

Æ ATOM EONS RESEARCH

RADIANCE- LUMINANCE THEORY

AND ALPHA WOLF EYES

A Deterministic Retina-to-IT Visual System
Without Learned Weights

Theory, implemented architecture, development evidence,
and preregistered scale validation

Atom McCree

/EoNs Research Laboratory / AtomEons Systems Laboratory

Radiance-Luminance Theory and Alpha Wolf Eyes

A Deterministic Retina-to-IT Visual System Without Learned Weights

Atom McCree

ÆoNs Research Laboratory / AtomEons Systems Laboratory, Marco Island, Florida, USA

Correspondence: a.mccree [at] gmail [dot] com

Preprint · Final manuscript copy · July 10, 2026

Abstract

We introduce Radiance-Luminance Theory (RLT), a photon-first account of visual perception, and Alpha Wolf Eyes (AWE-3), an implemented deterministic retina-to-inferotemporal visual pipeline containing no learned filter weights. Rather than treating perception as a mapping from pixel arrays into opaque embeddings optimized through gradient descent, RLT treats the sensor-observable spatiotemporal radiometric field as primary evidence and recognition as the recurrence of stable structural relationships across changes in illumination, scale, orientation, and viewpoint. AWE-3 applies fixed transformations for linear-light recovery, tone adaptation, CAT02 chromatic adaptation, foveated log-polar canonicalization, scotopic processing, retinal ganglion-inspired channel decomposition, opponent-color encoding, parallel parvocellular, magnocellular, and koniocellular processing, orientation- and scale-selective V1 responses, V2 contour and texture-boundary measurements, V4 curvature and color-shape measurements, and an 80-dimensional inferotemporal identity representation. The complete canonical representation contains 412 deterministic output units while preserving the original input as immutable evidence. On a 47-class development corpus containing 282 observations across approximately six lighting conditions per class, the current system correctly classified 275 observations (97.52%; Wilson 95% confidence interval approximately 95.0%-98.8%). Selected few-shot configurations reached 100% development accuracy with four to five reference observations per class. Because scoring and feature combinations were repeatedly evaluated on this corpus, these values are developmental evidence rather than an unbiased estimate of generalization. RLT further proposes multi-encoding consensus and a sparse relational Pattern Engine built from local cylindrical or phase-toroidal structural descriptors. We specify a preregisterable 10,000-class, one-million-observation experiment designed to estimate few-shot saturation, pathological-illuminant robustness, identity-code collision behavior, open-set rejection, scaling laws, reproducibility, and complete storage cost.

Keywords: deterministic computer vision; photon-first perception; zero learned weights; retinal computation; few-shot recognition; color constancy; log-polar vision; visual invariance; inferotemporal representation; reproducible vision systems

STATUS OF EVIDENCE

The retina-to-IT AWE-3 pipeline, IT-80 representation, saccadic sampling, and exemplar recognition are implemented. The PhotonKnot descriptor, constellation graph, multi-codec consensus, and 10,000-class validation are prospective and are not required for the reported development results.

Contents

- [1. Introduction](#)
- [2. Scope, terminology, and claims](#)
- [3. Radiometric and biological foundations](#)
- [4. Implemented AWE-3 architecture](#)

5. Development evidence

6. Radiance-Luminance Theory

7. Multi-encoding consensus

8. Pattern Engine and PhotonKnots

9. Relation to prior work

10. Preregistered scale validation

11. Limitations and falsification conditions

12. Conclusion

References

1. Introduction

Modern computer vision is dominated by learned feature hierarchies. Images are mapped through parameterized networks into latent representations, and recognition is inferred from learned decision surfaces. This paradigm has produced major practical gains, but its intermediate states are usually not physical measurements and its behavior depends on optimization data, objective design, and model scale (Yamins and DiCarlo, 2016; Radford et al., 2021).

Biological vision begins from a different constraint: finite photon arrival at a sensor whose circuitry decomposes the signal before semantic interpretation. Rods can respond to individual photon absorptions under controlled conditions; retinal circuits create multiple parallel representations; the lateral geniculate nucleus preserves partially segregated contrast, color, and temporal pathways; V1 neurons exhibit orientation and spatial-frequency selectivity; V2 and V4 combine boundaries, contours, curvature, and color; and inferotemporal cortex supports increasingly invariant object representation (Baylor, Lamb, and Yau, 1979; Roska and Werblin, 2001; Hubel and Wiesel, 1962; Pasupathy and Connor, 2002; DiCarlo, Zoccolan, and Rust, 2012).

Radiance-Luminance Theory asks whether a useful retina-to-identity hierarchy can be built from fixed, auditable transformations rather than learned filter weights. Alpha Wolf Eyes is the current engineering realization of that question. Its contribution is not the invention of every component in isolation. Fixed filters, log-polar sampling, color constancy, cortex-like hierarchies, HMAX, and invariant scattering all have substantial prior histories. The contribution is the integration of radiometrically grounded acquisition, deterministic retinal and cortical-style decomposition, a compact identity representation, few-shot exemplar registration, evidence preservation, and receipt-based reproducibility within one operational system.

