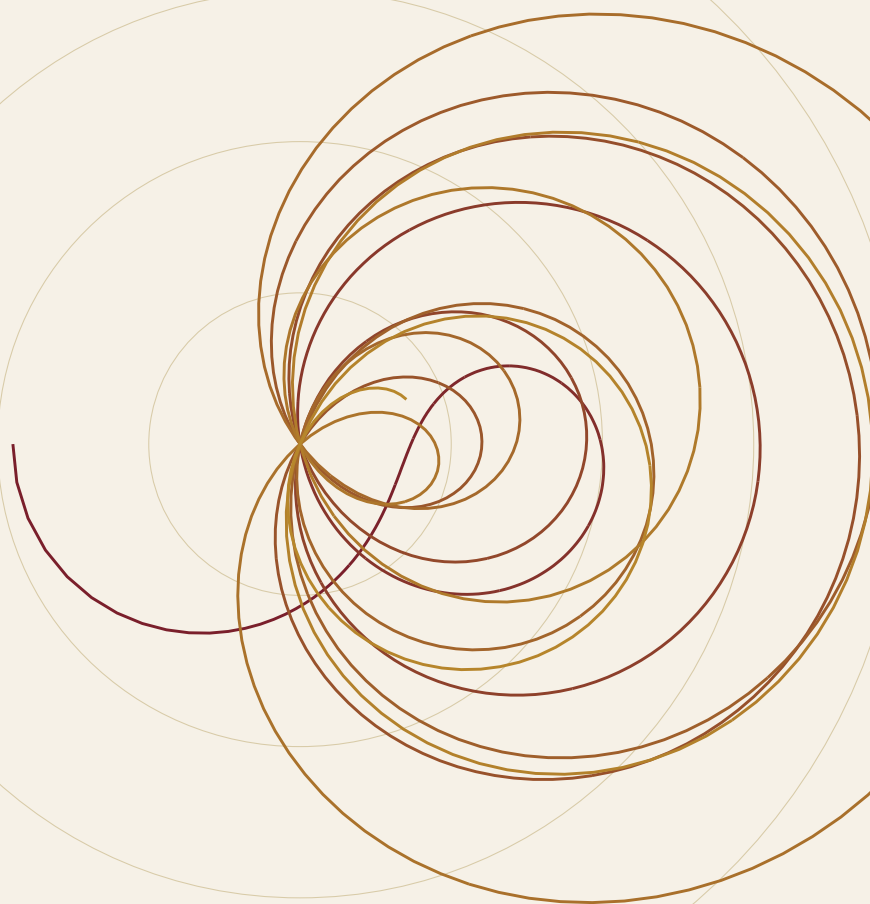


DRCC-Width-Zeta (DWZ-V2.2)

*A Methodological Demonstration of Reconstruction-Width Operationalization
Using the Riemann Zeta Function as a Numerical Case Study*

A Transparent, Auditable, and Limitation-Aware Case Study



The Zeta Spiral: $\zeta(0.5 + it)$ for $t \in [0, 60]$ on the Critical Line

Reza Hesamiy

DWZ-V2.2 • Audit-Strengthened Revision, August 9, 2026

DRCC-Width-Zeta (DWZ-V2.2)

A Methodological Demonstration of Reconstruction-Width Operationalization

Using the Riemann Zeta Function as a Numerical Case Study

Reza Hesamiy

DWZ-V2.2 – Audit-Strengthened Revision, August 9, 2026

Abstract

This manuscript presents a methodological demonstration of how an abstract reconstruction-width framework can be converted into a checkable numerical workflow. The Riemann zeta function in the critical strip is used as a deliberately well-understood case study, not as a setting in which a new or competitive zeta algorithm is claimed. The central research question is therefore not how to compute $\zeta(s)$ most efficiently, but what information, definitions, surrogates, error criteria, validation steps, and limitations are needed to operationalize the DRCC condition

$$0 < d_{\text{rec}}^{\text{DRCC}}(X) \leq d_{\text{frag}}^{\text{DRCC}}(X) < \infty$$

in a concrete numerical setting. Throughout this manuscript, “width” refers to the numerical truncation depth W_{trunc} , unless another manuscript-local quantity ($W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ or the register width ω_{reg}) is explicitly named.

The case study follows a reusable six-stage workflow: (1) declare the abstract quantity and computational task; (2) choose an observable finite fragment; (3) specify an admissible reconstruction model; (4) define a measurable stability and error criterion; (5) validate the result against independent references and standard methods; and (6) record the scope, resource separation, and unresolved limitations. For the zeta instance, the naive partial sum diverges in the critical strip, whereas a first-order Euler–Maclaurin correction yields a convergent, width-dependent approximation. The correction is classical and is not a mathematical contribution of this work. Its role is to provide the minimal transparent reconstruction model through which the general workflow can be examined end to end.

The resulting numerical threshold $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ is estimated empirically, only as an order-of-magnitude guide, not an individually reliable predictor, by $\widehat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon) \approx (C(t)/\varepsilon)^{2/3}$. The fit is calibrated on the first 50 zeros, assessed by bootstrap uncertainty, and checked via a proxy-based out-of-sample comparison on zeros #51–#100 (a different measurement procedure than the calibration fit, see §3.2). The aggregate trend transfers, but individual prediction remains weak ($R_{\text{oos}}^2 \approx 0.08$), consistent with this order-of-magnitude status. Across 100 reference zeros, the absolute position deviation is approximately 1.5×10^{-6} – 1.3×10^{-5} ; 98 cases use known starting values, while two are located through blind scans. The hardening study is reported with point-biserial correlation, logistic regression, confidence intervals, and a power analysis, and no statistically defensible spacing relationship is found in the available $n = 20$ case-level sample.

Benchmarking confirms, qualitatively and under repeated measurements in the available (not fully documented) environment, that the first-order DWZ realization has no algorithmic advantage. A single-evaluation benchmark shows a two-to-three-orders-of-magnitude disadvantage relative to higher-order Euler–Maclaurin and Riemann–Siegel calculations, confirmed by a repeated-evaluation check at one reference zero, while a separate, non-workload-matched cumulative workflow comparison yields an approximately 30-fold slowdown. These negative results are part of the methodological contribution: they separate structural state width, numerical truncation depth, and runtime, and prevent the DRCC vocabulary from being mistaken for a claim of numerical superiority.

The contribution is therefore a transparent and auditable operationalization recipe, together with one numerical instance documented to the extent of the retained record and one exact finite combinatorial operationalization, not a new zeta-function method, not a proof of universal DRCC advantage, and not a contribution to the Riemann Hypothesis.

Keywords: methodological demonstration, reconstruction-width operationalization, DRCC, Riemann zeta function, Euler–Maclaurin summation, numerical validation, auditable workflow.

Contents

Core Equation Register	3
1 Introduction	4
1.1 Starting Point	4
1.2 DRCC Origin and Scope	4
1.3 Contribution of This Work	6
1.4 DRCC Contribution, Equation Traceability, and Counterfactual Role	7
2 Methods	9
2.1 The Basic Problem of the Naive Width Definition	9
2.2 Variant A: Euler-Maclaurin Correction (the DWZ Method)	12
2.3 Variant B: Eta-Based Alternative	12
2.4 Variant C: Second-Order Euler-Maclaurin Correction	12
2.5 Order Degeneracy and Model Dependence	13
2.6 Packaging and Reduction Consistency: What It Does and Does Not Establish . .	23
2.7 Preliminary RP-Tensor Scoring Framework for Reconstruction Variants	24
2.8 Operationalizing $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$	27
2.8.1 Conditional Numerical Operationalization	29
2.8.2 Analytical Origin of the Observed $W^{-3/2}$ Rate	32
2.8.3 Local Calibration of $C(t)$	36
2.8.4 A General Operationalization Workflow	38
2.9 Verification Strategies	41
2.10 Detailed Worked Example: Determination of Zero #1	45
3 Results	48
3.1 Comparison of the Repair Variants	48
3.2 $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$	49
3.3 Blind and Hybrid Zero Search	51
3.4 Hardening and Pattern Hypothesis	53
3.5 Numerical Reproduction and Reference Check for 100 Zeros	55
3.6 Gap Verification	59
4 Discussion	59
4.1 Scope and Limitations	59
4.2 Detailed Discussion	60
4.3 Why DRCC, Given That an Analytic Reconstruction Formula Already Converges	60
5 DRCC Application Assessment of the DWZ Method	64
5.1 Purpose of the Application Assessment	64
5.2 Mapping and Resource Separation, in Brief	65
5.3 Operational Scope Conditions	65
5.4 Net Assessment	66

Core Equation Register

Eq.	Equation	Section
(1)	$\zeta(s) = \sum_{n=1}^{\infty} n^{-s}$ (Dirichlet series)	§1.1
(2)	Euler product $\prod_p (1 - p^{-s})^{-1}$	§1.1
(3)	Riemann Hypothesis (RH)	§1.1
(4)	DRCC stability condition $0 < d_{\text{rec}}^{\text{DRCC}}(X) \leq d_{\text{frag}}^{\text{DRCC}}(X) < \infty$	§1.2
(5)	naive partial sum $\zeta_W(s, W)$	§2.1
(6)	fragment identification $d_{\text{frag}}^{\text{DWZ}}(X; W) := W_{\text{trunc}}$	§2.1
(7)	intermediate form $0 < d_{\text{rec}}^{\text{DRCC}}(X) \leq W < \infty$	§2.1
(8)	$\zeta_{\text{DWZ}}(s, W)$ – Variant A / model M_1	§2.2
(9)	$\eta_W(s, W), \zeta_W^\eta(s, W)$ – Variant B	§2.3
(10)	$\zeta_{\text{DWZ}}^{(2)}(s, W)$ – Variant C / model M_2	§2.4
(11)	$W_{\min}(M; s, \varepsilon)$	§2.5
(12)	$810 < 9000$ (M_2 vs. M_1)	§2.5
(13)	$W_{\min}^*(s, \varepsilon)$	§2.5
(14)	register width $\omega_{\text{reg}}(\tau_{\text{out}}, \pi_{\text{seq}}; \mathfrak{M}_{\text{seq}})$	§2.5
(15)	finite-width non-degeneracy $e_k^{(M_1)}(W) > 0$ (Corollary 1)	§2.5
(16)	RP-Tensor selection score $S_\alpha(\xi)$	§2.7
(17)	RP-Tensor score restricted to the evaluable criterion set J_{avail}	§2.7
(19)	$W_{\text{stab}}^{(M_1)}(t, \varepsilon)$, stable-threshold definition	§2.8
(20)	$\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon)$, empirical estimator (calibrated on zeros #1–50; weak individual out-of-sample accuracy on #51–100)	§2.8
(21)	convergence hypothesis for Proposition 2	§2.8.1
(34)	proven, conditional local rate bound $ \rho_W - \rho \leq (2K_\rho/ \zeta'(\rho))W^{-\sigma_U-1}$ (conditional on Prop. 3's hypotheses, incl. the uniform remainder bound K_ρ ; specializes to $W^{-3/2+\delta}$ on the critical line, where $\sigma_U = \frac{1}{2} - \delta$)	§2.8.2
(36)	heuristic $\delta \downarrow 0$ width law, derived informally from the rigorous fixed- δ bound of Corollary 2 – not itself a rigorous sufficiency statement	§2.8.2
(40)–(48)	general operationalization workflow and Operational DRCC Record $\mathfrak{D}$	§2.8.4
(62)	$N(T)$, Riemann-von Mangoldt counting function	§2.9
(63)	Newton iteration rule $s_{n+1} = s_n - f(s_n)/f'(s_n)$	§2.10
(64)	$ s^* - \rho_1 = 1.66 \times 10^{-5}$ (worked example, zero #1)	§2.10
(66)	$r_{\text{pb}} = -0.189$, point-biserial correlation (hardening test, $n = 20$)	§3.4

This register lists the load-bearing equations of the construction; it is a curated selection, not an exhaustive index of

every numbered equation in the manuscript. Equation numbers not listed are intentionally omitted from this curated register, not missing from the manuscript.

1 Introduction

Framing. Before turning to the zeta function itself, it is worth stating plainly what kind of contribution this paper is and is not. It is not a proposal for a faster or more accurate way to compute $\zeta(s)$ – standard high-efficiency methods such as the Riemann-Siegel formula already address this numerical task in their natural regime, and Section 3 confirms this directly rather than disputing it. It is a case study documented to the extent of the retained record: it takes one specific abstract framework (DRCC, with its stability condition (4)) and works through, in as much checkable detail as the retained record allows, what it takes to instantiate that framework on one concrete, well-understood numerical problem – including every place where the instantiation is partial, where a naive first attempt fails, and where the resulting numbers turn out less favorable than an established alternative. The value of this exercise, if any, lies in the transparency and reproducibility of that process, not in the competitive performance of its output.

1.1 Starting Point

The Riemann zeta function is defined for $\text{Re}(s) > 1$ by

$$\zeta(s) = \sum_{n=1}^{\infty} \frac{1}{n^s}, \quad (1)$$

with the Euler product representation

$$\zeta(s) = \prod_{p \text{ prime}} \frac{1}{1 - p^{-s}}, \quad (2)$$

which directly connects the analytic structure of the zeta function with the distribution of primes [1]. The analytic continuation of $\zeta(s)$ to $\mathbb{C} \setminus \{1\}$ has, in the critical strip $0 < \text{Re}(s) < 1$, the nontrivial zeros. The *Riemann Hypothesis* (RH) is the universal statement that all these zeros lie on the critical line, i.e. that no single-zero exception exists:

$$\text{RH} : \quad \forall \rho \in \mathbb{C} : (\zeta(\rho) = 0 \text{ and } 0 < \text{Re}(\rho) < 1) \implies \text{Re}(\rho) = \frac{1}{2}. \quad (3)$$

This is an open conjecture, not a proven theorem. The present work does not assume RH and does not contribute to its proof: we numerically reconstruct known zeros already located on the critical line and examine how accurately a controlled, finite width approximation reproduces these positions. Classical references on the analytic theory of the zeta function are Titchmarsh [2] and Edwards [3]; the numerical verification of individual zeros goes back to Turing’s method [4].

1.2 DRCC Origin and Scope

DRCC (*Dimensional Reduction via Controlled Combinatorics*, Hesamiy [7]) replaces an infinite or high-dimensional structure X by a controlled, finite *width* W , capturing the resulting reconstruction loss through a stability condition,

$$0 < d_{\text{rec}}^{\text{DRCC}}(X) \leq d_{\text{frag}}^{\text{DRCC}}(X) < \infty, \quad (4)$$

not yet numerically instantiated for the zeta case at the outset of this work (without thereby claiming anything about their status within the cited DRCC theory itself). Written this way, as in

the cited framework, $d_{\text{rec}}^{\text{DRCC}}(X)$ and $d_{\text{frag}}^{\text{DRCC}}(X)$ appear to depend on X alone; this manuscript’s own concretization below depends additionally on the declared model, t , ε , and W , and is therefore given its own, argument-explicit notation $d_{\text{frag}}^{\text{DWZ}}(X; W)$ and $d_{\text{rec}}^{\text{DWZ}}(X; t, \varepsilon)$ once it becomes formal (Proposition 2), rather than reusing the undecorated symbols above as if they denoted the same, single-argument quantity. The goal of this work is a conditional numerical instantiation of the form of inequality (4) for the concrete case of the Riemann zeta function – we call the resulting method the *DWZ method* (DRCC-Width-Zeta method; methodically the *Euler-Maclaurin width reconstruction method*, Section 2).

Terminological note: throughout this work, “width” W consistently refers to the numerical truncation depth W_{trunc} , a manuscript-internal quantity distinct from the register width ω_{reg} ; this question is tested empirically in Section 2.5 and made precise formally in Lemma 1 and Corollary 1 (Section 2.5) – not by Proposition 2, which instead concerns the conditional operationalization of W_{stab} (Section 2.8.1). Table 1 makes this commitment verifiable rather than merely asserted.

Table 1: Traceability: DRCC/DWZ terms and their role in this work.

DRCC term	Role in this work	Status
$d_{\text{rec}}^{\text{DRCC}}(X)$, $d_{\text{frag}}^{\text{DRCC}}(X)$ (stability condition)	Motivates the width idea; inequality (4) as the starting point	Conditionally instantiated by numerical surrogates in Proposition 2; no identity with the abstract quantities is established
Register width ω_{reg} (this manuscript’s own definition, Eq. (14))	Structurally proven in Section 2.5 (Lemma 1)	Proven $O(1)$ for the sequential zeta accumulation; identification with W_{trunc} is <i>not</i> shown
Numerical accuracy parameter ε (this manuscript’s own, e.g. in $W_{\text{min}}(M; s, \varepsilon)$, Eq. (11))	Used in Section 2.5, Test B (model comparison)	Manuscript-specific parameter
Width W	Operationalized as the numerical truncation depth W_{trunc} of the partial sum	Independent quantity, newly defined in this work

In short: the title and framework name the theoretical origin. What this work shows is a *numerical case study and formal partial operationalization* of the DRCC starting idea – not a proof that the abstract stability condition and the numerical truncation width are the same quantity. Section 2.5 does not merely assert this point but formally separates the two quantities within the declared reconstruction model (Lemma 1 and the subsequent instantiations): for the natural reconstruction tasks of the sequential zeta accumulation, the register width is provably constant ($\omega_{\text{reg}} = O(1)$), while the required truncation depth, as quantified by $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$, grows without bound as the target accuracy is tightened (under the hypotheses of Corollary 1) – not W_{trunc} itself, which remains a free truncation variable rather than a function of ε . DRCC thus supplies the motivation and conceptual frame for the stability condition under study, not a proven structural identity with the numerical method; this work accordingly treats DRCC consistently as context, not as the mathematical foundation of the numerical results.

On methodological novelty. The Euler-Maclaurin summation formula itself is a long-established standard technique of analysis [2, 3] and is not claimed here as a new mathematical invention. The contribution of this work lies in its application – the explicit, minimal (order-1)

operationalization of the abstract DRCC stability condition for the zeta function in the critical strip – and in the empirical characterization of the resulting procedure: convergence order, model-dependent truncation depth, reproduction across 100 zeros (with the empirical fit for $C(t)$ remaining based on the first 50 zeros), and a systematic comparison against a higher-order Euler-Maclaurin expansion and the Riemann-Siegel formula (Section 2.5, Test C), which shows that the minimal order-1 variant carries an efficiency disadvantage of two to three orders of magnitude relative to established methods. The present method is thus not intended as a faster alternative to Riemann-Siegel, but as the most directly traceable, minimal route for translating the DRCC condition into a checkable number on a concrete, nontrivial example.

1.3 Contribution of This Work

The contribution of this manuscript is a reusable operationalization workflow and one numerical instance of that workflow, documented to the extent of the retained record. It does not claim a new summation formula, a faster zeta computation, or a result concerning the Riemann Hypothesis.

Why the zeta function is a critical stress test. The Riemann zeta function is not chosen because the present work seeks a new or competitive algorithm for evaluating $\zeta(s)$. It is chosen because the critical strip provides a particularly discriminating test of the DRCC operationalization workflow.

The defining finite fragment,

$$F_W(s) = \sum_{n=1}^W n^{-s},$$

does not converge to $\zeta(s)$ in the critical strip as $W \rightarrow \infty$ (Section 2). Thus, merely increasing the numerical fragment size does not solve the reconstruction problem – unlike a problem where the naive fragment already converges, for which the reconstruction model would be vacuous. A successful operationalization must introduce a genuinely nontrivial reconstruction operator.

At the same time, the required reconstruction is independently known: the Euler-Maclaurin formula provides a classical analytic correction whose convergence can be checked against established reference values. This combination makes the zeta case unusually useful as a stress test. The naive fragment fails, the reconstructed quantity converges, and the correctness of the reconstruction mechanism can be evaluated without requiring acceptance of DRCC itself – it is the workflow’s own claims, not the correctness of the underlying mathematics, that remain the object under test.

The resulting test therefore separates three logically distinct questions:

fragment adequacy, reconstruction adequacy, resource width.

The first is answered negatively by the divergent partial sum. The second is answered positively, within the declared numerical model, by the Euler-Maclaurin reconstruction. The third yields a different and important negative result: as Lemma 1 shows, for the natural sequential accumulation model considered in Section 2.5, the structural register width satisfies $\omega_{\text{reg}} = O(1)$ – whereas the required stable truncation threshold $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ grows as the required tolerance is tightened (under the hypotheses of Corollary 1); W_{trunc} itself remains a free truncation variable, not a function of ε . Consequently,

$$\omega_{\text{reg}} \neq W_{\text{trunc}}.$$

This separation is not a defect to be hidden. It is a property of this specific reconstruction task, not a general feature asserted of the DRCC framework (Remark 1), and it identifies a limitation of the present case study that directly informs the design of future DRCC tests: the next informative test of the workflow is a problem family for which the structural width scales with problem size rather than remaining constant, as noted in Section 6.

The zeta case is therefore valuable precisely because it produces both a successful reconstruction result and an unfavorable structural-width result: it tests whether the operationalization framework can distinguish success from failure instead of merely reproducing favorable examples.

The contribution has four parts.

1. **Operational translation.** The abstract DRCC condition is translated into a declared fragment, reconstruction model, error functional, stable accuracy-relative depth, and operational record.
2. **Resource separation.** Numerical truncation depth, accuracy-relative reconstruction depth, structural active-state width, and runtime are explicitly separated rather than treated as one notion of “width.”
3. **Empirical and statistical assessment.** The manuscript reports convergence tests, comparisons among reconstruction variants, repeated timing measurements, bootstrap uncertainty, out-of-sample prediction, a worked zero-reconstruction example, a 100-zero reference check, and a power analysis for the hardening hypothesis – with explicitly stated limitations (e.g. weak out-of-sample prediction, $R_{\text{os}}^2 \approx 0.08$, and a null result for the hardening spacing relationship) rather than a claim of general validation.
4. **Transferable recipe with explicit limits.** A general workflow is extracted from the zeta instance in Section 2.8.4, and carried out a second time, in exact finite form, for graph three-coloring on C_4 (Appendix B.2). Only the six-stage workflow itself is claimed to be transferable across these two instances; no third application and no universal DRCC advantage are asserted.

The scientific result is therefore methodological: the manuscript shows what must be supplied, measured, and qualified before an abstract reconstruction principle becomes a reproducible numerical claim. The zeta calculation is the demonstration vehicle, not the claimed novelty.

1.4 DRCC Contribution, Equation Traceability, and Counterfactual Role

The Euler-Maclaurin correction underlying the DWZ method is classical (Section 1.2); it would exist, and would converge, with or without DRCC. This raises a fair question, one a reviewer is entitled to ask directly: what, concretely, did DRCC contribute to this manuscript? This section answers that question for each traceable piece of the construction, rather than asserting the answer in general terms. For every DRCC-motivated equation used in this work, we state: (1) which abstract DRCC idea it concretizes, (2) how it was concretized in the DWZ construction, (3) where its effect becomes mathematically or numerically visible, and (4) what would be missing from this manuscript without that step – understood historically, as a statement about how this particular manuscript was actually developed, not as a claim that no other route to the same numerical results could exist.

Three different kinds of success. Collapsing these into a single “DRCC worked” statement would blur three distinct claims that deserve to be kept separate.

Methodological success. DRCC did not supply the Euler-Maclaurin formula. What it supplied was the demand for a checkable chain from the abstract condition (4) to a reproducible numerical protocol: $d_{\text{frag}}^{\text{DRCC}}(X) \rightsquigarrow W$, then $d_{\text{rec}}^{\text{DRCC}}(X) \rightsquigarrow W_{\text{stab}}^{(M_1)}(t, \varepsilon)$, then the conditional conclusion $W \geq W_{\text{stab}}^{(M_1)}(t, \varepsilon) \Rightarrow e_k^{(M_1)}(W) < \varepsilon$ (Proposition 2). Turning an abstract reconstruction requirement into a testable numerical statement is the methodological contribution.

Negative, but important, success. The DRCC framework raised the specific question of whether the numerical truncation width W_{trunc} coincides with the register width ω_{reg} . The answer, proven

Table 2: DRCC equation register: traceability from the abstract DRCC condition to its concretizations in this manuscript.

DRCC idea	DWZ concretization	Where the effect is visible	Without this step
Abstract stability condition $0 < d_{\text{rec}}^{\text{DRCC}}(X) \leq d_{\text{frag}}^{\text{DRCC}}(X) < \infty$	Eq. (4), §1.2; Table 1	Frames the entire manuscript structure around a width/accuracy pair, not just a convergence question	No formal reason to track a fragmentation–reconstruction pair at all; convergence alone would suffice
Fragment identification $d_{\text{frag}}^{\text{DWZ}}(X; W) := W_{\text{trunc}}$ (numerical surrogate, not an intrinsic property of X)	Eq. (6), §2.1	Naive fragment tested and shown to diverge (Table 3), motivating the correction of §2.2	Naive divergence would read as an isolated numerical fact, not as the failure of a specific, declared fragment
Stable reconstruction depth $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$	Eq. (19), §2.8	Used throughout Section 3 as the operational accuracy threshold; drives the empirical estimator below	No formal accuracy-relative depth concept – only an ad hoc “large enough W ”
Empirical estimator $\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon) \approx (C(t)/\varepsilon)^{2/3}$	Eq. (20), §2.8	Closed-form order-of-magnitude guide (not an individually reliable predictor, see §3.2), tested out-of-sample on zeros #51–100 (Section 3)	No closed-form width-prediction formula; every t would require a fresh numerical search
Conditional formal instantiation	Proposition 2, §2.8.1	Turns the DRCC→DWZ transition into a stated, checkable conditional proof, not an assertion	The DRCC→DWZ link would remain prose and motivation only, as the first reviewer round correctly criticized
Register-width lemma $\omega_{\text{reg}} = O(1)$ (proven), $\neq W_{\text{trunc}}$	Lemma 1, §2.5	Separates runtime length from active state width; load-bearing for Section 4 and the Conclusion	The naive identification $W_{\text{trunc}} = \omega_{\text{reg}}$ would stand unchallenged and unexamined
General operationalization template and Operational DRCC Record $\mathfrak{D}$	Eqs. (40)–(48), §2.8.4	States, honestly and narrowly, what a second instantiation of this template would require	DWZ would read as a one-off numerical trick with no stated, reusable structure for future instances

in Lemma 1 rather than merely observed, is that it does not: $\omega_{\text{reg}} = O(1)$ is proved within the declared $\mathfrak{M}_{\text{seq}}$ register model, while the required truncation depth $W_{\text{stab}}^{(M_1)}(t, \varepsilon) \rightarrow \infty$ as $\varepsilon \rightarrow 0$ (Corollary 1) – W_{trunc} itself remains a free variable, not a function of ε . Ruling out a naive identification is not a failure of the framework – it is a genuine separation result,

$$\text{state width} \neq \text{truncation depth} \neq \text{runtime},$$

that this manuscript would very plausibly not have investigated without the DRCC vocabulary prompting the question in the first place.

Numerical success of the operationalization. The concretization chosen here produces checkable numbers: the naive fragment diverges; M_1 converges; M_2 reduces the smallest tested grid width by roughly a factor of 11 in the five reported comparisons (Eq. (12), 9000 vs. 810, W_{test} values, not proven minima); the observed exponent near $3/2$ is analytically motivated by, and locally bounded through, the first omitted Euler-Maclaurin term (Section 2.8.2); 100 known zeros are numerically reproduced with absolute deviation 1.5×10^{-6} – 1.3×10^{-5} , two of them found blindly; and the same comparison honestly shows Riemann-Siegel to be substantially faster. The claim this supports is precise:

DRCC’s role here is operationalization and resource separation,

not

DRCC produced a faster zeta method.

The second statement is not made anywhere in this manuscript; the first one is what Table 2 traces in detail.

What would remain without this DRCC step. Stripped of the DRCC framing, this manuscript’s numerical core would still exist as roughly: *a numerical study of a first-order Euler-Maclaurin approximation for reproducing known zeros of the Riemann zeta function*. That hypothetical paper would still contain equation (8), convergence tests, the zero search, and the Riemann-Siegel comparison. Based on how this manuscript actually developed, it would not contain: the starting condition (4); the $d_{\text{frag}}/d_{\text{rec}}^{\text{DRCC}}$ distinction; the stable reconstruction threshold (19); the empirical stable-width estimator (20); Proposition 2 as a formal operationalization; the separation of W_{trunc} from ω_{reg} ; the Operational DRCC Record (48); the general template of Section 2.8.4; or the RP-Tensor question of systematic model selection (Section 2.7). This is a statement about the development history and content of this specific manuscript, not a claim that DRCC is the only conceivable route to any of these individual ideas.

Closing statement. The Euler-Maclaurin formula is classical. DRCC’s contribution in this manuscript is the explicit operationalization of an abstract reconstruction condition, the definition of a measurable reconstruction depth, the separation of model choice from structural resource, and an explicit operational record – traced end to end in Table 2 rather than asserted in the abstract. The intended reading is

DRCC question $\rightarrow$ DWZ concretization $\rightarrow$ measured result $\rightarrow$ counterfactual without DRCC,
a complete accounting of DRCC’s role in this manuscript, not a stronger claim standing in place of one.

2 Methods

2.1 The Basic Problem of the Naive Width Definition

Throughout this manuscript, $\mathbb{N} := \{1, 2, 3, \dots\}$ (the positive integers, excluding 0), the convention used wherever a width, index, or threshold is declared to lie in $\mathbb{N}$ below.

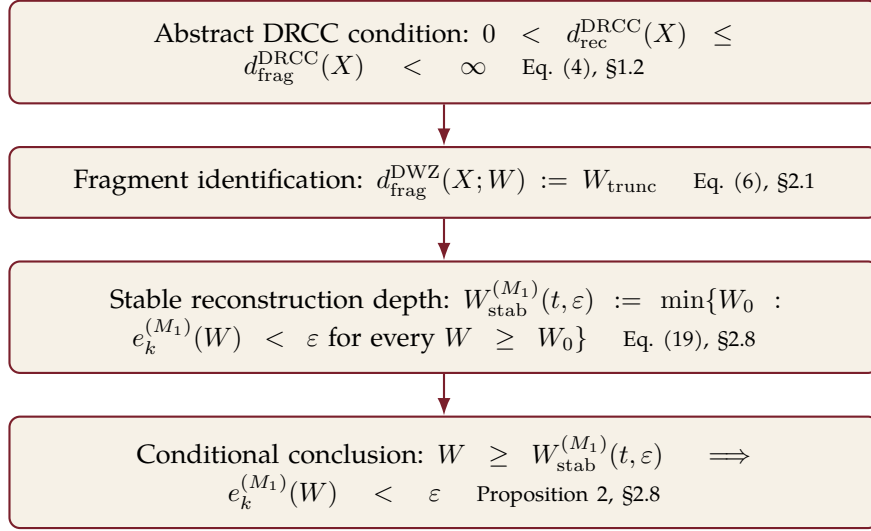
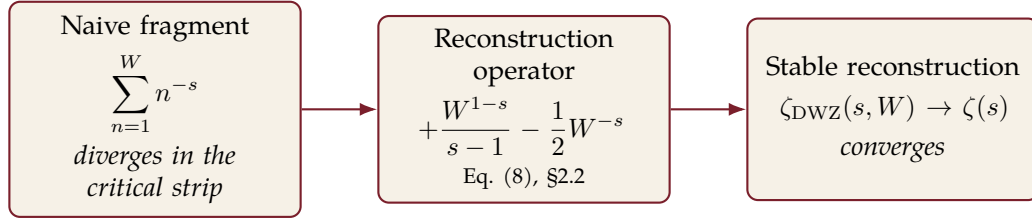
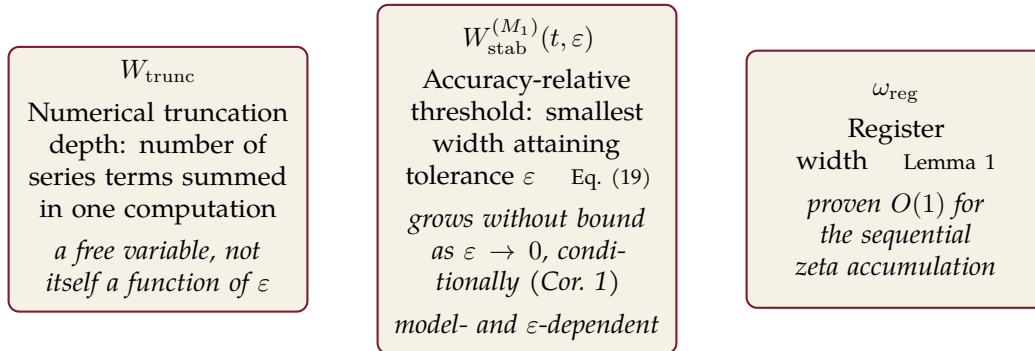


Figure 1: OPA-DRCC-01: From abstract condition to measurable width – the core traceability chain in one diagram.



W	naive $ \zeta_W(s, W) - \zeta(s) $	DWZ $ \zeta_{DWZ}(s, W) - \zeta(s) $
10	2.47×10^{-1}	3.85×10^{-2}
100	7.09×10^{-1}	1.18×10^{-3}
1000	2.24×10^0	3.73×10^{-5}
5000	5.00×10^0	3.33×10^{-6}

Figure 2: OPA-DRCC-02: Failure, repair, success, with real values at $s = 0.5 + 14.13i$ taken directly from Table 3 (naive column) and Table 13 (DWZ column) – the naive fragment diverges while the reconstructed function converges, on the same points.



$W_{\text{trunc}} \not\equiv W_{\text{stab}}^{(M_1)}(t, \varepsilon)$, $W_{\text{stab}}^{(M_1)}(t, \varepsilon) \not\equiv \omega_{\text{reg}}$, $W_{\text{trunc}} \not\equiv \omega_{\text{reg}}$ (pairwise non-identical resource notions – but not all pairs equally unrelated: $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ is itself a specific, achievable threshold value of the free variable W_{trunc} (equation (19)), whereas ω_{reg} is a register width of a different kind entirely; see Lemma 1 and Proposition 2)

Figure 3: OPA-DRCC-03: Three different widths that this manuscript is careful never to conflate.

