberative horizon is not bridged by any evidence in this paper.

*Remark 2.1* (What this does and does not license). It licenses the claim that a nervous system maintains a predictive model of its own next states and acts on the prediction rather than on the measurement. It does not license any claim about the length of a planning horizon, and a reader who takes the two together is making an architectural argument that the evidence here does not supply.

## 3 An acceptable region, and correction only on departure

*Seldon* holds that a frame may represent a region of acceptable futures (Definition 11.1) and intervene only when its projection leaves that region (Proposition 11.2). Motor control implements something with this shape, and the implementation is attested twice over — once measured and once derived.

The first is structural. Partitioning the variability of a practised movement into components that do and do not affect the task outcome shows that variance within the task-equivalent subspace is left alone while variance that would change the outcome is corrected — the *uncontrolled manifold* (Scholz and Schöner, 1999). The nervous system is not stabilising a trajectory; it is stabilising a region and permitting drift inside it.

The second is theoretical and derived rather than assumed. Under optimal feedback control the resulting policy exhibits what its authors name the *minimal intervention* principle: deviations from the average trajectory are corrected only when they interfere with task performance (Todorov and Jordan, 2002).

*Remark 3.1* (Three cautions on the correspondence). The principle is *minimal* intervention, not minimum: its criterion is task relevance, not magnitude, and the same analysis permits deviations to be *magnified* where doing so reduces cost. It is therefore not a smallest-sufficient-push result and should not be cited as one. And the controller does not project forward at decision time: the horizon is consumed offline in a backward pass, leaving an online policy that is a single operation on a state estimate. An architecture that re-projects at each step is closer to receding-horizon control, which Todorov and Jordan (2002) treat as a separate approach. And where that principle’s antecedent fails the region is not enough on its own: *Seldon* states the condition (Corollary 12.4), in which a frame remains inside every bound it holds while its capacity to leave has already gone, so a discipline reading acceptability alone has nothing to fire on.

The idea that acting well means keeping a projection inside a region is older than either result. Driving was described in 1938 in terms of a *field of safe travel*, an explicitly spatial region of acceptable trajectories within which the driver steers (Gibson and Crooks, 1938), a formulation whose computational descendants still maintain a perceived risk field below a threshold (Kolekar et al., 2020).

## 4 The acceptable region is Simon’s claim, not a simplification of it

Satisficing is usually reported as accepting the first option above a threshold, which is a scalar cutoff on one dimension. That is not what the original says. The aspiration level adapts — the balance between what is needed and what is available is maintained by *raising and lowering* aspiration levels — and the criterion is conjunctive across needs, a path being sought that permits satisfaction at some specified level of *all* of them (Simon, 1956). Simultaneous minima on every need dimension describe a region in a multi-dimensional space, and the fixed scalar version is explicitly the disposable simplification of a toy organism.

The formal development followed: aspiration adaptation theory treats goals as a *vector* with an aspiration on each component, adjusted against feasibility (Selten, 1998), and non-compensatory conjunctive models — a cutoff per attribute with no trade-offs between them — have been fitted to choice process data in preference to compensatory alternatives (Einhorn, 1970; Stüttgen et al., 2012). That the aspirations adapt to the environment rather than sitting fixed is itself measured (Caplin et al., 2011).

*Remark 4.1* (Where the region is genuinely absent). Several familiar constructs look like acceptable regions and are not. A reference point is one point on one axis with a kink, below which outcomes are treated continuously differently rather than categorically (Kahneman and Tversky, 1979); safety-first criteria minimise the probability of falling below a level, which is a half-line (Roy, 1952); and aspiration in the risk-preference literature, while a genuine second criterion rather than a reshaped utility function, remains a scalar cutoff (Lopes and Oden, 1999). The distinction between a threshold on a scalar and a region in several dimensions is exactly where an account can go wrong: the best-known attempt to explain accident rates by a single target level of risk failed, and failed because it collapsed a field to one scalar on one variable (Wilde, 1982; Evans, 1986).