This manuscript reports the implemented AWE-3 architecture and its development evidence, then presents RLT as the theory that organizes the implementation and motivates the next Pattern Engine. The distinction is deliberate: implemented findings are reported as findings; proposed mechanisms are reported as hypotheses; the million-observation experiment is presented as a prospective scale test.

2. Scope, terminology, and claims

AWE-3 is described as a system with zero learned weights, not as a parameter-free system. Fixed Gabor frequencies, CAT02 matrices, resampling kernels, thresholds, foveal exponents, channel definitions, and scoring coefficients are parameters in the ordinary mathematical sense. They are specified by design and are not estimated through back-propagation or gradient descent on a visual training corpus.

The term photon-first does not imply that standard decoded video preserves the original photon arrivals from a physical scene. It means that the system treats the measurable radiometric signal and its deterministic projections as primary evidence, preserves the source observation, and avoids replacing the evidence with an opaque learned embedding at the beginning of the pipeline.

The present recognition memory uses compact IT-80 observations and exemplars. The proposed Pattern Engine would add relational identity memory above that representation. Consequently, the current experimental results test the AWE-3 eye and exemplar recognizer; they do not yet validate the full constellation-graph theory.

Figure 2 | Evidence boundary

IMPLEMENTED	PILOT EVIDENCE	PROSPECTIVE
• Linear light / CAT02 • Log-polar canonical • Retinal-12 + axis-15 • LGN P/M/K • V1-V4 fixed projections • IT-80 + saccades	• 5/5 optometric suite • 7/7 lighting test • 18/19 extended test • 275/282 at 47 classes • N=4-5 best case • 192K search exhausted	• 10K × 100 validation • Open-set rejection • Codec consensus • PhotonKnot descriptor • Constellation graph • Human-scale Pattern Engine

Evidence status. Implemented architecture, development evidence, and prospective mechanisms are intentionally separated.

3. Radiometric and biological foundations

3.1 Sensor-observable radiometric field

Let $R(x,y,t,\lambda)$ denote spectral radiance incident through the optical path. A digital sensor does not directly measure an ideal scene luminance field. For channel c , the measured response is more accurately represented as an integration over spectral radiance, optical transmission, channel quantum efficiency, exposure, and noise:

$$E_c(x,y,t) = \int R(x,y,t,\lambda) T_{optics}(\lambda) Q_c(\lambda) d\lambda + \eta_c(x,y,t) \quad (1)$$

Photometric luminance is a particular human-observer projection using the photopic luminosity function $V(\lambda)$, not a universal camera measurement:

$$Y(x,y,t) = 683 \int R(x,y,t,\lambda) V(\lambda) d\lambda \quad (2)$$

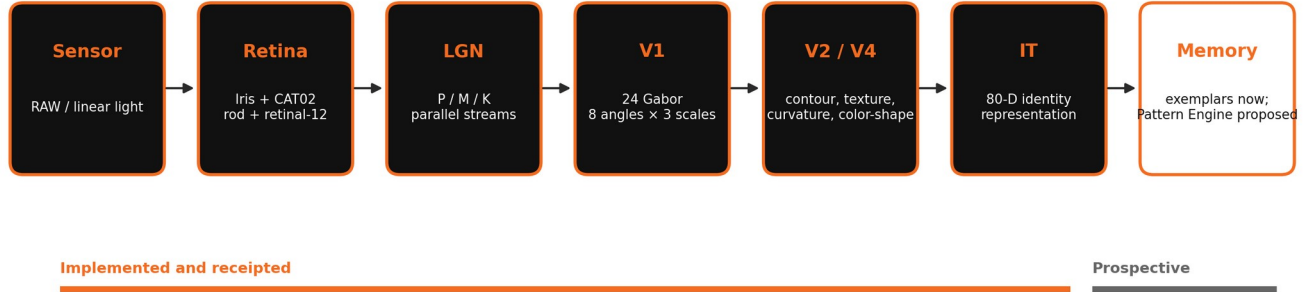
RLT therefore treats the sensor-observable radiometric field as the primary record. Linear luminance, opponent-color, scotopic, gradient, phase, and motion fields are derived measurements. When the source is decoded RGB rather than RAW sensor data, the input is explicitly described as a display-referred or scene-referred approximation according to available metadata.

3.2 Retinal and cortical constraints

The design takes biological organization as an engineering prior rather than as a claim of literal neural equivalence. Its retinal channels are motivated by parallel retinal representations and ganglion-cell diversity (Roska and Werblin, 2001; Baden et al., 2016; Masland, 2012). Its P/M/K split reflects partially segregated primate visual pathways (Kaplan and Shapley, 1986; Casagrande, 1994; Hendry and Reid, 2000). Its V1 stage uses fixed oriented filters consistent with classical Gabor descriptions of simple-cell receptive fields (Jones and Palmer, 1987). Its V2 and V4 stages measure contour, boundary, curvature, concavity, complexity, and color-shape conjunctions motivated by physiological studies (von der Heydt, Peterhans, and Baumgartner, 1984; Pasupathy and Connor, 1999, 2002). The IT stage is an engineered identity code, not a simulation of macaque IT firing.