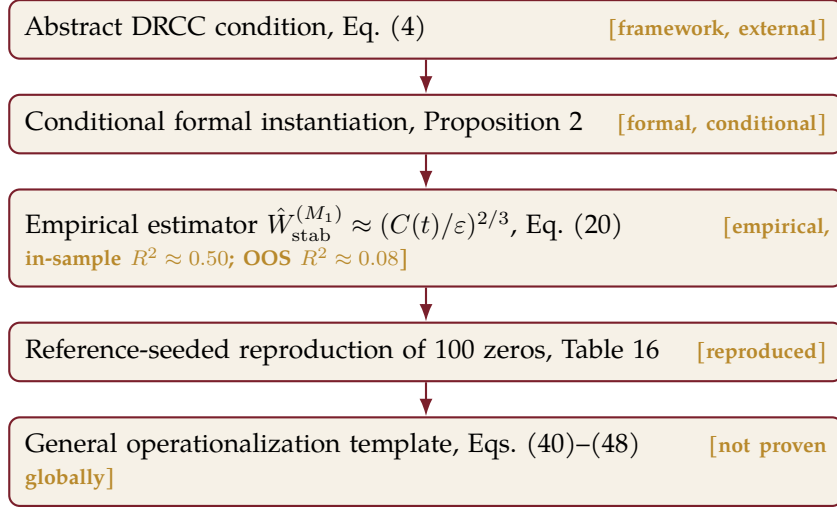


Figure 4: OPA-DRCC-04: Evidence ladder – each rung is labeled with its actual epistemic status; the ladder does not get stronger by omitting the labels on its weaker rungs.

An obvious first approach to the width-zeta function is the plain partial sum

$$\zeta_W(s, W) := \sum_{n=1}^W \frac{1}{n^s}, \quad W \in \mathbb{N}, \quad 0 < W < \infty, \quad (5)$$

To avoid notational inconsistency, the fragmentation quantity of equation (4) is introduced argument-explicitly from the outset, as its own numbered definition, distinct from the partial-sum formula of equation (5) that it references. Since X is fixed while W acts as a free variable, the fragmentation degree is defined as a function of both parameters,

$$d_{\text{frag}}^{\text{DWZ}}(X; W) := W_{\text{trunc}} \quad (\text{the } W \text{ of equation (5)}), \quad (6)$$

as the first concretization of equation (4) (the justification for this concretization is discussed in Section 2.5 and in the Conclusion).

Notational note. An earlier version of this manuscript wrote this, for simplicity, as $W := d_{\text{frag}}(X)$. This implicit notation is fully discarded here: since W is free while X is fixed, this quantity cannot literally be a function of X alone. The explicit form $d_{\text{frag}}^{\text{DWZ}}(X; W)$ used throughout this manuscript from this point on rigorously ensures that the functional dependence on the free variable W is represented correctly in all subsequent derivations, including at the one point (Proposition 2) where this is formally load-bearing. Here W is itself a free parameter that may be evaluated at any value; the *required* truncation depth needed to reach a declared accuracy (formalized below as W_{min} and, later, W_{stab}) depends on the chosen model, the accuracy ε , the point s , the search strategy, and the computational precision, and is not an intrinsic property of X itself. This series converges as $W \rightarrow \infty$ only for $\text{Re}(s) > 1$. In the critical strip, $\zeta_W(s, W)$ diverges as W grows (Table 3) instead of approaching $\zeta(s)$ – equation (5) is therefore unsuitable as a reconstruction approach in the critical strip.

At this point the inequality

$$0 < d_{\text{rec}}^{\text{DRCC}}(X) \leq W < \infty \quad (7)$$

remained incomplete: $d_{\text{rec}}^{\text{DRCC}}(X)$ was still undetermined.

Table 3: Divergence of the naive partial sum at $s = 0.5 + 14.13i$.

W	$ \zeta_W(s, W) - \zeta(s) $
10	2.47×10^{-1}
50	5.04×10^{-1}
100	7.09×10^{-1}
500	1.58×10^0
1000	2.24×10^0
5000	5.00×10^0

2.2 Variant A: Euler-Maclaurin Correction (the DWZ Method)

The Euler-Maclaurin summation formula – a classical, eighteenth-century result of analysis and not a contribution of this work (see the explicit disclaimer in Section 1.2, “On methodological novelty”) [2, 3] – provides a correction that makes ζ_W convergent even in the critical strip:

$$\zeta_{\text{DWZ}}(s, W) = \sum_{n=1}^W \frac{1}{n^s} + \frac{W^{1-s}}{s-1} - \frac{W^{-s}}{2} \quad (8)$$

with $\lim_{W \rightarrow \infty} \zeta_{\text{DWZ}}(s, W) = \zeta(s)$ also holding for $0 < \text{Re}(s) < 1$. The asymptotic error induced by the first omitted Euler-Maclaurin correction term is derived in Section 2.8.2.

2.3 Variant B: Eta-Based Alternative

As a comparison method, Variant B uses the Dirichlet eta function, which already converges for $\text{Re}(s) > 0$:

$$\eta_W(s, W) = \sum_{n=1}^W \frac{(-1)^{n-1}}{n^s}, \quad \zeta_W^\eta(s, W) = \frac{\eta_W(s, W)}{1 - 2^{1-s}}. \quad (9)$$

2.4 Variant C: Second-Order Euler-Maclaurin Correction

Variant A (8) uses only the first correction term of the Euler-Maclaurin expansion. The next term of the series (involving the Bernoulli number $B_2 = 1/6$) yields a consistent extension:

$$\zeta_{\text{DWZ}}^{(2)}(s, W) = \sum_{n=1}^W \frac{1}{n^s} + \frac{W^{1-s}}{s-1} - \frac{W^{-s}}{2} + \frac{s}{12} W^{-s-1} \quad (10)$$

We henceforth refer to equation (8) as model M_1 (first-order) and equation (10) as model M_2 (second-order), counting retained correction terms. This “order” label should not be confused with the Bernoulli-order index m used in the general Euler-Maclaurin parametrization of Section 2.5, where M_1 corresponds to $m = 0$ and M_2 to $m = 1$: the model labels count correction terms starting at one, the index m counts Bernoulli terms beyond the endpoint correction starting at zero, and the two conventions are offset by one by construction, not by inconsistency.

Methodological design principle. The purpose of this work is not to identify the numerically most efficient reconstruction model, but to demonstrate that the DRCC stability condition (Section 2) can already be operationalized by the smallest non-trivial Euler-Maclaurin reconstruction. To be explicit about scope: the Euler-Maclaurin summation formula itself is classical eighteenth-century analysis, not a contribution of this work (Section 1.2); what is discussed below concerns solely which correction order this manuscript adopts as its baseline, not the formula’s origin,

validity, or generality. “Smallest non-trivial” is not a stylistic preference here: the zero-order case, the naive partial sum of equation (5), diverges throughout the critical strip (Table 3), so M_1 is the first member of the Euler-Maclaurin correction family at which convergence – and hence a non-trivial reconstruction in the DRCC sense – is achieved at all. For this reason, M_1 is adopted as the primary reconstruction model throughout the manuscript. M_2 , and in principle any higher Euler-Maclaurin order, are included as comparative models within this same family, to quantify the improvement in width and accuracy obtainable through additional correction terms (Section 2.5), rather than to replace the methodological baseline established by M_1 . This scope is deliberately restricted to the Euler-Maclaurin family: M_η (Variant B, Section 2.3) is a structurally different construction and is retained in the comparison for a separate reason, not as a further point on this order hierarchy. Consequently, the finding of Section 2.7 that M_2 dominates M_1 on the two RP-Tensor criteria for which data exist does not contradict the choice of M_1 as baseline: that finding quantifies precisely the improvement this design principle anticipates from moving to a higher order, and the choice of M_1 is stated here as a declared methodological decision that sits outside, and precedes, the RP-Tensor ranking – not as its outcome.

2.5 Order Degeneracy and Model Dependence

Two targeted tests were carried out to clarify whether the summation order ω or the reconstruction model M is the actual lever controlling the DWZ width.

Test A – Order invariance. The same 500 terms n^{-s} ($s = 0.5 + 1001.3495i$) were summed in ascending, descending, and mixed order. At 30-digit floating-point precision, the observed ordering effects remained at rounding level, 5.6×10^{-31} , and did not change any tested width threshold – expected, since a finite sum is permutation-invariant in exact arithmetic. For the declared scalar accumulation model tested here, arithmetic term permutation does not change the exact finite sum, so minimization over orderings ω is degenerate for this model: the summation order is *not* the effective optimization variable. This test is confined to the specific scalar accumulation model used in this manuscript; it does not establish ordering degeneracy for reconstruction models in general, where different reconstruction orderings can in principle induce different persistent states or transition costs.

Test B – Model comparison. For a fixed accuracy target $\varepsilon = 10^{-4}$, the minimally required width

$$W_{\min}(M; s, \varepsilon) := \min\{W \in \mathbb{N} : |\zeta_M(s, W) - \zeta(s)| < \varepsilon\} \quad (11)$$

was defined (with $W_{\min}(M; s, \varepsilon) := \infty$ if this set is empty). It was determined for M_1 and M_2 at five zeros distributed across the tested range using a grid search, with grid step $h_{M_1} = 200$ and $h_{M_2} = 10$ (see scripts) (Table 4). The table values are thus the smallest width found on the grid used, W_{test} , not necessarily the exact minimum over every $W \in \mathbb{N}$.

Table 4: Smallest width found on the grid used for $\varepsilon = 10^{-4}$, model M_1 (Variant A) vs. M_2 (Variant C). W_{test} , not $W_{\min}$: grid values, not proof of the exact minimum.

Zero #	t	$W_{\text{test}}(M_1)$	$W_{\text{test}}(M_2)$	Factor
650	1001.3495	9000	810	11.11
675	1031.8332	9200	830	11.08
700	1062.9154	9400	850	11.06
725	1093.4409	9600	870	11.03
749	1122.4586	9600	890	10.79

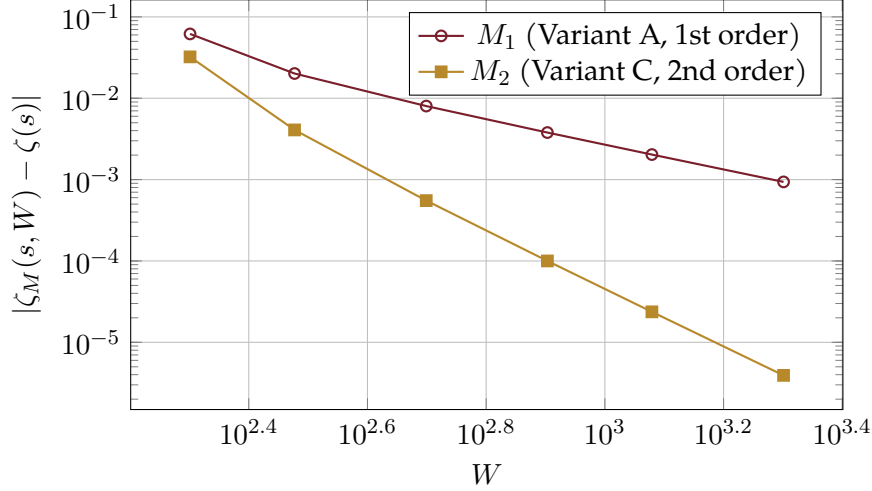


Figure 5: Convergence of M_1 and M_2 towards $\zeta(s)$ at zero #650 (real computed data).

In the five tested cases, the reduction factor ranged between 10.79 and 11.11 (Table 4); no general asymptotic constant is derived from this. What is reproducibly demonstrated for these five points is only the comparison of the values found on the grid:

$$W_{\text{test}}(M_2; s, \varepsilon) < W_{\text{test}}(M_1; s, \varepsilon), \quad \text{concretely } 810 < 9000 \quad (12)$$

– not a proven statement about the exact minima $W_{\min}$ themselves, which could turn out slightly smaller on a finer grid. The model optimization

$$W_{\min}^*(s, \varepsilon) := \min_{M \in \mathfrak{M}_{\text{adm}}} W_{\min}(M; s, \varepsilon), \quad \mathfrak{M}_{\text{adm}} := \{M_1, M_2\} \quad (13)$$

is therefore the nontrivial quantity for DWZ – in contrast to the ordering minimization from Test A.

Test C – comparison with standard methods (B_{eval}). Tests A and B compare only variants *within* the DWZ framework. To provide the comparison against established standard methods that the reviewer critique rightly requested, M_1 (Variant A, DWZ, order 1) is additionally compared against (a) a fifth-order Euler-Maclaurin expansion (M_{K5} , generally $\zeta_{\text{EM}}^{(m)}(s, W) = \sum_{n=1}^W n^{-s} + \frac{W^{1-s}}{s-1} - \frac{W^{-s}}{2} + \sum_{j=1}^m \frac{B_{2j}}{(2j)!} (s)_{2j-1} W^{-s-2j+1}$ with $(s)_k = s(s+1) \cdots (s+k-1)$ the rising factorial and B_{2j} the Bernoulli numbers. In this general- m formula, Variant A (M_1) uses $m = 0$ (only the endpoint correction $\frac{W^{1-s}}{s-1} - \frac{W^{-s}}{2}$, no Bernoulli term), Variant C (M_2) uses $m = 1$ (the endpoint correction plus one Bernoulli term), and M_{K5} uses $m = 5$ (the endpoint correction plus five Bernoulli terms) – the remainder term is not separately bounded) and (b) the Riemann-Siegel formula, evaluated via the mpmath functions $\vartheta(t)$ (siegeltheta) and $Z(t)$ (siegelz), with $\zeta(\frac{1}{2} + it) = e^{-i\vartheta(t)} Z(t)$. At a fixed accuracy target $\varepsilon = 10^{-6}$: for M_1, M_{K5} we determine the tested truncation depth W required, directly measured; for Riemann-Siegel we instead report the standard asymptotic main-sum term estimate $\sim \sqrt{t/2\pi}$, not a measured count of internal operations for the specific mpmath siegelz implementation used – and, in both cases, the computation time of a *single evaluation* at that minimal term count (30-digit mpmath arithmetic; see the code and data availability note at the end of this manuscript regarding the status of the underlying scripts). We denote this single-evaluation benchmark type B_{eval} throughout, to distinguish it explicitly from the cumulative, non-workload-matched benchmark B_{workflow} reported later (Section 3, “Runtime against an established reference method”) – the two measure different things and are not interchangeable, a distinction stated here in one place

because the two figures otherwise appear close together in the text and could be conflated by a reader skimming for a single “the” benchmark factor. Regarding the absolute time values reported in Table 5 and Figure 6: these individual entries remain single measurements (no median over repeated runs, no warm-up, no confidence intervals) on the hardware/software environment available in this sandbox (processor, Python, and `mpmath` version details are not separately specified in the manuscript text; see the code and data availability note at the end of this manuscript). Single measurements in the millisecond range can be affected by system noise; taken in isolation, the reported order-of-magnitude factors (roughly two to three orders of magnitude) would therefore be only a preliminary, exploratory tendency, not an exact, repeatedly validated benchmark constant – the word “robust” is deliberately avoided for these single-shot table entries, since it would imply a statistical stability that a single, unreplicated measurement cannot support on its own. The paragraph immediately below reports an independent repeated-measurement check that supersedes this limitation for the qualitative conclusion, while the single-shot entries in the table are left unchanged for transparency and reproducibility of the original run.

Repeated-measurement check. As requested in review, we repeat the single-shot measurement above to check whether its qualitative conclusion survives replication. At zero #1 ($t = 14.1347$), M_1 ($W=11,498$) was timed over 15 repeated evaluations, M_{K5} ($W=10$) and Riemann-Siegel (via `mpmath`’s `siegeltheta/siegelz`) over 30 repeated evaluations each, on the sandbox environment used for this revision: $M_1 = 146.9 \pm 2.2$ ms ($n = 15$), $M_{K5} = 0.343 \pm 0.009$ ms ($n = 30$), Riemann-Siegel = 0.53 ± 0.21 ms ($n = 30$) (mean $\pm$ standard deviation). The standard deviations are small relative to the means for M_1 and M_{K5} (coefficient of variation below 2% for both), but noticeably larger for Riemann-Siegel (coefficient of variation below 40%; at these sub-millisecond durations, this is likely attributable to scheduler and system noise and general short-duration timing variability – a plausible but unmeasured hypothesis, not independently confirmed here – rather than to the timer’s own resolution, which is well below this scale on typical systems). Because of this asymmetry, the two resulting ratios should not be read with the same precision: $M_1/M_{K5} \approx 430\times$, backed by low variability on both sides, is a comparatively precise figure, whereas $M_1/\text{Riemann-Siegel} \approx 280\times$ carries the wider uncertainty of the Riemann-Siegel measurement and should be read as an order-of-magnitude estimate (a few hundredfold) rather than a precise benchmark factor. Both ratios are nonetheless consistent in order of magnitude with the single-shot figures in Table 5 (measured in a different execution environment/run, with hardware not separately documented for either run – see Appendix B.0 – hence the absolute values differ slightly but not the conclusion). The qualitative two-to-three-orders-of-magnitude disadvantage of M_1 is thus confirmed under repetition, not merely asserted from a single run; this does not retroactively justify calling the original single-shot table “robust” – that wording remains avoided – but it does show that repeating the measurement would not have overturned the conclusion.

Benchmark environment. The exact processor model, operating-system version, Python version, `mpmath` version, and process-isolation conditions were not recorded together with the historical timing run reported above.

The recorded benchmark parameters were a working precision of 30 decimal digits and the use of `time.perf_counter` as timer.

Consequently, the reported absolute wall-clock measurements should be interpreted as environment-specific historical observations rather than values filled in after the fact. The reproducible conclusion is restricted to the observed order-of-magnitude ranking under the stated numerical protocol, not hardware-independent wall-clock constants.

Open Arb benchmark. A matched comparison against Arb or an equivalent ball-arithmetic implementation has not yet been executed. No Arb runtime, speedup, slowdown, or ratio is inferred from the existing measurements. The benchmark remains open until all methods are measured under the same precision, target accuracy, hardware, process isolation, warm-up, repetition, and stopping rules; the planned protocol and the still-unfilled comparison table are deferred to Appendix B.3 rather than shown here, since no measurements exist yet to report.

Table 5: Systematic comparison: required terms and computation time for $\varepsilon = 10^{-6}$, against M_1 (DWZ), M_{K5} (fifth-order Euler-Maclaurin), and Riemann-Siegel.

Zero #	t	$W(M_1)$	Time(M_1)	$W(M_{K5})$	Time(M_{K5})	Terms/Time (R.-S.)
1	14.1347	11,498	0.151 s	10	0.0004 s	1 / 0.0016 s
25	88.8091	42,695	0.635 s	43	0.0008 s	3 / 0.0009 s
50	143.1118	55,504	0.742 s	73	0.0011 s	4 / 0.0058 s
650	1001.3495	206,087	2.775 s	466	0.0060 s	12 / 0.0156 s

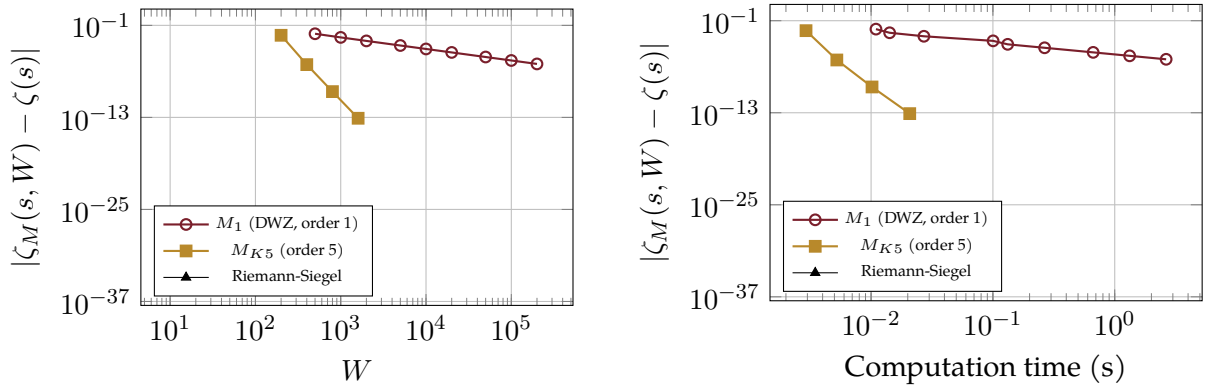


Figure 6: At zero #650 ($t = 1001.3495$): left, error against number of terms W ; right, error against computation time; for M_1 , M_{K5} , and Riemann-Siegel (shown as a single benchmark point because the DWZ truncation parameter W is not Riemann-Siegel’s control parameter in this comparison).

The result is unambiguous and is reported here without embellishment: the Riemann-Siegel formula reaches the same accuracy with an asymptotic main-sum term estimate of 1–12 terms (not a directly measured internal operation count, as noted above) and consistently under 16 ms, while M_1 (DWZ) requires, depending on t , between 11,500 and over 200,000 terms and 0.15–2.8 s – an efficiency disadvantage of roughly two to three orders of magnitude. Simply extending to five Euler-Maclaurin correction terms (M_{K5}) already reduces this disadvantage drastically (factor ≈ 442 –1150 in W , factor ≈ 378 –794 in time relative to M_1 , computed directly from the four rows of Table 5), but is itself no substitute for Riemann-Siegel and is numerically more delicate: the Euler-Maclaurin series is asymptotic, not convergent, and the correction terms can blow up in magnitude if W is too small relative to $|s|$, before the series converges cleanly for sufficiently large W (observed numerically; not shown in Table 5, whose W values already lie in the stable regime). M_1 reduces this particular high-order asymptotic-instability risk through its simpler correction form – it retains only the first correction term, so it cannot exhibit the specific blow-up behavior possible when a higher-order asymptotic series is truncated at too small a W – although it does not eliminate all finite-width numerical or truncation effects in general, at the cost of requiring substantially more terms. The scientific contribution of this manuscript therefore does not lie in an efficiency advantage over established methods – no such advantage is claimed anywhere – but in the explicit, minimal operationalization of the DRCC

stability condition through a single correction order, together with its empirical characterization (Section 3).

Comparative complexity analysis (time, state, domain). Beyond the empirical comparison of Test C, the comparison can also be conducted at the level of asymptotic complexity – with due care about which of the following statements are actually supported by this work and which are not.

The four rows are thus supported to different degrees: the time comparison is a direct measurement (Table 5), the state comparison a formal proof (Lemma 1), while the domain of validity is only the portion actually tested in this work – a claim of validity over the entire complex plane would not be supported by our own tests and is not made here. No comparative test on stability near Gram points was carried out; an advantage of Kahan compensated summation over Riemann-Siegel at such points is therefore not claimed – Kahan summation compensates floating-point rounding error over long addition chains, which is a different phenomenon from the structural sign degeneracy of $Z(t)$ near Gram points, a property of the analytic function itself rather than of floating-point arithmetic; conflating the two would be misleading.

The defensible core of this comparison is thus: Riemann-Siegel uses an $O(\sqrt{t})$ main-sum length versus the empirically observed $O(t^{0.68})$ tested truncation requirement for DWZ at fixed ε – a term-count comparison, not by itself a runtime claim; the wall-clock benchmarks of Table 5 separately show a large runtime advantage for Riemann-Siegel. For DWZ, in contrast, it is *proven*, not merely observed, that the register width stays constant ($\omega_{\text{reg}} = O(1)$, Lemma 1), independently of the truncation cost. This contrast – a measured time disadvantage against proven register-width constancy – is the differentiated claim actually supported here; broader formulations (validity over the entire complex plane, Gram-point robustness via Kahan summation) are deliberately not adopted, since they go beyond what is actually tested in this work.

Register width is not a runtime notion. An important conceptual clarification follows from this analysis: register width and runtime complexity are distinct notions that cannot be traded off against each other. The argument in Lemma 1 rests solely on the fact that the terms k^{-s} are deterministically reconstructible from the index k and the parameter s and need not be held simultaneously; the main sum of the Riemann-Siegel formula (terms $n^{-1/2} \cos(\vartheta(t) - t \ln n)$) has the same property. This is not, however, a general claim about every implementation of the Riemann-Siegel formula, but only about a reconstruction model that shares the same assumptions: under a reconstruction model with the same assumptions about deterministic recomputability, a constant register width $\omega_{\text{reg}} = O(1)$ appears plausible for the Riemann-Siegel main sum as well, under the task τ_{out} ; a corresponding formal proof, however, is *not* carried out in this work, and the claim is therefore not asserted, but stated as a plausible, open corollary. If it held, an important consequence would follow: two algorithms can share the same register-width class ($O(1)$) and still differ by orders of magnitude in runtime, exactly as the times measured in this section and in Section 3 show. The register-width finding for DWZ reported above is thus not a claim of an advantage over Riemann-Siegel, but a claim about the structure of the DWZ reconstruction itself, regardless of whether Riemann-Siegel shares the same property.

Conceptual caveat. The identification $d_{\text{frag}}^{\text{DWZ}}(X; W) := W_{\text{trunc}}$ introduced in Section 2 (equation (6)) denotes a numerical truncation depth: the number of explicitly summed series terms. This is not automatically the same as the register width ω_{reg} defined below, which counts the number of components that must be *simultaneously kept active* at an intermediate stage. The following lemma makes this conjecture precise and provable for the most natural reconstruction-model case.

Table 6: Comparative complexity: Riemann-Siegel vs. DWZ (M_1). The marker at the end of each row indicates how the corresponding claim is supported.

Criterion	Riemann-Siegel	DWZ (M_1)
Time per evaluation	$O(\sqrt{t})$ terms (established result)	$O(W_{\text{trunc}})$ per construction; empirically the required truncation depth $W_{\min}(M_1; s, \varepsilon) \sim t^{0.68}$ at fixed ε (W_{trunc} itself remains the free variable over which this minimum is taken) (regression fit over the four measured points in Table 5, $R^2 = 0.999$; here ε is the function-value tolerance of $W_{\min}(M; s, \varepsilon)$, whereas the exponent $2/3$ of equation (20) is fitted to the zero-position error $e_k^{(M_1)}(W)$. Numerically close to the $2/3$ -type exponent appearing in equation (20), but obtained under a different error functional; no direct identification is claimed) – <i>measured</i>
State width	$O(1)$ (or $O(\log t)$ for the bit width of the index/loop register, standard convention) – <i>plausible under an analogous streaming model; not formally instantiated here</i> (see “Register width is not a runtime notion” below)	$\omega_{\text{reg}} = O(1)$, independent of W_{trunc} – <i>proven</i> (Lemma 1)
Domain of validity	established primarily on the critical line $\text{Re}(s) = \frac{1}{2}$	tested in this work for the critical strip $0 < \text{Re}(s) < 1$; the underlying Euler-Maclaurin expansion is classically continuable analytically to a larger domain, but this is not independently tested here – <i>not tested</i>
Numerical stability near Gram points	documented in the literature as sensitive (sign changes of $Z(t)$) [3]	no comparative test performed in this work – <i>no claim</i>

Register-model instantiation. This manuscript introduces its own, self-contained register-counting quantity for the sequential zeta task. Let $\mathfrak{M}_{\text{seq}}$ denote the reconstruction-model family consisting of implementations that process terms $k = 1, \dots, N$ in ascending order and hold, at each step, exactly the quantities required to compute the next partial sum and nothing else. Register counting here uses an idealized *real-RAM-style scalar-storage convention*: each retained component of the persistent state holds exactly one value of the underlying numeric domain (real or complex), for the sole purpose of counting how many such simultaneously-held values are required – not a claim about finite-precision exact representability, which is addressed separately below (“Word-count versus bit-space”). For a realization $\mathcal{A} \in \mathfrak{M}_{\text{seq}}$ with persistent logical state $\mathcal{Z}_i$ held at step i , define the *register width*

$$\omega_{\text{reg}}(\tau_{\text{out}}, \pi_{\text{seq}}; \mathfrak{M}_{\text{seq}}) := \min_{\mathcal{A} \in \mathfrak{M}_{\text{seq}}} \max_i |\mathcal{Z}_i|, \quad (14)$$

a quantity defined directly and entirely within this manuscript, in terms of persistent logical state registers of a declared implementation family. Lemma 1 below proves a statement about ω_{reg} alone, self-contained and not dependent on any externally cited formalism.

Lemma 1 (Register width of sequential zeta accumulation). *Let τ_{out} be the task “output only the final value S_N ”, and let π_{seq} be the ordering that processes the terms $k = 1, \dots, N$ in ascending index order. Then*

$$\omega_{\text{reg}}(\tau_{\text{out}}, \pi_{\text{seq}}; \mathfrak{M}_{\text{seq}}) = O(1)$$

independent of N – a constant number of persistent logical registers in a fixed-precision word model, not constant total bit-space as $N \rightarrow \infty$ (see the “Word-count versus bit-space” paragraph below for the $\Theta(\log N)$ index-width and precision-growth caveats).

Proof. The series $S_N = \sum_{k=1}^N k^{-s}$ is evaluated via the recursion $S_k = S_{k-1} + f(k, s)$ with $f(k, s) = k^{-s}$. At time $k-1$, S_{k-1} has been computed; determining the successor state S_k requires, by the recursion, only the evaluation of $f(k, s)$, with input quantities: the previous partial sum S_{k-1} , the running index k , and the fixed parameter s . Hence a sufficient persistent state at step k is

$$\mathcal{Z}_k = \{S_{k-1}, k, s\}, \quad |\mathcal{Z}_k| = 3.$$

(This exhibits a witness realization achieving $|\mathcal{Z}_k| = 3$; the lemma’s $O(1)$ claim below needs only this upper bound, not a proof that no smaller persistent state exists, so “minimal” is not claimed here.) Each element of $\mathcal{Z}_k$ occupies a fixed number c of registers or memory words, independent of N . Since $f(k, s)$ is evaluated directly at step k and immediately accumulated into S_k , there is no informational dependence on earlier terms $f(k-1, s), f(k-2, s), \dots$ – these are overwritten after addition and removed from the persistent state. For an error-corrected variant (Kahan summation), the state grows only by one constant correction variable c_{k-1} ,

$$\mathcal{Z}_{k,\text{Kahan}} = \{S_{k-1}, k, s, c_{k-1}\}, \quad |\mathcal{Z}_{k,\text{Kahan}}| = 4.$$

In both cases $|\mathcal{Z}_k|$ is constant and independent of the truncation index N . This realization is admissible for τ_{out} under the sequential ordering π_{seq} (every step is deterministically executable from $\mathcal{Z}_k$, and S_N is fully reconstructed at the end), so it witnesses $\omega_{\text{reg}}(\tau_{\text{out}}, \pi_{\text{seq}}; \mathfrak{M}_{\text{seq}}) \leq |\mathcal{Z}_k| = O(1)$ by the definition of ω_{reg} (equation (14)) as a minimum over realizations in $\mathfrak{M}_{\text{seq}}$. This proves the lemma’s register-width statement. Since $\omega_{\text{reg}} = O(1)$ independent of N , this shows that the truncation index W_{trunc} (resp. N) has no effect on the register width of this realization; under τ_{out} it acts exclusively as a scaling factor for the loop’s running time, not for the register width. $\square$

Corollary 1 (Structural–Numerical Width Separation). *Fix a reference zero $\rho_k = \frac{1}{2} + it_k$ and suppose, in addition to the convergence hypothesis (21) for this t_k , that*

$$e_k^{(M_1)}(W) > 0 \quad \text{for every finite } W \in \mathbb{N}, \quad (15)$$

i.e. no finite truncation width reproduces ρ_k exactly (observed without exception for the tested pairs; not analytically proven that no finite W could ever do so). Then

$$\omega_{\text{reg}}(\tau_{\text{out}}, \pi_{\text{seq}}; \mathfrak{M}_{\text{seq}}) = O(1) \quad (\text{unconditionally, Lemma 1})$$

while

$$W_{\text{stab}}^{(M_1)}(t_k, \varepsilon) \longrightarrow \infty \quad \text{as } \varepsilon \downarrow 0 \quad (\text{conditionally on (21) and (15)}),$$

hence $W_{\text{stab}}^{(M_1)}(t_k, \varepsilon) \neq \omega_{\text{reg}}$.

Proof. The first claim is Lemma 1, proved unconditionally above. For the second claim, suppose toward a contradiction that $W_{\text{stab}}^{(M_1)}(t_k, \varepsilon)$ does *not* diverge as $\varepsilon \downarrow 0$: then there is a bound $W^* \in \mathbb{N}$ and a sequence $\varepsilon_n \downarrow 0$ with $W_{\text{stab}}^{(M_1)}(t_k, \varepsilon_n) \leq W^*$ for every n . By the defining property (19) of $W_{\text{stab}}^{(M_1)}(t_k, \varepsilon_n)$, this means $e_k^{(M_1)}(W) < \varepsilon_n$ for every $W \geq W_{\text{stab}}^{(M_1)}(t_k, \varepsilon_n)$; in particular, since $W^* \geq W_{\text{stab}}^{(M_1)}(t_k, \varepsilon_n)$, this gives $e_k^{(M_1)}(W^*) < \varepsilon_n$ for every n . Letting $n \rightarrow \infty$, $\varepsilon_n \rightarrow 0$, so $e_k^{(M_1)}(W^*) \leq 0$. Since $e_k^{(M_1)}(W^*) \geq 0$ by definition (it is a distance), this forces $e_k^{(M_1)}(W^*) = 0$, contradicting the non-degeneracy hypothesis (15) at the finite width W^* . Hence no such finite bound W^* exists, i.e. $W_{\text{stab}}^{(M_1)}(t_k, \varepsilon) \rightarrow \infty$ as $\varepsilon \downarrow 0$. $\square$

This corollary should not be read more strongly than stated: it is conditional on the same numerically-observed, analytically-unproven hypotheses already carried by Proposition 2 and by the defining set of equation (19), plus the additional explicit non-degeneracy hypothesis (15) introduced here. What it adds over Proposition 2 is not a stronger convergence claim, but the precise logical reason that ω_{reg} and $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ cannot be the same quantity: one is bounded independently of the target accuracy, the other is not, under the same hypotheses already in use elsewhere in this manuscript.