## 5 Resolution falls with horizon

*Seldon* treats resolution (Definition 5.2) and horizon as separate axes under a shared budget, spent where a distinction would change an operative decision (Hypothesis 8.1) and characteristically degrading as horizon extends. There is a measured version of this with a functional form. In interval timing the standard deviation of an estimate is proportional to the interval being

estimated, so the coefficient of variation stays constant as the interval grows (Gibbon, 1977; Rakitin et al., 1998; Zeki and Balçi, 2019). Absolute uncertainty therefore scales with distance: the further out, the coarser, in a fixed ratio. The result is fifty years old, holds across species, and is not in dispute.

*Remark 5.1* (A stronger candidate that does not survive). There is an obvious alternative anchor, and it should be named because a reader will otherwise supply it. Construal level theory holds that the further an object is from direct experience the more abstract its representation (Trope and Liberman, 2010), which is resolution falling with distance stated for cognition generally rather than for timing. A meta-analysis of over two hundred experiments reported a reliable medium effect (Soderberg et al., 2015). But a registered multi-laboratory replication across four distance types, with nearly twelve thousand participants, found effects of  $d \approx 0.03\text{--}0.08$  with the social effect reversed, and a bias-corrected re-analysis of the existing literature finds strong evidence of publication bias (Calderon et al., 2026; Maier et al., 2022). The theory’s authors contest the replication (Liberman and Levit-Mor, 2026); the registered report appeared in a peer-reviewed journal in 2026, while the bias-corrected re-analysis remains a preprint. So the accurate statement is that the effect is in serious dispute rather than that it is refuted — and either way it cannot carry weight. The timing result above is narrower and secure, and is used instead.

## 6 Horizon is plural and frame-relative

This section is one of the places in this programme where the frame-relativity is *measured* rather than inherited, which is worth marking before the evidence rather than after. Graphic Equalisation’s founding axiom (Axiom -1) requires it — every represented quantity is held in some frame, so a horizon is one frame’s and there is no view from nowhere at which the horizon is. That is an entailment and carries no empirical risk on its own. What follows does: the axiom predicts that a system built of many frames should exhibit many horizons rather than one, and that is a claim about brains which could have come out otherwise. It does not.

A single agent does not have one horizon. Cortical regions differ systematically in the window over which they integrate information, from tens of milliseconds in sensory areas to many seconds in higher-order ones (Hasson et al., 2008, 2015), and predictive representations differ in reach along the same gradient, with anterior prefrontal cortex showing the largest predictive horizons and posterior hippocampus the smallest (Brunec and Momennejad, 2022). Different subsystems hold different horizons over the same world, which is *Seldon’s* frame-relativity (Remark 5.3) appearing within one skull rather than between agents.

## 7 The manifold is indexed to the body

The strongest support in this paper is for the claim that different agents hold different manifolds of one situation, with no frame-free manifold that theirs approximate. The boundary between a stairway that can be climbed and one that cannot is not a property of the stairway: it scales to the observer’s leg length and matches biomechanical prediction (Warren, 1984), and the boundary between an aperture that can be walked through and one requiring rotation sits at a constant ratio of aperture to shoulder width, recovered from body-scaled optical information (Warren and Whang, 1987). The same physical situation has measurably different boundaries for different bodies, and there is no body-free version of the boundary for either of them to be wrong about.

*Remark 7.1* (The extension is not free). This is the perception of immediately available action, not a multi-step future. That the same frame-relativity holds of extended horizons is an argument *Seldon* makes and this literature does not supply.

*Remark 7.2* (Saliency follows the aperture). The doorway result just given is also the smallest setting for a claim *Seldon* makes about apertures generally (Remark 5.4): an aperture is a shape and not a size, so what passes is settled by the fit between two forms rather than by a comparison of magnitudes, resolution is spent where the fit is in doubt, and saliency therefore tracks fit rather than magnitude. The response to a poor fit in the doorway is not more looking but *turning* — a reshaping of the passer to match — which is that whole mechanism, in a doorway.