4. Implemented AWE-3 architecture

Figure 1 | AWE-3 implemented visual pathway and the boundary to prospective memory



AWE-3 pipeline. The first six blocks are implemented; the relational Pattern Engine is prospective.

The current system is implemented in Bun as a deterministic local pipeline. The canonical path is:

CANONICAL PATH		
RAW input -> linearization -> Reinhard tone adaptation -> CAT02 cone-space illuminant division -> 256×256 log-polar canonicalization with Catmull-Rom interpolation and foveal bias 1.6 -> 64×64 scotopic pathway -> retinal-12 -> axis bundle -> opponent channels -> LGN P/M/K streams -> V1 -> V2 -> V4 -> IT-80 -> multi-fixation saccades.		
Stage	Role	Status
Acquisition and linearization	Preserve source; reverse transfer function where known; establish linear-light working representation	Implemented
Iris / adaptation	Reinhard-style dynamic-range handling and CAT02 chromatic adaptation	Implemented
Canonicalization	256×256 foveated log-polar field; bicubic Catmull-Rom sampling; bias 1.6	Implemented
Rod pathway	64×64 scotopic structure channel	Implemented
Retinal-12	Twelve deterministic channel families inspired by retinal parallel processing	Implemented
Axis bundle	Fifteen scalar structural axes contributing to the canonical signature	Implemented
LGN	Parvocellular, magnocellular, and koniocellular-style streams	Implemented
V1	24 Gabor responses: eight orientations across three spatial scales	Implemented
V2	Contour, texture-boundary, and cross-orientation suppression measurements	Implemented
V4	Curvature, concavity, complexity, and color-shape coupling	Implemented

IT	80-dimensional identity representation used by the present recognizer	Implemented
Saccades	Multiple deterministic fixation samples	Implemented
PhotonKnot / graph	Local cylindrical or phase-toroidal structural node and relational identity graph	Proposed

Table 1. Implemented and prospective components of AWE-3 and RLT.

4.1 Output representation and evidence preservation

Each canonical observation emits 412 deterministic output units, including the 80-dimensional identity vector. The source observation is preserved rather than overwritten. The earlier phrase '2888% derived information' is replaced here by representational coordinate expansion: the system deterministically expands a compact source fixture into a larger set of interpretable measurements. It does not create new Shannon information absent from the source evidence.

4.2 Runtime and reproducibility

Bun is the governing runtime, reference implementation, experiment dispatcher, and receipt producer. Numeric kernels operate over typed arrays and deterministic iteration order. Scale experiments use independent operating-system processes addressed by PROC_RANK and PROC_WORKERS, immutable output shards, and a single deterministic reducer. This process model supports crash isolation, resumability, explicit memory reclamation, and rank-level audit receipts. Optional future acceleration may use self-owned WebAssembly or a stable native interface, but the Bun implementation remains the reference truth implementation.

5. Development evidence

AWE-3 has been evaluated on synthetic and controlled visual fixtures intended to stress spatial resolution, contrast, chromatic response, motion, depth, class identity, and illumination variation. All results in this section are development results. The same 47-class cache was used during extensive scoring and feature-combination exploration; therefore, the 47-class number is not an unbiased held-out estimate.

Evaluation	Result	Interpretation
Optometric suite	5/5 pass	Spatial, contrast, chromatic, motion, and depth fixtures
Photon-print fidelity	100.0%	Raw input preserved unaltered for every fixture
7 classes × 6 lighting conditions	7/7 classes	Development test
19 classes × 6 lighting conditions	18/19 classes	Starry Night under neon was the reported failure
47 classes, 282 observations	275/282 = 97.52%	Wilson 95% CI approximately 95.0%-98.8%
Few-shot ceiling on current cache	N=4-5 reached 100% best case	Selected development configuration, not independent validation
Search effort	192,000 scoring/weight experiments	No further gain at this corpus scale

Table 2. Current AWE-3 development evidence.

5.1 Recognition result and confidence interval

The 47-class cache contains 282 observations. The current locked development configuration correctly classified 275 observations:

$$\hat{p} = 275 / 282 = 0.9752 \quad (3)$$

The corresponding Wilson 95% interval is approximately 0.950-0.988. This interval quantifies binomial uncertainty conditional on the tested cache. It does not correct for model-selection bias introduced by repeated evaluation on the same observations.

