

Auditing Constraint Asymmetry in Latency-Based Resolution Tests

A Latency-Curve Diagnostic for Neutrality versus Confound-Induced Delay

Jeremy C. Jones

HoldingLight LLC — ORCID 0009-0007-2515-3774 — universalcollapse.com

Series: Universal Collapse Theory — Methods Paper • Version v1.0 • 2026-05 • CC BY 4.0

Abstract

Principle. Latency-based resolution tests are ubiquitous across the behavioral, biological, and physical sciences: drift-diffusion modeling of perceptual decisions (Ratcliff & McKoon 2008; Ratcliff et al. 2016), evidence-accumulation models more broadly (Forstmann et al. 2016; Evans & Wagenmakers 2020), first-passage-time models of biological commitment, and Kramers-escape pictures of metastable transition (Kramers 1940; Hänggi, Talkner & Borkovec 1990). Standard practice fits the drift parameter μ to the same response-time distribution it then helps explain. This is appropriate for model description and within-model comparison but not, by itself, a test of an external constraint-asymmetry claim: the resulting per-condition μ cannot serve as independent evidence that constraint neutrality is what produced the observed delay.

Result. We propose a latency-curve diagnostic with a four-axis falsifier comparison plus a fifth confound case. Independently specify constraint asymmetry ΔK and a functional form $\mu = \beta \cdot \Delta K$; estimate two identifiable objects from non-circular channels: the latency ceiling $C = a^2/\sigma^2$, which has units of time under the chosen DDM scale convention, and the drift-link parameter $\lambda = a\beta/\sigma^2$, which fixes the curve width in audited- ΔK units. Where an external evidence-state channel is available, separate (a, σ) estimates can be obtained, but (C, λ) suffice for the four-axis test. Compare four predicted features of the latency curve against the closed-form prediction $E[\tau](\Delta K) = C \cdot \tanh(\lambda \Delta K)/(\lambda \Delta K)$: ceiling at $\Delta K = 0$ against C ; near-ceiling curvature / fractional-drop scale — the bias scale $|\Delta K|_{-\rho} \approx \sqrt{(3\rho)/|\lambda|}$ at which the curve drops by a fraction ρ of the ceiling; symmetry around $\Delta K = 0$; full-curve shape. A fifth axis — the non-decision-time check, comparing raw RT against decision time T_{er} -corrected per condition — catches conditions where stimulus encoding, motor execution, or reporting delay shifts vary across ΔK . Distinct confounds produce distinct deviation patterns; the protocol is demonstrated on synthetic drift-diffusion cases including a circular-fit failure mode in which fitting μ from RT and choice under an assumed scale convention partially but does not cleanly disambiguate confounds. Application to perceptual-decision iEEG data is specified for execution conditional on the dataset providing speeded RT, choice, and a near-threshold ΔK ladder. The result is operational, not philosophical: it characterizes when a latency-versus-asymmetry curve is evidence of constraint-asymmetry-driven resolution under stated independence and scale conditions.

Keywords: drift-diffusion; first-passage time; constraint asymmetry; perceptual decision; circular fitting; reproducibility; non-decision time; scale identifiability; drift-link parameter; neutrality.

Review target

The review target is split into two acceptance levels, paralleling the structure suggested by external review.

Specification acceptance

Acceptance of v1.0 of this Methods Paper means the reviewer agrees that the proposed audit structure is well-defined: independently specify ΔK and a functional form $\mu = \beta \cdot \Delta K$ (or pre-registered f); estimate β by holdout, choice-only at mid-to-high $|\Delta K|$, or pre-registered specification; estimate the ceiling $C = a^2/\sigma^2$ and drift-link $\lambda = a\beta/\sigma^2$ from channels not using the latency feature being tested; where externally calibrated, obtain separate (a, σ) estimates additionally; report non-decision time T_{er} per condition; compare the audited- $\Delta K = 0$ ceiling, near-ceiling curvature / fractional-drop scale, symmetry around $\Delta K = 0$, full-curve shape, and decision-time-versus-raw-RT consistency against the NDR closed-form (Jones 2026d) prediction; test named confound alternatives. The structural specification can be accepted independently of whether the diagnostic outperforms standard DDM goodness-of-fit on any specific dataset.

Performance acceptance

Acceptance of the protocol's performance — a stronger claim than acceptance of the specification — requires the frozen synthetic demonstrations and a downstream real-data application that satisfies the dataset prerequisites listed in §5. v1.0 presents the synthetic demonstrations; real-data execution against an actual dataset is the work of a future T16 demonstration paper.

Rejection criterion

The protocol should be revised or rejected if ΔK , $C = a^2/\sigma^2$, $\lambda = a\beta/\sigma^2$, z_0 , non-decision time, or the $\Delta K \rightarrow \mu$ mapping cannot be estimated without circular use of the latency curve being explained, or if controlled synthetic cases fail to distinguish neutrality-induced delay from the five named confounds (boundary expansion, reduced effective noise, starting-point bias, hidden asymmetry, non-decision-time shift) and from circular per-condition drift fitting. Acceptance does not require that all confounds be disambiguated from RT/choice data alone; some require external scale-setting or external T_{er} calibration, and the protocol's sharpness is conditional on those external channels being available.

Stack placement

This paper is a Universal Collapse Theory Methods Paper for S_2 applications. Records Across Nature, Life, and Mind (Jones 2026a) defines records as the persistence layer and states the latency signature S_2 ("weak effective bias $\rightarrow$ longer expected resolution time") as a portable empirical claim. The Structuralization of Empiricism (Jones 2026b) treats S_2 as a stabilization test for resolution under near-symmetric constraint fields. The Update Integrity Standard (Jones 2026c) requires an independence audit before S_2 latency curves can support a Level 3 claim. Neutrality Delays Resolution (NDR; Jones 2026d) supplies the formal drift-diffusion bound: Theorem 1 establishes that under one-dimensional drift-diffusion with constant drift, constant noise, fixed symmetric absorbing boundaries, and start point $z = 0$, expected first-passage time $E[\tau]$ is strictly maximized at $\mu = 0$ with diffusion-only ceiling a^2/σ^2 . Auditing Independence in Multi-Channel Measurement (Jones 2026e) is the parallel S_1 Methods Paper this paper takes as architectural template. The present paper translates the S_2 commitments into a field-deployable latency-curve diagnostic and is a citable bridge between the formal lemma and downstream T16 empirical demonstrations. Formal pair: TN- S_2 (Neutrality Delays Resolution). Together they constitute the S_2 formal-protocol pair — the paired Technical Note proves the formal bound; this Methods Paper translates it into a deployable audit protocol. It is a Methods Paper, not a standard.

1. Introduction

Resolution latency under near-symmetric alternatives is a recurring observable in the behavioral, biological, and physical sciences. Perceptual decisions slow down at near-threshold stimulus values; cell-fate commitment is delayed when conflicting morphogen gradients near-balance; chemical reactions exhibit Kramers escape times that grow exponentially in the barrier-to-noise ratio (Kramers 1940; Hänggi, Talkner & Borkovec 1990). The qualitative regularity — weakly differentiating constraint fields produce longer expected resolution times — is what the Universal Collapse Theory library calls the S_2 stabilization signature.

Standard practice for fitting evidence-accumulation models to behavioral data is well established: hierarchical drift-diffusion (Wiecki, Sofer & Frank 2013; Vandekerckhove, Tuerlinckx & Lee 2011), fast-dm (Voss & Voss 2007), the Ratcliff toolboxes (Ratcliff & McKoon 2008; Ratcliff et al. 2016), and EZ-DDM-style closed-form fits (Wagenmakers, van der Maas & Grasman 2007) all infer drift, threshold, and starting-point parameters jointly from observed RT distributions. Posterior-predictive checks, deviance-information criteria, and Vincentile plots assess whether the fitted model reproduces the RT distribution. These procedures are appropriate for model description, parameter estimation under stated scale conventions, and within-family model comparison. They are not, by themselves, an independent test of an external constraint-asymmetry claim: a fitted μ that reproduces the RT distribution does not constitute evidence that stimulus or condition coding caused the inferred drift.

This is the methodological gap the present note addresses. Two distinct generating processes can produce nearly indistinguishable RT distributions and indistinguishable per-condition DDM fits:

- (i) Genuine constraint-asymmetry-driven resolution. The drift μ tracks an independently specified ΔK (stimulus coherence, morphogen ratio, contrast level, evidence weight). As $\Delta K \rightarrow 0$, the latency curve approaches the diffusion-only ceiling $C = a^2/\sigma^2$ with the closed-form tanh shape parameterized by $\lambda = a\beta/\sigma^2$.
- (ii) Confound-induced delay. Boundary expansion, reduced effective noise, starting-point bias, hidden asymmetry, or non-decision-time shift produces lengthened RTs that per-condition DDM fits absorb into adjusted parameter estimates. Per-condition goodness-of-fit can match the genuine case under plausible parameter substitutions, but the latency-versus- ΔK curve scales differently and — critically — violates the closed-form ceiling, curvature, symmetry, and shape predictions when $(\Delta K, C, \lambda, z_0, T_{er})$ are specified independently of the RT data.

The diagnostic proposed here is a comparison of the observed latency curve against the closed-form prediction under independently specified $(\Delta K, \beta, C)$ and externally measured $(z_0, T_{er}, \text{optionally } \sigma \text{ separately})$. Five observable features admit independent comparison: ceiling, near-ceiling curvature / fractional-drop scale, symmetry, full-curve shape, and decision-time-versus-raw-RT consistency. Confounds produce distinct, predictable deviation patterns across these axes.