Whether human frames do this is not established here, and the claim is modest: manifolds of this kind appear extant in human frames and severely limited, running short horizons at coarse resolution, which is what the timing result of Section 5 describes on one axis. Nothing here shows that saliency is *allocated by* fit rather than merely correlating with it, and that distinction is what an experiment would have to settle.

The same indexing to the body appears one level down, in how the eye is built and run: free viewing proceeds at roughly three to four fixations a second (Rayner, 1998), vision is suppressed across each saccade (Matin, 1974), and cone density peaks at the foveal centre and falls by more than an order of magnitude across the central few degrees (Curcio et al., 1990) — all measured and uncontested, and together the measured form of *Seldon*’s eye remark (Remark 8.2) and of Graphic Equalisation’s one-write remark (Remark 7.7). What this literature does not settle is the accounting those remarks turn on: whether visual cost scales with the number of fixations rather than with the number of elements resolved within one is the test that remark commits to, and the rate and the gradient are reported here while the bill is not.

## 8 Bounded search, and how it stops

Measured planning depth is shallow. Where it has been fitted directly it runs to roughly three steps, and subjects trade depth against precision under a near-constant workload, preferring to reduce depth rather than compute less precisely (Snider et al., 2015); deep problems are approached by concatenating shallow searches rather than by deepening them, with subgoal selection chosen so as to nearly minimise the cost of computing values (Huys et al., 2015).

How the search stops is less comfortable for any account that wants a cost-benefit rule. Branches are pruned reflexively on encountering a large loss, and the pruning persists when it is made overwhelmingly counterproductive (Huys et al., 2012). The same subjects show near-optimal subgoal selection and a hard-wired affective guillotine in the same task (Huys et al., 2015). The honest reading is that stopping is heuristic in mechanism and only approximately rational in aggregate.

*Remark 8.1* (Shallow search is not a shallow forward model). It would be a mistake to infer from three-step search that the forward representation is impoverished, and there is a better-supported alternative. People appear to hold a precomputed, horizon-discounted predictive map of expected future occupancy — flexible when rewards change, insensitive when the sequence of states changes — rather than either searching a tree or caching a policy (Momennejad et al., 2017). Such a map is much closer to a manifold with fields over it than tree search is, and it explains shallow measured search without impoverishing the representation. The caveat is that it is a compressed occupancy prediction rather than a branching structure with explicit alternatives, so it under-delivers on reachable futures in the plural.

## 9 Within-frame quality, and why it does not transfer

Forecasting quality inside a frame can be very high, and it is not bought by searching further. Replicating the classical result across skill levels finds that stronger players search *faster*, not deeper, with pattern recognition rather than search the principal discriminator (de Groot, 1965;

Connors et al., 2011); the sharpest demonstration is that a grandmaster’s rated strength falls only slightly when playing six boards simultaneously against tournament-length thinking time (Gobet and Simon, 1996a). Depth differences with skill are not zero (Campitelli and Gobet, 2004), so recognition should be called the principal discriminator rather than the only one.

The anticipation is genuinely prospective where it has been isolated: skilled performers respond to advance kinematic information before the event they are anticipating occurs, with the effect established across dozens of studies and hundreds of effect sizes (Müller et al., 2006; Mann et al., 2007), and the advantage is eliminated when attention is forced to switch between cues rather than held on the informative ones (Müller et al., 2010). This is the psychological side of the case *Seldon* calls the closest operational instance of its whole assembly it has found (Section 13): the batsman reading a bowler’s arm and the soaring pilot computing whether reach remaining covers the distance to goal are the same within-frame forecasting over a domain with recurrent structure to hold, run at two grains and two horizons.

It also does not transfer. Far transfer from trained domains shrinks as study design improves and vanishes against active controls (Sala and Gobet, 2017a,b; Melby-Lervåg et al., 2016). The mechanism predicts this rather than merely tolerating it: a store of domain-structured patterns has nothing to decode a foreign domain against. The complementary case is forecasting in a domain without stable recurrent structure, where experts do no better than simple statistical models or informed crowds (Forecasting Collaborative, 2023) — though within the same literature domain expertise still predicts accuracy, and calibration improves with training and aggregation (Mellers et al., 2014). Where structure exists the map forms; where it does not, expertise buys nothing.