Remark 1 (Separation Result, Not a Deficiency). *Lemma 1 is best read as a positive separation result rather than as a limitation of the case study. It establishes that two quantities this manuscript is careful never to conflate – the register width ω_{reg} on one hand, and the numerical truncation depth W_{trunc} resp. the stable threshold $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ on the other – are provably distinct for the sequential zeta accumulation: the former is $O(1)$, constant in the truncation index, while the latter grows without bound as the target accuracy $\varepsilon \rightarrow 0$ (Corollary 1, conditional on the non-degeneracy hypothesis stated there). This is not a coincidental byproduct of an easy problem; it is the manuscript’s central structural finding, precisely because it shows that “more terms summed” and “more state carried” are not the same kind of resource for this reconstruction task. The finding is specific to the sequential-accumulation realization studied here, not a claim that $\omega_{\text{reg}} = O(1)$ for every DRCC instantiation; a problem family for which the structural width instead scales with problem size would be the natural next test of the same separation question (Section 6).*

Word-count versus bit-space. The quantity actually shown to be $O(1)$ by this lemma is the number of persistent logical state components (equivalently, registers or memory words) in a fixed-precision word model – $|Z_k| = 3$ (resp. 4 with Kahan compensation), independent of N . This is not the same statement as $O(1)$ bit-space complexity: the index k requires $\Theta(\log N)$ bits to represent for k up to N , and the precision needed for the complex accumulator S_{k-1} may itself grow with the target accuracy ε or with N , depending on the numerical representation chosen (as in the multi-precision `mpmath` arithmetic used throughout the numerical experiments of this work, where working precision is fixed by hand, not derived from a bit-complexity analysis). $\omega_{\text{reg}} = O(1)$ in this manuscript therefore means: constant *word count* in a fixed-precision model, not constant total bit-space as $N \rightarrow \infty$. Table 6 already notes the standard $O(\log t)$ bit-width convention for the loop index in passing; this paragraph makes the same point explicit at the level of the lemma itself, where it is most load-bearing.

Scope of the lemma. In one sentence: Lemma 1 is not sufficient to establish the original, stronger headline identification $\omega_{\text{reg}} = W_{\text{trunc}}$, but it remains a valid and non-trivial structural result about the sequential reconstruction model in its own right – it formally separates runtime length from active state width, which the naive identification had conflated. More precisely: Lemma 1 shows that, for the natural task “output only the final value”, the identification $\omega_{\text{reg}} = W_{\text{trunc}}$ does *not* hold – ω_{reg} is constant under τ_{out} , while $W_{\text{trunc}} = N$ grows without bound. This direction of the originally open question is thus now proven, not merely conjectured. Whether there exists a stricter, separately declared task (e.g., requiring that every term remain individually reconstructable or auditable afterward) under which an active interface of size $\Theta(W)$ is genuinely required remains a separate, open construction question; it is neither claimed nor shown here. What remains defensible in addition is the numerically demonstrated statement (12): a stronger analytic reconstruction model reduces the required DWZ truncation width by more than a factor of ten.

Even under a stricter task: the reconstruction dilemma. One might suspect that a stricter task τ_{audit} – “every term must remain individually reconstructable afterward” – automatically forces an active interface of size $\Theta(W)$. This is not the case, however, as long as the declared reconstruction model permits recomputation as an admissible reconstruction operation: since each term k^{-s} is a pure, deterministic function of k and s , its value can be recomputed at any time from the pair (k, s) without ever needing to be stored in between. The entire information content of the sequence of terms is thus exhausted by the compact descriptor (N, s) ; a linear state width could only be justified if the sequence of terms admitted no compact description – which is evidently not the case for an analytic function such as k^{-s} , since it is fully determined by the closed-form map $(k, s) \mapsto k^{-s}$. (This paragraph deliberately states this intuition without invoking a formal notion of Kolmogorov complexity – such a formalization, including its dependence on the representation of s and the required output accuracy, is not carried out here; the load-bearing, formal argument of this section remains the direct register-count construction developed in the following paragraph, not this informal remark.) Two prefixes of the accumulation that agree on the running sum S_i , the index i , and the parameter s can answer any future admissible query – including an audit query for any already-processed term $j \leq i$, by recomputing $f(j, s) = j^{-s}$ on demand – identically, regardless of the specific earlier history that produced them; terms already processed but remain recomputable at any time can therefore be omitted from the persistent state.

Direct register count for τ_{audit} . Consider the sequential zeta accumulation under the stricter task τ_{audit} (“every term must remain individually reconstructable afterward”), still processed in ascending order π_{seq} . An admissible realization needs to answer, at every step i , both the running-sum query (as under τ_{out}) and any audit query for a term $j \leq i$. Retaining exactly

$$\mathcal{Z}_i = \{S_i, i, s\}, \quad |\mathcal{Z}_i| = 3$$

(one additional register for a Kahan correction raises this to 4) suffices for both: the running sum is answered directly from S_i , and an audit query for any $j \leq i$ is answered by recomputing $f(j, s) = j^{-s}$ from j (supplied by the query itself) and the retained s – no raw term $f(1, s), \dots, f(i, s)$ needs to be stored. Hence

$$\omega_{\text{reg}}(\tau_{\text{audit}}, \pi_{\text{seq}}; \mathfrak{M}_{\text{seq}}) \leq 3 = O(1),$$

extending Lemma 1 to the stricter task τ_{audit} , under the reconstruction model admitting recomputation as a valid operation. What remains open is exclusively the question of whether there exists a task under which this compression fails – such a task is neither constructed nor excluded here.

A further caveat: determinism limits the strength of this compression. For fixed s and the declared ascending order, the entire sequence of terms $f(j, s) = j^{-s}$, $j = 1, \dots, N$, is a deterministic function of j and s alone: there is essentially only one reachable prefix history per stage i (namely $r_i = (f(1, s), \dots, f(i, s))$, fully determined once s and i are fixed), not a large family of distinct reachable histories being quotiented down to a small representative set. Part of what makes $|\mathcal{Z}_i| = 3$ small here is therefore this determinism and lack of branching in the task, not solely a non-trivial compression argument over many genuinely distinct reachable states. This does not invalidate the register-width count $|\mathcal{Z}_i| = 3$ itself, but it tempers how much should be read into it as evidence of a general compression phenomenon.

Conceptual separation. We separate the required truncation depth $W_{\min}(M; s, \varepsilon)$, Equation (11) (the depth needed to reach accuracy ε ; W_{trunc} itself remains the free variable this minimum ranges over), from the register width ω_{reg} , which counts the maximum number of simultaneously active components required for future state continuations within the declared reconstruction model. $\omega_{\text{reg}} = O(1)$ holds for the sequential accumulation of the zeta approximation, including under τ_{audit} – since the terms are deterministically reconstructable from k and s . It is worth stressing further that $\omega_{\text{reg}} = O(1)$ for τ_{audit} concerns only the *storage width* at any given moment, not the *time cost* of an audit request itself: each individual recomputation of a term k^{-s} takes constant time, but a full audit of all N terms requires $\Theta(N)$ such recomputations and hence $\Theta(N)$ total time across all requests. Storage width, reconstruction time per request, and total time for all requests are thus three distinct quantities; only the first is captured by $\omega_{\text{reg}} = O(1)$.

Cost model assumptions for the audit task. The instantiation above relies on four modeling choices that are stated here explicitly rather than left implicit. First, the parameter s is treated as encoded at a fixed working precision, chosen once for the whole computation and independent of the term index k or the truncation length N ; a precision that must itself grow with N or with k would not change the register *count* in $\mathcal{Z}_k$, but would affect the size of each register (see the word-count-versus-bit-space discussion above), which the $O(1)$ claim does not address. Second, a recomputed term $f(k, s) = k^{-s}$ is assumed to reproduce the same value, to the same working precision, as if it had been retained from storage; no separate error analysis for recomputed versus stored terms is carried out here, and none is needed for the state-width claim itself, which concerns which quantities must be held, not how accurately they are computed. Third, evaluating $f(k, s)$ for a single k is treated as an elementary, constant-time operation at fixed working precision; this is the standard convention used throughout this manuscript’s runtime comparisons (Section 2.5), but it hides a real, precision-dependent cost – evaluating k^{-s} to p digits of precision via standard multiprecision exponentiation costs more than $O(1)$ bit operations in p (and, for very large k , in $\log k$ as well) – which is a statement about bit-complexity, not about the register count $|\mathcal{Z}_k|$ that Lemma 1 addresses. Fourth, and consequently, none of these choices affects the lemma’s actual claim: $\omega_{\text{reg}} = O(1)$ is, throughout this section, a statement about the number of persistent logical state components (a small, fixed set of registers), not about the bit-size of each register or the time cost of computing their contents – both of which may grow with precision and are explicitly out of scope for this particular claim.

Algorithmic illustration (register invariance). The following Python program realizes the reconstruction constructed in the proof of Lemma 1 using Kahan compensated summation, making the register invariance of ω_{reg} directly observable:

```

1 def naive_partial_sum_kahan(s: complex, N: int) -> complex:
2     """
3     Computes the NAIVE partial sum (equation (5)), not the DWZ-corrected
4     function (equation (8)), via sequential Kahan-compensated accumulation.
5     For 0 < Re(s) < 1 this partial sum diverges as N grows (Table 3); the

```

```

6   example below with s=2 therefore deliberately lies outside the critical
7   strip and serves only to illustrate register behavior, not to validate
8   the DWZ method in the critical strip.
9   Space complexity (state width): O(1) -- constant
10  Time complexity (truncation length): O(N) -- linear
11  """
12  S = 0.0 # accumulator (running sum)
13  c = 0.0 # compensation term for rounding error (Kahan)
14  for k in range(1, N + 1):
15      term = k ** (-s)
16      y = term - c # subtract previous error from new term
17      t = S + y # add to accumulator
18      c = (t - S) - y # compensation estimate for lost low-order bits
19      S = t # update state
20  return S
21
22  s_value = 2.0
23  truncation = 1_000_000
24  result = naive_partial_sum_kahan(s_value, truncation)
25  print(f"Partial sum({s_value}) approximation: {result}")

```

Regardless of whether $N = 10$ or $N = 10^9$ is chosen, the program allocates the same number of registers at every point in time ($S, c, k, s, \text{term}, y, t$) – no memory array of length N is ever created; each term exists only transiently in a register and is immediately overwritten after processing. Of these seven registers, $\{S, c, k, s\}$ form the *persistent* active state $\mathcal{Z}_{k,\text{Kahan}}$ defined in Lemma 1 (carried forward from step k to step $k+1$, $|\mathcal{Z}_{k,\text{Kahan}}| = 4$); term, y , and t are, by contrast, purely *transient* intermediate values that are computed and discarded within a single step and are not carried forward to the next step – they therefore do not count toward the persistent state $\mathcal{Z}_k$ from Lemma 1. The total of seven simultaneously allocated registers is a property of *this particular implementation* and a valid upper bound ($\omega_{\text{reg}} \leq 7$ for this code), but not a tighter figure than the bound proven in the lemma for the persistent state, namely 3 (without Kahan) or 4 (with Kahan). The variables c, y, t reduce accumulated floating-point summation error even for very large N (Kahan compensation does not guarantee full numerical accuracy and does not change the truncation error), without the algorithmic state width ever growing beyond $O(1)$. For $s = 2$ and $N = 10^6$, the program returns $1.6449330668\dots$ against the known limit $\pi^2/6 = 1.6449340668\dots$, a deviation of about 10^{-6} – consistent with the expected truncation error of a directly truncated Dirichlet series at this N (numerically verified). The algorithmic implementation thus illustrates the register-count argument of Lemma 1 on one concrete run: the register width remains strictly at $\omega_{\text{reg}} = O(1)$ throughout, while W_{trunc} (here controlled via the loop limit N) scales exclusively the running time, not the register width. A single run of one implementation illustrates, but does not itself add further proof to, the already-proven lemma.

2.6 Packaging and Reduction Consistency: What It Does and Does Not Establish

A reviewer asked for an explicit justification, in DRCC terms, of why adding the correction term $C_W(s)$ in (8)/(10) to the naive fragment $F_W(s)$ should count as an admissible “reconstruction” in the sense of Hesamiy [7]. We give the minimal packaging formalization relevant to that question, while making explicit from the outset that it establishes only packaging consistency, not reconstruction admissibility, analytic validity, or numerical correctness.

Definition 1 (Packaging operator). *Let $D \subseteq \mathbb{C}$ denote the domain of the complex parameter s under consideration (e.g. the critical strip). For a fixed correction term $C_W(s)$ – equal to the full correction added to $F_W(s)$ in the boxed formula, i.e. $C_W(s) = \frac{W^{1-s}}{s-1} - \frac{W^{-s}}{2}$ for M_1 (equation (8)), resp.*

$C_W(s) = \frac{W^{1-s}}{s-1} - \frac{W^{-s}}{2} + \frac{s}{12}W^{-s-1}$ for M_2 (equation (10)), so that $F_W(s) + C_W(s) = \zeta_{\text{DWZ}}(s, W)$ identically – define

$$S_W : D \longrightarrow \mathcal{Y}_{\text{pack}}, \quad S_W(s) := (F_W(s), C_W(s)), \quad \mathcal{Y}_{\text{pack}} := \mathbb{C} \times \mathbb{C},$$

with $F_W(s) = \sum_{n=1}^W n^{-s}$ the retained fragment (equation (5)). The packaging codomain is deliberately notated $\mathcal{Y}_{\text{pack}}$ rather than Ω , to avoid any visual clash with the reconstruction fiber $\Omega_{C_4}(f)$ of Appendix B.2, which denotes an unrelated construction. Note that S_W takes the parameter $s \in D$ as its argument, not the value $F_W(s)$ – this avoids any ambiguity in case F_W fails to be injective on D .

Definition 2 (Reduction operator). Define $R_W : \mathcal{Y}_{\text{pack}} \rightarrow \mathbb{C}$ by $R_W(F, C) := F$.

Proposition 1 (Packaging Consistency Identity for M_1 and M_2 – a Tautological Bookkeeping Fact, Not a Reconstruction Result). $R_W \circ S_W = F_W$ as functions $D \rightarrow \mathbb{C}$.

Proof. For every $s \in D$, $R_W(S_W(s)) = R_W(F_W(s), C_W(s)) = F_W(s)$, directly from Definitions 1 and 2. $\square$

What this does not show. The proof of Proposition 1 uses only the defining equation of R_W and never uses the specific value or analytic origin of $C_W(s)$. The identity $R_W \circ S_W = F_W$ therefore holds for *every* choice of correction term whatsoever, including an analytically incorrect correction, for example $C_W \equiv 0$ or $C_W \equiv 42$ on D . As formalized here, the packaging-consistency identity is a tautological consistency check on the packaging $(F, C) \mapsto F$: it confirms only that the reduction operator was defined consistently with the packaging operator, not that the reconstruction is numerically sound or even sensible. It therefore cannot, by itself, serve as a criterion for choosing between M_1 , M_2 , or any other additively-corrected model – that evidence is supplied separately, by the convergence properties of (8)/(10) and by the empirical results of Section 3, not by Proposition 1 itself.

A second limitation concerns Variant B. $\zeta_W^\eta(s, W) = \eta_W(s, W)/(1 - 2^{1-s})$ is a quotient by the s -dependent factor $1 - 2^{1-s}$; defining $C_W^\eta := \zeta_W^\eta - F_W$ would algebraically give $\zeta_W^\eta = F_W + C_W^\eta$, but this rewriting would make Proposition 1 tautologically applicable without independently motivating C_W^η , unlike the Euler-Maclaurin corrections of M_1 , M_2 . No such packaging is constructed for M_η here, so its packaging-consistency status remains uninstantiated. This matters for the selection procedure of Section 2.7.

2.7 Preliminary RP-Tensor Scoring Framework for Reconstruction Variants

Section 2.5 shows that M_1 , M_η , and M_2 differ measurably in accuracy and required width but stops short of a declared selection rule. We formalize the selection problem, and then check, criterion by criterion, how much of it the manuscript’s own data currently support – rather than presenting a completed optimization that has not actually been carried out.

Formal setup. Let $V_{\text{crit}} = \{c_{\text{err}}, c_{\text{run}}, c_{\text{depth}}, c_{\text{width}}, c_{\text{stab}}\}$ and, for each model $M_\alpha \in \{M_1, M_\eta, M_2\}$ and evaluation context $\xi = (s, \varepsilon)$ – deliberately not denoted τ , to avoid confusion with the computational tasks $\tau_{\text{out}}, \tau_{\text{audit}}, \tau_{\text{grad}}$ used elsewhere in this manuscript, since ξ here is a numerical evaluation point, not a task – let $T_{\alpha j}(\xi)$ be the raw value of criterion j for model α . With admissibility indicator $a_\alpha \in \{0, 1\}$, assigned only after the numerical and analytic validity requirements of the declared model have been checked – not from the packaging-consistency identity of Section 2.6 alone, which by itself supports a_α for no model – and weights $w_j \geq 0$, $\sum_j w_j = 1$, define, piecewise to avoid the ill-defined expression $0 \cdot \infty$ that a single combined formula could produce

for an inadmissible model with $a_\alpha = 0$, where the criterion terms $\tilde{T}_{\alpha j}$ need not themselves be well-defined,

$$S_\alpha(\xi) = \begin{cases} \sum_j w_j \tilde{T}_{\alpha j}(\xi), & a_\alpha = 1, \\ +\infty, & a_\alpha = 0, \end{cases} \quad M^*(\xi) \in \arg \min_\alpha S_\alpha(\xi), \quad (16)$$

with $\tilde{T}_{\alpha j}$ the min-max normalization of $T_{\cdot j}$ over the admissible candidate set (conventionally 0 if $T_j^{\max} = T_j^{\min}$, since a tied criterion then discriminates no candidate), and where “ $\in$ ” rather than “ $=$ ” is used because $\arg \min$ is in general set-valued (ties are possible).

Data availability check. Before choosing any weights, we check which cells of the 3×5 criterion table are actually backed by data already reported in this manuscript.

Table 7: Availability of RP-Tensor criteria for M_1 , M_η , M_2 in this manuscript’s own results. “Measured” = direct table entry; “extrapolated” = derived from measured points via a fitted or stated asymptotic law; “structural” = a definitional count, not an empirical measurement; “–” = not evaluated anywhere in this work.

Criterion	M_1	M_η	M_2
c_{err} (accuracy at fixed W)	measured	measured	measured
c_{width} (W at fixed ε)	measured	extrapolated	measured
c_{run} (wall-clock time)	measured	–	–
c_{depth} (EM correction order)	structural (1)	n/a (not an EM series)	structural (2)
c_{stab} (near Gram points)	–	–	–
Packaging-consistency identity status (§2.6)	established	open	established (same trivial argument)

To state this without hedging: no data exist in this work for $c_{\text{run}}(M_\eta)$, $c_{\text{run}}(M_2)$, or c_{stab} for any of the three models – these are not scores that remain to be discussed or interpreted, they are entries this manuscript has not measured. Two of five criteria are entirely empty for two of the three models (c_{run} , measured only for M_1 against M_{K5} /Riemann-Siegel in Table 5; c_{stab} , explicitly not tested for any model near Gram points, Section 2.5), one criterion is a definitional count rather than a measurement, and M_η ’s admissibility itself is unresolved. The last row of Table 7 was relabeled to make explicit that it reports only whether the tautological packaging-consistency identity of Section 2.6 holds for each model – as stated already in that section, this status is not itself the admissibility indicator a_α of equation (16), since a_α must additionally be assigned from the numerical and analytic validity requirements of the declared model, which the packaging identity alone does not establish. A genuine five-criterion, three-model RP-Tensor cannot currently be computed. What can honestly be computed is a pairwise comparison restricted to the two empirically measured performance criteria, c_{err} and c_{width} – at the two reference points where those numbers were actually measured. c_{depth} is deliberately excluded here despite being numerically defined (Table 7: 1 for M_1 , 2 for M_2), since it is a structural

model-description count, not a measured performance quantity, and is not defined on the same scale for M_η .

Table 8: Illustrative pairwise normalized comparison on the two available criteria (not a genuine five-criterion, three-model RP-Tensor selection result, per the preceding paragraph), computed from Table 4 (width, $t = 1001.3495$, $\varepsilon = 10^{-4}$), Figure 5 (error at common $W = 2000$, same t), and Table 13 (error, $t = 14.1347$; width at $\varepsilon = 10^{-4}$ obtained by log-log interpolation for M_1 , resp. extrapolation via the fitted local exponent -0.4997 for M_η , consistent with the stated law $O(W^{-\text{Re}(s)})$). $\tilde{T} \in \{0, 1\}$: min-max normalization over a two-element set is necessarily binary and reports only which model is better, not by how much; the ratio column preserves that information.

Comparison	Criterion	M_1	other	Ratio (disadvantage)
M_1 vs. M_2 , $t=1001.35$	$W_{\text{test}}(\varepsilon=10^{-4})$	9000	810	$11.1 \times$
M_1 vs. M_2 , $t=1001.35$	error at $W=2000$	9.37×10^{-4}	3.92×10^{-6}	$239 \times$
M_1 vs. M_η , $t=14.13$	$W(\varepsilon=10^{-4})$	≈ 517	$\approx 4.46 \times 10^6$	$8633 \times$ (favors M_1)
M_1 vs. M_η , $t=14.13$	error at $W=5000$	3.33×10^{-6}	2.98×10^{-3}	$895 \times$ (favors M_1)

With equal weights over the two available criteria ($w_{\text{err}} = w_{\text{width}} = 0.5$), and setting $a_\alpha = 1$ for both models in each pairwise comparison purely to make the min-max normalization well-defined for this illustrative computation – not as an assertion that $a_{M_\eta} = 1$ is established, which remains open (Section 2.6) – the restricted score (17) below ($J_{\text{avail}} = \{c_{\text{err}}, c_{\text{width}}\}$, avoiding a sum over the otherwise-undefined remaining criteria) gives $S_{M_1} = 1$, $S_{M_2} = 0$ at $t = 1001.35$ (so $M^* = M_2$), and $S_{M_1} = 0$, $S_{M_\eta} = 1$ at $t = 14.13$ (so $M^* = M_1$ conditional on $a_{M_\eta} = 1$, which this computation assumes for illustration rather than establishes).

Honest reading of these two computations. On the only two criteria for which this manuscript reports real numbers, M_2 dominates M_1 decisively at $t \approx 1001$, and M_1 dominates M_η decisively at $t \approx 14$. Combining these into a single three-way ranking additionally assumes – without separate verification – that the ordering transfers across t , since the two comparisons were not measured at a common (s, ε) . At the context where M_1, M_2 were actually compared, the restricted two-criterion score selects M_2 over M_1 ; no three-model $M^*(\xi)$ is established here, since M_η was scored at a different ξ . This directly contradicts reading M_1 as the output of an already-completed RP-Tensor optimization. M_1 is retained as the primary model in this manuscript for a different, already stated reason – minimality and direct traceability of the construction (Section 2) – not because a computed selection procedure ranked it first on error, width, runtime, or stability. That reason stands on its own, but it is not one of the five declared criteria in V_{crit} . Either a sixth, explicitly weighted criterion for structural simplicity should be added to equation (16), or the choice of M_1 should be stated plainly as an expository decision that sits outside the RP-Tensor ranking rather than as its result.

Formalizing which criteria are actually evaluable. To make the partiality of the score in equation (16) explicit rather than implicit, restrict the sum to the evaluable criterion set $J_{\text{avail}} \subseteq J_{\text{full}} = \{c_{\text{err}}, c_{\text{run}}, c_{\text{depth}}, c_{\text{width}}, c_{\text{stab}}\}$,

$$S_\alpha(\xi; J_{\text{avail}}) = \begin{cases} \sum_{j \in J_{\text{avail}}} w_j \tilde{T}_{\alpha j}(\xi), & a_\alpha = 1, \\ +\infty, & a_\alpha = 0, \end{cases} \quad (17)$$

where the weights are renormalized on J_{avail} , $w_j \mapsto w_j / \sum_{\ell \in J_{\text{avail}}} w_\ell$, so that $\sum_{j \in J_{\text{avail}}} w_j = 1$ and the restricted score stays on the same $[0, 1]$ scale as equation (16), using otherwise exactly

the same notation as that equation (S_α a function of the evaluation context ξ , with α indexing the model M_α , and $\tilde{T}_{\alpha j}$ the same min-max-normalized criterion values defined there), now made explicit in the evaluable criterion set J_{avail} , reported alongside every score it produces. Deliberately, J_{avail} is not restricted to directly measured entries: it also admits entries obtained by a stated, checkable extrapolation law, as long as that provenance is disclosed. For the present manuscript,

$$J_{\text{avail}} = \{c_{\text{err}}, c_{\text{width}}\},$$

where, per Table 7, $c_{\text{err}}(M_1)$, $c_{\text{err}}(M_2)$, and $c_{\text{width}}(M_1)$, $c_{\text{width}}(M_2)$ are measured directly, while $c_{\text{width}}(M_\eta)$ is only extrapolated via a fitted local exponent, not directly measured – a distinction J_{avail} itself does not encode and that any consumer of equation (17) should recover from Table 7 before treating the three models as evaluated on equal footing. c_{run} , c_{depth} , c_{stab} are undefined/not measured or extrapolated for at least one of M_1 , M_η , M_2 , not merely omitted from discussion. A ranking computed from equation (17) is a statement about J_{avail} only; it is not upgraded to a full five-criterion ranking by adding more models or more t -values without also measuring c_{run} , c_{depth} , and c_{stab} for the same models, the same t , the same ε , and the same benchmark protocol.

A caveat on min-max normalization with few candidates. With exactly two candidate models – the case for every row of Table 8, which is pairwise – min-max normalization necessarily maps every criterion to $\{0, 1\}$ regardless of the actual size of the gap: a 10% disadvantage and a 10,000% disadvantage both become $\tilde{T} = 1$. With three candidates (the full set $\mathcal{M}_{\text{cand}} = \{M_1, M_2, M_\eta\}$), this collapse to $\{0, 1\}$ applies only to the best- and worst-scoring model on a given criterion; the remaining, middle-ranked model’s normalized value can fall anywhere in the open interval $(0, 1)$ and is not itself forced to 0 or 1. In either case, Table 8’s ratio column ($11\times$ to $8633\times$ depending on criterion and reference point) carries information the normalized score alone discards. Reporting the raw ratio alongside the binary score, or normalizing against a fixed external reference (e.g. a target-accuracy budget) rather than min-max over the candidate set itself, would make the RP-Tensor more informative than a bare ranking.

Remaining work before the RP-Tensor is complete. c_{run} and c_{stab} have no data for M_η or M_2 in this manuscript. Completing them requires two further numerical experiments – a runtime benchmark for M_η and M_2 at matched accuracy targets, and a Gram-point stability test across all three models – together with an independently stated analytic/numerical admissibility criterion for M_η , before a_{M_η} can be assigned rather than left open; an optional packaging-consistency construction (Section 2.6) would address bookkeeping symmetry with M_1, M_2 but, per that section, would not by itself establish $a_{M_\eta} = 1$. Until then, equation (16) is a declared selection framework, not yet a completed calculation, and the two-criterion computation above should be read accordingly: as evidence that M_2 currently outperforms M_1 on the measured axes, not as a full justification for using M_1 going forward.

2.8 Operationalizing $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$

We denote by $\rho_k = \frac{1}{2} + it_k$ the k -th reference zero (`mp.zetazero(k)`) and by $s_k^{(M)}(W)$ the zero numerically determined by model M and width W that is assigned to ρ_k according to the formal assignment rule below. If $\zeta_M(\cdot, W)$ has more than one zero near ρ_k , which particular zero is found also depends on the declared search/refinement procedure (coarse scan, fine scan, Newton refinement from a stated starting point, Section 2.10), not on M and W alone; $s_k^{(M)}(W)$ is used throughout as shorthand for the zero produced by that declared procedure, given the formal assignment rule $k^* = \arg \min_j |s_j^{(M)}(W) - \rho_j|$ (nearest reference zero; this same rule underlies the evaluation of the hybrid hardening runs in Section 3, where it produces the neighbor hits reported there whenever zero spacing is small); the corresponding position error is $e_k^{(M)}(W) :=$

$|s_k^{(M)}(W) - \rho_k|$ (to distinguish it from the function-value error $E_\zeta(s, W) := |\zeta_M(s, W) - \zeta(s)|$, used in Table 3 and Figures 5–6; Figure 12 below instead plots the zero-position error $e_k^{(M)}(W)$ itself). The numerical measurement of zero convergence (Section 3, model M_1) yielded the empirical power law

$$e_k^{(M_1)}(W) \approx C(t_k) W^{-p}, \quad p \approx 1.4997\text{--}1.5010, \quad (18)$$

determined numerically, not proven exactly as $3/2$. The formally defined stable truncation threshold is thus a conditional exact quantity, numerically approximated on the tested grid, defined, in the same *stable* form used in Section 2.8.4, as an integer minimum over widths beyond which the error remains below tolerance, and – since $e_k^{(M_1)}(W)$ is only defined relative to a known reference zero ρ_k – operationally only at reference ordinates $t = t_k$, $k = 1, \dots, 100$, from Table 16 (not yet at an arbitrary $t \in \mathbb{R}$; the continuous estimator for general t follows only via equation (20) below),

$$W_{\text{stab}}^{(M_1)}(t, \varepsilon) := \min\{W_0 \in \mathbb{N} : e_k^{(M_1)}(W) < \varepsilon \text{ for every } W \geq W_0\}, \quad t = t_k, \quad (19)$$

where $e_k^{(M_1)}(W)$ by definition measures the distance to the already-known reference zero ρ_k , so the notation $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ used throughout this manuscript is, at this exact (non-estimator) level, implicitly restricted to $t \in \{t_1, \dots, t_{100}\}$ unless stated otherwise. $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ is thus a *post-hoc* error or calibration measure against an already-known zero, not a standalone criterion by which an as-yet-unknown zero could be blindly located – a blind search (Section 2.9) instead uses the scan minimum of $|\zeta_{\text{DWZ}}|$, not $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ itself. The statement that $W_{\text{stab}}^{(M_1)}$ "operationalizes reconstruction" (Section 2.8.1) thus refers, at this stage, to accuracy against known reference values, not to a procedure for finding unknown zeros. Here $t := t_k$ is to be set when $\rho_k = \frac{1}{2} + it_k$ is the corresponding reference zero – equation (19) itself is thus, at first, only evaluated at the discrete points $t = t_k$ for which a reference zero and hence an error curve $e_k^{(M_1)}(W)$ exist, not directly defined for arbitrary real t . The stable form of the definition builds the requirement that accuracy, once reached, persists directly into $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ itself, rather than leaving it as a separate assumption; whether the defining set is nonempty for the tested pairs, however, still rests on a numerical observation, not a proof. By the definition of a limit, convergence alone, $e_k^{(M_1)}(W) \rightarrow 0$ as $W \rightarrow \infty$, already suffices to guarantee that such a W_0 exists for every $\varepsilon > 0$ – monotonicity is not required for this conclusion, although it is, in fact, also observed in the tested range (Figure 12), as a strictly stronger property than what the argument here actually needs. This convergence itself is numerically observed for the tested pairs, not analytically proven for arbitrary W : a finite computation can be consistent with $e_k^{(M_1)}(W) \rightarrow 0$ but cannot, on its own, verify the universal statement "for all $W \geq W_0$ " that the definition of $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ in equation (19) requires. The stable minimum is therefore treated as well-defined for the tested (t_k, ε) pairs only conditionally on this numerical evidence, in the same sense made precise and conditional later in Proposition 2 – not as an independently established fact at this point in the text. For an arbitrary t not coinciding with a reference zero, the zero-position error $e_k^{(M_1)}(W)$ underlying equation (19) is itself not defined, so there is no exact extension of equation (19) to such t ; what follows is an interpolated/extrapolated proxy estimate, not a definitional extension. This proxy is achieved only through the continuous relationship $C(t)$ fitted below, not through equation (19) itself. The power law (18) provides an estimate for this: in general, $e(W) \approx C(t)W^{-p}$ implies the form $W \approx (C(t)/\varepsilon)^{1/p}$, and only under the – numerically plausible but unproven – assumption that $p = 3/2$ exactly does this yield the special exponent $2/3$:

$$\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon) = \left\lceil \left(\frac{C(t)}{\varepsilon} \right)^{1/p} \right\rceil \stackrel{p \approx 3/2}{\approx} \left\lceil \left(\frac{C(t)}{\varepsilon} \right)^{2/3} \right\rceil \approx W_{\text{stab}}^{(M_1)}(t, \varepsilon), \quad (20)$$

with $C(t) \approx 0.0272t + 1.043$ (linear fit over the first 50 zeros #1–#50 from Table 16, $n = 50$; Table 16 itself has since been extended to 100 zeros, but the $C(t)$ fit still covers only the first 50, see Table 14; slope 0.0272 ± 0.0040 , intercept 1.043 ± 0.375 , 95% confidence intervals). The fit explains only part of the scatter ($R^2 \approx 0.50$); neither a square-root, nor a logarithmic, nor a pure power-law form in t fits noticeably better ($R^2 \approx 0.49, 0.47$, and 0.48 respectively). $C(t)$ increases significantly with t (the trend is highly significant at $n = 50, p \ll 0.001$), but the simple linear form explains only about half of the observed scatter – the other half likely depends on properties not captured by t alone. This $R^2 \approx 0.50$ is an in-sample figure from the calibration zeros #1–50; equation (20)’s individual-zero predictive accuracy out-of-sample, on the independent zeros #51–100, is markedly weaker than this in-sample fit quality suggests (Section 3) – it should be read as a calibrated estimator with weak individual out-of-sample predictive power, not as a validated predictive formula. Equation (19) is a *definition*; equation (20) is a derived, *empirical estimator* – the two are deliberately not treated here as exactly identical. This provides a closed-form empirical surrogate associated with the operational workflow of the originally abstract inequality (4), not a validated predictive formula, given the weak individual out-of-sample performance ($R_{\text{os}}^2 \approx 0.08$) noted above. For an untested t (not among the reference ordinates $t_1, \dots, t_{100}$), $\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon)$ is only a starting-width estimator and does not certify a zero or a stable threshold there, given the weak individual out-of-sample predictive power noted above; it should not be read as asserting that " $W \geq \hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon)$ means the zero at that t is determined to within ε ". The quantity $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ is a model-dependent numerical truncation threshold. It is motivated by, but deliberately given different notation from, the original, abstract DRCC reconstruction degree $d_{\text{rec}}^{\text{DRCC}}(X)$ from the abstract DRCC condition in Section 1.2, Equation (4) – this notational separation marks a concrete interpretation for the zeta case, not a proven identity, and is intended to prevent exactly the notational overloading that a shared symbol would otherwise invite.