5.2 Storage accounting

An IT-80 vector stored as float32 requires 320 bytes. At five observations per identity, the raw vector payload is 1,600 bytes per identity. A 5 GiB raw-vector budget therefore holds approximately 3,355,443 identities at $N=5$. At $N=100$, the same budget holds approximately 167,772 identities, not one million. These are upper bounds that exclude identity keys, class metadata, lighting metadata, provenance, indexing structures, checksums, graph edges, and allocator overhead.

$$Capacity(N) = 5,368,709,120 / (320 N) \quad (4)$$

STORAGE LAW

All future capacity claims must report raw vectors, metadata, search index, evidence references, and Pattern Engine graph overhead separately. Quantization is permitted only after nearest-neighbor ordering, accuracy, and collision margins are receipted against float32.

6. Radiance-Luminance Theory

6.1 Central axiom

RLT's central axiom is that visual identity should be grounded in recurring structure within the spatiotemporal radiometric field. Recognition is not defined as retrieving a label attached to a pixel array. It is defined as reactivation of a stable structural account that explains multiple observations as transformations of one persistent cause.

6.2 Objecthood as internal predictability

An operational object hypothesis can be expressed through comparative dependence: a candidate region is object-like when its internal measurements predict one another more strongly than they predict measurements outside the region. The exact estimator may use mutual information, conditional code length, normalized compression distance, or deterministic structural agreement. The equation is a research criterion rather than a claim that the current AWE-3 pipeline computes exact mutual information over continuous light fields.

$$InternalDependence(O) > BoundaryDependence(O, \Omega \setminus O) \quad (5)$$

6.3 Identity as description-length advantage

Identity is formalized more carefully through minimum description length. Let L_i be observation i , M_i a separate model for that observation, M one persistent identity model, and T_i the transformation needed to explain L_i from M . The identity advantage is:

$$\Delta_{ID} = \sum_i DL(L_i / M_i) - [DL(M) + \sum_i DL(L_i / M, T_i)] \quad (6)$$

A positive and stable delta indicates that one persistent identity plus transformations explains the observations more compactly than unrelated models. This captures the intuitive claim that a familiar face is not stored as one

frozen image: many views, expressions, illuminants, and resolutions are economically explained by one recurring cause.

6.4 Six-layer hypothesis

1. **Photon acquisition.** Capture the best available sensor-observable radiometric record and preserve the source evidence.
2. **Canonical linear-light representation.** Remove known transfer functions; estimate adaptation and illuminant state; expose confidence.
3. **Local structural decomposition.** Measure luminance, opponent color, orientation, phase, texture, radial structure, motion, and symmetry.
4. **Temporal persistence.** Track which local structures remain coherent across motion and changing observation conditions.
5. **Transformation-invariant coding.** Represent rotation, scale, illumination, and temporal phase as explicit transformations rather than category-specific cases.
6. **Identity reactivation.** Recognize a persistent entity from a weighted constellation of matching structures and relations.

7. Multi-encoding consensus

Multi-encoding consensus is a proposed calibration and invariance procedure. It does not reconstruct the original physical photons that struck the camera sensor. Instead, it estimates source structure that is stable across codec, bitrate, resizing, chroma-subsampling, and rendering transforms applied to a common master.

$$F_j = S_{\text{codec}} + C_j + R_j + N_j \quad (7)$$

Here F_j is decoded version j , S_{codec} is codec-stable source structure, C_j is codec-specific residue, R_j contains resizing and color-transfer effects, and N_j contains noise and alignment error. A robust consensus estimator may operate over precisely aligned multiscale structural fields:

$$\hat{S}_{\text{codec}} = \text{RobustConsensus}(F_1, F_2, \dots, F_m) \quad (8)$$

The simplest implementation is a robust median or trimmed estimator over aligned feature fields. Alignment should use deterministic timebase reconciliation, phase correlation, enhanced correlation coefficient optimization, feature registration, or optical flow as appropriate. Structural similarity can evaluate agreement after alignment, but SSIM is not itself a complete registration procedure.

The key experiment compares recognition from a single source encoding against recognition from a consensus prototype built from independently encoded copies. Improvement on held-out encodings would support the claim that codec diversity can function as a non-learning invariance teacher.

8. Pattern Engine and PhotonKnots

The Pattern Engine is the proposed memory layer above AWE-3. It should consume IT-80, selected retinal and cortical intermediates, fixation metadata, illuminant state, source provenance, recognition margin, and references to preserved evidence. It should not duplicate the eye or infer semantic labels.

8.1 Cylindrical and toroidal PhotonKnots

Around a stable anchor (x_0, y_0, t_0) , a local descriptor samples the radiometric and derived channel fields in log-polar coordinates. Angular position is inherently circular, whereas ordinary short time is a bounded interval. The default domain is therefore cylindrical:

$$K(r, \theta, \tau, c), \quad \theta \in S^1, \quad \tau \in [-T, T] \quad (9)$$

A toroidal domain is justified when time is converted into phase for periodic or phase-normalized motion:

$$(\theta, \varphi - t) \in S^1 \times S^1 = T^2 \quad (10)$$

A co-prime traversal can sample angular, radial, and temporal indices without repeatedly aligning the same cross-sections. For example, 31 angular positions, 17 radial shells, and seven phase offsets provide a long repeat period. The exact tuple is a design parameter and requires ablation; co-primality alone does not prove recognition invariance.