*Remark 9.1* (One control is weaker than usually reported). The random-position control is often cited as showing that expert advantage vanishes without domain structure. It does not fully vanish: strong players retain a smaller but real advantage on random material, and it generalises across fields (Gobet and Simon, 1996b; Sala and Gobet, 2017c). The non-transfer claim above rests on the far-transfer meta-analyses, not on this control.

## 10 Blind spots, reshaping, and the growth of the dimension set

Three further components have implemented counterparts.

*Blind spots are real and are not introspectively accessible.* People report understanding complex phenomena with far greater precision and depth than they possess, and discover this only on attempting explanation (Rozenblit and Keil, 2002); attended observers miss salient unexpected events entirely (Simons and Chabris, 1999); and adults who can distinguish their own knowledge from another’s routinely fail to deploy that distinction online (Keysar et al., 2003). This supports *Seldon*’s treatment of a blind spot (Definition 9.6) as a distinct deficit from coarseness (Remark 9.2). It does not by itself rule out a mechanism in which a frame audits its own dimension set from inside: three findings that people do not are not a demonstration that nothing could. The impossibility claim is Graphic Equalisation’s (Proposition 14.5); the psychology above is consistent with it and offers no counterexample to it.

*Reshaping rather than navigating is implemented.* Committing in advance so as to remove a future option is a documented and modelled behaviour (Rachlin and Green, 1972; Ainslie, 1975), and the broader category of changing one’s situation rather than choosing within it is treated as a distinct class of self-control (Duckworth et al., 2016).

*The dimension set can grow.* Perceptual learning includes attentional weighting, differentiation and unitisation, and people demonstrably learn to separate dimensions that were previously fused (Goldstone, 1998; Goldstone and Steyvers, 2001). A represented dimension set is therefore not fixed, which *Seldon* requires of its represented dimensions (Definition 5.1) and does not otherwise evidence.

## 11 What the evidence does not support

Four things should be stated plainly, since a grounding paper that reports only the matches is not worth reading.

*The composite is not evidenced.* Every component above is implemented somewhere, but in different systems, at timescales separated by three orders of magnitude, and no study assembles them. That a nervous system holds a Seldon manifold is not shown here and is not claimed.

*Whether the machinery is required is unresolved.* Patients with hippocampal amnesia have been reported both as unable to construct coherent imagined experiences (Hassabis et al., 2007) and as unimpaired at imagining the future (Squire et al., 2010), with a third position arguing that the verbal measures used cannot adjudicate between them (Miloyan et al., 2019). The dispute is live. It is compatible with *Seldon*’s treatment of the manifold as optional machinery (Remark 1.5), but it should be presented as unresolved rather than as settled favourably.

*Forecasting error is on magnitude, not structure — and this is a reframing.* People systematically overestimate the intensity and duration of their future emotional reactions (Gilbert et al., 1998; Wilson and Gilbert, 2005). Much of the effect appears to be an artefact of how the question is interpreted (Levine et al., 2012), and when measured in daily life people predict correctly *which* events will make them feel better or worse while erring on absolute levels (Moeck et al., 2026). A manifold whose shape is right and whose field magnitudes are miscalibrated is what this describes. That is a legitimate reading and it is still a reading, offered as such.

*The per-agent axes are not established.* *Seldon* allows different frames to differ independently in horizon, resolution and represented dimensions (Remark 5.3). No study assembled here measures those in one sample, and at least one natural axis does not separate — imagery vividness and richness of future simulation covary rather than dissociating, with higher visual-imagery capacity predicting more sensory and contextual detail in imagined future events as well as in remembered past ones (D’Argembeau and Van der Linden, 2006). Separable per-agent axes are a synthesis, not a finding.

## 12 Falsifiability commitments

*On the tags.* Each commitment below is marked [SHAPE] or [DETAIL]. A [DETAIL] failing is a revision: the account survives with that mechanism replaced. A [SHAPE] failing costs *range* rather than a mechanism, because what fails is something the substrate is claimed to require or to permit — so the account does not hold where it claimed to, which is the more serious of the two and is still a boundary rather than an annihilation.