2.8.1 Conditional Numerical Operationalization

The preceding sections motivate the concretization of inequality (4) through prose and numerical examples. The following proposition makes this step explicit and conditionally provable – as a direct answer to the legitimate criticism that the DRCC→DWZ transition has so far only been asserted, not demonstrated – given one stated, numerically (not analytically) supported hypothesis.

Proposition 2 (Conditional Surrogate Consistency of the Numerical Concretization). *Let $t = t_k$ for one of the tested reference zeros $\rho_k = \frac{1}{2} + it_k$ ($k = 1, \dots, 100$ from Table 16, for which $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ is defined per equation (19)), $\varepsilon > 0$ fixed, and define the manuscript-specific instantiations*

$$d_{\text{frag}}^{\text{DWZ}}(X; W) := W_{\text{trunc}} \quad (\text{equation (6)})$$

$$d_{\text{rec}}^{\text{DWZ}}(X; t, \varepsilon) := W_{\text{stab}}^{(M_1)}(t, \varepsilon) \quad (\text{equation (19)})$$

to be the concretizations chosen in this work, with W_{trunc} used throughout as a numerical surrogate for the abstract, representation-independent quantity $d_{\text{frag}}^{\text{DRCC}}(X)$, not as a proof that the two coincide. These are surrogates in a stronger sense than a mere naming choice: $d_{\text{rec}}^{\text{DRCC}}(X)$ in the cited framework is a reconstruction dimension/degree, whereas $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ here is a numerical truncation threshold – different mathematical types of object, not merely different notations for the same kind of quantity; “:=” below fixes a name for the DWZ-specific threshold, it does not assert that a type-compatible realization of $d_{\text{rec}}^{\text{DRCC}}(X)$ has been constructed. A note on the “:=” above: the DWZ superscript together with the explicit $(X; W)$ / $(X; t, \varepsilon)$ arguments marks $d_{\text{frag}}^{\text{DWZ}}$ and $d_{\text{rec}}^{\text{DWZ}}$ as new, manuscript-local quantities – the “:=” only fixes notation for these new symbols, exactly as any proposition may freely name a quantity it introduces. It is not an assertion that $d_{\text{frag}}^{\text{DWZ}}$ or $d_{\text{rec}}^{\text{DWZ}}$ is representation-independent or proven identical to the abstract $d_{\text{frag}}^{\text{DRCC}}(X)$ and $d_{\text{rec}}^{\text{DRCC}}(X)$ of the cited framework that they instantiate; that non-identity is the point of

the surrogate language above, and is why the DWZ-superscripted, argument-explicit symbols used here are kept distinct from the DRCC-superscripted abstract quantities of equation (4) used throughout this manuscript; an undecorated $W := d_{\text{frag}}(X)$ notation was used in an earlier version of this manuscript and has since been fully replaced by the argument-explicit form (Section 2). Assume that

$$e_k^{(M_1)}(W) \longrightarrow 0 \quad (W \rightarrow \infty) \quad (21)$$

for this reference zero – a hypothesis that is numerically observed for the tested pairs (Figure 12), not analytically proven. Then, for every finite width $W \geq W_{\text{stab}}^{(M_1)}(t, \varepsilon)$, the DWZ-specific numerical surrogate satisfies an inequality of the same order form as the DRCC stability condition (4) – not an instantiation of $d_{\text{rec}}^{\text{DRCC}}(X)$ or $d_{\text{frag}}^{\text{DRCC}}(X)$ themselves, which remain a different mathematical type, per the surrogate discussion above –

$$0 < W_{\text{stab}}^{(M_1)}(t, \varepsilon) \leq W < \infty,$$

and moreover

$$e_k^{(M_1)}(W) < \varepsilon.$$

Proof. By the definition of convergence, hypothesis (21) already implies that, for every $\varepsilon > 0$, there exists $W_0 \in \mathbb{N}$ with $e_k^{(M_1)}(W) < \varepsilon$ for every $W \geq W_0$; monotonicity of $e_k^{(M_1)}(W)$ is not required for this step, although it is, in fact, also observed numerically in the tested range (Figure 12), as a strictly stronger property than what is used here. Hence, under hypothesis (21), the defining set $\{W_0 \in \mathbb{N} : e_k^{(M_1)}(W) < \varepsilon \text{ for every } W \geq W_0\}$ is nonempty for every $\varepsilon > 0$ in this reference-zero case, exactly as required by the defining condition later formalized in Proposition 4; this is a logical consequence of the hypothesis as granted, distinct from the separate, already-noted question of how well hypothesis (21) itself is empirically supported (numerically observed for the tested pairs, not analytically proven, as stated above). Granting this, $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ is well-defined and finite as the least element of this nonempty subset of $\mathbb{N}$. For every finite $W \geq W_{\text{stab}}^{(M_1)}(t, \varepsilon)$, the defining property of that least element directly gives $e_k^{(M_1)}(W) < \varepsilon$, and the inequalities $0 < W_{\text{stab}}^{(M_1)}(t, \varepsilon) \leq W < \infty$ then follow from the ordering of W_0 and W in $\mathbb{N}$, as in the proof of Proposition 4. $\square$

Remark 2 (Practical Conditions for Proposition 2). *Proposition 2 is conditional on hypothesis (21), stated abstractly as a convergence assumption. In a concrete numerical application, this hypothesis becomes checkable – not proven, but checkable – once the following are in place, mirroring the structure already used in the proof of Proposition 3 (Section 2.8.2):*

- an isolated, simple reference zero ρ (so that a local neighborhood U containing no other zero of ζ can be declared);
- a zero-free boundary curve ∂U for the target function on that declared neighborhood, so that $m_U := \min_{\partial U} |\zeta| > 0$ is well-defined;
- a local Euler–Maclaurin remainder estimate that holds uniformly on $\overline{U}$ (the constant K_ρ of equation (28));
- a threshold W_0 beyond which the approximation error on ∂U stays strictly below m_U , so that Rouché’s theorem (Step 0 of the proof of Proposition 3) applies and yields a unique $\rho_W \in U$; and
- a numerically checkable tolerance condition, verified to remain stable for all tested W beyond W_0 – not merely at a single W – consistent with the stable-minimum form already built into the definition of $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ (equation (19)) rather than a one-off “first hit” criterion.

These five conditions do not constitute a proof that hypothesis (21) holds for a given t ; they specify what a full local verification of Proposition 3’s hypotheses would require. The present numerical record checks

only part of this list: Table 16 documents the resulting position deviations $e_k^{(M_1)}(W=10^4)$ for the tested reference zeros, not a per-zero record of m_U , the uniform remainder bound K_ρ , the Rouché threshold W_0 , or a direct check of stability beyond W_0 – so it is evidence consistent with the hypothesis holding, not a documented verification of all five conditions above for each zero. A further, distinct assumption – likewise declared rather than proven here – underlies the use of $s_k^{(M_1)}(W)$ throughout this manuscript: that the declared search/refinement procedure (starting near the known reference ρ_k and converging under Newton refinement, Section 2.10) locates precisely the Rouché-guaranteed unique zero $\rho_W \in U$ of Proposition 3, i.e. $s_k^{(M_1)}(W) = \rho_W$ for the tested W , rather than some other, unrelated zero of $\zeta_{\text{DWZ}}(\cdot, W)$ outside U . This is a standard and plausible property of Newton refinement from a sufficiently close starting point, but it is not separately proven or numerically audited against the Rouché disk U in this manuscript; the small, consistently converging position deviations reported in Table 16 are evidence for, not a proof of, this identification. They identify exactly what would need to be reestablished for any new t or any other reconstruction model before Proposition 2 could be applied to it. This connects directly to the six-stage operationalization workflow of Section 2.8.4: satisfying this checklist is what “validate the result against independent references” (stage 5) concretely requires for the present model M_1 .

Reading the quantifiers in plain terms. To remove any ambiguity: Proposition 2 is an “IF-THEN” statement about numerical proximity to an *already known* reference zero ρ_k , not a statement about zero location or about the critical line. It does not say “IF the error bound holds THEN the zeros lie on the critical line”; the critical-line membership of ρ_k is given as input (from `mp.zetazero`), not derived. What the proposition says is: IF the numerically observed convergence hypothesis (21) holds for a given reference zero, THEN a finite width $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ exists such that every W at or above it brings the DWZ reconstruction within ε of that already-given reference position. The Riemann Hypothesis is neither assumed nor addressed by this statement.

Scope and limits of the proposition. Proposition 2 shows that the concretizations chosen in this work satisfy the *form* of the original DRCC inequality (4) without contradiction – the DRCC→DWZ transition is thus formal at this point, not merely motivational. The additional conclusion $e_k^{(M_1)}(W) < \varepsilon$ for every $W \geq W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ rests entirely on hypothesis (21) (convergence $e_k^{(M_1)}(W) \rightarrow 0$), which is numerically observed but not analytically proven for the tested pairs; it is not a global guarantee for every $t > 0$, and it is conditional, not unconditional, exactly as the proposition’s title now states. Note, further, that the central inequality $W_{\text{stab}}^{(M_1)}(t, \varepsilon) \leq W$ follows here directly from the choice $W \geq W_{\text{stab}}^{(M_1)}(t, \varepsilon)$: the proposition primarily establishes the *consistency* of the notation so chosen with the form of equation (4), not any further nontrivial property derived from DRCC theory itself – more fitting labels would be “consistency of the numerical concretization” or “conditional numerical instantiation”; we retain “operationalization” here only because it continues the usage introduced in Section 1.2. The proposition does *not* show that $d_{\text{frag}}^{\text{DWZ}}(X; W)$ corresponds to the abstract quantity $d_{\text{frag}}^{\text{DRCC}}(X)$, or that $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ corresponds to the abstract quantity $d_{\text{rec}}^{\text{DRCC}}(X)$, or that $\omega_{\text{reg}} = W_{\text{trunc}}$ – the latter naive identification has been excluded for the declared sequential model by Lemma 1 and Corollary 1. The proposition is thus deliberately modest: it secures the logical consistency of the chosen concretization, not its identity with the original, abstract DRCC quantities.

Relationship of $\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon)$ (Eq. (20)) to the DRCC condition (Eq. (4)). Since this question recurs in review, the chain is stated explicitly and in one place here: there is *no* direct relationship between the estimator $\hat{W}_{\text{stab}}^{(M_1)}$ and equation (4), only a two-link chain of two results of different kinds. First, $\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon)$ (Eq. (20)) approximates the exact quantity $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ (Eq. (19)) – this is a purely *empirical* result from a linear fit over $n = 50$ zeros ($R^2 \approx 0.50$, see above); it is not a proven equality or a convergence statement, but an estimate with substantial unexplained

scatter. Second, the exact quantity $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ itself – together with $d_{\text{frag}}^{\text{DWZ}}(X; W) := W_{\text{trunc}} -$ satisfies the form of equation (4); this is Proposition 2, a *formally proven, but explicitly conditional* (on the numerically observed convergence hypothesis (21)) and purely structural result (see the paragraph above: not an identity with the abstract DRCC quantities). Composed, this means: $\hat{W}_{\text{stab}}^{(M_1)} \xrightarrow{\text{empirical, unproven}} W_{\text{stab}}^{(M_1)}(t, \varepsilon) \xrightarrow{\text{formal, conditional, Proposition 2}} \text{form of Eq. (4)}$. The two links are of different epistemic quality, and neither transfers its proof status to the other: the estimator remains a numerical approximation regardless of the fact that the underlying exact quantity can be formally embedded in the DRCC form.

2.8.2 Analytical Origin of the Observed $W^{-3/2}$ Rate

The numerical experiments reported in Section 3 show an exponent close to $3/2$ for the displacement of reconstructed zeros on the critical line.

This exponent is not arbitrary.

It is consistent with the first omitted Euler–Maclaurin correction term in the reconstruction model M_1 .

Let

$$s = \sigma + it, \quad 0 < \sigma < 1,$$

and recall the first-order DWZ approximation,

$$\zeta_{\text{DWZ}}(s, W) = \sum_{n=1}^W n^{-s} + \frac{W^{1-s}}{s-1} - \frac{1}{2}W^{-s}. \quad (22)$$

The next Euler–Maclaurin correction term – already present in this work as the additional term of Variant C, equation (10) – is

$$\frac{s}{12}W^{-s-1}. \quad (23)$$

Accordingly, for fixed s and sufficiently large W , the function-value error has the asymptotic form

$$\zeta(s) - \zeta_{\text{DWZ}}(s, W) \stackrel{\text{heuristic}}{=} \frac{s}{12}W^{-s-1} + O_s(W^{-\sigma-3}), \quad (24)$$

where the notation $O_s(\cdot)$ indicates that the implied constant may depend on the fixed complex parameter s , and the “heuristic” label over the equality sign is not decorative: this is the standard Euler–Maclaurin asymptotic heuristic (the truncation error is of the same order as the first neglected term), not a fully rigorous remainder bound in the sense of Titchmarsh’s contour-integral treatment of the periodic Bernoulli remainder [2]; no such rigorous bound – for either the displayed leading term or the stated $O_s(W^{-\sigma-3})$ order of what follows it – is derived in this work.

Taking absolute values yields

$$|\zeta(s) - \zeta_{\text{DWZ}}(s, W)| = O_s(W^{-\sigma-1}). \quad (25)$$

On the critical line,

$$\sigma = \frac{1}{2},$$

and therefore

$$|\zeta(s) - \zeta_{\text{DWZ}}(s, W)| = O_s(W^{-3/2}). \quad (26)$$

Equations (25) and (26) follow from (24) by taking absolute values and are therefore heuristic in the same sense, not independently rigorous; Proposition 3 below does not rely on them directly but instead states an analogous uniform error bound as an explicit hypothesis of its own, motivated by, but not derived from, this heuristic argument. This explains analytically why an exponent near $3/2$ appears in the numerical experiments. We verified this explanation numerically (not merely asserted it): for the reconstruction error at zero #1 ($t = 14.1347$), $t = 100$, and $t = 1000$, the ratio of the actual measured error $|\zeta_{\text{DWZ}}(s, W) - \zeta(s)|$ to the predicted magnitude $|s|/12 \cdot W^{-3/2}$ approaches 1 as W grows large relative to t (e.g. at $W = 1000$ the ratio is approximately 1.0000 for $t = 14.1347$, 1.0002 for $t = 100$, and 1.0171 for $t = 1000$ – the remaining deviation at $t = 1000$ illustrates that the asymptotic regime is reached more slowly when t is large relative to W ; for $t = 1000$ the ratio is still 4.35 at $W = 50$, confirming that the asymptotic regime requires W large relative to t , consistent with the truncation depths actually used elsewhere in this work).

It remains necessary, however, to distinguish the function-value error

$$E_{\zeta}(s, W) = |\zeta_{\text{DWZ}}(s, W) - \zeta(s)|$$

from the zero-position error

$$e_k(W) = \left| s_k^{(M_1)}(W) - \rho_k \right|.$$

The following proposition gives the local connection between them.

Proposition 3 (Local Transfer from Function Error to Zero Error). *Let*

$$\rho$$

be a simple zero of ζ , so that

$$\zeta(\rho) = 0, \quad \zeta'(\rho) \neq 0. \quad (27)$$

Since the zeros of ζ are isolated, fix an open disk $U = B(\rho, \delta_0)$ for some $\delta_0 > 0$ (held fixed, independent of W) such that $\overline{U}$ is compact, contains no zero of ζ other than ρ , and satisfies $\overline{U} \subset \mathbb{C} \setminus \{1\}$ – automatically achievable by taking δ_0 small enough, since ρ is a nontrivial zero with $0 < \text{Re}(\rho) < 1$ strictly, hence bounded away from the pole at $s = 1$; this ensures ζ and $\zeta_{\text{DWZ}}(\cdot, W)$ (which shares the same pole at $s = 1$ through its $W^{1-s}/(s-1)$ term) are both holomorphic on a neighborhood of $\overline{U}$, as required for Rouché’s theorem below – such a δ_0 exists precisely because ρ is isolated among the zeros of ζ ; this isolation hypothesis is essential and is stated explicitly here because the argument below would not otherwise exclude ρ_W drifting toward a different zero of ζ inside U . Taking U to be a disk, rather than a general open set, ensures its boundary ∂U is a single circle, on which Rouché’s theorem below applies directly, without recourse to an auxiliary sub-disk.

Assume that the approximation error is uniformly bounded on the closure $\overline{U}$ (not merely on the open disk U itself, so that the bound below is available on the boundary circle ∂U for the Rouché step) by

$$\sup_{s \in \overline{U}} |\zeta_{\text{DWZ}}(s, W) - \zeta(s)| \leq K_{\rho} W^{-\sigma_U - 1}, \quad \sigma_U := \inf_{s \in \overline{U}} \text{Re}(s), \quad (28)$$

for a constant $K_{\rho} > 0$ and all sufficiently large W ; since $\text{Re}(s)$ is not constant on $\overline{U}$, the exponent uses the worst-case (infimum) real part over $\overline{U}$ (equal to the infimum over the open disk U itself, since Re is continuous and U is an open disk with closure $\overline{U}$), not $\text{Re}(\rho)$ itself. Assume in addition $\sigma_U > -1$, so that $W^{-\sigma_U - 1} \rightarrow 0$ as $W \rightarrow \infty$ – automatically satisfied whenever ρ is a nontrivial zero of ζ (so $0 < \text{Re}(\rho) < 1$) and δ_0 is small enough, the only case this manuscript applies the proposition to; the hypothesis is stated explicitly here because the proposition itself is not restricted to nontrivial zeros. Unlike an earlier version of this proposition, the existence of an approximation zero $\rho_W \in U$ is not assumed here:

it is derived below, via Rouché's theorem, from the uniform error bound and the isolation hypothesis on U alone.

Then, for all sufficiently large W , $\zeta_{\text{DWZ}}(\cdot, W)$ has exactly one zero ρ_W in U (existence and uniqueness), $\rho_W \rightarrow \rho$ as $W \rightarrow \infty$, and moreover

$$|\rho_W - \rho| = O_\rho(W^{-\sigma_U-1}). \quad (29)$$

In particular, for ρ on the critical line and a fixed open disk $U = B(\rho, \delta)$ (any $\delta > 0$ small enough that $\overline{U}$ contains no other zero of ζ), one has $\sigma_U = \frac{1}{2} - \delta$ – not $\frac{1}{2}$ exactly: an open neighborhood of a point on the line $\text{Re}(s) = \frac{1}{2}$ necessarily contains points with $\text{Re}(s) < \frac{1}{2}$, so σ_U can be made arbitrarily close to, but not equal to, $\frac{1}{2}$ by shrinking δ . This gives

$$|\rho_W - \rho| = O_\rho(W^{-3/2+\delta}) \quad (30)$$

for every fixed $\delta > 0$ small enough for $U = B(\rho, \delta)$ to satisfy the isolation hypothesis above.

Recovering the sharp exponent $3/2$ exactly (rather than $3/2 - \delta$ for every fixed $\delta > 0$) is not established by this proposition; it would require a shrinking-neighborhood argument (e.g. $U = U_W$ with $\delta = \delta_W \rightarrow 0$ at a controlled rate) or a Rouché-type argument, neither of which is carried out in this work. The numerically observed exponent close to $3/2$ (Section 3) is consistent with, but not rigorously pinned to, this bound in the $\delta \rightarrow 0$ limit.

Proof. Step 0: existence and uniqueness of ρ_W , via Rouché. Since $\overline{U}$ contains no zero of ζ other than ρ , and ρ is not on the boundary circle ∂U , the boundary is a compact set on which the continuous function ζ is nowhere zero, so

$$m_U := \min_{s \in \partial U} |\zeta(s)| > 0.$$

By equation (28) (which holds on $\overline{U}$, and hence on $\partial U \subset \overline{U}$), for all sufficiently large W ,

$$\sup_{s \in \partial U} |\zeta_{\text{DWZ}}(s, W) - \zeta(s)| \leq K_\rho W^{-\sigma_U-1} \rightarrow 0 \quad (W \rightarrow \infty), \quad (31)$$

so there is W_0 such that, for all $W \geq W_0$,

$$\sup_{s \in \partial U} |\zeta_{\text{DWZ}}(s, W) - \zeta(s)| < m_U \leq \min_{s \in \partial U} |\zeta(s)|.$$

By Rouché's theorem, $\zeta_{\text{DWZ}}(\cdot, W)$ and ζ then have the same number of zeros, counted with multiplicity, inside U , for every $W \geq W_0$. Since ζ has exactly one zero in U (namely ρ , which is simple), it follows that $\zeta_{\text{DWZ}}(\cdot, W)$ has exactly one zero ρ_W in U , for every $W \geq W_0$; this is the required existence and uniqueness statement, obtained directly on U itself rather than inferred from a smaller sub-disk.

Because $\zeta_{\text{DWZ}}(\rho_W, W) = 0$,

$$\zeta(\rho_W) = \zeta(\rho_W) - \zeta_{\text{DWZ}}(\rho_W, W), \quad \text{so} \quad |\zeta(\rho_W)| \leq K_\rho W^{-\sigma_U-1} \rightarrow 0. \quad (32)$$

Step 1: $\rho_W \rightarrow \rho$. The Rouché argument of Step 0 applies verbatim with any disk $D(\rho, \eta)$, $0 < \eta < \delta_0$, in place of U : $D(\rho, \eta) \subset U$ inherits the isolation property from U (its closure contains no zero of ζ other than ρ either, since $\overline{D(\rho, \eta)} \subseteq \overline{U}$), so the same argument gives a threshold $W_0(\eta)$ such that, for all $W \geq W_0(\eta)$, $\zeta_{\text{DWZ}}(\cdot, W)$ has exactly one zero in $D(\rho, \eta)$. Since $\zeta_{\text{DWZ}}(\cdot, W)$ has, by Step 0, exactly one zero ρ_W in the (possibly larger) set $U \supseteq D(\rho, \eta)$, that zero and the one located in $D(\rho, \eta)$ must coincide, so $\rho_W \in D(\rho, \eta)$ for all $W \geq W_0(\eta)$. Since $\eta > 0$ was arbitrary, $\rho_W \rightarrow \rho$ as $W \rightarrow \infty$.

Step 2: rate. Since ρ is a simple zero, the local Taylor expansion gives

$$\zeta(\rho_W) = \zeta'(\rho)(\rho_W - \rho) + O(|\rho_W - \rho|^2). \quad (33)$$

By Step 1, $|\rho_W - \rho| \rightarrow 0$, so for W large enough the quadratic remainder in equation (33) is dominated by the linear term; this domination is now justified, rather than assumed, because Step 1 already established $\rho_W \rightarrow \rho$ independently. Since $\zeta'(\rho) \neq 0$, combining equations (32) and (33) gives

$$|\rho_W - \rho| \leq \frac{2K_\rho}{|\zeta'(\rho)|} W^{-\sigma_U - 1} \quad (34)$$

for W sufficiently large. This proves the stated order. $\square$

Remark 3 (Status of the Exponent). *The proposition explains the observed exponent under explicit local assumptions: (1) the reference zero is simple; (2) U is a fixed neighborhood of ρ containing no other zero of ζ ; (3) the Euler–Maclaurin remainder estimate holds uniformly in U . Under these assumptions, existence and uniqueness of an approximation zero $\rho_W \in U$ for all sufficiently large W is now proved via Rouché’s theorem (Step 0 of the proof above), not assumed as a hypothesis, and $\rho_W \rightarrow \rho$ is likewise proved (Step 1).*

On the critical line, this proposition rigorously delivers the exponent $3/2 - \delta$ for every fixed, sufficiently small $\delta > 0$, not the sharp exponent $3/2$ itself – no open neighborhood of a point on $\text{Re}(s) = \frac{1}{2}$ can be confined to $\{\text{Re}(s) \geq \frac{1}{2}\}$, so $\sigma_U = \frac{1}{2}$ exactly is not attainable by a fixed U . The numerical fits in Section 3 support the exponent being close to $3/2$ in practice, but they do not replace these assumptions by a global theorem for all nontrivial zeros, nor do they establish the sharp exponent analytically. The exponent $3/2$ is therefore analytically motivated and locally conditional, approached but not exactly attained by the bound proved here, rather than inferred solely from regression.

Remark 4 (Analytic Reading of the Empirical Coefficient $C(t)$). *Equation (34) gives, at $\sigma_U = \frac{1}{2} - \delta$, the bound $|\rho_W - \rho| \leq (2K_\rho/|\zeta'(\rho)|)W^{-3/2+\delta}$. Comparing this to the empirical fit $e_k(W) \approx C(t_k)W^{-p}$ (Section 3) suggests a local dependence of the proportionality constant, not an analytic identification of it – the proved exponent is $3/2 - \delta$ while the fitted exponent is $p \approx 3/2$, and these two exponents are not shown to coincide, so the two coefficients need not denote the same quantity in a strict sense:*

$$C(\gamma) \propto \frac{K_\rho}{|\zeta'(\rho)|}, \quad \rho = \frac{1}{2} + i\gamma, \quad (35)$$

stated here as a qualitative scaling relation, up to the fixed factor of 2, the unresolved constant of proportionality, and the δ -gap between the proved and the sharp exponent. This is offered as a qualitative reading of what $C(t)$ may represent, not as a derivation of its value: K_ρ is itself an unspecified bound on the local Euler–Maclaurin remainder (equation (28)), not computed in closed form here, so equation (35) is not used to predict $C(t)$ numerically anywhere in this manuscript. It does, however, suggest a concrete, testable explanation for why $C(t)$ varies with t beyond the fitted linear trend of Section 3 ($R^2 \approx 0.50$): $|\zeta'(\rho)|$ is known to fluctuate substantially from one nontrivial zero to the next, so part of the unexplained scatter in $C(t)$ may reflect this local conditioning factor rather than a missing functional form in t alone. Testing this explicitly – for instance by regressing the fitted $C(t_k)$ against an independent estimate of $|\zeta'(\rho_k)|$ – is not carried out in this manuscript and is left as a specific, checkable extension of Section 3’s open scatter question.

Corollary 2 (Width Sufficient for a Target Local Tolerance). *Fix $\rho = \frac{1}{2} + i\gamma$, $\delta \in (0, \frac{1}{2})$, and a target tolerance $\varepsilon > 0$. Under the hypotheses of Proposition 3, bound (34) holds only for W already at or above the unspecified, ρ -dependent threshold from that proposition’s proof, denoted $W_{\rho, \delta}^{(0)}$ below, where the quadratic Taylor remainder is dominated by the linear term; the derivation below is therefore conditional on W already lying in that regime, not a guarantee that the resulting bound is attained starting from an arbitrary, possibly small, W . Requiring the proven bound (34) to satisfy $|\rho_W - \rho| \leq \varepsilon$ gives, by direct inversion,*

$$W \geq \left(\frac{2K_\rho}{\varepsilon |\zeta'(\rho)|} \right)^{\frac{1}{3/2-\delta}} = \left(\frac{2K_\rho}{\varepsilon |\zeta'(\rho)|} \right)^{\frac{2}{3-2\delta}},$$

which alone does not guarantee W already lies in the regime where (34) holds; combined with $W_{\rho,\delta}^{(0)}$, the fully sufficient condition is instead $W \geq \max\{W_{\rho,\delta}^{(0)}, (2K_\rho/(\varepsilon|\zeta'(\rho)|))^{2/(3-2\delta)}\} \Rightarrow |\rho_W - \rho| \leq \varepsilon$. And, as a formal limiting heuristic only – not a rigorous limit, since K_ρ is not shown to be controlled uniformly as $\delta \downarrow 0$ (see the discussion following this corollary) –

$$W \gtrsim \left(\frac{2K_\rho}{\varepsilon|\zeta'(\rho)|} \right)^{2/3} \quad (\delta \downarrow 0, \text{ heuristic}). \quad (36)$$

This is a direct algebraic consequence of the already-proven local bound, not an independent claim – it merely restates equation (34) as a width-sufficiency condition. Two things should not be conflated with it. First, equation (36) concerns the local, analytically bounded quantity $K_\rho/|\zeta'(\rho)|$ of Proposition 3 – an unspecified, un-evaluated bound on the local Euler–Maclaurin remainder – not the independently fitted global coefficient $C(t)$ of equation (20); Remark 4 already cautions that these two constants are, at best, related by a qualitative proportionality, not an identity. Second, the exponent $2/3$ in equation (36) is reached only in the limit $\delta \downarrow 0$ of the rigorously proven exponent $3/2 - \delta$; for any fixed $\delta > 0$ actually usable in Proposition 3, the sufficient width exponent is the slightly larger $2/(3 - 2\delta)$, not exactly $2/3$. Equation (36) therefore shows that a $2/3$ -type width law is the correct asymptotic consequence of the proven, conditional local rate bound, in the same qualitative sense that Remark 4 already uses to read $C(t)$ – it sharpens *why* an exponent near $2/3$ is analytically expected here, without upgrading the independently fitted exponent $p \approx 3/2$ of equation (20) to a proven statement, and without implying that $K_\rho/|\zeta'(\rho)|$ and $C(t)$ can be substituted for one another numerically. A third caveat applies to the limit itself: K_ρ is the uniform Euler–Maclaurin remainder bound on $U = B(\rho, \delta_0)$, and the disk radius δ_0 is exactly what sets $\sigma_U = \frac{1}{2} - \delta$; letting $\delta \downarrow 0$ therefore also shrinks U , and this manuscript does not establish how K_ρ itself behaves as U shrinks. The passage to equation (36) is accordingly stated formally/heuristically, suppressing this possible δ -dependence of the local constant, not as a rigorous limit with K_ρ controlled uniformly in δ .

2.8.3 Local Calibration of $C(t)$

The global fit of Section 3 reports a single linear trend $C(t) = (0.0272 \pm 0.0040)t + (1.043 \pm 0.375)$ across all 50 fitted zeros, which explains only about half of the observed scatter ($R^2 \approx 0.50$) and performs markedly worse out of sample ($R_{\text{os}}^2 \approx 0.08$). This global fit alone gives a practitioner no way to calibrate an estimate for a new, uncharacterized target t_* beyond the coarse linear trend. A locally weighted alternative is declared here to make this actionable; only the simplest instance of it (the unweighted mean, illustrated in a single worked example below) is numerically evaluated in this manuscript, not the general kernel-weighted family or a systematic bandwidth study.

For a target t_* , let $t_{j_1}, \dots, t_{j_m}$ be m reference points from Table 14 nearest to t_* , with fitted coefficients $C(t_{j_1}), \dots, C(t_{j_m})$. Define the locally weighted estimator

$$\hat{C}_{\text{loc}}(t_*) = \frac{\sum_{k=1}^m K_h(t_* - t_{j_k}) C(t_{j_k})}{\sum_{k=1}^m K_h(t_* - t_{j_k})}, \quad (37)$$

for a declared kernel K_h of bandwidth h , the simplest admissible choice being the unweighted mean of the m nearest known values (the uniform kernel $K_h \equiv 1$ in equation (37)). The median of the same m values,

$$\hat{C}_{\text{loc}}^{\text{med}}(t_*) := \text{median}\{C(t_{j_1}), \dots, C(t_{j_m})\},$$

is reported alongside as a robustness comparison below; it is a separate, non-kernel estimator, not a special case of the kernel-weighted average (37) for any choice of K_h . Propagate the chosen estimate into the width estimate via equation (20) in place of the global $C(t)$:

$$\widehat{W}_{\text{stab,loc}}^{(M_1)}(t_*, \varepsilon) = \left(\frac{\hat{C}_{\text{loc}}(t_*)}{\varepsilon} \right)^{2/3} \quad (38)$$

– a real-valued proxy, deliberately not integer-rounded here, since equation (38) is used below only as an intermediate quantity for the interval (39); if used directly as a required integer truncation width, it should be rounded up, $\lceil \widehat{W}_{\text{stab,loc}}^{(M_1)}(t_*, \varepsilon) \rceil$, in the same convention as equation (20). Because $\widehat{C}_{\text{loc}}(t_*)$ is itself uncertain, a local uncertainty band $C_L(t_*) \leq \widehat{C}_{\text{loc}}(t_*) \leq C_U(t_*)$ (for instance from the spread of the m neighboring values, or a local bootstrap analogous to Section 3) should be propagated, monotonically in C , to a (likewise real-valued) width interval

$$\left(\frac{C_L(t_*)}{\varepsilon} \right)^{2/3} \leq \widehat{W}_{\text{stab,loc}}^{(M_1)}(t_*, \varepsilon) \leq \left(\frac{C_U(t_*)}{\varepsilon} \right)^{2/3}, \quad (39)$$

rounded to integers only at the point of practical use, analogously. This declared local-calibration procedure is expected to track genuine local variation in $C(t)$ (including, per Remark 4, variation attributable to $|\zeta'(\rho)|$) better than a single global linear fit, since it uses only nearby reference points rather than the entire tested range.

A worked example at a held-out target. To make this concrete, equations (37)–(39) are evaluated once here, at target $t_* = t_{51} = 146.0010$ (zero #51, Table 16), using $m = 5$. Because Table 14 is fitted only over #1–50, and t_* exceeds every t_k in that table, its 5 nearest reference points are simply its last 5 entries:

k	46	47	48	49	50
t_k	134.7565	138.1160	139.7362	141.1237	143.1118
$C(t_k)$	3.993	3.652	6.457	5.874	3.200

With the unweighted-mean kernel (the simplest admissible choice declared above), $\widehat{C}_{\text{loc}}(t_*) = 4.635$; the median of the same 5 values is 3.993, noticeably lower, reflecting the right-skewed spread of this particular neighborhood. A local case-resampling bootstrap on these 5 values, using the same resampling principle as the global bootstrap (Section 3: resampling with replacement, 10,000 resamples, disclosed seed, Appendix B.0) but a different statistic – the local arithmetic mean, refit on each resample, rather than an OLS regression – gives the 95% percentile interval $[C_L(t_*), C_U(t_*)] = [3.539, 5.779]$. Propagated via equation (39) at $\varepsilon = 10^{-6}$,

$$\widehat{W}_{\text{stab,loc}}^{(M_1)}(t_*, 10^{-6}) \approx 27,800, \quad [W_L, W_U] \approx [23,225, 32,205].$$

Zero #51 was excluded from the calibration set (it lies outside #1–50 entirely), so its already-tabulated deviation $e_{51}(W=10^4) = 3.75 \times 10^{-6}$ (Table 16) provides a genuine held-out check: applying the same fixed-exponent convention used in the out-of-sample check above, $C_{\text{actual}}(t_{51}) = e_{51}(10^4) \cdot (10^4)^{1.5} = 3.750$, which falls inside $[C_L, C_U]$, near its lower edge, giving the fixed-exponent proxy $W_{\text{proxy}} \approx 24,137$ – likewise inside $[W_L, W_U]$. This proxy is itself derived from a single error value at $W = 10^4$ under the assumed exponent $p = 1.5$, not an independently measured W_{stab} obtained by testing multiple widths against tolerance ε ; it is used here only as the held-out comparison point for the calibration estimate, not as a ground-truth stable width. For comparison, the global fit’s prediction at the same point is $C_{\text{pred}}(t_{51}) = 0.0272 \cdot 146.0010 + 1.043 = 5.014$, i.e. $W_{\text{global}} \approx 29,296$: for this single held-out zero, the local mean-based estimate (27,800) lies closer to the proxy value (24,137) than the global prediction (29,296) does, and the local interval brackets the proxy value while remaining narrower than the full scatter of the global fit.