The contribution of this paper is methodological:

- (1) A latency-curve protocol applicable to drift-diffusion-natural domains, with the same four-step structure used in Auditing Independence (Jones 2026e) for S_1 , extended with an explicit non-decision-time audit.
- (2) A confound catalog with predicted curve-shape signatures: boundary expansion, reduced effective noise, starting-point bias, hidden asymmetry, non-decision-time shift, and the partial-circular-fit failure mode.
- (3) A model-conditional diagnostic (sharper to the extent that β, C, λ, z_0 , and T_{er} are estimated independently of the latency curve). The diagnostic is sharp when (β, C, λ) come from holdout or non-latency channels and when (z_0, T_{er}) come from external calibration or distributional fits respecting their identifiability constraints. It is a model-conditional comparison otherwise.

The sharpness comes from comparison against a prediction computed without using the latency curve being explained. This breaks the circularity that affects diagnostics calibrated against the same data they purport to validate. The methodological problem is structurally identical to the one Bland and Altman (1986) flagged for method-comparison correlation in clinical measurement: a goodness-of-fit summary computed on the same data the model was tuned to is not evidence about the data-generating process.

The paper is organized as follows. Section 2 states the formal results: the closed-form prediction (§2.1), the four diagnostic axes plus the non-decision-time check (§2.2), the DDM scale-identifiability

constraint (§2.3), the circular-fit trap (§2.4), and a comparison to existing diagnostics (§2.5). Section 3 specifies the protocol. Section 4 demonstrates the protocol on six synthetic cases including the circular-fit failure. Section 5 specifies the application to perceptual-decision iEEG data, conditional on dataset prerequisites. Sections 6–7 discuss applicable domains and limitations. Section 8 provides a reporting template. Section 9 places the protocol in the broader UCT context. Appendix A provides the simulator core; the full plotting script is archived as the companion file `methods_s2_v10_figures.py`.

2. Theory: latency curves under independent ΔK identification

2.1 The closed-form prediction

Setup. Let Z_t be a one-dimensional drift-diffusion process

$$dZ_t = \mu dt + \sigma dW_t, \quad Z_0 = 0,$$

with drift $\mu \in \mathbb{R}$, noise scale $\sigma > 0$, and W_t standard Brownian motion. The process resolves at $\tau = \inf\{t \geq 0 : |Z_t| = a\}$ where $a > 0$ is the absorbing boundary. The drift μ represents the effective constraint bias, with $\mu = 0$ corresponding to full neutrality.

Theorem 1 (NDR; Jones 2026d). Under the above setup with $\mu \in \mathbb{R}$, $\sigma > 0$, absorbing boundaries at $\pm a$, and $Z_0 = 0$, the expected first-passage time admits the closed form

$$E[\tau] = (a / \mu) \tanh(a\mu / \sigma^2) \text{ for } \mu \neq 0, \quad E[\tau | \mu = 0] = a^2 / \sigma^2.$$

$E[\tau]$ is a strictly decreasing function of $|\mu|$, with maximum a^2 / σ^2 at $\mu = 0$. Full proof from the Kolmogorov backward equation is in Jones 2026d.

Two identifiable objects. It is convenient to introduce two parameters that organize the diagnostic axes. The *latency ceiling* C , defined as

$$C := a^2 / \sigma^2,$$

has units of time under the chosen DDM scale convention (with σ conventionally fixed to set the evidence-state scale, C is in the same time units as τ) and fixes the height of the latency curve at neutrality. Under the constraint-coding map $\mu = \beta \cdot \Delta K$, the

drift-link parameter λ is defined as

$$\lambda := a\beta / \sigma^2,$$

with units inverse to those of ΔK , fixing the curve width in audited- ΔK units. The latency curve as a function of audited asymmetry then admits the compact form

$$E[\tau](\Delta K) = C \cdot \tanh(\lambda \Delta K) / (\lambda \Delta K), \quad E[\tau | \Delta K = 0] = C.$$

For a sharp audit, C and λ must be estimated from channels not using the latency feature being tested: C from speed-accuracy manipulation, cross-dataset calibration, or external evidence-state scale information; λ from independently specified β , held-out choice/accuracy data, or mid-to-high- $|\Delta K|$

choice-proportion information. If C is estimated from the same near-neutrality ceiling that Axis 1 later tests, or λ from the same latency-curve width that Axis 2 later tests, the result is descriptive curve fitting rather than a falsifier.

Starting-point extension. For $Z_0 = z_0$ with $|z_0| < a$, the zero-drift expected first-passage time admits the simpler closed form $E[\tau \mid \mu = 0, z_0] = (a^2 - z_0^2)/\sigma^2$. This makes starting-point bias analytically transparent: a nonzero z_0 depresses the zero-drift ceiling by z_0^2/σ^2 relative to the centered case. The full μ -dependent expression for $z_0 \neq 0$ (Karatzas & Shreve 1991, §2.8) is also even-symmetry-breaking, providing an independent check on Axis 3.

2.2 The four diagnostic axes plus non-decision-time check

The closed-form expression in §2.1 admits four observable features that can be compared independently against fitted values from data, plus a fifth check on raw RT versus decision time. Each axis tests a different facet of the model fit, and each has different scale-convention requirements (§2.3).

Axis 1: Ceiling at $\Delta K = 0$

At $\Delta K = 0$ the prediction is $E[\tau \mid 0] = C = a^2/\sigma^2$. This is the identifiable ratio — a single number, not a separate (a, σ) estimate. Independent estimation of C is possible from speed-accuracy tradeoff manipulations within the dataset, from across-dataset calibration, or from non-latency neural readouts. Observed ceiling values systematically above predicted C indicate effective boundary expansion or noise reduction (not separable from RT/choice alone, see §2.3); values systematically below indicate hidden asymmetry, starting-point bias (where $E[\tau \mid 0, z_0] = (a^2 - z_0^2)/\sigma^2 < C$), or decision-time-versus-RT confusion. The Axis 1 comparison uses the pre-specified audited $\Delta K = 0$ condition, not the empirical maximum over conditions, since hidden asymmetry can shift the peak away from $\Delta K = 0$ while preserving its height (§4.5). The empirical maximum and its ΔK location are reported separately as a peak-shift diagnostic feeding Axis 3.

Axis 2: Near-ceiling curvature / fractional-drop scale

Axis 2 tests the independently estimated drift-link parameter $\lambda = a\beta/\sigma^2$, equivalently the bias scale $1/|\lambda| = \sigma^2/(a|\beta|)$. The ceiling C fixes the height of the curve; λ fixes its width in audited- ΔK units. Under $E[\tau](\Delta K) = C \cdot \tanh(\lambda \Delta K)/(\lambda \Delta K)$, the function $h(x) = \tanh(x)/x$ admits the expansion $h(x) = 1 - x^2/3 + O(x^4)$, so the curve is even in ΔK and the first derivative at $\Delta K = 0$ is zero; the diagnostic is the second derivative $(-2/3) \cdot C \cdot \lambda^2$ at the maximum, equivalently the bias scale at which the curve drops below C by a stated fraction. For a fractional decrement ρ of the ceiling — i.e., $E[\tau] \leq (1 - \rho) \cdot C$ — the bias scale is

$$|\Delta K|_{-\rho} \approx \sqrt{(3\rho) / |\lambda|}.$$

This is the operational quantity to compare: when Axis 1 passes, the ΔK bias level at which the empirical latency curve has dropped 100 ρ % below the predicted ceiling C — equivalently, below the audited- $\Delta K = 0$ value within confidence. When Axis 1 fails, the same calculation is informative as a shape diagnostic but is no longer an independent fractional-drop falsifier and should be reported

both relative to predicted C and relative to the observed audited- $\Delta K = 0$ value. Fitted $|\Delta K|_p$ inconsistent with the right-hand-side prediction (where λ comes from independent estimation) diagnoses misspecification of any of (a, β, σ) . Crucially, Axis 2 is sharp only when λ , or the components defining it, are estimated by holdout, choice-only, or external channels; absent independent identification of λ , the test reduces to model-conditional curve fitting. NDR (Jones 2026d, Theorem 1c) gives the absolute-decrement form $\mu_\epsilon \approx (\sigma^3/a^2)\sqrt{(3\epsilon)}$; the fractional form here expresses the same content in inverse drift-link units.

Axis 3: Symmetry around $\Delta K = 0$ (scale-invariant, but not confound-free)

Under $z_0 = 0$ (centered start point), and with response-mapping symmetric across $\pm\Delta K$ conditions, $E[\tau]$ is an even function of μ : $E[\tau](\mu) = E[\tau](-\mu)$. Asymmetry of the empirical latency curve around $\Delta K = 0$ — specifically, $E[\tau(+\Delta K)] / E[\tau(-\Delta K)]$ departing from unity across audited ΔK levels — diagnoses starting-point bias $z_0 \neq 0$ or hidden asymmetry in the constraint-coding channel.

Axis 3 has a methodological property the other axes lack: it is invariant under the DDM scale convention. Whatever convention is used to fix σ , the symmetry ratio $E[\tau(+\Delta K)] / E[\tau(-\Delta K)]$ depends only on the symmetry of the underlying generating process, not on the choice of measurement scale. Axis 3 is therefore the cleanest of the falsifiers from a scale-identifiability perspective, but it is not confound-free: response-side asymmetries (e.g., motor-side mappings, response-key biases, unequal stimulus encoding delays for + vs. $-\Delta K$ conditions, or condition-asymmetric T_{er}) can also distort the symmetry ratio. A systematic $\pm\Delta K$ latency asymmetry indicates $z_0 \neq 0$ or hidden asymmetry only after response-mapping symmetry and T_{er} asymmetry have been ruled out by external evidence (e.g., counterbalanced response keys, motor-task control conditions, or symmetric encoding-time markers).