1. [SHAPE] **A predictive forward model.** Fails if the attenuation of self-produced sensation does not track prediction error, or if prediction is not separable from control.
2. [SHAPE] **Region rather than target.** Fails if task-irrelevant variability is corrected as strongly as task-relevant variability, i.e. if there is no uncontrolled subspace.
3. [DETAIL] **Resolution scaling with distance.** Fails if the coefficient of variation of temporal estimates is not approximately constant across intervals.
4. [DETAIL] **Salience tracks fit, not magnitude.** Fails if resolution is allocated by the magnitude of a represented quantity rather than by its fit against what remains acceptable. The test is an allocation manipulation that sets fit and magnitude against each other — fit held while magnitude varies, and the reverse — not a correlation between measured salience and measured fit, which the account and its rivals both predict.
5. [SHAPE] **Horizon is plural within one agent.** Fails if integration windows and predictive horizons across cortical regions collapse to one value, or if a single horizon suffices to model

within-agent prediction across timescales (Hasson et al., 2008; Brunec and Momennejad, 2022).

6. [SHAPE] **Body-indexed boundaries.** Fails if perceived action boundaries do not scale with the relevant bodily dimension, or if one boundary fits all bodies.
7. [DETAIL] **Search is shallow and stops reflexively.** Fails if fitted planning depth is not bounded at a few steps, or if pruning is cost-sensitive where it has been made overwhelmingly counterproductive.
8. [SHAPE] **Structure-dependent forecasting.** Fails if within-domain forecasting advantage is independent of the domain’s recurrent structure, or if it transfers to unrelated domains under active controls.
9. [DETAIL] **Blind spots are not introspectively auditable.** Fails if a frame can be shown to enumerate its own missing dimensions from inside.
10. [DETAIL] **The dimension set grows.** Fails if perceptual learning never separates dimensions that were previously fused.
11. [DETAIL] **Affective forecasting errs on magnitude, not on structure.** Fails if, in daily-life measurement, people misorder which events will make them feel better or worse as often as they misjudge how much (Moeck et al., 2026) — which would make the manifold’s shape wrong and not only its field magnitudes. Section 11 offers that reframing as a reading; this is where it becomes a commitment.

*Scope note (not a falsifier).* This paper would be over-claiming if any section were read as evidence that a nervous system holds a Seldon manifold, or as validation of a construct that does not require it.

## References

- George Ainslie. Specious reward: a behavioral theory of impulsiveness and impulse control. *Psychological Bulletin*, 82(4):463–496, 1975.
- Sarah-Jayne Blakemore, Chris D. Frith, and Daniel M. Wolpert. Spatio-temporal prediction modulates the perception of self-produced stimuli. *Journal of Cognitive Neuroscience*, 11(5):551–559, 1999.
- Iva K. Brunec and Ida Momennejad. Predictive representations in hippocampal and prefrontal hierarchies. *Journal of Neuroscience*, 42(2):299–312, 2022.
- Sofia Calderon, Erik Mac Giolla, Karl Ask, et al. Effects of psychological distance on mental abstraction: a registered report of four tests of construal-level theory. *Advances in Methods and Practices in Psychological Science*, 9(2):25152459251401177, 2026.
- Guillermo Campitelli and Fernand Gobet. Adaptive expert decision making: skilled chess players search more and deeper. *ICGA Journal*, 27(4):209–216, 2004.
- Andrew Caplin, Mark Dean, and Daniel Martin. Search and satisficing. *American Economic Review*, 101(7):2899–2922, 2011.
- Michael H. Connors, Bruce D. Burns, and Guillermo Campitelli. Expertise in complex decision making: the role of search in chess 70 years after de Groot. *Cognitive Science*, 35(8):1567–1579, 2011.