Figure 3 | PhotonKnot sampling domains (prospective Pattern Engine)

General observation: cylinder $S^1 \times [-T, T]$

Phase-normalized periodic observation: torus $S^1 \times S^1$



Prospective PhotonKnot geometry. General temporal windows form a cylinder; periodic or phase-normalized observations can form a torus.

8.2 Constellation recognition

A stored identity M is proposed as a sparse weighted graph of recurring local structures and observed relationships. A new observation O activates a partial subgraph. Recognition combines node agreement, relation agreement, and contradiction penalties:

$$Score(M, O) = \sum_{v \in match} w_v + \sum_{e \in match} w_e - \lambda \pi(M, O) \quad (11)$$

Relations may include spatial-neighbor, moves-with, transforms-to, persists-as, occludes, reappears-as, and co-activates-with. Rare discriminative structures receive greater weight than ubiquitous structures. Graph topology is learned only in the ordinary memory sense of registering observations; the perceptual filters remain fixed. This distinction should be explicit whenever the system is described as zero learned weights.

9. Relation to prior work

9.1 Classical and biologically inspired vision

Marr's computational theory, Marr-Hildreth edge detection, Canny edges, steerable filters, SIFT, and log-polar robotic vision established that substantial visual structure can be extracted through deterministic operators (Marr, 1982; Marr and Hildreth, 1980; Canny, 1986; Freeman and Adelson, 1991; Lowe, 2004; Traver and Bernardino, 2010). HMAX and subsequent cortex-like models established hierarchical alternation between selectivity and tolerance as a practical recognition strategy (Riesenhuber and Poggio, 1999; Serre et al., 2007).

AWE-3 differs through its radiometric front end, explicit retinal/LGN streams, compact IT code, evidence preservation, and operational few-shot registration.

9.2 Invariant scattering

Wavelet scattering is the closest mathematical precedent for fixed-filter hierarchical invariance. Scattering networks cascade wavelet convolutions, modulus nonlinearities, and averaging to obtain translation invariance and deformation stability without learned convolutional filters (Mallat, 2012; Bruna and Mallat, 2013). Published scattering classifiers generally add fitted classifiers or dimensionality reduction. RLT shares the commitment to fixed transformations but emphasizes radiometric adaptation, biological pathway decomposition, active fixation, and prospective relational identity memory.

9.3 Learned ventral-stream models

Goal-driven deep networks are strong predictive models of portions of sensory cortex and achieve high benchmark performance (Yamins and DiCarlo, 2016). RLT does not deny those results. It tests a different engineering proposition: whether a constrained, fixed, interpretable hierarchy can achieve useful identity recognition and few-shot registration without learning its perceptual filters. AWE-3 should therefore be compared against both classical descriptors and modern learned baselines under matched data and compute conditions.

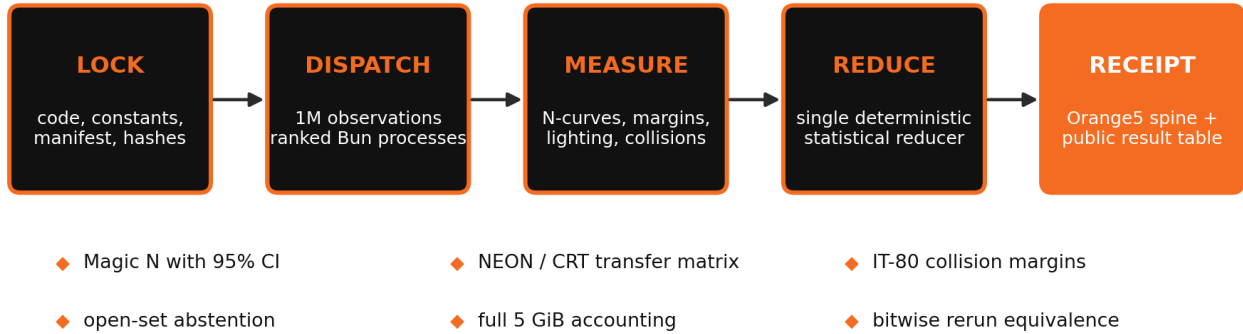
9.4 Radiance fields and computational photography

Neural radiance fields and Gaussian splatting represent scene-dependent radiance for view synthesis (Mildenhall et al., 2020; Kerbl et al., 2023). Their objective is scene reconstruction or rendering and their representations are optimized from multiple views. RLT instead concerns deterministic perceptual decomposition and recurring identity recognition. The shared radiance vocabulary should not be interpreted as algorithmic equivalence.