Boundary expansion and reduced effective noise both preserve symmetry; they cannot be detected by Axis 3 alone. Axis 3 specifically tests starting-point bias, hidden asymmetry, response-side asymmetry, and T_{er} asymmetry, jointly.

Axis 4: Full-curve shape

The complete observed $E[\tau]$ vs $|\Delta K|$ curve admits comparison against the closed-form $C \cdot \tanh(\lambda\Delta K)/(\lambda\Delta K)$ shape using a pre-registered residual statistic, calibrated by parametric bootstrap, subject-level bootstrap, or posterior predictive simulation under the independently fixed $(C, \lambda, z_0, T_{er})$ values. Large calibrated residuals diagnose either formalism inadequacy (drift-diffusion is the wrong model — collapsing bounds, urgency signals, time-varying drift) or unmeasured confound. This axis catches confound combinations that individually pass Axes 1–3 — e.g., boundary expansion that is partially compensated by reduced effective noise to keep the ceiling at the predicted value while distorting the curve width.

Axis 5: Non-decision-time check

Standard DDM decomposes total response time as $RT_{total} = T_{er} + \tau_{decision}$, where T_{er} is non-decision time absorbing stimulus encoding, motor execution, and reporting delay (Ratcliff & McKoon

2008; Ratcliff & Tuerlinckx 2002). If T_{er} varies across ΔK conditions — because stimulus encoding is harder at low coherence, because motor execution is influenced by response certainty, or because task switching introduces a condition-dependent overhead — raw RT will appear lengthened in conditions where decision time itself was unaffected. The latency curve fit to raw RT will then exhibit ceiling and curve-shape distortions that have nothing to do with constraint asymmetry.

Raw RT equals decision time plus an additive $T_{er}(\Delta K)$ term. If T_{er} is constant across conditions, raw and corrected curves differ only by a vertical offset; if T_{er} varies with ΔK , the raw curve is shape-distorted as well as offset, and the corrected curve should be used for Axes 1–4. The diagnostic check is to estimate $T_{er}(\Delta K)$ per condition (from RT-distribution shape via the Ratcliff & Tuerlinckx 2002 protocol, from explicit motor-task control conditions, or from neural markers of stimulus encoding versus response onset where available) and audit the protocol on $RT_{total} - T_{er}(\Delta K)$. If the four-axis comparison passes on T_{er} -corrected decision time but fails on raw RT, the confound is non-decision-time shift; if it fails on both, the confound is decision-time-internal. Reporting both (corrected and uncorrected) latency curves is the standard.

2.3 DDM scale identifiability

In standard DDM fitting, exactly one parameter must be fixed by convention to set the measurement scale of the evidence state (Ratcliff & McKoon 2008; Ratcliff et al. 2016; HDDM documentation, Wiecki et al. 2013). The most common convention sets the diffusion scale $\sigma = 1$ (or $\sigma = 0.1$ in some toolboxes); a is then estimated in those units. Under this convention, the identifiable objects from RT/choice data are the ratio $C = a^2/\sigma^2$ (the ceiling) and the drift-link $\lambda = a\beta/\sigma^2$ (the curve width), not a , β , and σ separately.

This has direct consequences for the diagnostic. Axis 1 tests C — the identifiable ratio — against the observed ceiling. Axis 2 tests λ via the curvature / fractional-drop scale. Axis 4 tests the full-curve shape under the stated convention. Axis 3 is scale-invariant. Axis 5 is independent of scale convention but depends on T_{er} identifiability.

Boundary expansion and reduced effective noise are not uniquely separable from RT/choice data alone unless β , C , and λ are independently constrained strongly enough to fix the evidence-state scale. Axis 1 detects elevation of C ; Axes 2 and 4 add leverage through λ , but clean separation requires an external estimate of at least one of a , σ , or β . With independently known β , the joint pair (C, λ) constrains a and σ ; without independent β , the boundary-vs-noise ambiguity remains, since C , λ , and the choice-data calibration of β all sit on the same evidence-state scale. The protocol audits the identifiable quantities (C and λ under the stated convention; signed- μ recovery via EZ-DDM at mid-to-high $|\Delta K|$) and reports the residual scale ambiguity honestly. This is a real limitation, not a flaw, and stating it explicitly makes the protocol harder to reject.

2.4 The circular-fit trap

Standard practice in DDM fitting infers μ , a , and z_0 jointly from the same RT distribution that is then asked to validate the model. With a stated scale convention and joint use of RT and choice data (e.g.,

the EZ-DDM closed-form recovery of Wagenmakers et al. 2007: $\mu \approx (\sigma^2/(2a)) \log[P(\text{upper})/P(\text{lower})]$), the per-condition fit is no longer perfectly tautological — the two-channel joint fit constrains the parameters more tightly than either channel alone. But the fit still embeds the assumed scale convention. Under boundary-expansion or reduced-noise data with the wrong scale assumption, the per-condition ($E[\tau]$, μ_{fit}) pairs depart from the prediction curve in characteristic ways (§4.7), but the departure is not as clean as the honest comparison against independently audited ΔK (§4.7, Figure 2A).

The diagnostic proposed here breaks the residual circularity by requiring ΔK to be specified by a separate channel, distinct from both RT and choice data: stimulus coherence levels in random-dot motion paradigms (Britten et al. 1992; Roitman & Shadlen 2002; Palmer, Huk & Shadlen 2005), contrast levels in detection tasks, evidence ratios in token-collection tasks, morphogen ratios in developmental fate-timing studies, externally measured rate constants in chemical transition systems. The diagnostic also requires (C , λ , z_0 , T_{er}) to be identified independently of the latency curve being explained. Without independent identification of these quantities, the four-axis comparison reduces to per-condition curve-fitting under a stated convention, and the falsifier becomes a model-conditional diagnostic rather than a sharp test.

The structural problem here parallels the disjoint-subset trap in S_1 (Jones 2026e, §2.3): an audit that reuses the same data it is supposed to validate cannot discriminate between regimes that share that data. The S_1 case requires disjoint fragments; the S_2 case requires an independent ΔK channel and either holdout or non-latency estimation of β . Both audits fail trivially when this discipline is dropped, and both audits become model-conditional when partial discipline is maintained.

2.5 Diagnostic comparison

Standard DDM goodness-of-fit (deviance-information criteria, posterior predictive checks, Vincentile-plot agreement, parameter-recovery simulations) tests whether the fitted DDM reproduces the empirical RT distribution under a stated scale convention. The proposed diagnostic tests something different: whether the scaling of $E[\tau]$ with independently specified ΔK matches the closed-form prediction under independent identification of (β , C , λ , z_0 , T_{er}). The two are complementary, not substitutes. Posterior-predictive checks catch gross model misspecification (the wrong distributional family, the wrong tail structure); the four-axis-plus- T_{er} comparison catches confounds that preserve distributional structure but violate the scaling prediction. A clean dataset should pass both; a confounded dataset may pass posterior-predictive checks while failing the proposed comparison.

Method-comparison practice in clinical measurement (Bland & Altman 1986) makes the same structural argument: high correlation between two measurement methods does not imply agreement between them, because correlation measures co-variation while agreement measures absolute scaling. The proposed comparison is the resolution-time analogue: DDM goodness-of-fit measures within-condition distributional match, while the four-axis comparison measures across-condition scaling match against an external reference.

Statistical-testing approaches to drift-diffusion (Fudenberg, Newey, Strack & Strzalecki 2020) provide additional rigor for testing DDM-like models against alternatives within the model family. The present diagnostic is concerned with a specific external claim (μ tracks an independently specified ΔK) rather than with model selection per se; the two approaches address adjacent but distinct questions.

3. The protocol

The protocol has four steps plus an explicit scale-convention statement. The first three operate on data and on independent channels; the fourth is the falsifier comparison. The structure parallels Auditing Independence §3 (Jones 2026e), with the channel-characterization step replaced by constraint-asymmetry-channel specification, the effective-sample-size audit replaced by ceiling-and-drift-link audit, and an explicit non-decision-time substep added.

3.1 Step 1: Independent ΔK channel and β specification

Specify the constraint-asymmetry channel before any RT analysis begins. The ΔK channel must produce a quantitative ordering of conditions independently of the RT data it will be compared against. Acceptable channels include: stimulus parameters (motion coherence, luminance contrast, evidence count), externally measured neural drift-rate proxies (e.g., LIP firing-rate ramp slope under fixed stimulus conditions, where the proxy is calibrated against an independent dataset), externally measured physical or chemical rate constants, externally specified condition codings (with explicit pre-registration of ΔK assignments). Unacceptable channels include: μ values fit to the same RT data, condition codings derived post hoc from RT pattern, or any quantity computed from the response data.

Specify the functional form $\mu = f(\Delta K; \theta)$. The simplest and most defensible choice is the linear form $\mu = \beta \cdot \Delta K$, with β a single global scaling parameter. Estimate β from choice/accuracy data in held-out nonzero- ΔK conditions, preferably using mid-to-high $|\Delta K|$ levels that avoid both ceiling latency (where τ is dominated by diffusion variance and β is weakly identified) and saturated choice proportions (where $P(\text{upper})$ approaches 0 or 1 and the EZ-DDM logit estimate becomes unstable). Acceptable channels: (a) non-latency evidence such as neural drift-rate slopes calibrated on an independent dataset; (b) choice/accuracy data alone via the EZ-DDM closed form $P(\text{upper}) = 1/(1 + \exp(-2a\mu/\sigma^2))$ at mid-to-high- $|\Delta K|$ conditions; (c) cross-validation across held-out ΔK levels (fit β on $|\Delta K| \geq \Delta K_{\text{train}}$, test on $|\Delta K| < \Delta K_{\text{train}}$). Estimating β by latency-curve fit on the same data is acceptable only as a model-conditional exercise, not as a sharp test.