This is one worked instance, not a systematic evaluation, and should not be read as demonstrating that local calibration outperforms the global fit in general: a single held-out point cannot establish an out-of-sample R^2 , and both estimates here overshoot the proxy value by a comparable relative margin. What it does establish is that the declared procedure is executable end to end on real data already in this manuscript, with a disclosed random seed and a genuine, excluded

held-out target. Repeating this construction across all of #51–100 to obtain a local analogue of $R_{\text{oos}}^2 \approx 0.08$ – the systematic comparison the qualitative claim above still requires – remains the concrete extension of Section 3 left open here.

A conservative policy for an unseen t_* . Neither the global estimator $\hat{W}_{\text{stab}}^{(M_1)}$ (equation (20)) nor the local estimator $\hat{W}_{\text{stab,loc}}^{(M_1)}$ (equation (38)) of this section should be treated as a certified stopping width for a target t_* that was not among the tested reference zeros: both are empirical point predictions with demonstrated, substantial uncertainty ($R_{\text{oos}}^2 \approx 0.08$ for the global fit; a single held-out check, not a systematic validation, for the local estimator above), not proven bounds. The policy used throughout this manuscript for any such t_* is:

1. compute a point estimate $\hat{C}(t_*)$ – global (equation (20)) or local (equation (37)) – and the resulting width estimate;
2. treat this value exclusively as a *starting width for a search*, never as a certified threshold in its own right; and
3. validate the candidate width by the means actually available for t_* , which depends on whether an independent reference zero exists for it. If t_* coincides with, or is later matched to, a known reference zero ρ_k (so that $e_k^{(M_1)}(W)$ is defined at all, per the restriction noted earlier in this section), the candidate width may be checked post hoc against the condition of equation (19),

$$e_k^{(M_1)}(W) < \varepsilon \quad \text{for every tested } W \geq W_0,$$

increasing W_0 beyond the point estimate if the check fails there. For a genuinely unseen t_* with no such reference available, this check cannot be performed at all, since $e_k^{(M_1)}(W)$ is by definition undefined at such a point; the candidate is instead assessed only through the declared blind-scan residual and Newton-refinement criteria of Section 2.9 ($|\zeta_{\text{DWZ}}|$ scan minimum, refinement residual $\varepsilon_{\text{scan}}$), not through equation (19).

This introduces no new statistical method and no new claim about the distribution of $C(t)$; it states explicitly, as a single policy, what the definitions already in this section imply separately: that $\hat{C}(t_*)$ and any width estimate derived from it are search heuristics, while the actual grounding required by Proposition 2 comes from the direct error check of equation (19) when a reference zero is available, and from the blind-scan residual criteria otherwise – never from the fitted coefficient alone. Where it applies, the grid check of equation (19) is itself a validation on the finitely many tested widths, not a proof of the universal statement “for every $W \geq W_0$ ” in equation (19) – the same limitation already noted earlier in this section, where that equation is first defined, and in the hypothesis of Proposition 2 itself.

2.8.4 A General Operationalization Workflow

The reviewer-motivated general recipe extracted from the DWZ case is the following.

1. **Abstract quantity and task.** State the abstract DRCC quantity and the exact computational task τ before choosing a numerical surrogate.
2. **Observable fragment.** Define a finite, measurable fragment $F_W(X)$ and state whether W is structural, numerical, or merely a surrogate parameter.
3. **Reconstruction model.** Declare the admissible model family $\mathcal{M}$ and the model-specific reconstruction operator R_α . Model choice must precede the evaluation of the required depth.

4. **Stability criterion.** Specify an error functional $E_\alpha(X, W)$, a tolerance ε , and a stable threshold that requires the error to remain below tolerance for every larger width, not only at the first successful test point.
5. **Validation and comparison.** Test convergence, compare against independent references and established methods, replicate timing measurements where practical, and report uncertainty, out-of-sample behavior, and statistical power when inferential claims are made.
6. **Limitations and operational record.** Separate truncation depth, structural state width, bit-space, and runtime; then record the instance, task, model, error measure, tolerance, threshold, evidence status, and unresolved assumptions in one operational record.

This six-stage workflow is the generic methodological output of the paper. The equations below formalize it. Their existence does not imply that a useful model or a favorable runtime bound exists for every target object.

The zeta-function case is one concrete instance of a more general operational pattern. The essential DRCC contribution of the present work is not the Euler–Maclaurin formula itself. Rather, it is the explicit separation of finite fragment, reconstruction model, reconstruction width, and accuracy condition for a deterministic approximation process.

Let X be a target mathematical or computational object, and let

$$F_W(X) \tag{40}$$

be a finite approximation or fragment indexed by a numerical depth parameter $W \in \mathbb{N}$. Let

$$\mathcal{M} = \{M_\alpha : \alpha \in A\} \tag{41}$$

be a declared family of admissible reconstruction models. For each model M_α , define a reconstructed approximation

$$\hat{X}_{\alpha, W} = \mathcal{R}_\alpha(F_W(X)), \tag{42}$$

where $\mathcal{R}_\alpha$ is the model-specific reconstruction operator. Let

$$\mathcal{E}_\alpha(X, W) = d_Y(\hat{X}_{\alpha, W}, X) \tag{43}$$

be the reconstruction error in a declared metric space $(\mathcal{Y}, d_Y)$.

Definition 3 (Accuracy-Relative Reconstruction Depth). *For a target tolerance $\varepsilon > 0$, the stable accuracy-relative reconstruction depth of model M_α is*

$$W_{\text{stab}}^{(\alpha)}(X, \varepsilon) = \min \{W_0 \in \mathbb{N} : \mathcal{E}_\alpha(X, W) < \varepsilon \text{ for every } W \geq W_0\}, \tag{44}$$

provided that the defining set is nonempty. This is deliberately stronger than the first hitting time $\min\{W : \mathcal{E}_\alpha(X, W) < \varepsilon\}$: the latter only guarantees accuracy at that single W , whereas equation (44) guarantees it for every larger width as well, which is the property actually needed for $W_{\text{stab}}^{(\alpha)}(X, \varepsilon)$ to serve as a usable stability threshold rather than an isolated lucky value. For the concrete zeta case in this manuscript, decreasing error with increasing width is illustrated explicitly for zero #1 in Figure 12 and is consistent with the multi-width fits underlying Table 14; Table 16 itself records only a single deviation value per zero at one tested width and so does not, by itself, exhibit a sequence over W . No finite numerical experiment proves the eventual-accuracy property for every $W \geq W_0$ and, consistent with the caveat already stated in Section 2.8.1, it is not proven here for arbitrary W .

This definition separates three quantities that must not be conflated: W_{trunc} (the numerical truncation or iteration depth used in one computation), $W_{\text{stab}}^{(\alpha)}(X, \varepsilon)$ (the smallest depth beyond which the declared accuracy is attained and maintained under model M_α), and ω_{reg} (the register width of the declared reconstruction process). In general, no equality among these quantities is implied.

Proposition 4 (Conditional Operationalization Principle – a Threshold Existence Statement from Well-Ordering, Not a General Existence Result for Reconstruction Models). *Assume that, for a fixed object X , model M_α , and tolerance $\varepsilon > 0$, there exists a finite $W_0 \in \mathbb{N}$ such that*

$$\mathcal{E}_\alpha(X, W) < \varepsilon \quad \text{for every } W \geq W_0. \quad (45)$$

Then the stable accuracy-relative reconstruction depth $W_{\text{stab}}^{(\alpha)}(X, \varepsilon)$ is well-defined and finite, and for every finite $W \geq W_{\text{stab}}^{(\alpha)}(X, \varepsilon)$,

$$0 < W_{\text{stab}}^{(\alpha)}(X, \varepsilon) \leq W < \infty \quad (46)$$

and

$$\mathcal{E}_\alpha(X, W) < \varepsilon. \quad (47)$$

Proof. By assumption, the set $\{W_0 \in \mathbb{N} : \mathcal{E}_\alpha(X, W) < \varepsilon \text{ for every } W \geq W_0\}$ is nonempty; being a nonempty subset of $\mathbb{N}$, it has a least element, so $W_{\text{stab}}^{(\alpha)}(X, \varepsilon)$ exists, is positive, and is finite. For every finite width satisfying $W \geq W_{\text{stab}}^{(\alpha)}(X, \varepsilon)$, the defining property of that least element directly gives $\mathcal{E}_\alpha(X, W) < \varepsilon$, and the inequalities $0 < W_{\text{stab}}^{(\alpha)}(X, \varepsilon) \leq W < \infty$ follow directly from the ordering of W_0 and W in $\mathbb{N}$. $\square$

Remark 5 (Meaning and Limits of the General Principle). *Proposition 4 is a general consistency statement: given the (nontrivial, but modest) assumption that accuracy is eventually attained and maintained, the well-ordering of $\mathbb{N}$ guarantees a well-defined, finite, stable threshold; the mathematical content beyond this existence argument is deliberately kept minimal. It does not prove that a useful reconstruction model exists for every high-dimensional problem, it does not show that the required depth is polynomially bounded, and it does not identify the numerical depth with the register width. Its role is only to specify exactly which additional objects must be supplied – a fragment F_W , a model $\mathcal{M}$, an error measure $\mathcal{E}_\alpha$, and an eventually-maintained accuracy condition – before an abstract reconstruction condition such as (4) becomes a measurable computational statement. Whether such objects can be supplied usefully for a given problem is a separate, problem-specific question that this proposition does not answer.*

Definition 4 (Operational DRCC Record). *An operational DRCC record for a deterministic approximation process is the tuple*

$$\mathfrak{D} = \left(X, F_W, \mathcal{M}, \mathcal{R}_\alpha, \mathcal{E}_\alpha, \varepsilon, W_{\text{stab}}^{(\alpha)}, \omega_{\text{reg}}, T_\alpha \right), \quad (48)$$

where X is the target object; F_W is the finite fragment or approximation; $\mathcal{M}$ is the admissible reconstruction-model family; $\mathcal{R}_\alpha$ is the reconstruction operator; $\mathcal{E}_\alpha$ is the declared error measure; ε is the required tolerance; $W_{\text{stab}}^{(\alpha)}$ is the required numerical reconstruction depth; ω_{reg} is the register width; and T_α is the complete runtime or operation-count model.

The DWZ method instantiates the principal, problem-specific fields of this record concretely through

$$X = \zeta(s), \quad F_W(s) = \sum_{n=1}^W n^{-s}, \quad (49)$$

$$\mathcal{R}_{M_1}(F_W) = F_W + \frac{W^{1-s}}{s-1} - \frac{1}{2}W^{-s}, \quad (50)$$

and

$$\mathcal{E}_{M_1}(s, W) = |\zeta_{\text{DWZ}}(s, W) - \zeta(s)| \quad (51)$$

or, for zero reconstruction,

$$\mathcal{E}_{M_1}^{\text{zero}}(\rho_k, W) = \left| s_k^{(M_1)}(W) - \rho_k \right|. \quad (52)$$

The remaining fields of $\mathfrak{D}$ – the tolerance ε , the stable threshold $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ (equation (19)), the register width ω_{reg} (equation (14)), and the runtime model T_{M_1} (Section 4.3) – are each instantiated concretely elsewhere in this manuscript, but are not reassembled into a single explicit nine-tuple at this location. This is the only instance of the numerical record $\mathfrak{D}$ worked out in this manuscript; the C_4 graph-coloring instance of Appendix B.2 is exact and finite and uses a simpler operational record of its own (equation (76)), not $\mathfrak{D}$, since $\mathfrak{D}$'s fields (truncation depth, runtime decomposition, accuracy tolerance) are specific to iterative numerical approximation. The same operational template could, in principle, be applied to other deterministic iterative approximation problems with an explicit fragment, reconstruction operator, error metric, and declared stopping tolerance – candidates might include truncated integral transforms, finite spectral expansions, or iterative linear-system reconstruction – but no such second instance is worked out here, and we deliberately stop short of claiming that one exists for every, or even any, such problem. For each hypothetical application, the mathematical burden remains entirely problem specific: one would have to independently establish $\mathcal{E}_\alpha(X, W) \rightarrow 0$, derive or estimate $W_{\text{stab}}^{(\alpha)}(X, \varepsilon)$, separately determine ω_{reg} , and account for the complete computational cost $T_\alpha = T_{\text{build}} + T_{\text{reconstruct}} + T_{\text{verify}} + T_{\text{output}}$ – none of which is automatic, and none of which follows from Proposition 4 alone.

Remark 6 (What Is and Is Not Transferable). *The transferable insight of the DWZ case study is not a new summation formula, and it is not a proof that DRCC generalizes to other high-dimensional problems. What is transferable, and what this manuscript actually delivers, is a worked template of the operational sequence – finite fragment, declared reconstruction model, measurable reconstruction error, accuracy-relative reconstruction depth, and separate structural-width/runtime analysis – carried out completely and honestly for one concrete, nontrivial example. Whether this sequence can be carried out as completely for a different target object X is an open question specific to that object, not something this manuscript establishes. The zeta function is the first, and so far only, detailed numerical realization developed in this manuscript; the C_4 graph-coloring instance (Appendix B.2) carries out the same six-stage sequence exactly and finitely, but is not itself a numerical approximation problem and does not use $\mathfrak{D}$.*

2.9 Verification Strategies

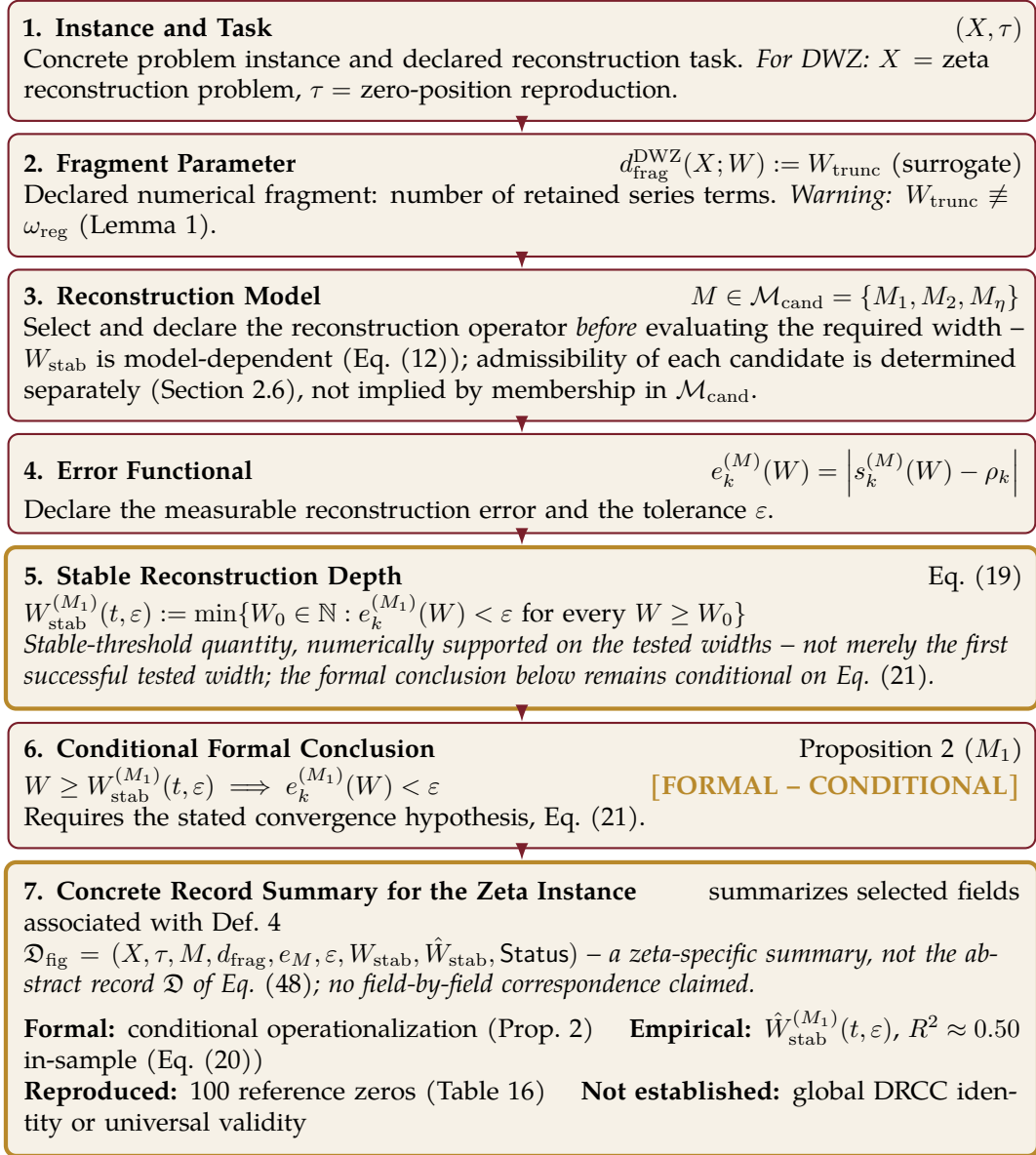
Declared Blind-Scan Algorithm

The blind zero-search procedure is defined independently of any reference-zero starting value. Let

$$I_{\text{scan}} = [t_{\min}, t_{\max}] \subset \mathbb{R}_{>0}, \quad s(t) = \frac{1}{2} + it,$$

and let $h_{\text{scan}} > 0$ be a fixed coarse-grid step. The scan grid is

$$\mathcal{G} = \{t_j = t_{\min} + jh_{\text{scan}} : j = 0, \dots, J\}, \quad J = \left\lfloor \frac{t_{\max} - t_{\min}}{h_{\text{scan}}} \right\rfloor. \quad (53)$$



What the record establishes	What the record does not establish
A documented and partially reproducible/checkable operationalization (some historical environment and blind-scan parameters not recorded, Appendix B.0) for one declared model, on one declared instance and task.	$W_{\text{trunc}} = \omega_{\text{reg}}$; global convergence for all t ; a faster zeta algorithm; or a result on the Riemann Hypothesis.

Figure 7: OPA-DRCC-05: Operational DRCC Record from an abstract reconstruction condition to a checkable numerical claim. The diagram separates the declared instance and task, fragment parameter, reconstruction model, error functional, stable reconstruction depth, conditional formal conclusion, and evidence status. The record is model-relative and does not identify numerical truncation depth with the register width.

At every grid point, the score

$$g_j = \left| \zeta_{\text{DWZ}} \left(\frac{1}{2} + it_j, W_{\text{scan}} \right) \right| \quad (54)$$

is evaluated. For each interior grid point $1 \leq j \leq J-1$ (the rule is undefined at $j = 0, J$, where g_{-1} or g_{J+1} would be required, so the two endpoints are never treated as candidates), a coarse candidate is retained exactly when

$$g_j < g_{j-1}, \quad g_j \leq g_{j+1}, \quad g_j < \theta_{\text{scan}}. \quad (55)$$

For every retained candidate, define

$$I_j = [t_j - h_{\text{scan}}, t_j + h_{\text{scan}}] \cap I_{\text{scan}}. \quad (56)$$

The refinement minimizes

$$t \mapsto \left| \zeta_{\text{DWZ}} \left(\frac{1}{2} + it, W_{\text{ref}} \right) \right| \quad (57)$$

over I_j using the declared one-dimensional refinement method, producing a scan-minimum ordinate $\hat{t}_j$; this is a starting candidate only, not yet a located zero of ζ_{DWZ} in $\mathbb{C}$. It stops when the interval width or successive-ordinate change is at most δ_t , when the residual is at most $\varepsilon_{\text{scan}}$, or when N_{ref} iterations have been reached.

Complex Newton refinement. From the starting point $s_0 = \frac{1}{2} + i\hat{t}$ (the scan minimum $\hat{t}$ as imaginary part, $\frac{1}{2}$ as the starting real part – not a constraint on the result), a complex Newton iteration of the same form as equation (63) (Section 2.10), with $f(s) := \zeta_{\text{DWZ}}(s, W_{\text{ref}})$, is applied in $\mathbb{C}$, *not* confined to the line $\text{Re}(s) = \frac{1}{2}$, stopping under three analogous but distinct criteria (a complex step-size tolerance $\delta_s > 0$, stopping when $|s_{n+1} - s_n| \leq \delta_s$ – not the one-dimensional ordinate tolerance δ_t of the scan step above, since a complex Newton step is a change in $\mathbb{C}$, not merely in an ordinate; residual $\varepsilon_{\text{scan}}$; or N_{ref} iterations), and producing the output $\hat{s} \in \mathbb{C}$ of the blind-scan protocol – consistent with the worked example and Table 16, where the resulting $\hat{s}$ generally satisfies $\text{Re}(\hat{s}) \neq \frac{1}{2}$ exactly. Two refined candidates are duplicates when

$$|\hat{s}_a - \hat{s}_b| < \delta_{\text{dup}}, \quad (58)$$

and only the smallest-residual member of a duplicate cluster is retained. Reference values are not used in scanning, candidate selection, refinement, stopping, or duplicate removal. After the candidate set has been frozen, a candidate $\hat{s}$ is assigned to the nearest reference ordinate, using its imaginary part $\hat{t} = \text{Im}(\hat{s})$,

$$k^* = \arg \min_k |\hat{t} - t_k^{\text{ref}}|, \quad (59)$$

equivalent here to assignment by the full complex distance $|\hat{s} - \rho_k^{\text{ref}}|$ used elsewhere in this manuscript (Section 2.10): since every reference zero $\rho_k^{\text{ref}} = \frac{1}{2} + it_k^{\text{ref}}$ shares the same real part $\frac{1}{2}$, the term $(\text{Re}(\hat{s}) - \frac{1}{2})^2$ is constant across k and does not affect the $\arg \min$, regardless of $\text{Re}(\hat{s})$ itself; the two ranking rules therefore agree exactly for this reference set, though not in general. The candidate is a validated hit only if the full complex distance to the assigned reference zero $\rho_{k^*}^{\text{ref}} = \frac{1}{2} + it_{k^*}^{\text{ref}}$ satisfies

$$|\hat{s} - \rho_{k^*}^{\text{ref}}| \leq \delta_{\text{val}}. \quad (60)$$

The formal protocol template is

$$\mathcal{B} = (t_{\min}, t_{\max}, h_{\text{scan}}, W_{\text{scan}}, \theta_{\text{scan}}, W_{\text{ref}}, \delta_t, \delta_s, \varepsilon_{\text{scan}}, N_{\text{ref}}, \delta_{\text{dup}}, \delta_{\text{val}}). \quad (61)$$

As the following table shows, the historical experiment did not retain a value for every field of this template.

Table 9: Status of the retained blind zero-search protocol parameters. Values not present in the retained manuscript record are explicitly marked as not recoverable rather than reconstructed or guessed.

Parameter	Retained record
General scan interval I_{scan}	not recoverable from retained logs
Coarse step h_{scan}	not recoverable for the complete experiment
Coarse width W_{scan}	2000 for the plotted zero-#9 scan
Candidate threshold θ_{scan}	not recoverable from retained logs
Refinement method	Newton refinement after scan selection
Refinement width W_{ref}	not recoverable as a frozen global parameter
Position tolerance δ_t (1D scan/refinement)	not recoverable from retained logs
Complex step tolerance δ_s (complex Newton)	not recoverable from retained logs
Residual tolerance $\varepsilon_{\text{scan}}$	not recoverable from retained logs
Maximum iterations N_{ref}	not recoverable from retained logs
Duplicate radius δ_{dup}	not recoverable from retained logs
Validation radius δ_{val}	not recoverable as a predeclared value
Reference source	<code>mpmath.zetazero</code> , after candidate freezing

The missing historical blind-scan parameters and any Arb runtime values are deliberately not invented here. This reflects the actual limitation of the underlying experiment, which recovered only two blind hits (#9, #11) and did not retain a fully frozen candidate set; reporting placeholder or reconstructed numbers in their place would misrepresent that limitation rather than disclose it.

The plotted record for zero #9 additionally preserves the grid $t \in [47.85, 48.15]$ with step 0.01 and $W = 2000$. This local record does not establish that the same values governed all blind scans. Consequently, Q3 is specified algorithmically here, but the historical experiment is not retroactively presented as a fully parameter-frozen benchmark.

We use four methodologically distinct tests:

1. **Known start:** Newton refinement of equation (8) starting from the known reference position (48 of the first 50 zeros #1–#50 in Table 16, all except #9 and #11; the further 50 zeros #51–#100 in the same table were all treated with a known start, see Section 3).
2. **Blind search:** scanning $|\zeta_{\text{DWZ}}|$ over a t -range without reference knowledge, Newton refinement of the best candidate (zeros #9, #11).
3. **Hybrid search:** analytic estimate (Riemann-von Mangoldt formula, Eq. 62) as the scan center, with a reduced window (zero #12, hardening #13–#41).
4. **Gap verification:** fine-grained scan between two known zeros to search for additional numerical candidates on the chosen grid (not a rigorous completeness proof, see Section 4).

Conceptually, three objects must be distinguished: the *scan minimum* (starting candidate from the one-dimensional scan along $\text{Re}(s) = \frac{1}{2}$), the *zero of the approximation* $s_k^{(M)}(W)$ (result of the complex Newton iteration in the two-dimensional search space), and the *reference zero* ρ_k of ζ itself. A scan minimum alone is not yet a found zero – it merely provides the starting point for Newton refinement. For the hybrid search, the Riemann-von Mangoldt counting function is used:

$$N(T) \approx \frac{T}{2\pi} \ln \frac{T}{2\pi} - \frac{T}{2\pi} + \frac{7}{8}. \quad (62)$$

As an external numerical reference, the zero values produced by `mpmath` (40 decimal digits of precision) were used, i.e. generated independently of the DWZ computation – not necessarily

independent of established zeta algorithms in the sense of a second, independently implemented piece of software or a certified zero count. For the first zeros, these values agree with published tables such as [4, 5].

2.10 Detailed Worked Example: Determination of Zero #1

To document the method to the extent of the retained record, we lay out here every single computational step for determining a specific zero – with all intermediate values. The first nontrivial zero $\rho_1 = 0.5 + 14.1347251417 \dots i$ serves as the example. All numerical values in this section were generated with `mpmath` at 30 decimal digits of precision and are expected to be reproducible with a standard `mpmath` installation at the same precision, though the exact `mpmath`/Python/OS version used for the historical run is not separately recorded (Appendix B.0).

Figure 8 shows the complete chain of steps at a glance, before the individual steps are documented in detail with numerical values below.

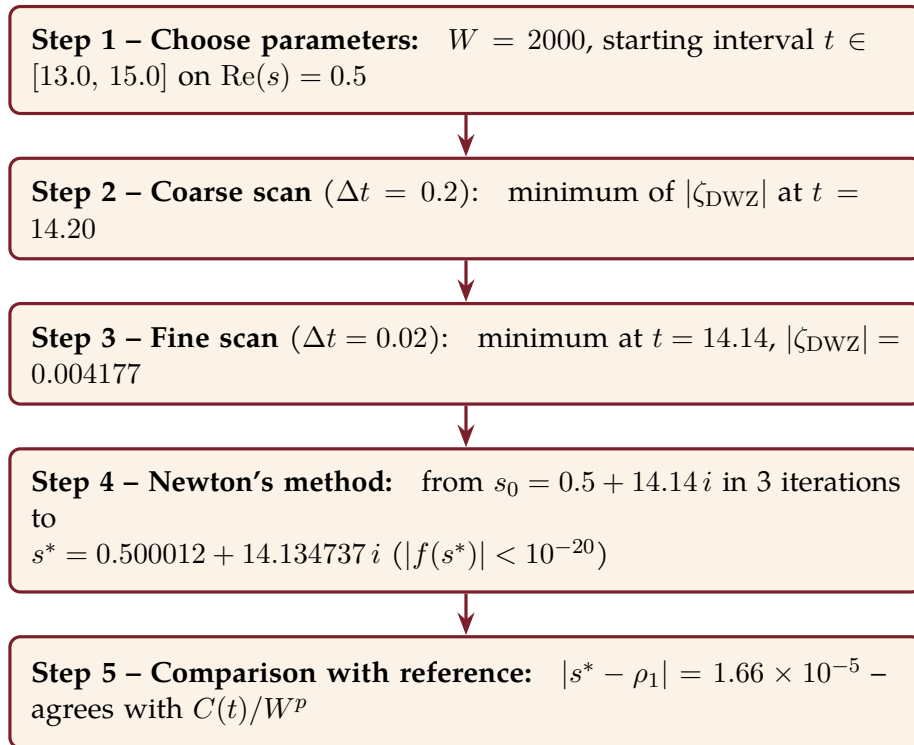


Figure 8: The chain of steps in DWZ zero determination, illustrated for ρ_1 : each arrow represents a fully documented computational step (details and all intermediate values below).

Step 1 – Choice of parameters. We choose the width $W = 2000$ and use $\zeta_{\text{DWZ}}(s, W)$ according to equation (8). A literature-informed coarse interval $t \in [13.0, 15.0]$ on the critical line $\text{Re}(s) = 0.5$ serves as the starting range of the search – chosen from the rough, published location of the first zero, not from its precise position, and not to be confused with the strictly blind scan protocol of Section 2.9, which uses no such prior information.

Step 2 – Coarse scan. We evaluate $|\zeta_{\text{DWZ}}(0.5 + it, 2000)|$ in steps of $\Delta t = 0.2$ (Table 10). The minimum lies at $t = 14.20$.

Step 3 – Fine scan. To refine the coarse scan, we sample the interval $t \in [14.00, 14.30]$ in steps of $\Delta t = 0.02$ (Table 11). The minimum now lies at $t = 14.14$, with $|\zeta_{\text{DWZ}}| = 0.004177$ – already

Table 10: Coarse scan, step size 0.2.

t	$ \zeta_{\text{DWZ}}(0.5 + it, 2000) $
13.00	0.791151
13.20	0.669990
13.40	0.540293
13.60	0.402573
13.80	0.257453
14.00	0.105639
14.20	0.052042
14.40	0.214601
14.60	0.380919
14.80	0.549793
15.00	0.719932

an order of magnitude smaller than in the coarse scan.

Table 11: Fine scan, step size 0.02.

t	$ \zeta_{\text{DWZ}}(0.5 + it, 2000) $
14.000	0.105639
14.020	0.090120
14.040	0.074543
14.060	0.058909
14.080	0.043219
14.100	0.027474
14.120	0.011675
14.140	0.004177
14.160	0.020081
14.180	0.036036
14.200	0.052042
14.220	0.068096
14.240	0.084198
14.260	0.100346
14.280	0.116540
14.300	0.132779

Step 4 – Newton’s method. Starting from the best candidate of the fine scan, $s_0 = 0.5 + 14.14i$, we refine the zero using Newton’s method for complex functions:

$$s_{n+1} = s_n - \frac{f(s_n)}{f'(s_n)}, \quad f(s) := \zeta_{\text{DWZ}}(s, 2000), \quad (63)$$

where $f'(s_n)$ is approximated via a central difference with step size $h = 10^{-6}$: $f'(s_n) \approx (f(s_n+h) - f(s_n-h))/(2h)$. Table 12 shows the complete iteration sequence.

Already after two iterations, $|f(s_n)|$ has fallen to 10^{-11} ; from $n = 3$ onward, s_n changes only in decimal places far beyond any practical relevance – over these reported iterations the residuals shrink at the rate expected of quadratic local convergence (each step roughly doubling the number of correct digits), consistent with, but not a formal proof of, quadratic convergence of this particular Newton iteration with its finite-difference-approximated derivative.

Table 12: Newton iteration starting from $s_0 = 0.5 + 14.14 i$.

n	s_n	$ f(s_n) $
0	$0.500000000 + 14.140000000 i$	4.18×10^{-3}
1	$0.500023320 + 14.134738876 i$	9.09×10^{-6}
2	$0.500012099 + 14.134736526 i$	4.32×10^{-11}
3	$0.500012099 + 14.134736526 i$	9.77×10^{-22}

Step 5 – Comparison with the reference. The point found,

$$s^* = 0.500012099 + 14.134736526 i,$$

is compared with the independently computed reference $\rho_1 = 0.5 + 14.134725142 i$ (mp.zetazero(1), 30 decimal digits):

$$|s^* - \rho_1| = 1.66 \times 10^{-5}. \quad (64)$$

This residual deviation is *not* an error of Newton’s method (which found the zero of the width-2000 approximation to the reported numerical residual, $|f(s_3)| = 9.77 \times 10^{-22}$ in Table 12, not exactly in the mathematical sense), but the observed finite-truncation deviation of $\zeta_{\text{DWZ}}(s, 2000)$ from $\zeta(s)$ at finite W (not a discretization error in the classical continuous-to-discrete sense; the underlying series is already discrete, and the deviation instead comes from truncating it at finite W) – measured here, not predicted in advance from a proven error interval for $W = 2000$. It agrees post hoc with the empirical power law derived in Section 2: with $C_1^{\text{local}} \approx 1.499$ and $p \approx 1.501$ – a local fit computed exclusively from the W -error values of this worked example (Figure 12), to be distinguished from (a) the systematic table entry $C(t_1) = 1.486$ from Table 14, based on the same W -grid as all other 49 rows of that table, and (b) the value of the global linear regression $C_{\text{reg}}(14.13) = 0.0272 \times 14.13 + 1.043 \approx 1.427$ from equation (20). All three values lie in the range 1.43–1.50 and arise from different fitting procedures (local, on a few W -points; systematic, per zero; global, over all 50 zeros); the differences between them are expected fit sensitivity, not a contradiction, but should not be left unremarked side by side – hence this clarification. With C_1^{local} and p , the power law (18) predicts an error of $C/W^p = 1.499/2000^{1.501} \approx 1.68 \times 10^{-5}$ – practically identical to the actually observed value 1.66×10^{-5} . This agreement is, however, an in-sample consistency check, not an independent validation: C_1^{local} and p were fitted directly from the error curve of this same zero #1 (Figure 12), so the prediction merely checks whether the fit curve reproduces its own data. An actual test independent of this local, single-zero calibration is provided by the other zeros in Table 16, none of which entered the fit of C_1^{local} and p above (this is a separate independence claim from the held-out test of the global $C(t)$ regression of equation (20), whose own calibration set is zeros #1–50 and whose independent test is correspondingly restricted to zeros #51–100, discussed in Section 3). Increasing W to 10^4 (as used in Table 16 for all 100 zeros) correspondingly reduces the deviation to $\approx 1.5 \times 10^{-6}$ (for the early, small- t entries of that table; for larger t it rises slightly to $\approx 1.3 \times 10^{-5}$, see Table 16).