- Christine A. Curcio, Kenneth R. Sloan, Robert E. Kalina, and Anita E. Hendrickson. Human photoreceptor topography. *Journal of Comparative Neurology*, 292(4):497–523, 1990.
- Arnaud D’Argembeau and Martial Van der Linden. Individual differences in the phenomenology of mental time travel: the effect of vivid visual imagery and emotion regulation strategies. *Consciousness and Cognition*, 15(2):342–350, 2006.
- Adriaan D. de Groot. *Thought and Choice in Chess*. Mouton, The Hague, 1965. Originally published in Dutch, 1946.
- Angela L. Duckworth, Tamar Szabó Gendler, and James J. Gross. Situational strategies for self-control. *Perspectives on Psychological Science*, 11(1):35–55, 2016.
- Hillel J. Einhorn. The use of nonlinear, noncompensatory models in decision making. *Psychological Bulletin*, 73(3):221–230, 1970.
- Leonard Evans. Risk homeostasis theory and traffic accident data. *Risk Analysis*, 6(1):81–94, 1986.
- J. Randall Flanagan, Philipp Vetter, Roland S. Johansson, and Daniel M. Wolpert. Prediction precedes control in motor learning. *Current Biology*, 13(2):146–150, 2003.
- Forecasting Collaborative. Insights into the accuracy of social scientists’ forecasts of societal change. *Nature Human Behaviour*, 7(4):484–501, 2023.
- John Gibbon. Scalar expectancy theory and Weber’s law in animal timing. *Psychological Review*, 84(3):279–325, 1977.
- James J. Gibson and Laurence E. Crooks. A theoretical field-analysis of automobile-driving. *American Journal of Psychology*, 51(3):453–471, 1938.
- Daniel T. Gilbert, Elizabeth C. Pinel, Timothy D. Wilson, Stephen J. Blumberg, and Thalia P. Wheatley. Immune neglect: a source of durability bias in affective forecasting. *Journal of Personality and Social Psychology*, 75(3):617–638, 1998.
- Fernand Gobet and Herbert A. Simon. The roles of recognition processes and look-ahead search in time-constrained expert problem solving. *Psychological Science*, 7(1):52–55, 1996.
- Fernand Gobet and Herbert A. Simon. Recall of random and distorted chess positions: implications for the theory of expertise. *Psychonomic Bulletin & Review*, 3(2):159–163, 1996.
- Robert L. Goldstone. Perceptual learning. *Annual Review of Psychology*, 49:585–612, 1998.
- Robert L. Goldstone and Mark Steyvers. The sensitization and differentiation of dimensions during category learning. *Journal of Experimental Psychology: General*, 130(1):116–139, 2001.
- Demis Hassabis, Dharshan Kumaran, Seralynne D. Vann, and Eleanor A. Maguire. Patients with hippocampal amnesia cannot imagine new experiences. *Proceedings of the National Academy of Sciences*, 104(5):1726–1731, 2007.
- Uri Hasson, Eunice Yang, Ignacio Vallines, David J. Heeger, and Nava Rubin. A hierarchy of temporal receptive windows in human cortex. *Journal of Neuroscience*, 28(10):2539–2550, 2008.
- Uri Hasson, Janice Chen, and Christopher J. Honey. Hierarchical process memory: memory as an integral component of information processing. *Trends in Cognitive Sciences*, 19(6):304–313, 2015.