3.2 Step 2: Latency curve construction

For each condition c with specified ΔK_c , estimate $E[\tau_c]$ from RT data using a robust mean-RT estimator with bootstrap confidence intervals (10^3 resamples standard). Plot $E[\tau]$ vs ΔK , preserving sign of ΔK (the curve should be plotted on signed ΔK , not folded to $|\Delta K|$, so Axis 3 can be evaluated). Constructing the latency curve uses RT data; fitting the closed-form prediction does not. The two are

separated explicitly so the comparison in Step 4 is between an empirical curve and a model-derived prediction, not between two fits to the same data.

3.3 Step 3: Scale convention, ceiling and drift-link audit, external channels

State the scale convention explicitly: which of (a, σ) is fixed and at what value. The standard choice in cognitive DDM practice is $\sigma = 1$ (or $\sigma = 0.1$; either is fine as long as it is stated). Estimate the ceiling $C = a^2/\sigma^2$ from speed-accuracy tradeoff manipulations, cross-dataset calibration, external evidence-state scale information, or a held-out high- $|\Delta K|$ training subset not used in the Axis 1–4 latency-curve test. If same-experiment behavioral fits are used without holdout, report the result as model-conditional rather than as a sharp audit. Estimate $\lambda = a\beta/\sigma^2$ either directly from held-out choice proportions, which identify $a\mu/\sigma^2 = \lambda\Delta K$ under the fixed-boundary DDM, or from an independently estimated β together with C under the explicitly stated scale convention. If β is not in the same evidence-state scale convention as C , λ is not identified without an additional estimate of a , σ , or a/σ^2 . Both estimates should come from channels not using the latency feature being tested in Step 4 (§2.1).

Where an external evidence-state calibration is available — e.g., a neural readout of effective trial-to-trial noise that fixes σ absolutely, or a speed-accuracy manipulation that fixes a explicitly — separate (a, σ) estimates can be obtained, and the boundary-expansion / reduced-effective-noise confounds become separately diagnosable. Where such calibration is not available, the protocol audits C as a single quantity, audits λ separately for additional leverage when β is independently constrained, and reports the residual ambiguity honestly per §2.3.

Estimate non-decision time T_{er} per condition. Acceptable estimators: minimum-RT or fast RT-quantile (in the spirit of Ratcliff & Tuerlinckx 2002), explicit motor-task control conditions, neural markers of stimulus encoding versus response onset, or an EZ-DDM joint fit that returns T_{er} per condition. Report both raw RT and T_{er} -corrected decision time; the four-axis comparison runs on the corrected curve.

Estimate starting-point z_0 from choice-proportion bias at $\Delta K = 0$ (where $P(\text{upper})$ departs from 0.5 only if $z_0 \neq 0$ and the zero- ΔK stimulus is externally validated as drift-neutral; absent such validation, the same choice-proportion shift may reflect hidden asymmetry rather than starting-point bias) or from explicit response-bias manipulations within the dataset. Report z_0 estimates with confidence intervals and disclose any external drift-neutrality validation step.

3.4 Step 4: Falsifier comparison

Compare five predicted features against fitted values:

- Axis 1 — Ceiling. Predicted: $C = a^2/\sigma^2$ from independent estimate, or $(a^2 - z_0^2)/\sigma^2$ if $z_0 \neq 0$. Observed: $E[\tau]$ at the pre-specified audited $\Delta K = 0$ condition, or a pre-registered fitted intercept at $\Delta K = 0$ if no exact zero condition is present, in either case computed on T_{er} -corrected decision time. Also report the empirical maximum and its ΔK location as a separate peak-shift diagnostic feeding Axis 3, since hidden asymmetry can yield a peak at

the correct ceiling height but at $\Delta K \neq 0$. Test: $|\text{observed} - \text{predicted}|$ within bootstrap confidence interval.

- Axis 2 — Near-ceiling curvature / fractional-drop scale. Predicted: $|\Delta K|_{\rho} \approx \sqrt{(3\rho)/|\lambda|}$ for stated fractional drop ρ (e.g., $\rho = 0.1$), with $\lambda = a\beta/\sigma^2$ from independent estimation. Observed: if Axis 1 passes, the empirical ΔK bias level at which the latency curve falls $100\rho\%$ below the predicted ceiling C — equivalently, below the audited- $\Delta K = 0$ estimate within confidence. If Axis 1 fails, report Axis 2 both relative to predicted C and relative to the observed audited- $\Delta K = 0$ value; the latter is descriptive shape information, not an independent fractional-drop falsifier. Test: ratio observed/predicted within confidence band; ratio < 1 diagnoses a higher effective $|\lambda|$ than the independent estimate.
- Axis 3 — Symmetry (scale-invariant). Predicted: $E[\tau](+\Delta K) / E[\tau](-\Delta K) = 1$ across audited ΔK levels. Observed: empirical ratio at each ΔK , plus the peak-shift quantity from Axis 1. Test: ratio confidence intervals consistent with unity. Systematic departure diagnoses $z_0 \neq 0$ or hidden asymmetry only after response-mapping symmetry and T_{er} asymmetry are independently ruled out; otherwise the asymmetry diagnosis remains ambiguous between these four sources. This axis requires no scale convention and no external (a, σ) channel; it does require external response-side and T_{er} audit.
- Axis 4 — Full-curve shape. Predicted: $E[\tau](\Delta K) = C \cdot \tanh(\lambda \Delta K) / (\lambda \Delta K)$ with (C, λ) fixed at independent values. Observed: residual statistic of the empirical curve against the closed form. Test: a pre-registered residual statistic, calibrated by parametric bootstrap, subject-level bootstrap, or posterior predictive simulation under the independently fixed $(C, \lambda, z_0, T_{\text{er}})$ values. Large calibrated residuals diagnose formalism inadequacy or unmeasured confound.
- Axis 5 — Non-decision-time check. Predicted: raw $RT = \text{decision time} + T_{\text{er}}(\Delta K)$ per condition; if T_{er} is constant across conditions, raw and corrected curves differ only by a vertical offset; if T_{er} varies with ΔK , the raw curve is shape-distorted as well as offset. Observed: raw- RT curve, T_{er} -corrected decision-time curve, and per-condition $T_{\text{er}}(\Delta K)$ estimate. Test: corrected curve passes Axes 1–4 while raw curve fails diagnoses non-decision-time shift; both failing diagnoses decision-time-internal confound.

Pass the diagnostic if all five axes are within confidence; flag for confound diagnosis if one or more axes deviate. The confound catalog in §4 maps deviation patterns to candidate confounds. Failure of multiple axes simultaneously indicates either complete model misspecification, that the constraint-asymmetry channel itself is invalid, or that multiple confounds are present and require disentanglement.

4. Synthetic demonstrations

Six synthetic cases demonstrate the protocol — a clean baseline plus five named confounds. All simulations use one-dimensional drift-diffusion with Euler discretization, $dt = 5 \times 10^{-4}$, 8,000 trials per ΔK level, and $\Delta K \in \{\pm 2.0, \pm 1.5, \pm 1.0, \pm 0.75, \pm 0.5, \pm 0.3, \pm 0.2, \pm 0.1, \pm 0.05, 0\}$. Reference parameters under the analyst's assumed model are $a = 1$, $\sigma = 1$ (stated convention), $z_0 = 0$, $T_{\text{er}} = 0$, with $\beta = 1$

such that $\Delta K = \mu$ in the simulations. The analyst's prediction is therefore $E[\tau] = (1/\mu) \tanh(\mu)$ with $C = 1.000$ and $\lambda = 1.000$. Reproduction code is in Appendix A and in the companion file `methods_s2_v10_figures.py`.

4.1 Clean DDM (passes all five axes)

Clean drift-diffusion with stimulus-coded $\Delta K = \mu$, $a = 1$, $\sigma = 1$, $z_0 = 0$, $T_{er} = 0$. Observed ceiling at $\Delta K = 0$: 1.024 (predicted 1.000; small upward bias is the finite-time-step boundary-detection bias of Euler simulation, present in NDR Figure 1 as well). Symmetry ratio: within 5% across audited ΔK levels. Curve-shape residual: consistent with the tanh prediction. T_{er} -corrected curve identical to raw RT (no offset). All five axes pass. Figure 1A.

4.2 Boundary-expansion confound

True $a = 1.3$ instead of analyst-estimated $a = 1.0$; $\sigma = 1$, $z_0 = 0$, $T_{er} = 0$, $\Delta K = \mu$. True $C = 1.69$; observed ceiling 1.732 (consistent with the true ratio modulo Euler bias). Analyst's predicted ceiling 1.00 violated by 73%. True $\lambda = 1.3$, so $|\Delta K|_{0.1} \approx \sqrt{0.3/1.3} \approx 0.42$ vs. analyst-predicted 0.55 — the curve narrows. Symmetry preserved (Axis 3 passes). Axes 1, 2, and 4 fail; Axes 3 and 5 pass. Diagnostic signature: ceiling overshoot with narrowed curve; symmetry preserved. Without an external evidence-state channel or independent β , this signature is degenerate with reduced effective noise (§4.3); see §2.3. Figure 1B.