10. Preregistered scale validation

The next decisive experiment is a locked 10,000-class by 100-observation evaluation, totaling one million observations. Its purpose is not to optimize the feature system. The code, feature weights, thresholds, split manifest, randomization seed, and environment must be frozen before evaluation. Any post-lock change creates a new experiment version.

Figure 4 | Preregistered 10,000-class validation logic



Prospective experiment flow. Independent Bun ranks write immutable shards; one reducer computes the final statistics and receipt.

10.1 Primary endpoints

Magic N. For $N=1$ through 100, compute top-1 accuracy, macro accuracy, worst-decile class accuracy, and bootstrap or Wilson confidence intervals. Magic N is the smallest N whose locked success criteria remain satisfied.

Cross-illuminant transfer. Report reference-lighting $\times$ query-lighting matrices. NEON and CRT-like conditions must be isolated rather than averaged into general lighting accuracy.

Collision behavior. Measure intra-class distance, nearest-impostor distance, classification margin, reciprocal nearest-neighbor collisions, and class-family confusion.

Open-set behavior. Include unknown identities and report false-accept rate, false-reject rate, threshold curves, and calibrated abstention.

Scaling law. Report performance at 47, 100, 250, 500, 1,000, 2,500, 5,000, and 10,000 classes to identify gradual decay or collision phase transitions.

Storage truth. Report vectors, metadata, provenance, database, search index, receipts, and graph overhead separately.

Reproducibility. Repeat selected shards and the final reducer; require bitwise identity or explicitly declared numeric tolerance.

10.2 Required baselines and ablations

The validation should include deterministic baseline descriptors such as raw canonical distance, histogram-based structure, SIFT-family aggregation where licensing and implementation permit, Gabor-only features, and wavelet scattering. A learned baseline may be included as context but should not define success. Ablations should remove or neutralize CAT02, foveation, retinal-12, P/M/K splitting, V2, V4, saccades, and selected IT components one at a time. The winner-plus-additive doctrine requires reporting both improvements and regressions.

11. Limitations and falsification conditions

- The current evidence corpus is small relative to general visual recognition benchmarks and was used during model development.
- The exact degree to which the engineered modules correspond to biological cell classes is limited; the architecture is biologically inspired rather than a cell-level simulation.
- Color constancy cannot recover chromatic information eliminated by a spectrally narrow illuminant or clipped sensor channel. The proper behavior in such cases is reduced confidence or abstention.
- No fixed feature system is immune to distribution shift. AWE-3 avoids learned-weight drift, but its assumptions can fail under unfamiliar optics, severe blur, extreme occlusion, adversarial patterns, or domains outside its design envelope.
- Multi-codec consensus can remove codec-specific residue only to the extent that encodings provide independent distortions and precise alignment. It cannot undo shared source mastering, camera, demosaicing, grading, or clipping.
- The proposed PhotonKnot and constellation graph have not yet been validated at 10,000-class scale.
- Human-level photon vision remains a roadmap target, not an achieved claim. Human vision includes active embodiment, binocular geometry, developmental learning, attention, memory, and task-dependent feedback beyond the present system.

FALSIFICATION CONDITIONS

RLT is weakened if the locked AWE-3 representation collapses as class count increases, if Magic N rises without saturation, if pathological illuminants cannot be handled through confidence-aware adaptation, if codec consensus does not outperform the best single encoding, or if the Pattern Engine fails to improve partial and cross-view recognition beyond compact exemplar memory.

12. Conclusion

Radiance-Luminance Theory reframes visual perception as deterministic measurement and recurrence of structured light evidence. Alpha Wolf Eyes provides an implemented test of the first half of that thesis: a fixed retina-to-IT hierarchy, operating without learned perceptual weights, can produce a compact identity representation and strong few-shot development recognition across controlled lighting variation.

The result should neither be understated nor overstated. AWE-3 is more than a speculative architecture: it is an operating deterministic visual system with receipted development results. It is not yet proof of human-level vision or universal recognition. The 10,000-class experiment is the boundary between an impressive pilot and a statistically defensible scale result. The Pattern Engine is the boundary between compact exemplar recognition and a persistent relational account of identity.

The immediate scientific program is therefore clear: freeze the eye, validate its scaling behavior, expose its failure geometry, and only then drive the toroidal/cylindrical Pattern Engine with real cortical observations. If the proposed predictions survive that sequence, RLT would establish a practical route to visual perception that is local, deterministic, interpretable, reproducible, and grounded in the measurable structure of light.

Declarations

Author contributions. Atom McCree conceived the theory, directed the architecture and experiments, and authored the manuscript.

Declaration of interests. The author declares no competing interests.

Ethics. The reported development fixtures do not involve human or animal subjects. Future evaluations involving identifiable people should use appropriately consented or licensed data.

Data and code availability. The implementation, receipts, and scale-dispatch tooling are maintained within the AtomEons / Orange5 research environment. A public reproducibility package is planned after the 10,000-class configuration and manifest are frozen. No proprietary methods or unpublished benchmark data are required to understand the theoretical claims in this manuscript.