Reproducibility. The complete computational procedure thus consists of exactly five steps: (1) fix the width W and starting interval, (2) coarse scan of $|\zeta_{\text{DWZ}}|$ on the critical line, (3) fine scan around the coarse-scan minimum, (4) Newton refinement with numerical derivative, (5) comparison with an independent reference or with the empirical power-law error estimate $C(t)/W^p$ (the fitted estimate, not the proven analytical bound of Section 2.8.2). The numerical procedure can be reconstructed from the information given here, formula (8), and a standard mpmath installation; exact bit-for-bit reproduction of the specific historical run reported above is not claimed (see the code and data availability note at the end of this manuscript).

3 Results

3.1 Comparison of the Repair Variants

This subsection compares two of the three reconstruction variants introduced in Section 2 – Variant A (the first-order Euler-Maclaurin correction, model M_1) and Variant B (the Dirichlet-eta-based alternative, which is not an Euler-Maclaurin correction) – against each other at the reference point $s = 0.5 + 14.13i$. Table 13 and Figure 9 report the resulting model-dependent reconstruction error $E_M(s, W) := |\zeta_M(s, W) - \zeta(s)|$ ($\zeta_M = \zeta_{\text{DWZ}}$ for Variant A, $\zeta_M = \zeta_W^\eta$ for Variant B – not the naive $\zeta_W(s, W)$ of equation (5)) across the same range of truncation widths W for both variants. For Variant A, the error decreases by several orders of magnitude more rapidly with increasing W than for Variant B over this range (a convergence-rate comparison, not a runtime comparison); this provides the empirical basis for preferring Variant A (M_1) over the eta-based Variant B in this comparison. The use of M_1 as this manuscript’s primary baseline throughout the remainder of the Results section remains the methodological design choice already stated in Section 1.2 (the minimal, directly traceable route for translating the DRCC condition into a checkable number), which this empirical comparison additionally supports but does not by itself establish. Figure 10 additionally shows that both the function-value error and the zero-position error for zero #1 decrease as W grows for Variant A, the numerical convergence behavior on which the operationalization of Section 2.8.1 relies.

Table 13: Model-dependent reconstruction error $E_M(s, W) = |\zeta_M(s, W) - \zeta(s)|$ for $s = 0.5 + 14.13i$: Variant A vs. Variant B.

W	Variant A (DWZ)	Variant B (Eta)
10	3.85×10^{-2}	8.24×10^{-2}
50	3.34×10^{-3}	2.99×10^{-2}
100	1.18×10^{-3}	2.11×10^{-2}
500	1.05×10^{-4}	9.42×10^{-3}
1000	3.73×10^{-5}	6.66×10^{-3}
5000	3.33×10^{-6}	2.98×10^{-3}

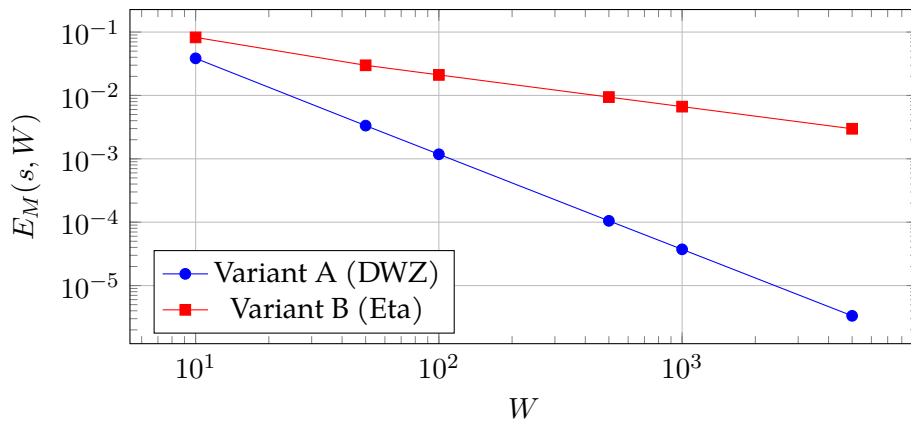


Figure 9: Variant A (DWZ) converges by orders of magnitude faster than Variant B, shown here at $s = 0.5 + 14.13i$ (empirically fitted exponent near $-3/2$ on the critical line for Variant A; Variant B scales as $\mathcal{O}(W^{-\text{Re}(s)})$). Elsewhere in the critical strip, the exponent for Variant A depends on $\text{Re}(s)$.

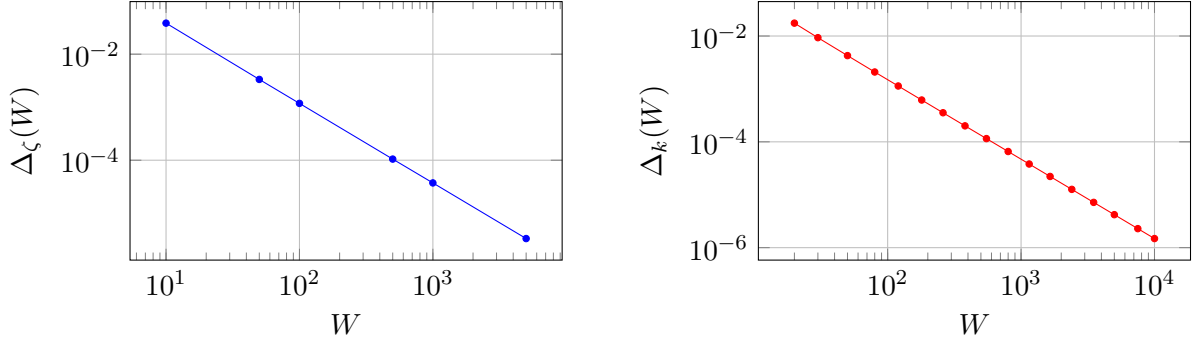


Figure 10: Left: convergence of the zeta error $\Delta_\zeta(W) = |\zeta_{\text{DWZ}}(s, W) - \zeta(s)|$ for $s = 0.5 + 14.13i$. Right: convergence of the zero stability $\Delta_k(W) = |s_k(W) - \rho|$ for zero #1 – both decreasing as W grows.

3.2 $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$

Table 14: Fitted $C(t)$ and exponent p per zero, over the first 50 zeros #1–#50 from Table 16 ($n = 50$, previously $n = 8$; Table 16 now spans 100 zeros, but this fit was not re-extended to the additional 50); the functional form of $C(t)$ is still considered preliminary, not conclusively proven (see Conclusion).

k	t	$C(t)$	p
1	14.1347	1.486	1.5000
2	21.0220	1.544	1.5002
3	25.0109	1.523	1.5002
4	30.4249	1.944	1.5000
5	32.9351	1.990	1.5002
6	37.5862	1.615	1.4998
7	40.9187	2.287	1.4999
8	43.3271	1.972	1.5002
9	48.0052	2.547	1.4998
10	49.7738	2.918	1.4998
11	52.9703	1.818	1.4999
12	56.4462	1.992	1.5003
13	59.3470	3.563	1.5003
14	60.8318	3.058	1.4997
15	65.1125	2.367	1.4998
16	67.0798	3.131	1.4999
17	69.5464	2.660	1.5004
18	72.0672	2.023	1.4999
19	75.7047	3.561	1.5005
20	77.1448	4.403	1.4998
21	79.3374	2.512	1.5003
22	82.9104	2.638	1.5004
23	84.7355	3.279	1.5004
24	87.4253	4.042	1.4999
25	88.8091	3.406	1.5002
26	92.4919	2.564	1.4999

continued on next page

Table 14 (continued)

k	t	$C(t)$	p
27	94.6513	5.345	1.5003
28	95.8706	4.727	1.5002
29	98.8312	2.354	1.5006
30	101.3179	2.673	1.4999
31	103.7255	3.998	1.5003
32	105.4466	4.681	1.4998
33	107.1686	2.993	1.4999
34	111.0295	5.848	1.5007
35	111.8747	6.907	1.5010
36	114.3202	3.745	1.5001
37	116.2267	3.112	1.5003
38	118.7908	2.552	1.5003
39	121.3701	4.427	1.5005
40	122.9468	6.608	1.5007
41	124.2568	4.474	1.5001
42	127.5167	2.704	1.5003
43	129.5787	4.609	1.5000
44	131.0877	4.742	1.5008
45	133.4977	4.881	1.5005
46	134.7565	3.993	1.5002
47	138.1160	3.652	1.5003
48	139.7362	6.457	1.4999
49	141.1237	5.874	1.5010
50	143.1118	3.200	1.5007

Regression fit: $C(t) = (0.0272 \pm 0.0040) t + (1.043 \pm 0.375)$; $R^2 \approx 0.50$; p range 1.4997–1.5010.

Bootstrap uncertainty for the $C(t)$ fit. As requested in review, we quantify the fit uncertainty beyond the analytic standard errors above via a case-resampling (pairs) bootstrap on the $n = 50$ $(t, C(t))$ points in Table 14 (10,000 resamples with replacement, ordinary least squares refit on each resample): slope 0.0272, 95% percentile CI [0.0205, 0.0343]; intercept 1.043, 95% CI [0.577, 1.502]; R^2 , 95% CI [0.329, 0.667]. The bootstrap intervals are consistent with, and slightly wider than, the analytic standard errors reported above, and confirm that the slope is reliably positive (the CI excludes zero) while the intercept and R^2 carry substantially more uncertainty – consistent with this fit being described throughout this work as preliminary rather than conclusively established (Section 4).

Out-of-sample check on zeros #51–100. As rightly requested in review, the $C(t)$ regression above was fitted exclusively on #1–50 and, prior to this check, had never been evaluated against the additional 50 zeros already present in Table 16. We supply this check here. For each zero $k = 51, \dots, 100$, a local value $C_{\text{actual}}(t_k)$ is obtained from the already-tabulated deviation $e_k(W=10^4)$ (Table 16) under the same fixed-exponent convention used to build Table 14, $C_{\text{actual}}(t_k) := e_k(10^4) \cdot (10^4)^{1.5}$, and compared against the out-of-sample prediction $C_{\text{pred}}(t_k) = 0.0272 t_k + 1.043$ from the original $n = 50$ fit – without refitting on the new data. This is a *proxy-based* out-of-sample check at the single fixed width $W = 10^4$ with a fixed exponent 1.5, not the same multi- W error-curve fit used to build Table 14 for #1–50; the two procedures for obtaining $C(t_k)$ are therefore not identical. The result: mean $C_{\text{actual}} = 6.32$ against mean $C_{\text{pred}} = 6.26$ (aggregate level in reasonable agreement), but at the level of individual zeros the out-of-sample fit is markedly weaker than the in-sample fit, $R_{\text{os}}^2 \approx 0.08$ (correlation ≈ 0.28) compared to $R^2 \approx 0.50$ in-sample,

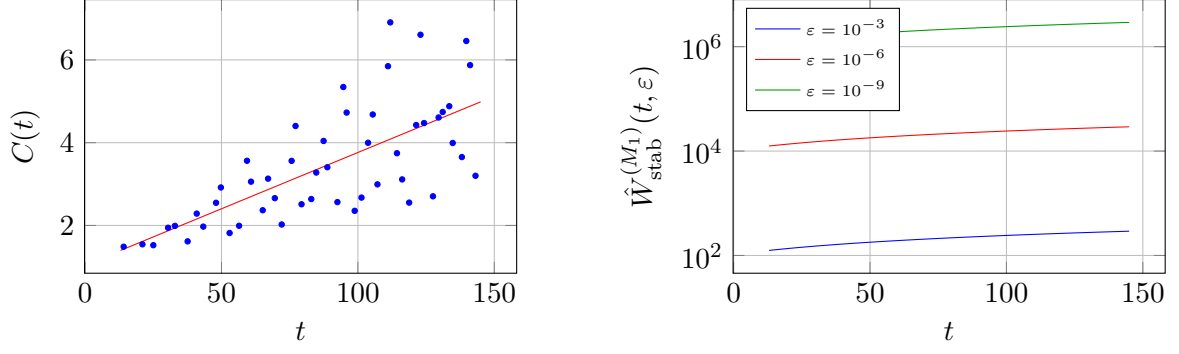


Figure 11: Left: fitted $C(t)$ values for the first 50 zeros #1–#50 (points) with linear trend (line); the trend is statistically significant but explains only about half of the scatter ($R^2 \approx 0.50$). Right: the resulting estimator function $\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon)$ (Eq. 20) for three tolerance levels, still based on the $n = 50$ fit (not extended to the additional 50 zeros in Table 16).

with mean absolute residual ≈ 2.06 and RMSE ≈ 2.57 . An independent regression fitted directly on #51–100 alone gives slope 0.0283 and intercept 0.896, close to the original 0.0272/1.043 – so the aggregate linear trend direction and magnitude approximately reproduce out of sample, but $C(t)$ cannot be reliably predicted for an individual, previously unseen zero from t alone. This confirms, with numbers rather than a general caveat, the concern raised in review: the $n = 50$ fit generalizes at the level of the overall trend, not at the level of a single prediction, and $\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon)$ (equation (20)) should be read accordingly – as an order-of-magnitude guide, not a precise, individually reliable estimate for a specific, previously unseen t .

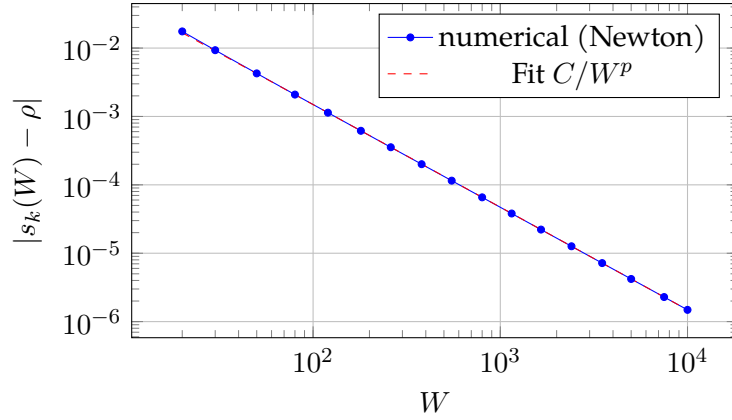


Figure 12: Error curve $|s_k(W) - \rho|$ for the first zero, with power-law fit ($C = 1.499, p = 1.501$).

3.3 Blind and Hybrid Zero Search

Table 15: Zeros found blindly (without reference knowledge).

Zero	Found	Reference	Difference
#9	$0.5000025 + 48.0051511i$	$0.5 + 48.0051508812i$	2.55×10^{-6}
#11	$0.4999983 + 52.9703208i$	$0.5 + 52.9703214777i$	1.82×10^{-6}

A direct Newton start from the analytic estimate (Eq. 62) for zero #12 ($T \approx 57.55$) diverged completely (values $> 10^{12}$ after a few steps). The hybrid strategy (estimate as scan center,

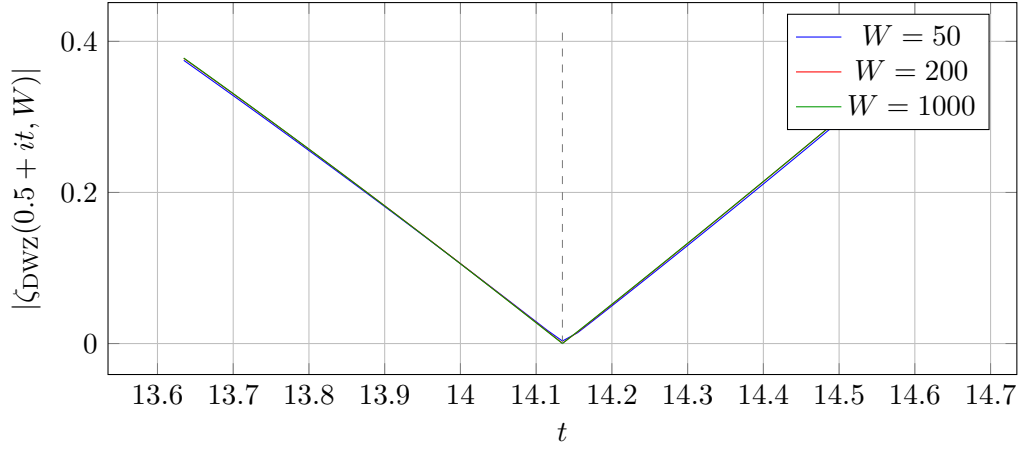


Figure 13: $|\zeta_{\text{DWZ}}(0.5 + it)|$ in Cartesian coordinates near $t = 14.1347$ (dashed line): at small W the minimum remains imprecise; as W grows, the location of the minimum approaches the reference ordinate t_1 without reaching it exactly at any of the finite widths shown.

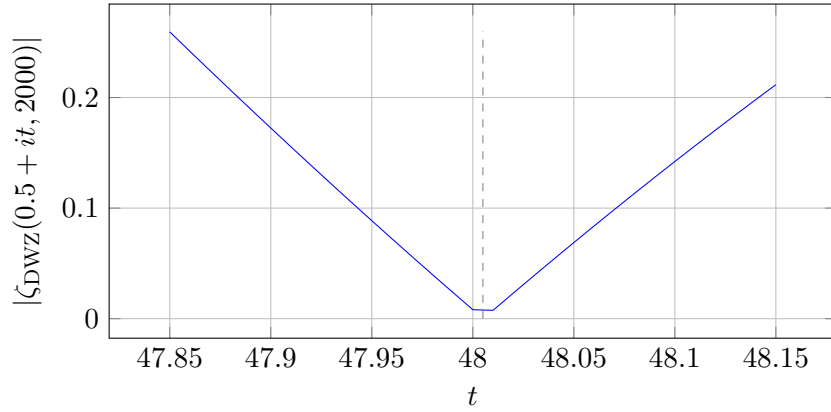


Figure 14: Scan of $|\zeta_{\text{DWZ}}(0.5 + it)|$ at $W = 2000$: on the plotted $\Delta t = 0.01$ grid, the smallest sampled value occurs at $t = 48.01$; the dashed line marks the subsequently refined candidate at $t \approx 48.005$ (Table 15), not the coarse-grid minimum itself.

window ± 4), by contrast, worked and required only 61 instead of 220 scan points, at the same absolute deviation (2.0×10^{-6}).

Distinguishing a found zero from a numerical artifact. Since the Riemann Hypothesis is not assumed, a candidate located by the blind scan (a local minimum of $|\zeta_{\text{DWZ}}(0.5 + it, W)|$) is, by itself, not sufficient evidence of a genuine zero. The validation step used throughout this work is comparison against the independent `mp.zetazero` reference (Riemann-Siegel-based, unrelated to the DWZ scan): both candidates #9 and #11 agree with their nearest reference zero to within 2.6×10^{-6} (Table above), which is treated here as strong evidence of correct identification rather than a formal proof; the blind scan is validated *against*, not instead of, an established high-precision method. A natural next question is the candidate false-positive fraction of the scan itself – a false discovery proportion among reported candidates, not a classical false-positive rate in the statistical sense (which would require a count of true negatives, not defined here), formally $\text{FDP}_{\text{cand}} = N_{\text{FP}}/N_{\text{final}}$ (equation (65) below, stated precisely once the relevant candidate sets have been defined). A complete audit requires the frozen sets

$$\mathcal{G}, \quad \mathcal{C}_{\text{raw}}, \quad \mathcal{C}_{\theta}, \quad \mathcal{C}_{\text{final}},$$

where $\mathcal{G}$ contains all grid points, $\mathcal{C}_{\text{raw}}$ all local minima before thresholding, $\mathcal{C}_{\theta}$ all threshold-passing candidates, and $\mathcal{C}_{\text{final}}$ the refined, deduplicated candidates. The corresponding counts are

$$N_{\text{scan}} = |\mathcal{G}|, \quad N_{\text{raw}} = |\mathcal{C}_{\text{raw}}|, \quad N_{\text{cand}} = |\mathcal{C}_{\theta}|, \quad N_{\text{final}} = |\mathcal{C}_{\text{final}}|.$$

After frozen post-hoc validation, one may report N_{TP} , N_{FP} , and N_{FN} , together with

$$\text{Precision} = \frac{N_{\text{TP}}}{N_{\text{TP}} + N_{\text{FP}}}, \quad \text{Recall} = \frac{N_{\text{TP}}}{N_{\text{TP}} + N_{\text{FN}}}, \quad \text{FDP}_{\text{cand}} = \frac{N_{\text{FP}}}{N_{\text{final}}}. \quad (65)$$

FDP_{cand} denotes a false discovery proportion among final candidates, not a classical false-positive rate. These sets were not preserved for the historical run, so none of these rates can be reconstructed honestly from the present files. The blind-scan experiment is therefore an exploratory existence demonstration, not a validated detection benchmark. Two reference zeros were recovered without reference-based initialization; no claim of general blind-search performance is made.

3.4 Hardening and Pattern Hypothesis

The term "hardening" here denotes a stress test of index assignment (cf. the formal assignment rule in Section 2), not a confirmation of exact target indices: as the following results show, the method reliably finds a known zero in the neighborhood, but only in part of the cases finds exactly the one targeted – the name should thus be read in the sense of "neighborhood stress test of index assignment," not as a general assurance.

The hybrid method was repeated for #13–#41 (29 cases): in all 29 cases, the point found agreed, within the numerical tolerance, with a known reference zero, but only in 10 cases with the exact one targeted – the remaining cases hit a neighboring zero. Plausible, but as shown below *not* statistically confirmed in this sample, is the conjecture that this is related to the zero spacing shrinking on average with t ($2\pi/\ln(t/2\pi)$).

A total of 29 zeros were subjected to the hybrid hardening procedure. However, complete case-level records (spacing, outcome, and precision error for each individual case) are available in the present manuscript only for the 20-case statistical subsample used in the point-biserial and logistic-regression analysis below (Figure 15). No complete 29-row reconstruction table is claimed here, and none of the statistics reported below are extrapolated, estimated, or otherwise derived to cover the remaining 9 cases.

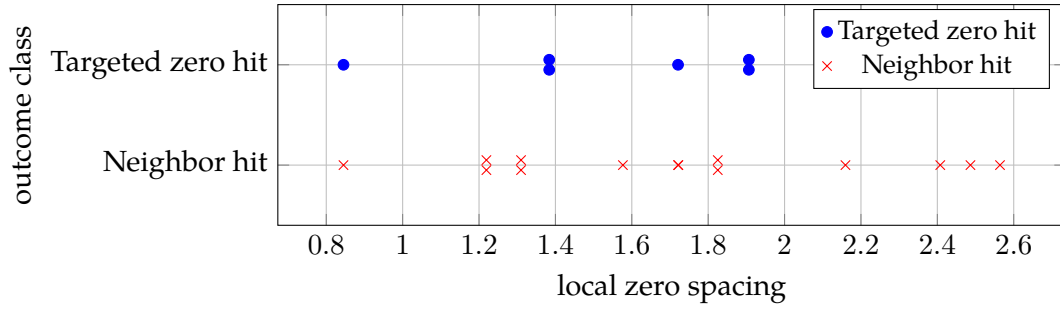


Figure 15: Strip plot of the 20 of 29 hardening runs for which individually documented case-level records are available. The remaining 9 were conducted under the same hardening protocol, but their individual case-level pairs are not available in the manuscript and are therefore not included in this plot. Each point is placed at its local zero spacing on the x -axis and its binary outcome class on the y -axis (small vertical jitter is added only to separate points sharing the same x -value within a class; it carries no quantitative meaning). This is the actual relationship entering the statistical tests below: local spacing (x) versus outcome class (targeted zero hit vs. neighbor hit). The descriptive Pearson coefficient $r = 0.055$ was previously reported over all 29 runs; however, only the 20 individually documented runs shown here support a point-biserial and logistic-regression analysis reported to the extent of this retained case-level record, reported in the text below.

A hypothesis derived from this ("smaller spacing $\Rightarrow$ more confusions") was first checked, as a simplification, against the Pearson coefficient for this sample ($n = 29$, $r = 0.055$, no p -value reported) – this $n = 29$ figure is reported here as a historical, non-reproducible preliminary observation only (see below for why it cannot be independently recomputed from this manuscript), not as a validated result in its own right. As rightly noted in review, for a binary outcome (targeted zero hit vs. neighbor hit) a point-biserial correlation or logistic regression is the methodologically appropriate tool. We supply this here: for the 20 of 29 hardening runs whose individual values are shown explicitly above (Figure 15; of which 6 have outcome "targeted zero hit" and 14 have outcome "neighbor hit"), the point-biserial correlation between local spacing and outcome – coded as a binary variable with 1 = "targeted zero hit" and 0 = "neighbor hit" (the same coding underlies the logistic regression below) – gives

$$r_{pb} = -0.189, \quad t(18) = -0.82, \quad p \approx 0.42, \quad 95\% \text{ CI} \approx [-0.58, 0.28], \quad (66)$$

and a logistic regression (outcome $\sim$ spacing) yields a slope coefficient of -0.93 (standard error 1.09; Wald $z = -0.85$, $p \approx 0.39$; 95% CI $\approx [-3.07, 1.21]$). Both confidence intervals comfortably straddle zero, which further illustrates the lack of significance. At $n = 20$ with only 6 events in the minority outcome class, standard Wald-based inference for logistic regression is known to be unreliable (separation and small-sample bias); Firth's penalized-likelihood or an exact logistic regression would be the more robust choice for this sample size and are noted here as a methodological improvement for future work, not recomputed in this manuscript. Accordingly, the point-biserial correlation above is treated as the primary, reproducible test for this sample; the Wald-based logistic regression is reported as an exploratory sensitivity check, not as an independent confirmation of equal statistical standing. Both tests are statistically non-significant, consistent with the previously reported observation that no defensible linear relationship is detected in this sample – now, however, with an effect size and a p -value rather than an unqualified correlation coefficient alone (a non-significant result does not itself confirm the absence of a relationship; see the power analysis below). Notably, though not to be over-interpreted given the lack of significance, the point estimate points, if anything, in the *opposite* direction from the original hypothesis (larger, not smaller, spacing is weakly associated with more, not fewer, confusions). For the remaining 9 hardening runs, no individually documented (t , outcome)

pairs are available in the manuscript; consequently, the Pearson coefficient $r = 0.055$ reported over all 29 cases at the outset cannot be independently recomputed from the manuscript alone – only the tests reported above, based on $n = 20$ (r_{pb} , logistic regression), are reproducible to the extent of the retained case-level record shown in Figure 15; the figure gives the exact spacing values as plotted coordinates but is a strip plot, not a numeric table, so exact recomputation of r_{pb} to the last reported digit requires reading those coordinates from the underlying plot data rather than a printed table. An analysis over all 29 points, like the previously announced increase in sample size, remains an open follow-up step.

Statistical power of this test. As requested in review, we report the statistical power of the $n = 20$ point-biserial test, computed via the standard Fisher z -transformation normal approximation for a two-tailed test at $\alpha = 0.05$. At $n = 20$, the power to detect a conventionally “medium” effect size ($r = 0.3$) is only 0.25; the power to detect an effect as large as the point estimate actually observed here ($r_{pb} = -0.189$) is just 0.12. Reaching 80% power to detect a medium effect ($r = 0.3$) would require $n \approx 85$; detecting the observed effect size itself at 80% power would require $n \approx 218$. The correct reading of this test is therefore not “no relationship exists,” but “this sample size could not have detected a small-to-medium relationship with reasonable probability even if one were present”; the non-significant result is compatible with a genuine absence of a relationship, but equally compatible with an underpowered test missing a real, moderate one. This sharpens, rather than contradicts, the qualitative conclusion already stated above.

3.5 Numerical Reproduction and Reference Check for 100 Zeros

To frame the following table: of the 100 zeros, 98 were treated with a known starting value (Newton’s method from the rounded `mp.zetazero` reference) and only 2 (#9, #11) were found blindly; 29 of these zeros (#13–#41) were additionally hardened via the hybrid method, of which only 10 of the 29 targeted indices were hit exactly (Section 3). The phrase “100 zeros numerically reproduced” should be read in this sense: predominantly as a local reproduction of already-known reference values, not as an independent discovery or standalone verification of the first 100 zeros.

Table 16: DWZ method: deviation from the reference position for $k = 1, \dots, 100$ (two blocks of 50 shown side by side, k and $k + 50$, to reduce the number of pages this table spans). Here $t_k := \text{Im}(\rho_k)$, and the deviation column is $|\text{DWZ} - \rho_k|$ in units of 10^{-6} .

k	t_k	dev.	k	t_k	dev.
1	14.1347	1.49	51	146.0010	3.75
2	21.0220	1.54	52	147.4228	4.36
3	25.0109	1.52	53	150.0535	7.72
4	30.4249	1.95	54	150.9253	7.95
5	32.9351	1.99	55	153.0247	3.01
6	37.5862	1.62	56	156.1129	4.20
7	40.9187	2.29	57	157.5976	7.53
8	43.3271	1.97	58	158.8500	6.20
9	48.0052	2.55	59	161.1890	3.74
10	49.7738	2.92	60	163.0307	3.30
11	52.9703	1.82	61	165.5371	4.06
12	56.4462	1.99	62	167.1844	5.67
13	59.3470	3.55	63	169.0945	9.17
14	60.8318	3.07	64	169.9120	7.15
15	65.1125	2.37	65	173.4115	4.14
16	67.0798	3.13	66	174.7542	6.18
17	69.5464	2.65	67	176.4414	5.53

Continued on next page

Table 16 (continued; $t_k = \text{Im}(\rho_k)$, deviation in units of 10^{-6})

k	t_k	dev.	k	t_k	dev.
18	72.0672	2.03	68	178.3774	5.27
19	75.7047	3.55	69	179.9165	4.28
20	77.1448	4.41	70	182.2071	3.03
21	79.3374	2.51	71	184.8745	10.30
22	82.9104	2.63	72	185.5988	12.50
23	84.7355	3.27	73	187.2289	5.50
24	87.4253	4.05	74	189.4162	3.13
25	88.8091	3.40	75	192.0267	5.78
26	92.4919	2.56	76	193.0797	6.73
27	94.6513	5.33	77	195.2654	5.71
28	95.8706	4.72	78	196.8765	7.52
29	98.8312	2.34	79	198.0153	5.23
30	101.3179	2.68	80	201.2648	4.96
31	103.7255	3.99	81	202.4936	7.56
32	105.4466	4.69	82	204.1897	8.19
33	107.1686	2.99	83	205.3947	6.11
34	111.0295	5.82	84	207.9063	3.84
35	111.8747	6.85	85	209.5765	4.24
36	114.3202	3.74	86	211.6909	5.33
37	116.2267	3.11	87	213.3479	8.53
38	118.7908	2.55	88	214.5470	8.09
39	121.3701	4.41	89	216.1695	4.08
40	122.9468	6.57	90	219.0676	4.42
41	124.2568	4.47	91	220.7149	12.50
42	127.5167	2.70	92	221.4307	11.40
43	129.5787	4.61	93	224.0070	7.05
44	131.0877	4.71	94	224.9833	6.49
45	133.4977	4.86	95	227.4214	3.68
46	134.7565	3.99	96	229.3374	5.14
47	138.1160	3.64	97	231.2502	13.00
48	139.7362	6.46	98	231.9872	13.18
49	141.1237	5.83	99	233.6934	4.51
50	143.1118	3.18	100	236.5242	4.93
Min. / mean / median / max. ($n = 100$, units 10^{-6}): 1.49 / 4.87 / 4.32 / 13.18					
Standard deviation ($n = 100$, units 10^{-6}): 2.59					

Summary statistics. In addition to the minimum, mean, and maximum, the median and standard deviation are computed directly from the 100 deviations reported in Table 16. These quantities provide measures of central tendency and dispersion for the complete numerical sample.