4.3 Reduced-effective-noise confound

True $\sigma = 0.7$ instead of analyst-estimated $\sigma = 1.0$; $a = 1$, $z_0 = 0$, $T_{er} = 0$, $\Delta K = \mu$. True $C = 2.04$; observed ceiling 2.081. Analyst's predicted ceiling 1.00 violated by 104%. True $\lambda = 1/0.49 \approx 2.04$, so $|\Delta K|_{0.1} \approx \sqrt{0.3/2.04} \approx 0.27$ vs. predicted 0.55 — the curve narrows much more than under boundary expansion. The full curve is above the prediction and sharper than the boundary-expansion case. Symmetry preserved. Axes 1, 2, and 4 fail; Axes 3 and 5 pass. Diagnostic signature: ceiling overshoot inversely proportional to σ^2 with sharper curve narrowing than boundary-expansion. Figure 1C. The (Axis 1) elevation alone is degenerate with boundary expansion; Axes 2 and 4 add leverage through λ (which scales as $1/\sigma^2$ vs. as a linearly), but clean separation requires an external estimate of one of (a, σ, β) per §2.3.

4.4 Starting-point bias

True $z_0 = 0.4$ instead of analyst-assumed $z_0 = 0$; $a = 1$, $\sigma = 1$, $T_{er} = 0$, $\Delta K = \mu$. The closed-form zero-drift prediction is $E[\tau \mid \mu = 0, z_0] = (a^2 - z_0^2)/\sigma^2 = (1 - 0.16)/1 = 0.84$. Observed at $\Delta K = 0$: 0.877 (consistent with 0.84 modulo Euler bias). The latency-curve maximum shifts away from $\Delta K = 0$ toward negative ΔK , since with $z_0 > 0$ a small negative drift compensates for the bias toward the upper boundary and maximizes residence time. The compensating drift depends on (a, σ, z_0, β) ; for the parameters here it is numerically near -0.4 but is not in general equal to $-z_0$. Symmetry ratio $E[\tau(+\Delta K)] / E[\tau(-\Delta K)]$ departs from unity across the audited range. Axis 1 fails (ceiling below predicted 1.00 by exactly z_0^2 in noise units); Axis 3 fails (asymmetric); Axes 2 and 4 fail with a

distinctive shape residual. Axis 5 passes (no T_{er} issue). Diagnostic signature: ceiling depressed by z_0^2/σ^2 with asymmetric curve. Figure 1D.

4.5 Hidden asymmetry

Generated with $\mu_{true} = \beta \cdot (\Delta K + \delta)$ where $\delta = 0.25$; fitted with assumed $\mu = \beta \cdot \Delta K$ (i.e., the analyst's coding underestimates the constraint-asymmetry origin). $a = 1$, $\sigma = 1$, $z_0 = 0$, $T_{er} = 0$. At audited $\Delta K = 0$, the true drift is 0.25, so observed $E[\tau \mid \text{audited } 0] \approx (1/0.25) \tanh(0.25) = 0.978$; observed 1.012 (consistent modulo Euler bias). The latency-curve maximum shifts to audited $\Delta K \approx -\delta = -0.25$ (marked by the vertical guide in Figure 1E), where the apparent and true drift cancel. This is precisely why Axis 1 uses the audited $\Delta K = 0$ condition rather than the empirical maximum: under hidden asymmetry the maximum can sit close to the predicted ceiling height but at the wrong ΔK . The audited-zero ceiling is depressed (here from $C = 1.00$ to about 0.98) while the peak-shift diagnostic (Axis 3) flags the offset. Axes 2, 3, and 4 fail; Axis 1 marginally passes on ceiling magnitude but the peak-shift diagnostic flags the confound. Axis 5 passes. Diagnostic signature: peak shifted off $\Delta K = 0$; audited-zero ceiling slightly depressed; symmetry violated. Figure 1E.

4.6 Non-decision-time shift

$T_{er}(\Delta K) = 0.4 + 0.15 \cdot |\Delta K|$ applied to RT_{total} ; analyst plots raw RT without subtracting T_{er} . $a = 1$, $\sigma = 1$, $z_0 = 0$, true $\mu = \Delta K$. At $\Delta K = 0$, observed raw RT = 1.419 (decision time 1.024 + T_{er} 0.4 $\approx$ 1.42). The raw-RT curve is offset by ≈ 0.4 plus a $|\Delta K|$ -dependent additive shift that mimics curve-shape deviation, since $T_{er}(\Delta K)$ is not constant across conditions. Axes 1, 2, and 4 appear to fail on raw RT but pass on T_{er} -corrected decision time. Axis 5 (the non-decision-time check) is what detects this confound: the residual between raw-RT and T_{er} -corrected curves matches the estimated $T_{er}(\Delta K)$ profile exactly. Diagnostic signature: raw-RT failures disappear under T_{er} correction. Figure 1F.

Figure 1. Confound discrimination via the latency curve (predicted under $a=\sigma=1$, ceiling $a^2/\sigma^2=1$)

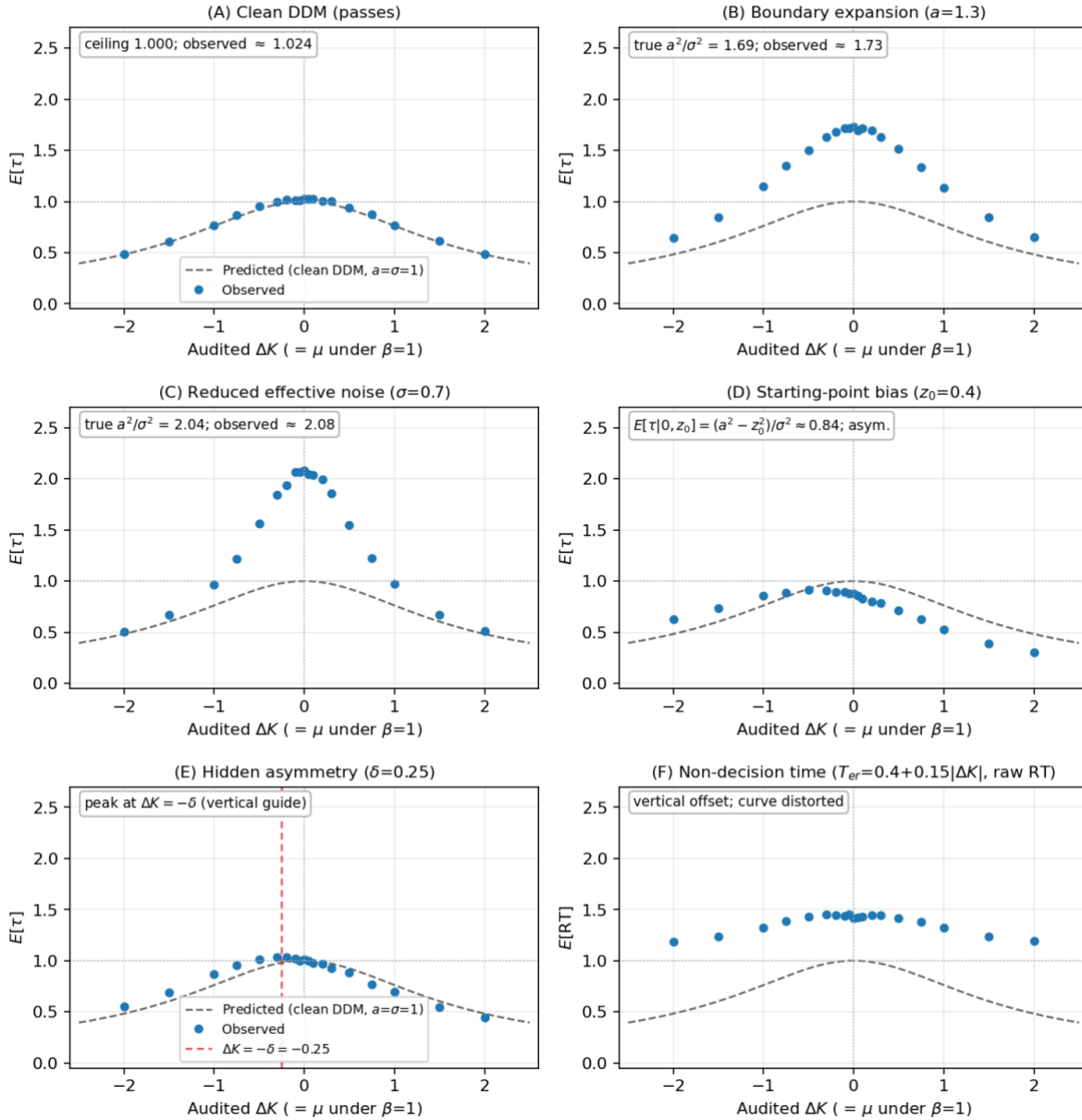


Figure 1. Confound discrimination via the latency curve. Each panel shows the predicted $E[\tau]$ curve under the analyst's assumed model (dashed grey, $a = \sigma = 1$, ceiling $C = 1$) against simulated observations (blue points, 8,000 trials per ΔK level). Panels: (A) clean DDM, (B) boundary expansion, (C) reduced effective noise, (D) starting-point bias, (E) hidden asymmetry (red vertical guide marks the predicted off-zero peak at audited $\Delta K = -\delta$), (F) non-decision time (raw RT). Per-panel parameters and pass/fail patterns are discussed in §4.1–4.6.

4.7 The circular-fit failure mode

Each of the three primary confound regimes in §4.2–4.3 plus the clean baseline produces RT distributions and choice proportions at each ΔK level. The honest comparison plots $E[\tau]$ versus the

audited ΔK (Figure 2A); the three regimes are clearly separable. The circular fit recovers signed μ per condition from choice proportions under the EZ-DDM closed-form $\mu \approx (\sigma^2/(2a)) \log[P(\text{upper})/P(\text{lower})]$ (Wagenmakers et al. 2007), with the analyst's assumed ($a = 1, \sigma = 1$). Plotting $E[\tau]$ versus this fitted μ (Figure 2B) shifts the regimes' placement relative to Panel A but does not collapse them onto the prediction. Under boundary-expansion data, choice proportions are driven by the true $a \cdot \mu / \sigma^2$ ratio; the analyst's assumed (a, σ) gives $\mu_{\text{fit}} = (\text{true } a) / (\text{assumed } a) \cdot \text{true } \mu$, and the $(\mu_{\text{fit}}, E[\tau_{\text{obs}}])$ points sit clearly above the prediction. Reduced-noise data behaves analogously.