Acknowledgment. The author acknowledges the accumulated work of physicists, neuroscientists, vision scientists, computer-vision researchers, and open-source runtime/tooling communities whose published methods make deterministic visual systems testable.

References

- Adelson, E. H., and Bergen, J. R. (1985). Spatiotemporal energy models for the perception of motion. *Journal of the Optical Society of America A*, 2(2), 284-299.
- Baden, T., Berens, P., Franke, K., Román Rosón, M., Bethge, M., and Euler, T. (2016). The functional diversity of retinal ganglion cells in the mouse. *Nature*, 529(7586), 345-350.
- Baylor, D. A., Lamb, T. D., and Yau, K.-W. (1979). Responses of retinal rods to single photons. *Journal of Physiology*, 288, 613-634.
- Buchsbaum, G. (1980). A spatial processor model for object colour perception. *Journal of the Franklin Institute*, 310(1), 1-26.
- Bruna, J., and Mallat, S. (2013). Invariant scattering convolution networks. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 35(8), 1872-1886.
- Callaway, E. M. (2005). Structure and function of parallel pathways in the primate early visual system. *Journal of Physiology*, 566(1), 13-19.
- Canny, J. (1986). A computational approach to edge detection. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 8(6), 679-698.
- Casagrande, V. A. (1994). A third parallel visual pathway to primate area V1. *Trends in Neurosciences*, 17(7), 305-310.
- Curcio, C. A., Sloan, K. R., Kalina, R. E., and Hendrickson, A. E. (1990). Human photoreceptor topography. *Journal of Comparative Neurology*, 292(4), 497-523.
- Debevec, P. E., and Malik, J. (1997). Recovering high dynamic range radiance maps from photographs. *Proceedings of SIGGRAPH*, 369-378.
- DiCarlo, J. J., Zoccolan, D., and Rust, N. C. (2012). How does the brain solve visual object recognition? *Neuron*, 73(3), 415-434.
- Fairchild, M. D. (2013). *Color Appearance Models* (3rd ed.). Wiley-IS&T.
- Field, D. J. (1987). Relations between the statistics of natural images and the response properties of cortical cells. *Journal of the Optical Society of America A*, 4(12), 2379-2394.
- Field, G. D., and Chichilnisky, E. J. (2007). Information processing in the primate retina: circuitry and coding. *Annual Review of Neuroscience*, 30, 1-30.
- Finlayson, G. D., and Trezzi, E. (2004). Shades of gray and colour constancy. *Color and Imaging Conference*, 37-41.
- Fossum, E. R., Ma, J., Masoodian, S., Anzagira, L., and Zizza, R. (2016). The quanta image sensor: every photon counts. *Sensors*, 16(8), 1260.
- Freeman, W. T., and Adelson, E. H. (1991). The design and use of steerable filters. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 13(9), 891-906.
- Freiwald, W. A., and Tsao, D. Y. (2010). Functional compartmentalization and viewpoint generalization within the macaque face-processing system. *Science*, 330(6005), 845-851.
- Gijsenij, A., Gevers, T., and van de Weijer, J. (2011). Computational color constancy: survey and experiments. *IEEE Transactions on Image Processing*, 20(9), 2475-2489.
- Han, J., Li, B., Mukherjee, D., Chiang, C.-H., Grange, A., Chen, C., et al. (2021). A technical overview of AV1. *Proceedings of the IEEE*, 109(9), 1435-1462.
- Hendry, S. H. C., and Reid, R. C. (2000). The koniocellular pathway in primate vision. *Annual Review of Neuroscience*, 23, 127-153.
- Horn, B. K. P., and Schunck, B. G. (1981). Determining optical flow. *Artificial Intelligence*, 17(1-3), 185-203.
- Hubel, D. H., and Wiesel, T. N. (1962). Receptive fields, binocular interaction and functional architecture in the cat's visual cortex. *Journal of Physiology*, 160(1), 106-154.
- Jensen, H. W., Marschner, S. R., Levoy, M., and Hanrahan, P. (2001). A practical model for subsurface light transport. *Proceedings of SIGGRAPH*, 511-518.