Runtime against an established reference method (B_{workflow}). For all 100 zeros, the actual runtime was measured in addition to the position deviation and compared against an established reference method: the zeros `mp.zetazero(k)` implemented in `mpmath`, which are determined internally via the Riemann-Siegel formula (evaluation of $\vartheta(t)$ and $Z(t)$ followed by sign-change isolation) and serve throughout this work as the independent reference position ρ_k – an implementation detail of the specific `mpmath` version used for the historical run (Appendix B.0, not recorded), not asserted to be version-independent. For $k = 1, \dots, 100$, the cumulative runtime (30-digit `mpmath` arithmetic, fixed width $W = 10^4$, Newton’s method with known start as in Section 3, identical protocol to Table 16) was

$$T_{\text{DWZ}}(k=1..100) = 135.03 \text{ s}, \quad T_{\text{R.-S.}}(k=1..100) = 4.55 \text{ s}, \quad \frac{T_{\text{DWZ}}}{T_{\text{R.-S.}}} \approx 29.7 \quad (B_{\text{workflow}}).$$

This factor ≈ 30 is the B_{workflow} figure; it is not the same benchmark as the B_{eval} figure of Test C (Section 2.5, roughly two to three orders of magnitude), and the two should not be quoted interchangeably as “the” DWZ/Riemann-Siegel runtime factor. The DWZ position deviation for these

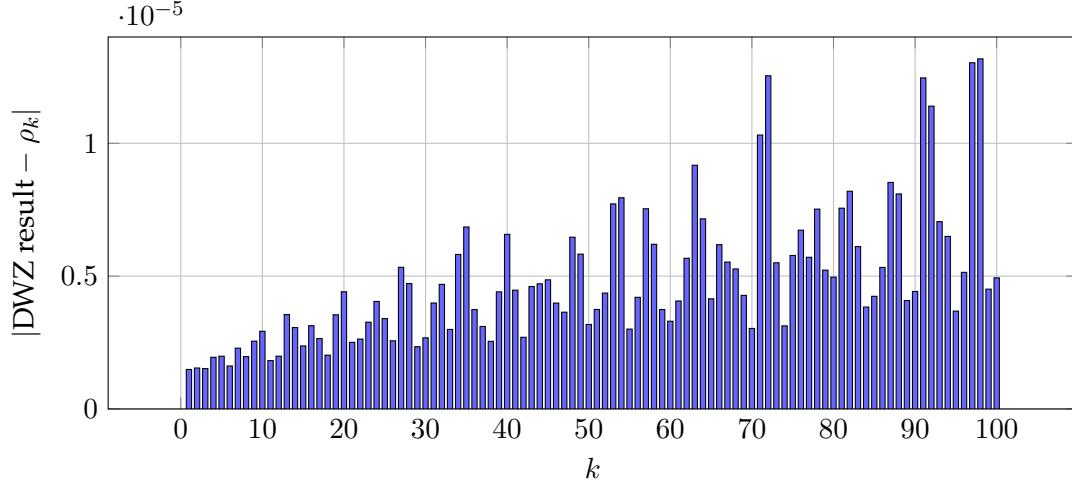


Figure 16: Absolute position deviation $e_k^{(M_1)}(W=10^4)$ for all 100 zeros: between 1.5×10^{-6} and 1.3×10^{-5} , with an increasing tendency for larger k (Pearson correlation between k and e_k : $r \approx 0.65$, $n = 100$, no formal significance test – consistent with $e_k^{(M_1)}(W) \approx C(t_k)W^{-p}$ at fixed W and growing t_k).

100 zeros ranges from 1.49×10^{-6} to 1.32×10^{-5} (mean 4.87×10^{-6} , Table 16), while `mp.zetazero` was evaluated at 30-digit working precision and is treated throughout as the numerical reference value (not as a competitor with its own separately quantified residual uncertainty – no certified error bound for `mp.zetazero` itself is derived or cited in this manuscript). The `mp.zetazero` reference workflow is thus, for these 100 zeros, faster (factor ≈ 30) and used as the higher-precision numerical reference in this manuscript – a result that confirms the efficiency gap already reported in Test C (Section 2.5), here however with a smaller factor: Test C measured the extreme case at much larger t (up to $t \approx 1001$) with a width adaptively optimized for $\varepsilon = 10^{-6}$, whereas the constant width $W = 10^4$ used here already suffices for the smaller t of this range ($t \in [14, 237]$) – the factor 30 is therefore not a contradiction of the factor 100–1000 from Test C, but the result of a deliberately different, here fixed, choice of width. Both measurements are snapshots on the hardware used and are not to be understood as absolute constants, but their relative order of magnitude is consistent across both protocols: the `mp.zetazero` reference workflow – which internally uses Riemann-Siegel but also includes its own isolation and search machinery, and is thus not identical to an isolated Riemann-Siegel implementation under matched input and stopping criteria, the comparison carried by B_{eval} /Test C and Table 5 – is markedly faster than the DWZ method on this measured task. No independent accuracy comparison between the two implementations is claimed here: DWZ accuracy in this manuscript is defined as deviation from the `mp.zetazero` reference itself (no certified error bound for `mp.zetazero` is derived or cited, §3.5), so comparing accuracy directly against that same reference would be circular; DWZ reproduces the reference positions with deviations of 1.49×10^{-6} to 1.32×10^{-5} (Table 16). It should also be noted that T_{DWZ} and $T_{\text{R.-S.}}$ do not measure exactly the same sub-task here: T_{DWZ} measures Newton *refinement* from an already-known starting value at fixed width $W = 10^4$, whereas `mp.zetazero(k)` determines the k -th zero entirely from scratch, including its own isolation and search steps. The measured factor ≈ 30 is thus a comparison of two different, each practically relevant, tasks (refinement from a known start vs. a full search), not a comparison of identical workloads; a strict algorithm comparison would additionally need to measure separately (a) pure function evaluation at the same s and (b) zero search from identical starting information with an identical stopping criterion – neither is carried out here.

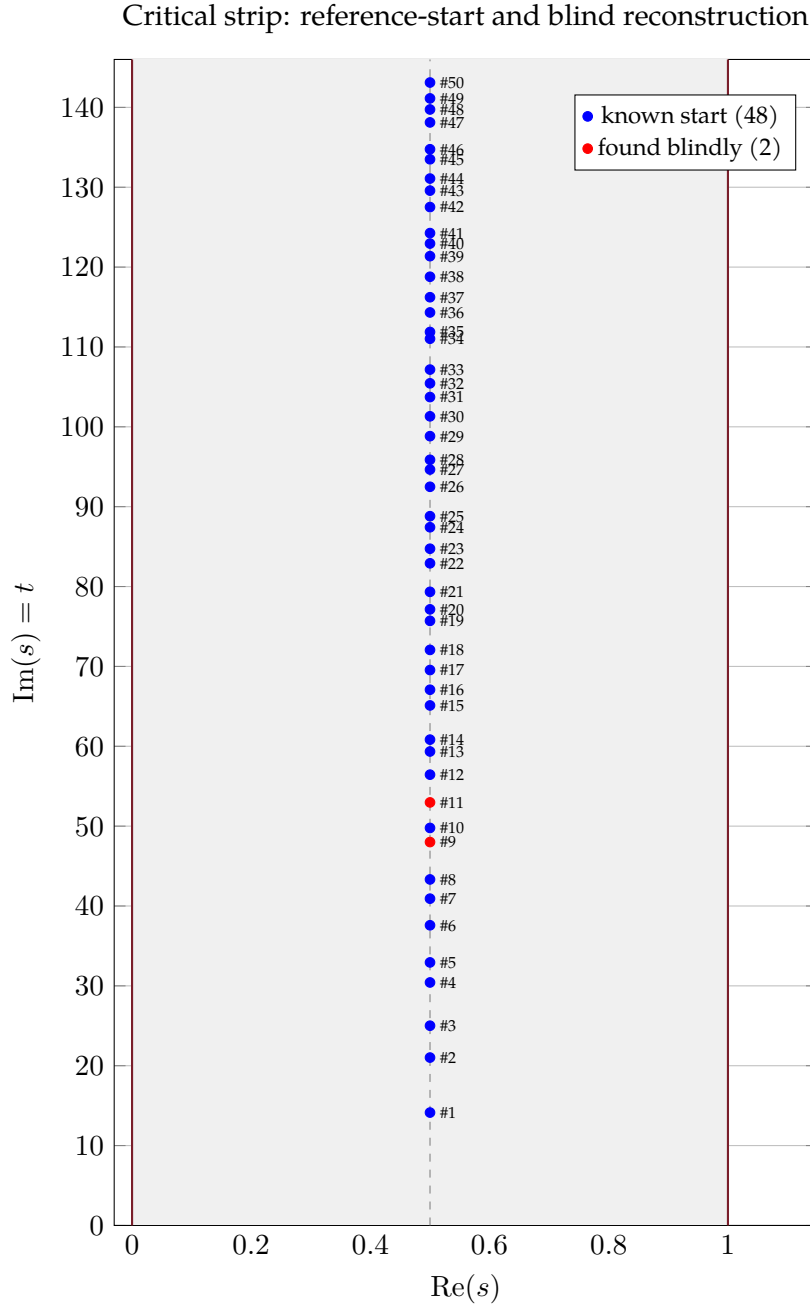


Figure 17: Critical strip $0 < \text{Re}(s) < 1$ (gray) with the first 50 of the now 100 zeros numerically reproduced by the DWZ method (Table 16; #51–#100 are not additionally plotted here for legibility – each point is individually labeled – but are listed in full in Table 16): blue = used as the known reference position for Newton refinement (48 of 50), red = found blindly, without prior knowledge of the exact position (2, #9 and #11) – 4 % of the 50 shown. For all 50, the result agrees with the reference to within $1.5\text{--}6.8 \times 10^{-6}$, regardless of the method used to find it (no point is unaccounted for); for #51–#100 the deviation ranges from 1.5×10^{-6} to 1.3×10^{-5} (Table 16). The additional hardening on 29 of these zeros (#13–#41, already included in the 100, Section 3) is not shown here. The figure distinguishes the numerical protocol used for each point (reference-start Newton refinement vs. blind scan) and does not represent an independent discovery of all plotted zeros.

3.6 Gap Verification

For nine neighboring pairs of zeros, the magnitude $|\zeta_{\text{DWZ}}(\frac{1}{2} + it, W)|$ at $W = 2000$ was sampled finely between the known positions (Table 17). In none of the nine intervals was an additional candidate with a sufficiently small ζ_{DWZ} magnitude observed on the chosen grid. This numerical check only shows that, on the critical line and within the grid used, no further candidate became visible – it is not a rigorous completeness proof: a zero between two grid points, a minimum shifted by finite W , or zeros away from the critical line are not thereby ruled out. Such a proof would require, for instance, an argument-principle count or a certified Turing-style zero count.

Table 17: Tested gaps between neighboring zeros.

Pair	Spacing	Result
#1–#2	6.89	no additional candidate detected on the tested grid
#2–#3	3.99	no additional candidate detected on the tested grid
#3–#4	5.41	no additional candidate detected on the tested grid
#4–#5	2.51	no additional candidate detected on the tested grid
#5–#6	4.65	no additional candidate detected on the tested grid
#6–#7	3.33	no additional candidate detected on the tested grid
#7–#8	2.41	no additional candidate detected on the tested grid
#22–#23	1.83	no additional candidate detected on the tested grid
#23–#24	2.69	no additional candidate detected on the tested grid

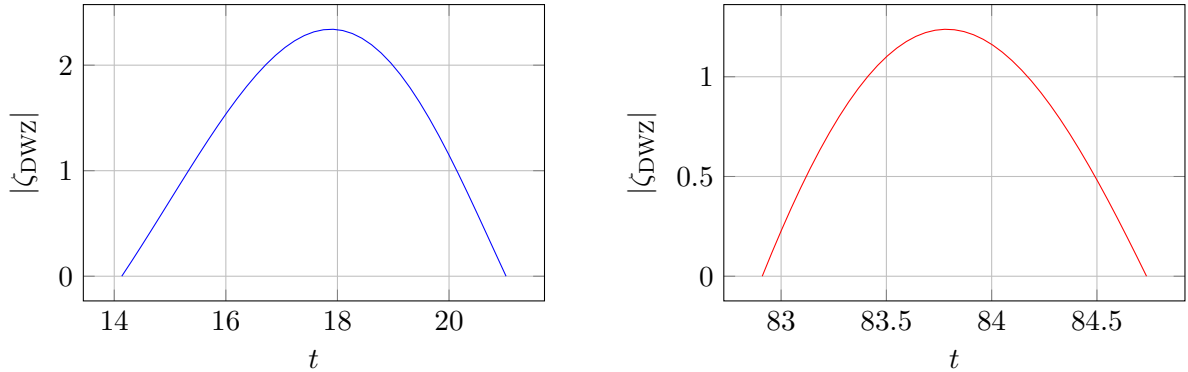


Figure 18: Two examples: largest tested spacing (#1–#2, left, spacing 6.89) and narrowest tested spacing (#22–#23, right, spacing 1.83) – in both cases only a single hump with one minimum at each edge, no additional candidate on the grid used here (not a proof of the absence of a zero, see text above).

4 Discussion

4.1 Scope and Limitations

The disclaimers in this work are scattered by necessity – each one is attached to the specific claim it qualifies – but their cumulative content is short and is collected here once, in full, as a single reference point.

What this work is not. It is not a contribution to the proof or disproof of the Riemann Hypothesis; the statement of RH (equation (3)) is not assumed, and the numerical experiments are restricted to zeros already known to lie on the critical line (Section 1.2). It is not a new mathematical discovery in Euler-Maclaurin summation, which is a classical, eighteenth-century

technique (Section 1.2, “On methodological novelty”). It is not a faster or more accurate way to compute $\zeta(s)$ than the established Riemann-Siegel formula, which this work’s own measurements show to be faster by roughly two to three orders of magnitude (Section 2.5, Test C). It is not a completeness proof for the zeros found or the gaps tested: the blind and hybrid searches (Section 3) and the gap verification (Section 3.6) each operate on a fixed numerical grid and cannot certify the absence of unfound zeros or additional candidates. It is not a claim that the numerical truncation width W_{trunc} equals the register width ω_{reg} ; Lemma 1 shows the opposite for the natural sequential model. It is not a claim that the DRCC operationalization template (Section 2.8.4) transfers generally to other problems; only one numerical instance (the zeta-function case) and one exact, finite combinatorial instance (C_4 graph three-coloring, Appendix B.2) of that template are worked out in this manuscript.

What this work is. It is a case study, documented to the extent of the retained record, demonstrating, end to end and with every approximation made explicit, how an abstract reconstruction-width stability condition can be turned into a checkable number on one concrete, nontrivial numerical problem, including a formal proof that runtime length and active reconstruction state width are logically separate quantities (Lemma 1) and an analytical account of the observed convergence rate (Section 2.8.2).

The remaining paragraphs of this section discuss specific results in more depth; they restate none of the above scope statements beyond what is needed for the point being made.

4.2 Detailed Discussion

The numerical reproduction of known zeros is a *validity demonstration of the method*, not a contribution to the Riemann Hypothesis itself. The reference values come from `mpmath` (Section 2); the DWZ method confirms that a controlled, finite width reconstruction is able to reproduce these positions with high, consistent accuracy – even where zeros lie close together.

Of the first 50 zeros treated in Table 16, 2 (#9, #11) were found blindly, without prior knowledge of the exact position (4 % of these 50); the remaining 48 were refined by Newton’s method starting from the known reference position. The further 50 zeros #51–#100 in the same table were all treated with a known start (Section 3); no separate blind-search test is available for them. The separate hardening on 29 of these zeros (#13–#41, already included in the 100) via the hybrid method at the same time reveals a real limitation of the simple hybrid search: index assignment can mismatch the targeted zero for a neighboring one. This limitation concerns *index assignment*, not the existence of the zeros found – all 29 hits agreed, within the stated numerical tolerance, with known reference zeros. The available $n = 20$ case-level sample does not establish that this limitation is caused by decreasing zero spacing: the point-biserial correlation and logistic regression reported in Section 3 are both statistically non-significant ($r_{\text{pb}} = -0.189$, $p \approx 0.42$; logistic slope $p \approx 0.39$), and the point estimate, if anything, points in the opposite direction from a simple “smaller spacing $\Rightarrow$ more confusions” hypothesis.

4.3 Why DRCC, Given That an Analytic Reconstruction Formula Already Converges

A direct objection deserves a direct answer. The Euler-Maclaurin correction (8) is a classical, self-contained analytic tool: it converges on its own, independently of any DRCC vocabulary, and Section 3 confirms that it does. This raises the question this subsection answers explicitly: *if the analytic formula already solves the numerical problem, what work is DRCC actually doing?*

The answer is not that DRCC improves the Euler-Maclaurin formula mathematically – it does not, and no such claim is made anywhere in this manuscript. The answer is a division of labor:

Euler-Maclaurin is the reconstruction tool; DRCC is the finite control framework.

Euler-Maclaurin supplies the analytic reconstruction operator – the rule that turns a finite partial sum into an estimate of $\zeta(s)$, Equation (8). DRCC supplies everything the formula itself does not decide: which finite fragment is actually processed; that the admissibility of a reconstruction variant among M_1, M_2, M_η for the declared task must be explicitly declared and tested rather than assumed (an admissibility question whose packaging-consistency component is established for M_1 and M_2 – Section 2.6 – while full admissibility additionally relies on the declared analytic and numerical validity conditions of §2.7–2.8, and remains altogether open for M_η); when the declared accuracy requirement remains satisfied for all larger widths; and which finite record is sufficient for the result to be reproduced by someone else. The building blocks for this division are already present in this manuscript – the stability condition (4), the stable reconstruction depth $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ of Equation (19), the separation of state width from truncation depth (Lemma 1), and the Operational DRCC Record (48) – but they have, until now, not been assembled explicitly into this single answer.

DRCC does not compute infinity. The zeta function itself has an infinite analytic definition (1), and the classical convergence statement $\lim_{W \rightarrow \infty} \zeta_{\text{DWZ}}(s, W) = \zeta(s)$ is a fact about that infinite limit. Nothing in the DRCC framework executes this limit. For the reference-calibrated zero-reproduction task actually considered in this manuscript – a tested reference zero ρ_k with $t = t_k$, so that $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ is defined per equation (19) – write

$$W_{\text{DRCC}}(t, \varepsilon) := W_{\text{stab}}^{(M_1)}(t, \varepsilon)$$

for the same stable-threshold quantity already defined in Equation (19) – the notation $W_{\text{DRCC}}(t, \varepsilon)$ is introduced here only to make its role explicit under this framing; it is not a second, independent definition, and t here inherits the same restriction as equation (19) itself: W_{stab} is a post-hoc quantity computed against an already-known reference value, not a quantity defined or computable for an arbitrary, unknown s without such a reference (Section 2.8.1). Once

$$W \geq W_{\text{DRCC}}(t, \varepsilon),$$

the process is declared complete for this reference-calibrated task: the *aggregate* contribution of the omitted tail $n > W_{\text{DRCC}}(t, \varepsilon)$ is represented, for the declared tolerance ε and within the tested model, by the Euler-Maclaurin correction term – not a claim that each individual omitted term n^{-s} is itself reconstructed or separately accounted for; Euler-Maclaurin bounds the sum of the omitted terms, not their individual values. This is a statement about the reference-calibrated task tested here, not a general claim that widths above threshold are redundant for every conceivable task τ . The actually executed reconstruction range is the finite set

$$\mathcal{I}_{\text{DRCC}}(t, \varepsilon) = \{1, \dots, W_{\text{DRCC}}(t, \varepsilon)\},$$

never the full infinite index set. This is the connection to the finite-control reading of DRCC referenced in [7]: for the reference-calibrated zero-reproduction task considered here, DRCC does not compute or approximate infinity – it selects a finite, empirically validated reconstruction sufficient for the declared tolerance within the tested model, and stops there.

Figure 19 answers the process question: what is actually executed, from the infinite object to the finite output. The converse question – what does Euler-Maclaurin alone provide, and what is added exclusively by DRCC? – is answered directly by Table 18.

Reusing the runtime decomposition of the Operational DRCC Record (Section 2.8.4), the honest runtime statement for this case study is

$$T_{\text{DWZ}} = T_{\text{EM}} + T_{\text{control}}, \quad T_{\text{control}} \geq 0, \quad \text{hence} \quad T_{\text{DWZ}} \geq T_{\text{EM}}.$$

This is not a measured gap: in the implementation used here, declaring a model, checking a tolerance, and assembling a record are computationally negligible next to evaluating the Euler-Maclaurin sum itself, so the measured T_{DWZ} in Table 5 already reflects, to good approximation,

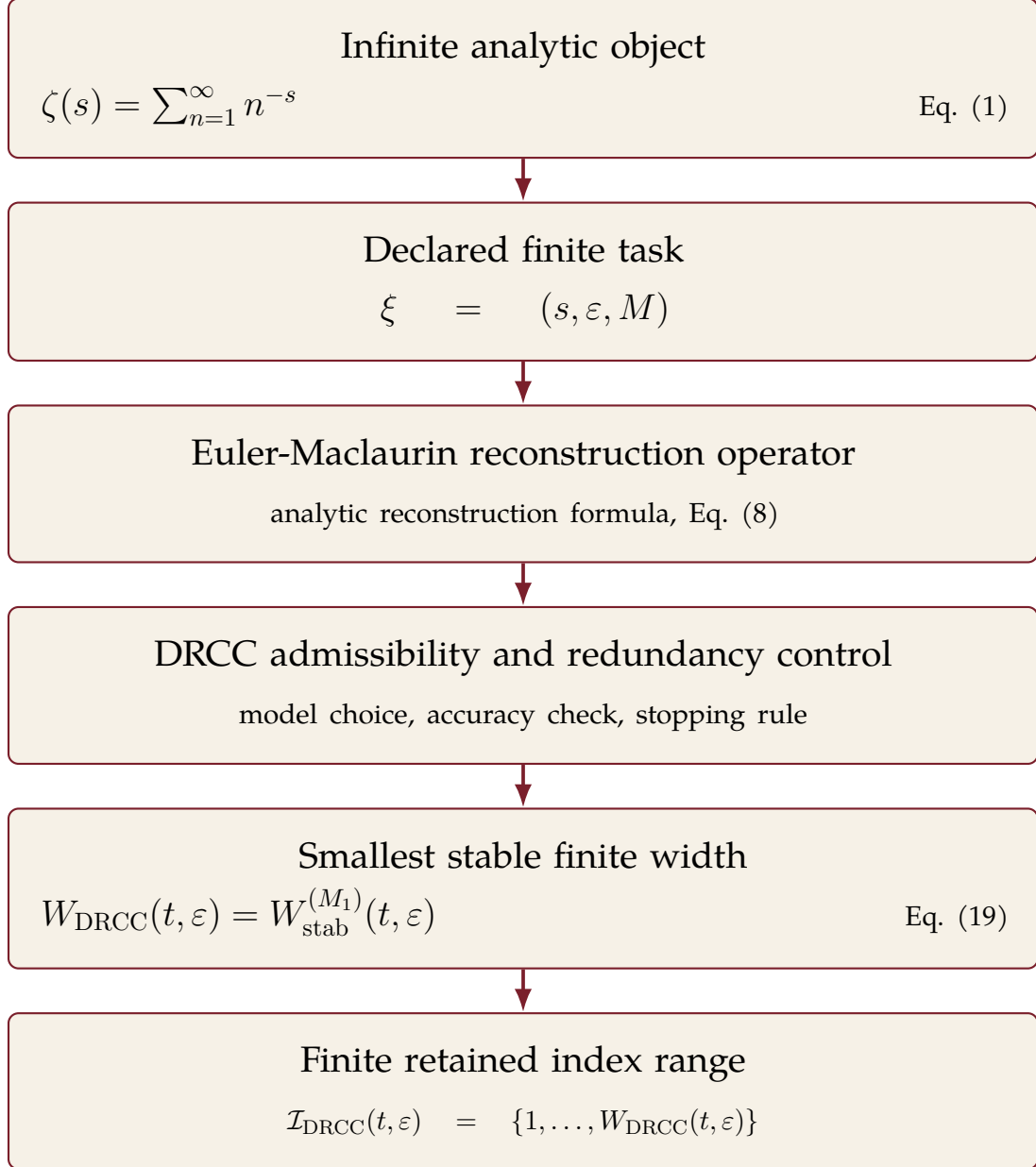


Figure 19: OPA-DRCC-06: The infinite analytic object $\zeta(s)$ is never processed as such. A declared finite task selects an analytic reconstruction operator (Euler-Maclaurin) and a finite control layer (DRCC), which together determine the smallest stable finite width and the finite, reproducible output actually computed.

Table 18: The bare Euler-Maclaurin formula, absent a declared a-priori stopping/tolerance protocol (left), versus Euler-Maclaurin under DRCC control (the DWZ method, right). Test B and Test C elsewhere in this manuscript do use numerical tolerance values ε for exploratory grid search of $W_{\min}$; the left column is not a claim that ε is never used anywhere in this manuscript, only that the bare formula, without the DRCC protocol, carries no predeclared stopping rule that is itself checked against a formally stated (though hypothesis-conditional) criterion, in contrast with Proposition 2’s conditionally formalized stopping criterion on the right, empirically checked on the tested grid, not an unconditional guarantee for every untested W . Not a limitation of classical Euler-Maclaurin theory, which can itself support declared remainder bounds and tolerances – DRCC does not alter or accelerate the analytic formula; it adds a declared, checkable control and documentation layer around it.

Aspect	Euler-Maclaurin as used without the manuscript’s declared control record	Euler-Maclaurin under DRCC control (DWZ)
Termination	No stopping point declared as a predeclared, formally checked protocol (see caption) – convergence left as the asymptotic statement $W \rightarrow \infty$	Explicit stopping rule $W \geq W_{\text{DRCC}}(t, \varepsilon) = W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ (Eq. (19)); a conditionally formalized stopping criterion, numerically supported on the tested widths
Tolerance	Used exploratively in this manuscript’s grid searches (Test B/C), but not checked against a formally stated criterion	ε fixed before computation and checked against $e_M(W)$ (Proposition 2), conditional on that proposition’s hypothesis
Model admissibility	One fixed formula, no declared alternative	Declared candidate set $\mathcal{M}_{\text{cand}} = \{M_1, M_2, M_\eta\}$, selected before evaluating the required width; admissibility of each candidate is assessed separately and remains open for M_η (Section 2.6)
Resource classification	No distinction between truncation depth and register width	$W_{\text{trunc}} \neq \omega_{\text{reg}}$ enforced (Lemma 1)
Reproducibility record	None declared as part of the method	Operational DRCC Record $\mathfrak{D}$ (Eq. (48))
Runtime	$T_{\text{EM}} := T_{\text{reconstruct}}$	$T_{\text{DWZ}} = T_\alpha = T_{\text{reconstruct}} + T_{\text{control}}$, $T_{\text{control}} := T_{\text{build}} + T_{\text{verify}} + T_{\text{output}} \geq 0$
Comparison with Riemann-Siegel	Not separately benchmarked in this manuscript	Mixed: M_1 markedly slower (two to three orders of magnitude); M_{K5} comparable to or faster (factor $1.1 \times - 5.3 \times$) at the same four measured points (Table 5; not a general ordering claim beyond those points). The separate $\approx 30 \times$ figure elsewhere refers to the cumulative B_{workflow} benchmark (Section 3.5), not this comparison. M_2 vs. Riemann-Siegel not separately benchmarked; DRCC neither closes nor claims to close this gap

T_{EM} alone. The inequality is stated for conceptual completeness – DRCC is not claimed to make the computation faster, only better specified – not because a measurable runtime cost of the control layer was observed.

Euler-Maclaurin computes; DRCC controls, terminates, separates, and records.

Euler-Maclaurin determines how the omitted analytic tail is reconstructed. DRCC determines which finite fragment is retained, requires the reconstruction model to be declared and its admissibility to be assessed before the width is interpreted, and determines, for the reference-calibrated task tested here, when the declared accuracy has been reached and when further computation becomes redundant *for that task*. Thus, the analytic formula performs the reconstruction, whereas DRCC controls the finite reconstruction process.

To be precise about what this does and does not claim: this is a statement about the *role* DRCC plays in this manuscript’s construction, not a claim that DRCC is mathematically necessary to derive or use the Euler-Maclaurin correction, and not a claim that no other finite-control vocabulary could describe the same stopping rule. The Euler-Maclaurin formula would converge with or without being described this way. What the DRCC framing adds is a declared, checkable protocol – fragment, model, tolerance, stopping condition, record – for exactly when and why the infinite analytic definition can be replaced by a finite computation, rather than leaving that decision implicit or ad hoc.

What the negative width result teaches about future case studies. The result $\omega_{reg} = O(1)$ is methodologically important because it prevents the zeta case from being used as evidence that register width explains the observed growth of numerical truncation depth.

In the present sequential reconstruction model, increasingly stringent accuracy requirements enlarge the required stable threshold $W_{stab}^{(M_1)}(t, \varepsilon)$, while W_{trunc} remains the free variable over which that threshold is defined; either way, tightening ε does not enlarge the active structural register required to carry the sequential state (Lemma 1). Thus, under the same modeling choice, the zeta experiment falsifies the simplest possible identification

$$\text{numerical depth} = \text{structural width}.$$

This negative result narrows the empirical question for future work, as already noted in Section 6: the next informative DRCC case studies should not merely involve large candidate spaces. They should involve reconstruction processes in which the size of the task-relevant active interface itself can vary with instance size or structural complexity, so that changes in ω_{reg} can be tested against changes in semantic-state growth, reconstruction cost, or total runtime – a test the present bounded-width instance cannot provide by construction.

In this sense, the zeta case performs an important filtering function: it demonstrates successful numerical reconstruction while showing, specifically for the sequential-accumulation realization studied here (Remark 1), that the chosen problem is structurally too narrow to validate a scaling law for register width.

5 DRCC Application Assessment of the DWZ Method

5.1 Purpose of the Application Assessment

The preceding sections establish, piece by piece, how the DWZ method builds a conditional surrogate consistent with the order form of the abstract DRCC stability condition (4) – not

an instantiation of the abstract DRCC quantities themselves, per the surrogate discussion of Section 2.8.1: the traceability chain and equation register of Section 1.4, the conditional surrogate result of Section 2.8.1, the resource separation of Lemma 1, and the role division of Section 4.3. This section does not repeat those derivations, and it does not restate the scope statements already collected in Section 4.1. It adds only what is not stated in full anywhere else in this manuscript: an explicit list of the boundary conditions under which the operationalization holds.

5.2 Mapping and Resource Separation, in Brief

For orientation, the mapping used throughout this manuscript is

$$\begin{aligned}
 & \underbrace{0 < d_{\text{rec}}^{\text{DRCC}}(X) \leq d_{\text{frag}}^{\text{DRCC}}(X) < \infty}_{\text{Eq. (4)}} \rightsquigarrow \underbrace{d_{\text{frag}}^{\text{DWZ}}(X; W) := W_{\text{trunc}}}_{\text{Eq. (6)}} \\
 & \rightsquigarrow \underbrace{W_{\text{stab}}^{(M_1)}(t, \varepsilon)}_{\text{Eq. (19)}} \rightsquigarrow \underbrace{W \geq W_{\text{stab}}^{(M_1)}(t, \varepsilon) \Rightarrow e_k^{(M_1)}(W) < \varepsilon}_{\text{Prop. 2, conditional}}
 \end{aligned}$$

with the three width-like quantities of this manuscript kept explicitly distinct,

$$W_{\text{trunc}} \not\equiv W_{\text{stab}}^{(M_1)}(t, \varepsilon), \quad W_{\text{stab}}^{(M_1)}(t, \varepsilon) \not\equiv \omega_{\text{reg}}, \quad W_{\text{trunc}} \not\equiv \omega_{\text{reg}}.$$

Full detail and traceability are in Table 2, Figures 1 and 3, and Lemma 1; none of it is reproduced here.

5.3 Operational Scope Conditions

The operationalization established in this manuscript holds only within stated boundaries, which this subsection collects explicitly.

- **Convergence hypothesis.** Proposition 2 requires $e_k^{(M_1)}(W) \rightarrow 0$ as $W \rightarrow \infty$ (Eq. (21)) to be assumed, not independently proven for every s in the critical strip.
- **Critical-line qualification.** The observed $W^{-3/2}$ error rate and its analytic origin (Section 2.8.2) are derived and verified on the critical line; away from it, the general heuristic dependence on $\sigma = \text{Re}(s)$ is written down in equation (25), but that equation is explicitly flagged there as a standard asymptotic heuristic, not a rigorous remainder bound, and is not independently validated numerically away from the critical line in this manuscript.
- **Empirical estimator scope.** The estimator $\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon)$ of Equation (20) is fitted on the first 50 of the 100 reproduced zeros, explains only about half the observed scatter ($R^2 \approx 0.50$), and its out-of-sample performance on zeros #51–100 is markedly weaker ($R_{\text{os}}^2 \approx 0.08$) than its in-sample fit.
- **Sample size for the pattern hypothesis.** The point-biserial test on the $n = 20$ hardening subset (Eq. (66)) is statistically non-significant and underpowered; roughly 85 observations would be required for 80 % power to detect an effect of $r \approx 0.3$, so the available sample supports neither confirming nor ruling out the proposed spacing-confusion relationship.
- **Instance scope.** The Operational DRCC Record $\mathfrak{D}$ (Eq. (48)) is worked out completely for one instance, the zeta reconstruction problem; no second numerical instance using $\mathfrak{D}$ is carried out in this manuscript (Remark 6). The C_4 graph-coloring instance (Appendix B.2) is a second, differently-formatted instance of the six-stage workflow, not a second instance of $\mathfrak{D}$.

5.4 Net Assessment

On accuracy, the operationalization succeeds within its stated scope (naive fragment diverges, M_1 converges, 100 known zeros reproduced to 1.5×10^{-6} – 1.3×10^{-5} , Table 16). On runtime, the honest result is unfavorable and is not softened here: M_2 needs roughly $11\times$ fewer terms than M_1 (Eq. (12), a measured term-count/width advantage, not a runtime comparison), M_{K5} a further ≈ 442 – $1150\times$ fewer terms than M_1 (Table 5), and M_1 remains markedly slower than Riemann-Siegel – roughly $94\times$ to $706\times$ across the four rows of Table 5, and independently $\approx 280\times$ (wider-uncertainty estimate) from the repeated-measurement check at zero #1 (Section 2.5); neither figure is to be confused with the separate $\approx 30\times B_{\text{workflow}}$ figure of Section 3.5); at the same four measured points, M_{K5} is comparable to or faster than Riemann-Siegel (roughly $1.1\times$ to $5.3\times$), so no claim of M_{K5} being slower than Riemann-Siegel is made here; no matched runtime benchmark of M_2 against Riemann-Siegel was carried out in this manuscript. Table 18 already lays out this division in full; DRCC does not close the performance gap and is not intended to.

Combined with Section 4.1, the resulting application-level judgment is: the DWZ case study establishes a concretization of the abstract DRCC condition (4) that is formally stated and conditional through Proposition 2, with its convergence premise numerically supported on the tested widths, together with a proven separation of truncation depth from register width ($\omega_{\text{reg}} = O(1)$, Lemma 1) and one fully worked Operational DRCC Record (48). It does *not* establish $W_{\text{trunc}} = \omega_{\text{reg}}$, a faster zeta algorithm, an unconditional convergence proof, a statistically defensible spacing-confusion effect, anything about the Riemann Hypothesis, or that the general template of Section 2.8.4 transfers universally beyond the two deliberately limited instances (Remark 6 and Appendix B) – the same non-claims already stated in Section 4.1, not repeated in expanded form here.