The methodologically honest reading is: the circular fit using both RT and choice constrains the parameters more tightly than RT alone, and confounded data is not perfectly tautologically explained — the $(\mu_{\text{fit}}, E[\tau])$ points still depart from the prediction in characteristic ways. But the departure is more compressed than in Panel A, and the signature of the underlying confound is harder to read off the circular-fit plot than off the audited- ΔK plot. The independent ΔK channel and stated scale convention together are what make the comparison sharp; the circular fit alone is a model-conditional check that complements but does not replace the four-axis comparison.

An audit that uses only the data twice (once to fit μ , once to validate the fit) remains weaker as evidence than an audit that uses an external ΔK channel. The diagnostic proposed in this paper requires ΔK to come from outside the RT/choice data stream. The same applies to the independent-channel discipline of S_1 (Jones 2026e, §2.3) and to method-comparison studies more broadly (Bland & Altman 1986).

Figure 2. The circular-fit trap (with signed- μ recovery)

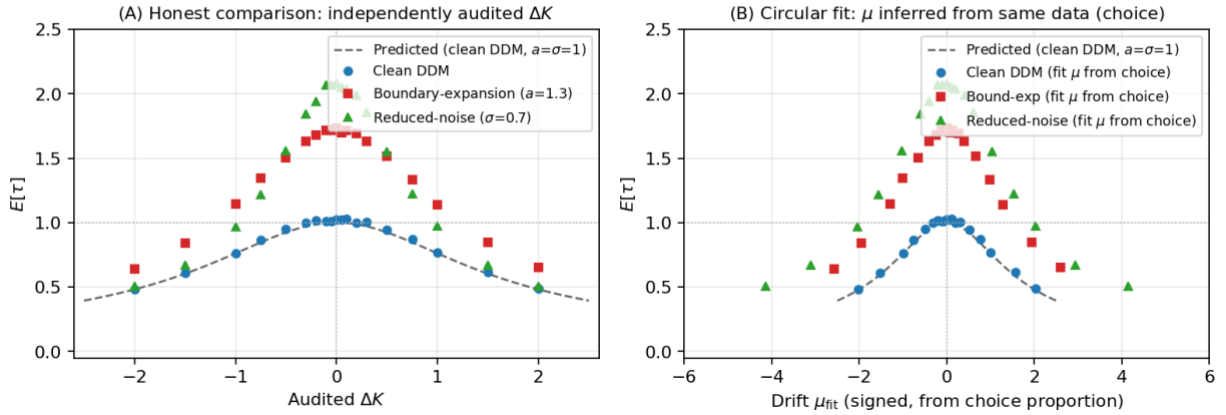


Figure 2. The circular-fit trap with signed- μ recovery via choice proportions. (A) Honest comparison plots mean decision time against independently audited ΔK ; clean DDM points follow the predicted curve, while boundary-expansion and reduced-noise confounds sit above it. (B) Circular comparison plots the same observations against μ_{fit} recovered from choice proportions under the analyst's assumed $a = \sigma = 1$. Confounded points shift along the x-axis but do not collapse onto the prediction. The circular fit is therefore a model-conditional check, not a sharp test.

5. Application: perceptual-decision iEEG (specification, conditional on dataset prerequisites)

We specify the application of the protocol to perceptual-decision iEEG data, paralleling the way Auditing Independence §5 (Jones 2026e) specifies application to Planck PR3 component-separated CMB maps. Full execution against a real dataset is the work of a downstream T16 demonstration paper; the present specification is the analysis plan, conditional on the dataset providing the required ingredients.

5.1 Dataset prerequisites

A dataset is a candidate application for this protocol only if it satisfies all of the following: (1) speeded response times with explicit time pressure or a stated speed-vs-accuracy regime; (2) recorded choice outcomes (which alternative was selected on each trial); (3) a near-threshold stimulus ladder spanning at least three ΔK levels including a near-zero condition; (4) both signs of ΔK available (or a paradigm in which the symmetry axis can be tested across other axes of the design); (5) an externally specified stimulus parameter (motion coherence, luminance contrast, evidence count, etc.) providing the ΔK channel; (6) sufficient trial count per condition for stable mean-RT estimates with bootstrap confidence intervals.

Datasets that lack any of these prerequisites cannot use this protocol as specified. Free-response paradigms without time pressure, paradigms that record only RT without choice, or paradigms that vary stimulus parameters in ways that do not deliver a clean ΔK ladder all require formalism extensions and are deferred.

5.2 Candidate dataset families

Random-dot motion paradigms (Britten et al. 1992; Roitman & Shadlen 2002; Palmer, Huk & Shadlen 2005) are the canonical fit: motion coherence $c \in \{0\%, 3.2\%, 6.4\%, 12.8\%, 25.6\%, 51.2\%\}$ provides a monotone, externally specified ΔK channel that is independent of RT. Equivalent paradigms (luminance-contrast detection, evidence-token accumulation, lexical-decision frequency manipulation) admit the same protocol provided they meet the prerequisites in §5.1.

Consciousness-theory iEEG corpora such as the COGITATE multi-center dataset (Cogitate Consortium 2025; see also the consortium’s adversarial-collaboration protocol, Melloni et al. 2023) are candidate applications only where the specific paradigm within the dataset includes speeded RT, choice outcomes, and an independently coded ΔK ladder suitable for evidence-accumulation modeling. Not every paradigm in such a corpus satisfies these prerequisites; the dataset must be screened against §5.1 before the protocol is applied.

5.3 Analysis plan (for a satisfying dataset)

(1) Pre-register the ΔK channel (stimulus parameter levels) and the functional form $\mu = \beta \cdot \Delta K$. (2) For each subject and ΔK level, estimate mean RT and choice proportion with bootstrap CIs. (3) Estimate β from mid-to-high- $|\Delta K|$ conditions (avoiding both ceiling latency and saturated choice proportions)

via either choice proportions alone or held-out cross-validation. (4) State the scale convention (e.g., $\sigma = 1$) and estimate $C = a^2/\sigma^2$ from speed-accuracy tradeoff manipulations within the dataset where available; otherwise from a single-convention DDM fit on high- $|\Delta K|$ trials. Compute $\lambda = a\beta/\sigma^2$ from β and C , both estimated from channels not using the latency feature being tested. (5) Estimate T_{er} per condition via the Ratcliff & Tuerlinckx 2002 protocol or motor-task control conditions. (6) Estimate z_0 from choice-proportion bias at $\Delta K = 0$, conditional on external validation that the zero- ΔK stimulus is drift-neutral. (7) Predict ceiling, $|\Delta K|_{-p} = \sqrt{(3\rho)/|\lambda|}$ for stated ρ (e.g., 0.1), and curve shape from the independent estimates. (8) Compare observed latency curve (on T_{er} -corrected decision time) against prediction across all five axes.

5.4 Falsifier

The clean S_2 interpretation for this paradigm fails if any axis deviates beyond its pre-specified confidence band. If the deviation pattern matches one of the named confounds in §4, the result is diagnosed as confounded rather than as evidence against the drift-diffusion formalism itself. If the deviation pattern does not match a named confound, the likely causes are formalism inadequacy, an invalid ΔK channel, or multiple interacting confounds. Combined-confound deviations (e.g., partial boundary expansion compensated by partial noise reduction; non-decision-time shift on top of starting-point bias) are diagnosable from the curve-width residual in Axis 4 and the corrected-vs-raw comparison in Axis 5, but with reduced sharpness and the explicit caveats of §2.3.

6. Discussion

6.1 Domains where the diagnostic applies directly

The drift-diffusion formalism is directly underwritten in perceptual decision-making (Ratcliff & McKoon 2008; Ratcliff et al. 2016; Forstmann, Ratcliff & Wagenmakers 2016) and in lexical decision and value-based choice paradigms with two-alternative response structure (Krajbich, Armel & Rangel 2010; Ratcliff 1978 for memory retrieval). In these domains the protocol applies as specified.

Closely related but requiring formalism adaptation: developmental fate-timing under stochastic threshold-crossing (cell-fate commitment with morphogen-ratio-coded ΔK , Gillespie-style stochastic simulation following Gillespie 1977 rather than drift-diffusion), phenotypic switching under fluctuating selection (Süel et al. 2007; Balázsi, van Oudenaarden & Collins 2011), absorbing-Markov-chain models in genotype-phenotype neutral-network traversal (Kemeny & Snell 1976; Ahnert 2017), and Kramers-escape regimes in chemical metastability (Kramers 1940; Hänggi, Talkner & Borkovec 1990). These domains are plausible candidate instances of the qualitative S_2 hypothesis but require their own Methods Paper for sharp falsification, since the closed-form drift-diffusion ceiling does not transfer directly.

6.2 The independent ΔK channel as the load-bearing step

Step 1 of the protocol — specifying ΔK from a channel separable from the RT/choice data — is the load-bearing step. Without it, the diagnostic collapses to model-conditional curve-fitting under a stated convention, and §4.7 applies: the fitted curves shift but do not cleanly disambiguate confounds. This parallels the independence-audit discipline in S_1 (Jones 2026e, §6.2): the formal lemma is not what does the evidential work in any real domain; the audit is.