- Quentin J. M. Huys, Neir Eshel, Elizabeth O’Nions, Luke Sheridan, Peter Dayan, and Jonathan P. Roiser. Bonsai trees in your head: how the Pavlovian system sculpts goal-directed choices by pruning decision trees. *PLoS Computational Biology*, 8(3):e1002410, 2012.
- Quentin J. M. Huys, Niáll Lally, Paul Faulkner, Neir Eshel, Erich Seifritz, Samuel J. Gershman, Peter Dayan, and Jonathan P. Roiser. Interplay of approximate planning strategies. *Proceedings of the National Academy of Sciences*, 112(10):3098–3103, 2015.
- Daniel Kahneman and Amos Tversky. Prospect theory: an analysis of decision under risk. *Econometrica*, 47(2):263–291, 1979.
- Boaz Keysar, Shuhong Lin, and Dale J. Barr. Limits on theory of mind use in adults. *Cognition*, 89(1):25–41, 2003.
- Konstantina Kilteni, Christian Houborg, and H. Henrik Ehrsson. Rapid learning and unlearning of predicted sensory delays in self-generated touch. *eLife*, 8:e42888, 2019.
- Sarvesh Kolekar, Joost de Winter, and David Abbink. Human-like driving behaviour emerges from a risk-based driver model. *Nature Communications*, 11:4850, 2020.
- Linda J. Levine, Heather C. Lench, Robin L. Kaplan, and Martin A. Safer. Accuracy and artifact: reexamining the intensity bias in affective forecasting. *Journal of Personality and Social Psychology*, 103(4):584–605, 2012.
- Nira Liberman and Or Levit-Mor. The effect of psychological distance on level of construal: what do we learn from the CLIMR failed replication? PsyArXiv preprint, 2026.
- Lola L. Lopes and Gregg C. Oden. The role of aspiration level in risky choice. *Journal of Mathematical Psychology*, 43(2):286–313, 1999.
- Maximilian Maier, František Bartoš, Minjeong Oh, Eric-Jan Wagenmakers, David Shanks, and Adam J. L. Harris. Adjusting for publication bias reveals that evidence for and size of construal level theory effects is substantially overestimated. PsyArXiv preprint, 2022.
- Derek T. Y. Mann, A. Mark Williams, Paul Ward, and Christopher M. Janelle. Perceptual-cognitive expertise in sport: a meta-analysis. *Journal of Sport and Exercise Psychology*, 29(4):457–478, 2007.
- Ethel Martin. Saccadic suppression: a review and an analysis. *Psychological Bulletin*, 81(12):899–917, 1974.
- Monica Melby-Lervåg, Thomas S. Redick, and Charles Hulme. Working memory training does not improve performance on measures of intelligence or other measures of “far transfer”. *Perspectives on Psychological Science*, 11(4):512–534, 2016.
- Barbara Mellers, Lyle Ungar, Jonathan Baron, Jaime Ramos, Burcu Gurcay, Katrina Fincher, Sydney E. Scott, Don Moore, Pavel Atanasov, Samuel A. Swift, Terry Murray, Eric Stone, and Philip E. Tetlock. Psychological strategies for winning a geopolitical forecasting tournament. *Psychological Science*, 25(5):1106–1115, 2014.
- R. Chris Miall, Lars O. D. Christensen, Owen Cain, and James Stanley. Disruption of state estimation in the human lateral cerebellum. *PLoS Biology*, 5(11):e316, 2007.
- R. Chris Miall and Daniel M. Wolpert. Forward models for physiological motor control. *Neural Networks*, 9(8):1265–1279, 1996.