6 Conclusion and Outlook

This work has presented a methodological demonstration of how the abstract DRCC stability condition $0 < d_{\text{rec}}^{\text{DRCC}}(X) \leq d_{\text{frag}}^{\text{DRCC}}(X) < \infty$ can be translated into a finite, measurable, and reviewable numerical workflow for one concrete problem. The Euler–Maclaurin correction is the classical analytic instrument used in the demonstration; it is not the claimed novelty and it is not made more efficient by the DRCC framing. Carrying out that operationalization required first showing that the naive width definition fails in the critical strip (divergent), then replacing the failed naive fragment by the classical first-order Euler-Maclaurin reconstruction (the DWZ method), which, at the tested point $s = 0.5 + 14.13i$ and the width values examined there, is considerably more efficient than the eta-based alternative (Table 13) – no general efficiency theorem over the entire critical strip is derived from this. For the zeta function, this yields a zeta-specific, model-dependent concretization of the originally abstract DRCC inequality $0 < d_{\text{rec}}^{\text{DRCC}}(X) \leq W < \infty$ (Equation (7), operationalized through Equations (19)–(20) and Proposition 2), *without* thereby showing or claiming $W_{\text{trunc}} = \omega_{\text{reg}}$ or $W_{\text{stab}}^{(M_1)}(t, \epsilon) = d_{\text{rec}}^{\text{DRCC}}(X)$ – the abstract, structural DRCC condition itself thus remains unchanged in its abstractness, and the general operational workflow of Section 2.8.4 is worked out numerically for the zeta instance and exactly, but only finitely, for the C_4 instance in Appendix B.

The principal conclusion is therefore deliberately narrow. DWZ shows that DRCC can organize a transparent chain from abstract quantity to observable surrogate, stability threshold, validation protocol, resource separation, and limitation record. It does not show that DRCC is required for the underlying mathematics, that it accelerates the computation, or that the same chain will succeed on another problem. Those distinctions are part of the result rather than qualifications added after it.

Open points for future work:

- A scaling-width case study, as the most important structural validation step remain-

ing for the present operationalization program. The two instances worked out in this manuscript are not directly comparable on this point: for the zeta accumulation, Lemma 1 proves a register width $\omega_{\text{reg}} = O(1)$; for the C_4 graph coloring (Appendix B.2), no such register (simultaneously-active-component) width was computed at all – only an $O(1)$ reconstruction fiber size ($d_{\text{rec}}(C_4, f) = 3$), a different quantity. What the two instances share is only that neither exhibits the fragment or reconstruction quantities growing with problem size in the tested case, not a common proven $\omega_{\text{reg}} = O(1)$ result. This is a property of these two specific reconstruction tasks, not evidence that the DRCC register width is generically constant. The decisive next test of the operationalization workflow is therefore a problem family P_n , with a declared task τ , ordering π , and realization family $\mathcal{M}$, for which the register width $\omega_{\text{reg}}(\tau, \pi; \mathcal{M})$ – abbreviated $\omega_{\text{reg}}(P_n)$ below once these are fixed for the family – provably scales with n rather than remaining bounded, so that the empirical width/runtime relationship established here can be examined outside the constant-width regime. Without such an example, the separation result of Remark 1 remains demonstrated only in the favorable case. A first, purely formal step in this direction is given in Appendix B.4: a minimal grid family for which the register width of equation (14), *under a fixed row-major ordering*, is proven to scale as $\Theta(\sqrt{N})$, establishing that the declared coordinate-wise scalar-register model admits scaling instances – not a scaling result for a globally minimized register width over all admissible orderings, and not a numerical or runtime-benchmarked case study, which remains the open task described here.

- The fit of $C(t)$ still rests on the first 50 zeros from Table 16 ($n = 50$ instead of $n = 8$, Table 14). The linear trend remains statistically significant but explains only about half of the scatter ($R^2 \approx 0.50$); square-root, logarithmic, and pure power-law forms in t do not perform better. Among these four tested functional forms in t alone, none explains the remaining scatter better than the linear fit, which points to at least one further explanatory quantity not captured here (plausible candidates include the local zero spacing $2\pi / \ln(t/2\pi)$ or properties of the Riemann-Siegel theta function). Table 16 itself has since been extended to 100 zeros (plain position deviation at fixed $W = 10^4$; no new per-zero, multi- W curve fit of $C(t)$ and p in the sense of Table 14 – distinct from the coarser, aggregate-level independent regression on #51–100 already reported in Section 3, slope 0.0283, intercept 0.896); extending the per-zero $C(t)/p$ fit to these additional 50 zeros and incorporating further covariates remain a separate, open step. It should also be examined whether the scaling observed here (Table 14, $t \lesssim 143$) also holds outside the tested numerical range ($t \approx 1000$ – 1123).
- Adaptive scan window width (proportional to $2\pi / \ln(t/2\pi)$), to avoid the index confusion observed during hardening.
- Connecting $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$ to the Riemann-Siegel formula for very large $\text{Im}(s)$.
- A stricter analytic error bound for the truncated-zero-sum reconstruction error of the Chebyshev prime-counting function, defined here for completeness as $E_p(W) := |\hat{\psi}_W(x) - \psi(x)|$, where $\psi(x)$ is the Chebyshev function and $\hat{\psi}_W(x)$ its reconstruction from the first W nontrivial zeros via the classical explicit formula [1, 3] – this quantity is introduced here only as a named target for future work and is not otherwise used, derived, or numerically evaluated in this manuscript.
- **Transferability to other Dirichlet series.** This point is specifically about the reconstruction-and-width mechanism, not about this manuscript’s methodology as a whole: the zero search, the analytic estimate of $\zeta'(\rho)$, the Rouché argument of Proposition 3, and the Riemann-Siegel benchmark comparisons are all specific to $\zeta(s)$ and are not claimed to transfer. What does not rely anywhere on properties specific to $\zeta(s) = \sum n^{-s}$ is narrower:

the construction used here (Euler-Maclaurin correction of a partial sum, a formally proven $O(1)$ register width over deterministically reconstructible terms). It should in principle carry over to other Dirichlet series $\sum a_n n^{-s}$ with sufficiently regular coefficients a_n – for instance Dirichlet L -functions or the eta function already used in Section 2 – provided an analogous Euler-Maclaurin correction can be derived for the series in question. This is stated here as a plausible direction for extension, not as a result: neither the convergence order nor the numerical efficiency of such a transfer has been tested in this work for any other series.

Appendix A: List of Symbols

Symbol	Meaning	First Use
s	complex variable, $s = \sigma + it$	§1.1
t	imaginary part of s (“height”)	§1.1
σ	real part of s , $\sigma = \text{Re}(s)$	§1.1
$\zeta(s)$	Riemann zeta function	Eq. (1)
ρ, ρ_k	nontrivial zero resp. k -th reference zero, $\rho_k = \frac{1}{2} + it_k$ (<code>mp.zetazero(k)</code>)	Eq. (3), §2.8
RH	Riemann Hypothesis (open conjecture, not proven here)	Eq. (3)
X	abstract DRCC object (here: the zeta function with its infinite sum)	§1.2
$d_{\text{rec}}^{\text{DRCC}}(X), d_{\text{frag}}^{\text{DRCC}}(X)$	reconstruction resp. fragmentation degree of the abstract DRCC stability condition [7]	Eq. (4)
W, W_{trunc}	width resp. numerical truncation depth (number of summed series terms)	§1.2, Eq. (5)
$\zeta_W(s, W)$	naive partial sum $\sum_{n=1}^W n^{-s}$	Eq. (5)
$\zeta_{\text{DWZ}}(s, W)$	DWZ function, Variant A, model M_1 (1st-order Euler-Maclaurin)	Eq. (8)
$\eta_W(s, W), \zeta_W^\eta(s, W)$	Dirichlet eta partial sum resp. Variant B	§2.3
$\zeta_{\text{DWZ}}^{(2)}(s, W)$	Variant C, model M_2 (2nd-order Euler-Maclaurin)	Eq. (10)
B_2	Bernoulli number, $B_2 = 1/6$	§2.4
M, M_1, M_2	reconstruction model in general resp. Variant A / Variant C	§2.4–2.5
ω	summation order (Test A, order invariance)	§2.5
ε	accuracy/tolerance threshold	Eq. (11)
$W_{\min}(M; s, \varepsilon)$	exact minimum over $\mathbb{N}$ (definition, Eq. (11)); the values actually reported in Table 4 are W_{test} , the smallest width found on a finite grid, not this exact minimum itself	Eq. (11)
$W_{\min}^*(s, \varepsilon)$	minimum of $W_{\min}(M; s, \varepsilon)$ over admissible models $\mathfrak{M}_{\text{adm}}$	Eq. (13)
$W_{\text{stab}}^{(M_1)}(t, \varepsilon)$	accuracy-relative truncation threshold, a numerical quantity distinct from the register width ω_{reg} (definition; numerically evaluated only at tested (t_k, ε) pairs)	Eq. (19)
$\hat{W}_{\text{stab}}^{(M_1)}(t, \varepsilon)$	empirical estimator for $W_{\text{stab}}^{(M_1)}(t, \varepsilon)$	Eq. (20)
$C(t), p$	fit coefficient resp. exponent of the empirical power law $e_k(W) \approx C(t_k)W^{-p}$	Eq. (20), Table 14
$s_k^{(M)}(W)$	zero numerically determined by model M and width W , assigned to ρ_k	§2.8
$e_k^{(M)}(W)$	position error $ s_k^{(M)}(W) - \rho_k $	§2.8
$E_\zeta(s, W)$	function-value error $ \zeta_M(s, W) - \zeta(s) $	§2.8
$N(T)$	Riemann-von Mangoldt counting function (analytic height estimate)	Eq. (62)

Continued on next page

Appendix A: List of Symbols (continued)

Symbol	Meaning	First Use
r	Pearson correlation coefficient (pattern-hypothesis test, hardening)	§3.4
$E_p(W)$	$:= \hat{\psi}_W(x) - \psi(x) $, reconstruction error of the Chebyshev function from W zeros via the explicit formula; a named target for future work only, not derived, used, or numerically evaluated in this manuscript	§6
$\mathfrak{D}$	operational DRCC record (tuple of fragment, model family, reconstruction operator, error measure, tolerance, depth, register width, runtime model); instantiated for the zeta case only	Def. 4, Eq. (48)
$\omega_{\text{reg}}(\tau, \pi; \mathfrak{M})$	register width: minimum, over admissible realizations in $\mathfrak{M}$, of the maximum persistent state size $ \mathcal{Z}_i $	Eq. (14)
$\mathfrak{M}_{\text{seq}}, \mathcal{Z}_k$	sequential-accumulation realization family resp. persistent state at step k	Lemma 1
ρ_W	numerically determined zero at width W (Proposition-3 notation), distinct from $s_k^{(M)}(W)$ used in §2.8	Prop. 3
$U, \overline{U}, \delta_0$	open disk $B(\rho, \delta_0)$ around a zero ρ used in the Rouché argument, resp. its closure	Prop. 3
K_ρ	uniform bound on the function-value error on $\overline{U}$ (Rouché hypothesis)	Prop. 3
σ_U	$:= \inf_{s \in \overline{U}} \text{Re}(s)$; assumed > -1 , automatically satisfied for nontrivial zeros	Prop. 3
m, G_m, N	grid side length, $m \times m$ grid graph, $N = m^2$ vertices (Appendix B.4)	App. B.4
$\tau_{\text{grad}}, S(G_m)$	grid-gradient task resp. grid-gradient energy $\sum_{\{u,v\} \in E_m} (x_u - x_v)^2$	App. B.4
$\mathfrak{M}_{\text{grid}}$	realization family for the grid-gradient sweep, analogous to $\mathfrak{M}_{\text{seq}}$	App. B.4

Appendix B: Technical Reproduction Record and Second Workflow Instantiation

B.0 Computational Environment Record

The fields expected of a complete, pinned computational environment record for the benchmarks of Section 2.5 are listed below; this record covers those timing and zero-reconstruction runs only, not the resampling procedure of Section 3 (separately addressed below). Fields not retained alongside the original timing runs are marked accordingly rather than filled with reconstructed or assumed values.

Field	Value
CPU	not recorded
RAM	not recorded
Operating system	not recorded
Python version	not recorded
mpmath version	not recorded
NumPy version	not recorded
SciPy version	not recorded
Repository commit	not recorded
Randomness used (Section 2.5 timing/reconstruction runs)	no
RNG seed (Section 2.5 timing/reconstruction runs)	not applicable
Bootstrap RNG / seed (Section 3 $C(t)$ -fit, 10,000-resample pairs bootstrap)	not recorded
Bootstrap RNG / seed (Section 2.8.3 worked example, local 10,000-resample bootstrap on $m = 5$ points)	Python random, seed 20260807
Precision setting	30 decimal digits
Number of repetitions	15 (M_1), 30 (M_{K5} , Riemann-Siegel) – Section 2.5
Warm-up protocol	not recorded for the historical run reported here
Timing function	<code>time.perf_counter</code>

No pseudo-random number generator was used in the reported benchmark or zero-reconstruction computations of Section 2.5; consequently, no RNG seed enters those results. This does not extend to the case-resampling (pairs) bootstrap used for the $C(t)$ -fit uncertainty in Section 3, which does draw random resamples; its RNG state and seed were not retained alongside the original run and are listed above as not recorded, rather than silently assumed to be covered by the “no randomness” statement for the timing benchmarks. This table is deliberately left with explicit “not recorded” entries rather than plausible-looking placeholders, consistent with Section 2.5’s benchmark-environment paragraph; a future archival release (see “Code and data availability” at the end of this manuscript) is the appropriate place to complete it with an actually pinned environment.

B.1 Blind-Scan Reproduction Record

The complete reproduction package for a future blind-scan rerun must retain $\mathcal{G}$, $\mathcal{C}_{\text{raw}}$, $\mathcal{C}_\theta$, and $\mathcal{C}_{\text{final}}$. For every final candidate it must store

$$\left(t_j^{\text{scan}}, \hat{t}_j, \left| \zeta_{\text{DWZ}} \left(\frac{1}{2} + i\hat{t}_j, W_{\text{ref}} \right) \right|, k_j^*, \left| \hat{t}_j - t_{k_j^*}^{\text{ref}} \right|, h_j \right),$$

where $h_j \in \{\text{TP}, \text{FP}\}$ is assigned only after the candidate set is frozen. The archive must additionally record the exact tuple $\mathcal{B}$ from equation (61), software versions, working precision, input and output file names, and the random-state record if any stochastic component is introduced. These materials do not exist as a complete retained archive for the historical two-hit experiment.

B.2 Second Workflow Instantiation: Exact Reconstruction in Graph Three-Coloring

This section applies the six-stage workflow of Section 2.8.4 to a finite problem that is independent of Euler–Maclaurin analysis.

Let $C_4 = (V, E)$ be the four-vertex cycle,

$$V = \{v_1, v_2, v_3, v_4\}, \quad E = \{\{v_1, v_2\}, \{v_2, v_3\}, \{v_3, v_4\}, \{v_4, v_1\}\},$$

with color set $K = \{1, 2, 3\}$. The raw assignment space has cardinality

$$d_{\text{frag,raw}}(C_4) = |K|^{|V|} = 3^4 = 81. \quad (67)$$

The correspondence between $d_{\text{frag,raw}}(C_4)$ here and W_{trunc} in the zeta case (Section 2) is only an analogy of workflow role – both are the raw-fragment-count field of the six-stage operationalization template – not a mathematical identity of type: $d_{\text{frag,raw}}(C_4)$ is a finite assignment-space cardinality, while W_{trunc} is a series-truncation index length; the two are not interchangeable or directly comparable in magnitude. A coloring $c : V \rightarrow K$ is admissible exactly when

$$c(u) \neq c(v) \quad \text{for every } \{u, v\} \in E. \quad (68)$$

The chromatic-polynomial identity $P(C_n, q) = (q-1)^n + (-1)^n(q-1)$ gives

$$d_{\text{frag,adm}}(C_4) = P(C_4, 3) = 2^4 + 2 = 18. \quad (69)$$

To state the fragment/fiber construction used below in the same formal terms as the rest of the manuscript, let $\Phi(C_4)$ denote the set of admissible colorings, $|\Phi(C_4)| = d_{\text{frag,adm}}(C_4) = 18$, and for a vertex subset $V_f \subseteq V$ let

$$F(C_4) := \bigcup_{V_f \subseteq V} \{f : V_f \rightarrow K\} \quad (70)$$

denote the space of partial colorings on all such subsets. Define the restriction (collapse) operator

$$R_{V_f} : \Phi(C_4) \rightarrow F(C_4), \quad R_{V_f}(c) := c|_{V_f}, \quad (71)$$

which maps a complete admissible coloring to its restriction on V_f . Fix the fragment $V_f = \{v_1, v_2\}$,

$$f(v_1) = 1, \quad f(v_2) = 2. \quad (72)$$

The admissible completions satisfy $c(v_3) \neq 2$, $c(v_4) \neq c(v_3)$, and $c(v_4) \neq 1$. A complete enumeration gives

$$(c(v_3), c(v_4)) \in \{(1, 2), (1, 3), (3, 2)\}, \quad (73)$$

so the reconstruction fiber, expressed via the restriction operator as

$$\Omega_{C_4}(f) := \{c \in \Phi(C_4) : R_{V_f}(c) = f\}, \quad (74)$$

is, by the enumeration above,

$$\Omega_{C_4}(f) = \{(1, 2, 1, 2), (1, 2, 1, 3), (1, 2, 3, 2)\}, \quad d_{\text{rec}}(C_4, f) = 3. \quad (75)$$

The exact operational record is therefore

$$81 \longrightarrow 18 \longrightarrow 3. \quad (76)$$

The first arrow removes assignments violating the edge constraints (a reduction by a factor of $81/18 = 4.5$); the second conditions the admissible solution space on the declared fragment (a further reduction by a factor of $18/3 = 6$), for a combined raw-to-fiber factor of $81/3 = 27$. These two arrows are *not* both instances of the restriction operator R_{V_f} : by construction, R_{V_f} (equation (71)) has domain $\Phi(C_4)$, the 18 admissible colorings, so R_{V_f} and the fiber definition (74) realize only the second arrow, $18 \rightarrow 3$. The first arrow, $81 \rightarrow 18$, is a separate, prior step – constraint filtering by the edge condition (68), applied before R_{V_f} is defined at all, not part of the collapse operator itself. Schematically,

$$\underbrace{81}_{\text{raw assignments}} \xrightarrow{\text{constraint filtering, Eq. (68)}} \underbrace{18}_{\text{admissible colorings } \Phi(C_4)} \xrightarrow{R_{V_f}^{-1}(f), \text{ Eq. (74)}} \underbrace{3}_{\text{reconstruction fiber}}.$$

A second, complementary view starts from the task as stated – reconstruct all admissible extensions of the *already fixed* fragment f on $V_f = \{v_1, v_2\}$ – in which case the immediate raw extension space consists only of the free vertices v_3, v_4 , of size $3^2 = 9$, not 81; the reduction $9 \rightarrow 3$ (a factor of 3, from the edge constraints on v_3, v_4 alone) is then the more directly relevant comparison to the reconstruction fiber, with $81 \rightarrow 18$ read as a separate, whole-graph normalization rather than a step in this narrower view. Both views are mathematically consistent – they simply start counting from different declared spaces (81: all colorings of all four vertices; 9: colorings of the two free vertices, fragment already fixed) – and neither replaces the other. The result is an exact finite reconstruction count, not a runtime theorem and not a proof that graph three-coloring is easy in general. To be explicit about which of the two upstream quantities plays the role of $d_{\text{frag}}^{\text{DRCC}}(X)$ in the abstract stability condition (4): it is $d_{\text{frag,raw}}(C_4) = 81$, the full, unconstrained assignment space, that plays the analogous fragmentation role played by W_{trunc} in the zeta case – an upper bound on the search space before any structural constraint is applied. The two are not the same *type* of quantity: W_{trunc} counts summed series terms (an index-range length), while $d_{\text{frag,raw}}(C_4)$ counts assignments in a state space (a cardinality of colorings); the analogy is at the level of the workflow role each plays, not an identity of the underlying mathematical objects. $d_{\text{frag,adm}}(C_4) = 18$ is an intermediate, edge-constrained quantity introduced for this combinatorial example specifically and has no counterpart in the zeta instantiation; it is not itself the $d_{\text{frag}}^{\text{DRCC}}(X)$ of equation (4), but a step on the way from $d_{\text{frag,raw}}$ to the reconstruction fiber size $d_{\text{rec}}(C_4, f)$.

The six workflow stages are instantiated as follows:

1. abstract quantity and task: reconstruct all admissible extensions of f ;
2. observable fragment (the fixed partial assignment f , not to be confused with the raw fragmentation-space size $d_{\text{frag,raw}}(C_4) = 81$ above): equation (72);
3. reconstruction model: complete colorings extending f ;
4. measurable criterion: equation (68);
5. exact validation by exhaustive enumeration of all 81 raw assignments;
6. operational record: equation (76).

Thus the workflow is demonstrated in two distinct settings: an empirical numerical approximation problem and an exact finite combinatorial reconstruction problem. This does not establish universal transferability, but it shows that the workflow is not tied exclusively to the analytic form

of the Euler–Maclaurin construction. To state this precisely: this appendix validates workflow portability only – the same six-stage template applies to both instances – and does not establish type-preserving portability of d_{frag} or d_{rec} themselves, which remain different mathematical objects ($d_{\text{frag,raw}}(C_4)$ a finite assignment-space cardinality, $d_{\text{rec}}(C_4, f)$ a reconstruction-fiber cardinality) in the two instances, not a single quantity transported unchanged between them.

B.3 Benchmark Status Record

The repeated M1/MK5/Riemann–Siegel measurements reported in the main text remain part of the evidence record. The Arb comparison is intentionally unfilled below. No numerical Arb claim may be made until a matched repeated benchmark has been executed and archived.

Table 19: Protocol template – no measurements yet. Planned matched repeated-evaluation benchmark against a modern ball-arithmetic baseline. Dashes denote measurements not yet executed; this table records the intended design, not a result.

Zero	t	M1 median	MK5 median	R.-S. median	Arb median
1	14.1347	–	–	–	–
25	88.8091	–	–	–	–
50	143.1118	–	–	–	–
650	1001.3495	–	–	–	–

The intended protocol uses five warm-up runs and at least thirty measured runs per method, reports the median and interquartile range as primary statistics, records mean and standard deviation secondarily, and fixes a common target accuracy $\varepsilon = 10^{-6}$ and working precision of 30 decimal digits or a declared bit-equivalent. The decisive methodological status is

The Arb benchmark remains open until real matched measurements exist.

Code and data availability. At the time of this revision, a complete executable code-and-data archive has not yet been deposited. The numerical procedures, parameter choices, and principal results are documented in the manuscript, while the consolidated scripts, raw output files, and environment specification remain to be prepared for archival release. References elsewhere in this manuscript to “accompanying scripts” describe the computational scripts used to produce the reported results, not a currently public, consolidated archive; they should not be read as a present reproducibility guarantee independent of this manuscript.

B.4 A Formal Scaling-Width Example (Grid-Gradient Sweep)

This is a self-contained formal register-width lemma in the style of Lemma 1, not a third full six-stage operationalization in the sense of Section 2.8.4 or Appendix B.2 (no fragment/model/criterion/validation/record instantiation is carried out for it).

Section 5 and Section 6 identify the absence of a problem family with a provably *scaling* structural width as the most important open point of this manuscript: both instances worked out above (the zeta accumulation, Lemma 1; the C_4 three-coloring, Appendix B.2) are bounded-width or non-width-computing cases. This subsection gives one minimal, self-contained family for which the register-width quantity of equation (14) is proven, not merely observed, to scale – built in exactly the same formalism as Lemma 1, but with the opposite conclusion.

The family. For $m \geq 2$, let $G_m = (V_m, E_m)$ be the $m \times m$ grid graph, $V_m = \{1, \dots, m\}^2$, $N = m^2$, with edges between horizontally and vertically adjacent vertices. Each vertex (i, j) carries an input

value $x_{i,j} \in \mathbb{R}$, presented once, in row-major order $\pi_{\text{row}}: (1, 1), (1, 2), \dots, (1, m), (2, 1), \dots, (m, m)$. The task τ_{grad} is to output the total grid-gradient energy

$$S(G_m) := \sum_{\{u,v\} \in E_m} (x_u - x_v)^2. \quad (77)$$

Let $\mathfrak{M}_{\text{grid}}$ denote the realization family analogous to $\mathfrak{M}_{\text{seq}}$ of Section 2.5: implementations that read the $x_{i,j}$ exactly once, in the declared order π_{row} , with no re-access to already-read input values except through explicitly retained persistent state $\mathcal{Z}_i$, and that correctly output $S(G_m)$ after the last value is read. As in Section 2.5, $|\mathcal{Z}_i|$ counts persistent scalar state components at a fixed working precision: one retained input value $x_{i,j} \in \mathbb{R}$ occupies exactly one such component, and lossless packing of several independent input values into a single retained component (e.g. via an injective pairing function) is not an admissible representation under this register-counting convention – the same convention already used, and stated explicitly, for $\mathfrak{M}_{\text{seq}}$. Formally, this is a *coordinate-wise scalar-storage model*: $\mathcal{Z}_i \in \mathbb{R}^r$ for some $r = |\mathcal{Z}_i|$, and each of the r coordinates may losslessly retain at most one independently-varying input scalar at a time (a coordinate may be a fixed function of one such retained scalar, but not an injective encoding of two or more of them jointly). This model assumption is what turns pairwise distinguishability of the retained values, established in the lower-bound proof below, into a lower bound on the number of *register components* $r = |\mathcal{Z}_i|$, rather than merely on the amount of information those components jointly encode; without it, a more compact sufficient encoding of the same information could not be ruled out a priori.

Lemma 2 (Register width of the grid-gradient sweep, under the scalar-register, no-lossless-packing convention).

$$\omega_{\text{reg}}(\tau_{\text{grad}}, \pi_{\text{row}}; \mathfrak{M}_{\text{grid}}) = m + O(1) = \Theta(\sqrt{N}).$$

Proof (upper bound). Exhibit $\mathcal{A}^* \in \mathfrak{M}_{\text{grid}}$: maintain an array $\text{buf}[1..m]$, a scalar left , and a running total S . While row 1 is read, set $\text{buf}[j] := x_{1,j}$ (no vertical edges yet) and accumulate horizontal terms into S via left . For each subsequent row $i \geq 2$ and column $j = 1, \dots, m$: upon reading $x_{i,j}$, add $(x_{i,j} - \text{buf}[j])^2$ (vertical edge to $(i-1, j)$) and, for $j \geq 2$, $(x_{i,j} - \text{left})^2$ (horizontal edge to $(i, j-1)$) to S ; then set $\text{left} := x_{i,j}$ and overwrite $\text{buf}[j] := x_{i,j}$ in place, so that buf now holds row i for use by row $i+1$. At every step, $|\mathcal{Z}_i| = |\text{buf}| + |\{\text{left}, S\}| = m + 2$. This realization is admissible for τ_{grad} under π_{row} , so $\omega_{\text{reg}}(\tau_{\text{grad}}, \pi_{\text{row}}; \mathfrak{M}_{\text{grid}}) \leq m + 2$. $\square$

Proof (lower bound). Fix any $\mathcal{A} \in \mathfrak{M}_{\text{grid}}$ and any row i with $2 \leq i \leq m$. Consider the instant immediately after $\mathcal{A}$ has read $x_{i,1}$. The values $x_{i-1,2}, x_{i-1,3}, \dots, x_{i-1,m}$ ($m-1$ values) are each required later – to compute the vertical edge term $(x_{i,j} - x_{i-1,j})^2$ once column j of row i is reached, for $j = 2, \dots, m$ – and, since $\mathcal{A} \in \mathfrak{M}_{\text{grid}}$ cannot re-read already-consumed input, each of these $m-1$ values must already be present in $\mathcal{A}$'s persistent state at this instant (none of these values can yet be discarded, because each is still required for a future vertical-edge term, and none will be re-presented by the input stream). To make this rigorous rather than merely suggestive: fix $j \in \{2, \dots, m\}$ and suppose, for contradiction, two input streams agreeing on every value read so far except $x_{i-1,j}$ (say $x_{i-1,j} = a$ versus $x_{i-1,j} = a' \neq a$) drive $\mathcal{A}$ to the *same* internal state at this instant. Both streams admit a common continuation in which (i, j) is later read with some value $x_{i,j} = b$. Writing $S_a(b)$ and $S_{a'}(b)$ for the correct total energy of the two continuations as functions of b (with a resp. a' held fixed at $x_{i-1,j}$), every term of $S(G_m)$ not involving $x_{i-1,j}$ or $x_{i,j}$ is common to both and cancels in the difference, every term involving $x_{i-1,j}$ but not $x_{i,j}$ (e.g. the horizontal edges of row $i-1$ incident to $x_{i-1,j}$, already accumulated by this instant) contributes a fixed constant $K(a) - K(a')$ independent of b , and the one term involving both, $(b-a)^2 - (b-a')^2 = (a'-a)(2b-a-a')$, is affine in b with nonzero slope $2(a'-a) \neq 0$; hence $S_a(b) - S_{a'}(b)$ is a non-constant affine function of b , so $S_a(b) \neq S_{a'}(b)$ for all but at most one value of b . So $\mathcal{A}$, being in the same state after the differing values and facing

the same subsequent input, cannot produce both correct outputs for such a b – a contradiction. Hence the internal state after reading $x_{i-1,j}$ must distinguish every value of $x_{i-1,j} \in \mathbb{R}$, and since $m - 1$ such values ($j = 2, \dots, m$) require pairwise-distinguishable information simultaneously, and the no-lossless-packing convention disallows encoding more than one of them into a shared component, each requires its own reserved component of $\mathcal{Z}_i$. Hence $|\mathcal{Z}_i| \geq m - 1$ at this instant for every admissible $\mathcal{A}$, so $\omega_{\text{reg}}(\tau_{\text{grad}}, \pi_{\text{row}}; \mathfrak{M}_{\text{grid}}) \geq m - 1$. $\square$

Combining both bounds, $m - 1 \leq \omega_{\text{reg}}(\tau_{\text{grad}}, \pi_{\text{row}}; \mathfrak{M}_{\text{grid}}) \leq m + 2$, proving the lemma; since $N = m^2$, this is $\Theta(\sqrt{N})$.

Small-instance check. Table 20 tabulates the upper-bound witness register size $m + 2$ against N and $\sqrt{N}$ for the first few instances of the family, making the linear-in- $\sqrt{N}$ growth concrete rather than only asymptotic.

Table 20: Grid-gradient register width against instance size (upper-bound witness of Lemma 2). ω_{reg} grows linearly in $\sqrt{N} = m$, in contrast to the constant $\omega_{\text{reg}} = O(1)$ of Lemma 1 for the zeta accumulation.

m	$N = m^2$	$\sqrt{N}$	Witness state size $ \mathcal{Z}_i = m + 2$
2	4	2	4
3	9	3	5
4	16	4	6
5	25	5	7
10	100	10	12

Scope of the register-width result. Exactly as in Section 2.5, this is a self-contained register-counting result about a declared implementation family $\mathfrak{M}_{\text{grid}}$. Lemma 2 gives the register width below for this specific, declared realization family *under the fixed row-major ordering* π_{row} :

$$\omega_{\text{reg}}(\tau_{\text{grad}}, \pi_{\text{row}}; \mathfrak{M}_{\text{grid}}) = \Theta(\sqrt{N})$$

This proves scaling register width for the declared row-major realization. It does *not* establish a $\Theta(\sqrt{N})$ lower bound for a register width minimized over all admissible orderings: a possibly more favorable ordering $\pi \neq \pi_{\text{row}}$ achieving a smaller register width is not excluded by this lemma alone. What the lemma *does* close, at the level of a formal existence proof, is the open point of Section 6 that the declared coordinate-wise scalar-register model admits scaling instances at all: a register width *can* scale with instance size under this framework for at least one declared ordering, so the constant width found for the zeta case (Lemma 1) is a property of that specific reconstruction task under its declared ordering, not an artifact of the register-width formalism itself. This example does *not* extend the zeta numerical case study: it is not connected to Euler–Maclaurin summation, no runtime or numerical-accuracy measurement is carried out for it here, and it is not offered as a fully worked empirical case study with an implemented, benchmarked reconstruction process – only as a minimal formal witness that the scaling regime this manuscript’s zeta instance does not test is not vacuous. A full empirical width/runtime case study on a scaling family, as called for in Section 6, remains future work.

References

- [1] B. Riemann, *On the Number of Prime Numbers Less Than a Given Quantity*, Monthly Reports of the Berlin Academy, 1859 (English title supplied for consistency).

- [2] E. C. Titchmarsh, *The Theory of the Riemann Zeta-Function*, 2nd ed., rev. D. R. Heath-Brown, Oxford University Press, 1986.
- [3] H. M. Edwards, *Riemann's Zeta Function*, Academic Press, 1974.
- [4] A. M. Turing, *Some calculations of the Riemann zeta-function*, Proc. London Math. Soc. (3) 3, 1953, pp. 99–117.
- [5] A. M. Odlyzko, *Tables of zeros of the Riemann zeta function*, available at https://www.dtc.umn.edu/~odlyzko/zeta_tables/.
- [6] F. Johansson et al., *mpmath: a Python library for arbitrary-precision floating-point arithmetic*, <http://mpmath.org> (the specific version used for the historical run reported here is not recorded, see Appendix B.0; a versioned citation is recommended for the next archival revision).
- [7] R. Hesamiy, *DRCC: Dimensional Reduction via Controlled Combinatorics*, Zenodo, Version V1.0, 2026, doi:10.5281/zenodo.20499433; see also aiXiv:260604.000001 (posted there under the title *DRCC: Dimensionality Reduction as Complexity Collapse – A Structural Framework for Controlled Reconstruction in Complex Search Spaces*). (This reference provides the theoretical motivation for the present work; it is cited here under its Zenodo title and DOI as the primary, permanent identifier.)
- [8] M. Rubinstein, *Computational methods and experiments in analytic number theory*, in: Recent Perspectives in Random Matrix Theory and Number Theory, London Math. Soc. Lecture Note Series 322, Cambridge University Press, 2005, pp. 425–506.