In practice, ΔK specification is often the part most resistant to formalization. Stimulus parameters are clean. Behavioral condition codings are less clean but still acceptable when pre-registered. Post hoc ΔK assignments derived from RT pattern are not acceptable. Composite ΔK channels (combining stimulus parameters with subject-rated difficulty) require explicit pre-registration of the combination rule. The protocol does not enforce a particular ΔK specification; it enforces that the specification be made before the latency curve is constructed.

6.3 Comparison to alternative diagnostics

Standard DDM diagnostics (posterior-predictive checks, deviance-information criteria, Vincentile plots, parameter-recovery simulations) test model-fit adequacy on the same RT distribution the model was fit to. The proposed diagnostic tests something distinct: scaling of the latency curve against an independent channel, under stated scale convention and with explicit T_{er} audit. Standard goodness-of-fit is appropriate for model description and within-family comparison; it is not by itself an independent test of an external constraint-asymmetry claim. The two are complementary, not substitutes.

Statistical tests of DDM-like models against alternatives (Fudenberg, Newey, Strack & Strzalecki 2020) provide additional rigor for testing model adequacy within the evidence-accumulation family. The present protocol addresses an adjacent question: not whether DDM fits the data, but whether a specific external claim about constraint asymmetry is supported by the latency curve's scaling. Both tools are useful; they answer different questions.

Method-comparison practice in clinical measurement (Bland & Altman 1986) makes the same structural argument: high correlation between two measurement methods does not imply agreement between them, because correlation measures co-variation while agreement measures absolute scaling. The proposed comparison is the resolution-time analogue.

7. Limitations

Formalism scope. The protocol is restricted to one-dimensional drift-diffusion with constant drift, constant noise, and fixed symmetric absorbing boundaries. Real DDM extensions — collapsing bounds (Hawkins et al. 2015), urgency signals (Cisek et al. 2009), time-varying drift, hierarchical-DDM with subject-level random effects — require formalism extension. The five-axis comparison can in principle accommodate these extensions: each extension produces a different closed-form (or numerically computed) prediction for the axes, and the diagnostic generalizes by replacing the tanh-

shaped prediction with the appropriate alternative. v1.0 development should formalize the collapsing-bounds extension as a high-priority extension target.

Scale separability. As discussed in §2.3, separating a from σ individually requires either an external evidence-state channel or independent constraint of β sharp enough to fix the scale via the joint pair (C, λ) . Boundary-expansion and reduced-effective-noise confounds are not uniquely separable from RT/choice data alone in the absence of such external constraint; the protocol audits the identifiable quantities and reports the residual ambiguity honestly.

Heterogeneous-condition and hierarchical settings. Where (a, σ, z_0, T_{er}) vary across conditions, per-condition independent estimates are required. Pooled estimates introduce an additional confound channel (variability across conditions absorbed into fitted μ). v1.0 development should formalize the heterogeneous case explicitly, paralleling the heterogeneous-fragments treatment in TN-S₁ (Jones 2026f).

Contaminant RTs and lapse trials. Real RT data contains contaminants (very fast guesses, very slow lapses, attention-failure trials) that distort mean-RT estimates and DDM parameter recovery (Ratcliff & Tuerlinckx 2002). The protocol assumes that contaminant trials are screened or modeled before the latency curve is constructed; v1.0 should formalize this preprocessing step.

Quantum-decision analogues. Quantum-cognitive models of decision (Pothos & Busemeyer 2009; Busemeyer & Bruza 2012) reframe the underlying state space and the resolution process; the drift-diffusion formalism does not directly apply. A separate Methods Paper would be required.

Kramers-escape and absorbing-Markov regimes. Developmental fate-timing, phenotypic switching, neutral-network traversal, and chemical metastability are S₂ candidate domains where the drift-diffusion ceiling does not transfer. The qualitative S₂ hypothesis may have analogues in these regimes, but the falsifier shape differs (e.g., Kramers-rate scaling $\exp(-E_b/k_{BT})$ for chemical metastability; expected-hitting-time bounds on absorbing-Markov chains for neutral-network traversal). These are deferred to a future Methods Paper.

8. Reporting template

Authors using this protocol should report the following items. Items marked (R) are required for the comparison to be evaluable; items marked (S) sharpen the diagnostic but are not strictly required.

- (R) The ΔK channel: source, scale, pre-registration status, dates of pre-registration.
- (R) The functional form $\mu = f(\Delta K; \theta)$ and how β (or other θ) was estimated: holdout, choice-only at mid-to-high $|\Delta K|$ (avoiding saturation), non-latency channel, or model-conditional.
- (R) Stated scale convention (e.g., $\sigma = 1, \sigma = 0.1$), independent estimate of $C = a^2/\sigma^2$ with method, source dataset, and confidence intervals, and independent estimate of $\lambda = a\beta/\sigma^2$. Estimates must come from channels not using the latency feature being tested. Where available, separate (a, σ) estimates from external channels.

- (R) Per-condition T_{er} estimates with method (Ratcliff-Tuerlinckx fit, motor-control condition, neural marker, EZ-DDM joint fit, etc.). Both raw-RT and T_{er} -corrected decision-time curves; explicit statement of whether $T_{er}(\Delta K)$ is constant or varies with ΔK .
- (R) z_0 estimate from choice-proportion bias at $\Delta K = 0$ with confidence interval, including disclosure of any external drift-neutrality validation step.
- (R) Mean $E[\tau_c]$ for each condition c with bootstrap CI (10^3 resamples standard).
- (R) Predicted ceiling C (or $(a^2 - z_0^2)/\sigma^2$ if $z_0 \neq 0$), predicted $|\Delta K|_\rho = \sqrt{(3\rho)/|\lambda|}$ for stated ρ , and predicted curve shape from independent estimates.
- (R) Five-axis comparison: ceiling at audited $\Delta K = 0$ (with empirical max + ΔK location reported separately as peak-shift diagnostic), fractional-drop scale, symmetry ratio, curve-shape residual, raw-vs-corrected RT consistency. Report each with confidence interval and pass/fail status. For Axis 3, also report response-side and T_{er} asymmetry audit results.
- (R) Confound diagnosis where any axis fails: which confound pattern in the catalog of §4 best matches the deviation, and why.
- (S) Posterior-predictive checks of the per-condition DDM fit, demonstrating that the fitted model reproduces the RT distribution. The five-axis comparison is complementary to this, not a replacement.
- (S) Sensitivity analysis: re-run the comparison under perturbations of $(C, \lambda, z_0, T_{er})$ within their confidence intervals. Report which axes pass robustly and which are marginal.

9. Broader context

S_2 is one of three portable empirical signatures in the Universal Collapse Theory library (Jones 2026a, 2026b). S_1 (redundancy implies consensus) has its formal support in TN- S_1 / Objectivity from Records (Jones 2026f) and its Methods Paper in Auditing Independence (Jones 2026e). The present paper is the Methods Paper for S_2 , paralleling that architecture. S_3 (constraint-stable updates produce hysteresis) remains in development.

The structural pattern across these Methods Papers is consistent: the technical note establishes a closed-form bound under load-bearing assumptions; the standard demands an audit of those assumptions before downstream claims can be made; the Methods Paper supplies the field-deployable audit protocol with synthetic demonstrations and a specified real-data application. The downstream T16 demonstration carries the domain-native execution. The Methods Paper is the bridge between the formal lemma and the field claim, and the audit it specifies is what does the evidential work in any real domain.

Appendix A. Reproduction code

Appendix A reproduces the simulator core used for Figures 1 and 2. The full plotting script, including panel layout, annotations, and the vertical guide on Panel E of Figure 1, is archived as the companion

file `methods_s2_v10_figures.py` (archived in this component). The script is about 500 lines, NumPy + Matplotlib only, with no SciPy dependency. The simulator tracks both first-passage time and choice (which boundary was hit), enabling signed- μ recovery in the circular-fit demonstration via the EZ-DDM closed form rather than numerical inversion.

```
import numpy as np

def simulate_dd(mu, sigma=1.0, a=1.0, z0=0.0, t_er=0.0,
               dt=5e-4, max_t=20.0, n_trials=8000, seed=0):
    """1D drift-diffusion with absorbing boundaries at +/- a.
    Returns (mean_RT, mean_decision_time, P(upper))."""
    rng = np.random.default_rng(seed)
    n_steps = int(max_t / dt)
    z = np.full(n_trials, float(z0))
    dt_arr = np.full(n_trials, max_t)
    upper = np.zeros(n_trials, dtype=bool)
    resolved = np.zeros(n_trials, dtype=bool)
    sqrt_dt = np.sqrt(dt)
    for t in range(n_steps):
        if resolved.all(): break
        incs = mu*dt + sigma*sqrt_dt*rng.standard_normal(n_trials)
        z[~resolved] += incs[~resolved]
        hit_up = (z >= a) & (~resolved)
        hit_lo = (z <= -a) & (~resolved)
        newly = hit_up | hit_lo
        dt_arr[newly] = (t+1) * dt
        upper[hit_up] = True
        resolved |= newly
    mean_dt = dt_arr.mean()
    return mean_dt + t_er, mean_dt, upper.mean()

def E_tau_theory(mu, sigma=1.0, a=1.0):
    if abs(mu) < 1e-9: return a*a / (sigma*sigma)
    return (a / mu) * np.tanh(a * mu / (sigma * sigma))

def mu_from_choice(p_upper, a=1.0, sigma=1.0):
    """Signed mu recovery from choice proportion (EZ-DDM)."""
    p = np.clip(p_upper, 1e-4, 1 - 1e-4)
    return (sigma**2 / (2*a)) * np.log(p / (1-p))

# Vary the simulation parameters per confound:
# (A) clean: a=1, sigma=1, z0=0, t_er=0
# (B) bound: a=1.3, sigma=1, z0=0, t_er=0
# (C) noise: a=1, sigma=0.7, z0=0, t_er=0
# (D) sp_bias: a=1, sigma=1, z0=0.4, t_er=0
# (E) hidden: a=1, sigma=1, z0=0, t_er=0; pass mu = DeltaK + 0.25
# (F) nondec: a=1, sigma=1, z0=0, t_er = 0.4 + 0.15*|DeltaK|
```