- Beyon Miloyan, Kimberley A. McFarlane, and Thomas Suddendorf. Measuring mental time travel: is the hippocampus really critical for episodic memory and episodic foresight? *Cortex*, 117:371–384, 2019.
- Ella K. Moeck, Karta K. Grewal, Amit Mehta, Katharine H. Greenaway, Peter Koval, and Elise K. Kalokerinos. Affective forecasting accuracy in everyday life. *Affective Science*, 7:306–318, 2026.
- Ida Momennejad, Evan M. Russek, Jin H. Cheong, Matthew M. Botvinick, Nathaniel D. Daw, and Samuel J. Gershman. The successor representation in human reinforcement learning. *Nature Human Behaviour*, 1(9):680–692, 2017.
- Sean Müller, Bruce Abernethy, and Damian Farrow. How do world-class cricket batsmen anticipate a bowler’s intention? *Quarterly Journal of Experimental Psychology*, 59(12):2162–2186, 2006.
- Sean Müller, Bruce Abernethy, Michael Eid, Rohan McBean, and Matthew Rose. Expertise and the spatio-temporal characteristics of anticipatory information pick-up from complex movement patterns. *Perception*, 39(6):745–760, 2010.
- Howard Rachlin and Leonard Green. Commitment, choice and self-control. *Journal of the Experimental Analysis of Behavior*, 17(1):15–22, 1972.
- Brian C. Rakitin, John Gibbon, Trevor B. Penney, Chara Malapani, Sean C. Hinton, and Warren H. Meck. Scalar expectancy theory and peak-interval timing in humans. *Journal of Experimental Psychology: Animal Behavior Processes*, 24(1):15–33, 1998.
- Keith Rayner. Eye movements in reading and information processing: 20 years of research. *Psychological Bulletin*, 124(3):372–422, 1998.
- A. D. Roy. Safety first and the holding of assets. *Econometrica*, 20(3):431–449, 1952.
- Leonid Rozenblit and Frank Keil. The misunderstood limits of folk science: an illusion of explanatory depth. *Cognitive Science*, 26(5):521–562, 2002.
- Giovanni Sala and Fernand Gobet. Does far transfer exist? Negative evidence from chess, music, and working memory training. *Current Directions in Psychological Science*, 26(6):515–520, 2017.
- Giovanni Sala and Fernand Gobet. Working memory training in typically developing children: a meta-analysis of the available evidence. *Learning & Behavior*, 45(4):414–421, 2017.
- Giovanni Sala and Fernand Gobet. Experts’ memory superiority for domain-specific random material generalizes across fields of expertise. *Memory & Cognition*, 45(2):183–193, 2017.
- John P. Scholz and Gregor Schöner. The uncontrolled manifold concept: identifying control variables for a functional task. *Experimental Brain Research*, 126(3):289–306, 1999.
- Reinhard Selten. Aspiration adaptation theory. *Journal of Mathematical Psychology*, 42(2–3):191–214, 1998.
- Herbert A. Simon. Rational choice and the structure of the environment. *Psychological Review*, 63(2):129–138, 1956.
- Daniel J. Simons and Christopher F. Chabris. Gorillas in our midst: sustained inattentional blindness for dynamic events. *Perception*, 28(9):1059–1074, 1999.

- Joseph Snider, Dongpyo Lee, Howard Poizner, and Sergei Gepshtein. Prospective optimization with limited resources. *PLoS Computational Biology*, 11(9):e1004501, 2015.
- Courtney K. Soderberg, Shannon P. Callahan, Annie O. Kochersberger, Elinor Amit, and Alison Ledgerwood. The effects of psychological distance on abstraction: two meta-analyses. *Psychological Bulletin*, 141(3):525–548, 2015.
- Larry R. Squire, Anna S. van der Horst, Susan G. R. McDuff, Jennifer C. Frascino, Ramona O. Hopkins, and Kristin N. Mauldin. Role of the hippocampus in remembering the past and imagining the future. *Proceedings of the National Academy of Sciences*, 107(44):19044–19048, 2010.
- Peter Stüttgen, Peter Boatwright, and Robert T. Monroe. A satisficing choice model. *Marketing Science*, 31(6):878–899, 2012.
- Emanuel Todorov and Michael I. Jordan. Optimal feedback control as a theory of motor coordination. *Nature Neuroscience*, 5(11):1226–1235, 2002.
- Yaacov Trope and Nira Liberman. Construal-level theory of psychological distance. *Psychological Review*, 117(2):440–463, 2010.
- William H. Warren. Perceiving affordances: visual guidance of stair climbing. *Journal of Experimental Psychology: Human Perception and Performance*, 10(5):683–703, 1984.
- William H. Warren and Suzanne Whang. Visual guidance of walking through apertures: body-scaled information for affordances. *Journal of Experimental Psychology: Human Perception and Performance*, 13(3):371–383, 1987.
- Gerald J. S. Wilde. The theory of risk homeostasis: implications for safety and health. *Risk Analysis*, 2(4):209–225, 1982.
- Timothy D. Wilson and Daniel T. Gilbert. Affective forecasting: knowing what to want. *Current Directions in Psychological Science*, 14(3):131–134, 2005.
- Mustafa Zeki and Fuat Balçi. A simplified model of communication between time cells. *Frontiers in Computational Neuroscience*, 12:111, 2019.