- Jones, J. P., and Palmer, L. A. (1987). An evaluation of the two-dimensional Gabor filter model of simple receptive fields in cat striate cortex. *Journal of Neurophysiology*, 58(6), 1233-1258.
- Kaplan, E., and Shapley, R. M. (1986). The primate retina contains two types of ganglion cells, with high and low contrast sensitivity. *Proceedings of the National Academy of Sciences*, 83(8), 2755-2757.
- Kerbl, B., Kopanas, G., Leimkühler, T., and Drettakis, G. (2023). 3D Gaussian splatting for real-time radiance field rendering. *ACM Transactions on Graphics*, 42(4).
- Land, E. H., and McCann, J. J. (1971). Lightness and retinex theory. *Journal of the Optical Society of America*, 61(1), 1-11.
- Livingstone, M. S., and Hubel, D. H. (1988). Segregation of form, color, movement, and depth: anatomy, physiology, and perception. *Science*, 240(4853), 740-749.
- Lowe, D. G. (2004). Distinctive image features from scale-invariant keypoints. *International Journal of Computer Vision*, 60(2), 91-110.
- Mallat, S. (2012). Group invariant scattering. *Communications on Pure and Applied Mathematics*, 65(10), 1331-1398.
- Marr, D. (1982). *Vision*. W. H. Freeman.
- Marr, D., and Hildreth, E. (1980). Theory of edge detection. *Proceedings of the Royal Society B*, 207(1167), 187-217.
- Masland, R. H. (2012). The neuronal organization of the retina. *Neuron*, 76(2), 266-280.
- Mildenhall, B., Srinivasan, P. P., Tancik, M., Barron, J. T., Ramamoorthi, R., and Ng, R. (2020). NeRF: Representing scenes as neural radiance fields for view synthesis. *European Conference on Computer Vision*.
- Olshausen, B. A., and Field, D. J. (1996). Emergence of simple-cell receptive field properties by learning a sparse code for natural images. *Nature*, 381(6583), 607-609.
- Palmer, S. E. (1999). *Vision Science: Photons to Phenomenology*. MIT Press.
- Pasupathy, A., and Connor, C. E. (1999). Responses to contour features in macaque area V4. *Journal of Neurophysiology*, 82(5), 2490-2502.
- Pasupathy, A., and Connor, C. E. (2002). Population coding of shape in area V4. *Nature Neuroscience*, 5(12), 1332-1338.
- Radford, A., Kim, J. W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., et al. (2021). Learning transferable visual models from natural language supervision. *International Conference on Machine Learning*.
- Reddy, B. S., and Chatterji, B. N. (1996). An FFT-based technique for translation, rotation, and scale-invariant image registration. *IEEE Transactions on Image Processing*, 5(8), 1266-1271.
- Reinhard, E., Stark, M., Shirley, P., and Ferwerda, J. (2002). Photographic tone reproduction for digital images. *ACM Transactions on Graphics*, 21(3), 267-276.
- Rieke, F., and Baylor, D. A. (1998). Single-photon detection by rod cells of the retina. *Reviews of Modern Physics*, 70(3), 1027-1036.
- Riesenhuber, M., and Poggio, T. (1999). Hierarchical models of object recognition in cortex. *Nature Neuroscience*, 2, 1019-1025.
- Roska, B., and Werblin, F. (2001). Vertical interactions across ten parallel, stacked representations in the mammalian retina. *Nature*, 410(6828), 583-587.
- Sandini, G., and Tagliasco, V. (1980). An anthropomorphic retina-like structure for scene analysis. *Computer Graphics and Image Processing*, 14(4), 365-372.
- Schwartz, E. L. (1977). Spatial mapping in the primate sensory projection: analytic structure and relevance to perception. *Biological Cybernetics*, 25(4), 181-194.
- Serre, T., Wolf, L., Bileschi, S., Riesenhuber, M., and Poggio, T. (2007). Robust object recognition with cortex-like mechanisms. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 29(3), 411-426.
- Shafer, S. A. (1985). Using color to separate reflection components. *Color Research and Application*, 10(4), 210-218.
- Simoncelli, E. P., and Olshausen, B. A. (2001). Natural image statistics and neural representation. *Annual Review of Neuroscience*, 24, 1193-1216.
- Sullivan, G. J., Ohm, J.-R., Han, W.-J., and Wiegand, T. (2012). Overview of the High Efficiency Video Coding standard. *IEEE Transactions on Circuits and Systems for Video Technology*, 22(12), 1649-1668.
- Tanaka, K. (1996). Inferotemporal cortex and object vision. *Annual Review of Neuroscience*, 19, 109-139.
- Traver, V. J., and Bernardino, A. (2010). A review of log-polar imaging for visual perception in robotics. *Robotics and Autonomous Systems*, 58(4), 378-398.
- von der Heydt, R., Peterhans, E., and Baumgartner, G. (1984). Illusory contours and cortical neuron responses. *Science*, 224(4654), 1260-1262.
- von Kries, J. (1902/1970). Chromatic adaptation. In D. L. MacAdam (Ed.), *Sources of Color Science*. MIT Press.
- Wandell, B. A. (1995). *Foundations of Vision*. Sinauer Associates.
- Wang, Z., Bovik, A. C., Sheikh, H. R., and Simoncelli, E. P. (2004). Image quality assessment: from error visibility to structural similarity. *IEEE Transactions on Image Processing*, 13(4), 600-612.
- Yamins, D. L. K., and DiCarlo, J. J. (2016). Using goal-driven deep learning models to understand sensory cortex. *Nature Neuroscience*, 19(3), 356-365.