References

- Ahnert, S. E. (2017). Structural properties of genotype-phenotype maps. *Journal of the Royal Society Interface*, 14(132), 20170275. <https://doi.org/10.1098/rsif.2017.0275>
- Balázsi, G., van Oudenaarden, A., & Collins, J. J. (2011). Cellular decision making and biological noise: from microbes to mammals. *Cell*, 144(6), 910–925. <https://doi.org/10.1016/j.cell.2011.01.030>

- Bland, J. M., & Altman, D. G. (1986). Statistical methods for assessing agreement between two methods of clinical measurement. *The Lancet*, 327(8476), 307–310.
[https://doi.org/10.1016/S0140-6736\(86\)90837-8](https://doi.org/10.1016/S0140-6736(86)90837-8)
- Britten, K. H., Shadlen, M. N., Newsome, W. T., & Movshon, J. A. (1992). The analysis of visual motion: a comparison of neuronal and psychophysical performance. *Journal of Neuroscience*, 12(12), 4745–4765. <https://doi.org/10.1523/JNEUROSCI.12-12-04745.1992>
- Busmeyer, J. R., & Bruza, P. D. (2012). *Quantum Models of Cognition and Decision*. Cambridge University Press. <https://doi.org/10.1017/CBO9780511997716>
- Cisek, P., Puskas, G. A., & El-Murr, S. (2009). Decisions in changing conditions: the urgency-gating model. *Journal of Neuroscience*, 29(37), 11560–11571.
<https://doi.org/10.1523/JNEUROSCI.1844-09.2009>
- Cogitate Consortium (2025). Adversarial testing of global neuronal workspace and integrated information theories of consciousness. *Nature*, 642, 133–142. <https://doi.org/10.1038/s41586-025-08888-1>
- Evans, N. J., & Wagenmakers, E.-J. (2020). Evidence accumulation models: current limitations and future directions. *The Quantitative Methods for Psychology*, 16(2), 73–90.
<https://doi.org/10.20982/tqmp.16.2.p073>
- Forstmann, B. U., Ratcliff, R., & Wagenmakers, E.-J. (2016). Sequential sampling models in cognitive neuroscience: advantages, applications, and extensions. *Annual Review of Psychology*, 67, 641–666. <https://doi.org/10.1146/annurev-psych-122414-033645>
- Fudenberg, D., Newey, W., Strack, P., & Strzalecki, T. (2020). Testing the drift-diffusion model. *Proceedings of the National Academy of Sciences*, 117(52), 33141–33148.
<https://doi.org/10.1073/pnas.2011446117>
- Gillespie, D. T. (1977). Exact stochastic simulation of coupled chemical reactions. *The Journal of Physical Chemistry*, 81(25), 2340–2361. <https://doi.org/10.1021/j100540a008>
- Hänggi, P., Talkner, P., & Borkovec, M. (1990). Reaction-rate theory: fifty years after Kramers. *Reviews of Modern Physics*, 62(2), 251–341. <https://doi.org/10.1103/RevModPhys.62.251>
- Hawkins, G. E., Forstmann, B. U., Wagenmakers, E.-J., Ratcliff, R., & Brown, S. D. (2015). Revisiting the evidence for collapsing boundaries and urgency signals in perceptual decision-making. *Journal of Neuroscience*, 35(6), 2476–2484. <https://doi.org/10.1523/JNEUROSCI.2410-14.2015>
- HDDM documentation. (n.d.). HDDM: hierarchical Bayesian estimation of the drift-diffusion model in Python (online documentation). <https://hddm.readthedocs.io/>
- Jones, J. C. (2026a). *Records Across Nature, Life, and Mind*. HoldingLight LLC.
<https://doi.org/10.17605/OSF.IO/7H6DY>
- Jones, J. C. (2026b). *The Structuralization of Empiricism*. HoldingLight LLC.
<https://doi.org/10.17605/OSF.IO/J4GZ9>

- Jones, J. C. (2026c). Update Integrity Standard (UIS v1.0). HoldingLight LLC.
<https://doi.org/10.17605/OSF.IO/DWM29>
- Jones, J. C. (2026d). Neutrality Delays Resolution: An Expected Resolution-Time Bound (UCT Technical Note v1.0). HoldingLight LLC. <https://doi.org/10.17605/OSF.IO/6WRQV>
- Jones, J. C. (2026e). Auditing Independence in Multi-Channel Measurement (UCT Methods Paper v1.0). HoldingLight LLC. <https://doi.org/10.17605/OSF.IO/7U8SK>
- Jones, J. C. (2026f). Objectivity from Records: An Exponential Consensus Bound (UCT Technical Note v1.0). HoldingLight LLC. <https://doi.org/10.17605/OSF.IO/6M7N3>
- Karatzas, I., & Shreve, S. E. (1991). *Brownian Motion and Stochastic Calculus* (2nd ed.). Springer.
<https://doi.org/10.1007/978-1-4612-0949-2>
- Kemeny, J. G., & Snell, J. L. (1976). *Finite Markov Chains*. Springer.
- Krajovich, I., Armel, C., & Rangel, A. (2010). Visual fixations and the computation and comparison of value in simple choice. *Nature Neuroscience*, 13(10), 1292–1298.
<https://doi.org/10.1038/nn.2635>
- Kramers, H. A. (1940). Brownian motion in a field of force and the diffusion model of chemical reactions. *Physica*, 7(4), 284–304. [https://doi.org/10.1016/S0031-8914\(40\)90098-2](https://doi.org/10.1016/S0031-8914(40)90098-2)
- Melloni, L., Mudrik, L., Pitts, M., Bendtz, K., Ferrante, O., Gorska, U., et al. (2023). An adversarial collaboration protocol for testing contrasting predictions of global neuronal workspace and integrated information theory. *PLOS ONE*, 18(2), e0268577.
<https://doi.org/10.1371/journal.pone.0268577>
- Palmer, J., Huk, A. C., & Shadlen, M. N. (2005). The effect of stimulus strength on the speed and accuracy of a perceptual decision. *Journal of Vision*, 5(5), 376–404.
<https://doi.org/10.1167/5.5.1>
- Pothos, E. M., & Busemeyer, J. R. (2009). A quantum probability explanation for violations of “rational” decision theory. *Proceedings of the Royal Society B*, 276(1665), 2171–2178.
<https://doi.org/10.1098/rspb.2009.0121>
- Ratcliff, R. (1978). A theory of memory retrieval. *Psychological Review*, 85(2), 59–108.
<https://doi.org/10.1037/0033-295X.85.2.59>
- Ratcliff, R., & McKoon, G. (2008). The diffusion decision model: theory and data for two-choice decision tasks. *Neural Computation*, 20(4), 873–922. <https://doi.org/10.1162/neco.2008.12-06-420>
- Ratcliff, R., Smith, P. L., Brown, S. D., & McKoon, G. (2016). Diffusion decision model: current issues and history. *Trends in Cognitive Sciences*, 20(4), 260–281.
<https://doi.org/10.1016/j.tics.2016.01.007>
- Ratcliff, R., & Tuerlinckx, F. (2002). Estimating parameters of the diffusion model: approaches to dealing with contaminant reaction times and parameter variability. *Psychonomic Bulletin & Review*, 9(3), 438–481. <https://doi.org/10.3758/BF03196302>

- Roitman, J. D., & Shadlen, M. N. (2002). Response of neurons in the lateral intraparietal area during a combined visual discrimination reaction time task. *Journal of Neuroscience*, 22(21), 9475–9489. <https://doi.org/10.1523/JNEUROSCI.22-21-09475.2002>
- Süel, G. M., Kulkarni, R. P., Dworkin, J., Garcia-Ojalvo, J., & Elowitz, M. B. (2007). Tunability and noise dependence in differentiation dynamics. *Science*, 315(5819), 1716–1719. <https://doi.org/10.1126/science.1137455>
- Vandekerckhove, J., Tuerlinckx, F., & Lee, M. D. (2011). Hierarchical diffusion models for two-choice response times. *Psychological Methods*, 16(1), 44–62. <https://doi.org/10.1037/a0021765>
- Voss, A., & Voss, J. (2007). Fast-dm: A free program for efficient diffusion model analysis. *Behavior Research Methods*, 39(4), 767–775. <https://doi.org/10.3758/BF03192967>
- Wagenmakers, E.-J., van der Maas, H. L. J., & Grasman, R. P. P. P. (2007). An EZ-diffusion model for response time and accuracy. *Psychonomic Bulletin & Review*, 14(1), 3–22. <https://doi.org/10.3758/BF03194023>
- Wiecki, T. V., Sofer, I., & Frank, M. J. (2013). HDDM: hierarchical Bayesian estimation of the drift-diffusion model in Python. *Frontiers in Neuroinformatics*, 7, 14. <https://doi.org/10.3389/fninf.2013.00014>

This paper is part of the Universal Collapse Theory library. For a reading guide and full architecture, visit universalcollapse.com/roadmap.

AI Disclosure

AI tools were used to assist with manuscript preparation. The underlying theory, arguments, and interpretive claims are the author's own, and the author takes full responsibility for the content.

Citation

Jones, J. C. (2026). Auditing Constraint Asymmetry in Latency-Based Resolution Tests: A Latency-Curve Diagnostic for Neutrality versus Confound-Induced Delay (UCT Methods Paper v1.0). HoldingLight LLC.