

What Kenneth Wilson Actually Did

A Renormalization Group Primer via the Ising Model

Sridhar Tayur

Tepper School of Business, Carnegie Mellon University

Tayur Musings on Physics

July 2026

Abstract

Kenneth Wilson won the 1982 Nobel Prize in Physics for the renormalization group (RG), a method that resolved a decades-old puzzle in the theory of phase transitions: why the critical exponents governing how order disappears at a critical point are the same number, reproducibly, across wildly different materials, while the best available theory (Landau’s) confidently predicted the wrong number. Popular accounts of this story tend to stop at metaphor: couplings that “run,” a universe that looks different depending on the scale you observe it from. This note does not stop there. Using the Ising model of a ferromagnet as the entry point we carry out real renormalization group calculations by hand: an exact decimation in one dimension, and a flagged approximate one in two dimensions, checked in both cases against known exact results. We connect the resulting divergent correlation length to current work on criticality as a quantum-sensing resource, and close by checking that the same running-coupling idea, applied to energy rather than length, correctly explains the direction (and much of the size) of the fine-structure constant’s growth with energy, and the sign that makes the strong coupling run the opposite way. This note has not attempted to reproduce Wilson’s full momentum-space construction for continuum field theories or his multi-coupling real-space flows; instead it has aimed to show, in the simplest solvable setting, the key move that distinguishes his work from Landau’s: treating couplings as scale-dependent variables whose flows and fixed points organize both critical phenomena and quantum field theories.

1 The Puzzle, Stated With a Real Number

1.1 Magnetism’s Old Mystery

Ask a physicist for the electron’s real charge, and you get roughly the same kind of answer a mathematician gives if you ask for the last digit of π : there isn’t one. Not because nobody has measured carefully enough, but because the question is the wrong shape. There is no final digit of π waiting to be found – there is only a never-terminating sequence of ever-closer rational approximations, and that sequence, in its entirety, *is* the answer. It will take us the length of this note to earn it properly (the payoff lands in Section 7), but that turns out to be almost exactly the right way to think about how a physical coupling – an electron’s charge, a magnet’s tendency to align – behaves once you take Kenneth Wilson’s ideas seriously.

The story, however, starts somewhere much more mundane: a piece of iron, and a question people had been asking about it since long before anyone knew what an atom was.

Magnetism defied explanation for centuries. People had known since antiquity that a lodestone attracts iron, and by the nineteenth century it was clear that heating a magnet strongly enough

makes the magnetism vanish – Pierre Curie documented this in 1895 – but nobody had a theory that predicted *how* the magnetism disappears, only that it does.

The first serious attempt at such a theory came from Pierre Weiss in 1907. His idea was simple and appealing: imagine that each atomic magnet feels an effective field produced by the average alignment of all its neighbors, so that a little bit of order reinforces itself and grows, below some critical temperature, while thermal jostling wins out and destroys it above that temperature. It is a good idea, and – as we will see over the course of this note – it turns out to be almost exactly the *wrong* good idea. The nature of that wrongness, in plain terms and well ahead of the math: Weiss’s picture treats every atom as if it only ever felt the *average* environment produced by its neighbors, never their actual, fluctuating, moment-to-moment values, and never any sense of how far away those neighbors are. Understanding why that omission matters, and what has to replace it, is what won Kenneth Wilson his Nobel Prize some seventy-five years later – and, as a bonus few readers see coming from a story that starts with iron and lodestones, the same idea Wilson worked out for magnets turned out, a few years later, to be just as important for particle physics, quietly setting the rules for how the fundamental forces themselves behave at different energies. We return to that connection in Section 6.

1.2 A Familiar Cousin: van der Waals

Before going further: most readers have already met a close cousin of this problem, quite possibly in a high school chemistry class – the van der Waals equation of state,

$$\left(P + \frac{a}{V^2}\right)(V - b) = RT, \quad (1)$$

the correction to the ideal gas law that accounts for the finite size of molecules (b) and their mutual attraction (a). That equation, from van der Waals’s 1873 doctoral thesis, predicts something that rarely gets emphasized in the chemistry-class version of the story: a specific pressure and temperature at which the distinction between liquid and gas disappears entirely – the liquid-gas critical point – in the same qualitative way a magnet’s order disappears at its Curie point.

This isn’t a coincidence or a loose analogy. Weiss’s theory of magnets and van der Waals’s theory of fluids are, when written out in full, the same mathematical theory wearing two different costumes – both special cases of a more general framework that Lev Landau put on a firm footing in 1937, entirely independently of either specific system. The same mathematics – Landau’s free-energy expansion, given in full in the next subsection – describes both, and the exponent we are about to compute is, quite literally, the same number in both systems: water and steam becoming indistinguishable at the critical point, and iron losing its magnetism at the Curie point, are the same phenomenon underneath.

1.3 Landau’s Guess, and a Wrong Number

Heat a piece of iron above a critical temperature $T_c \approx 1043\text{ K}$ and its permanent magnetism disappears entirely. Cool it back down through T_c and the magnetization M reappears continuously from zero, growing as

$$M \propto (T_c - T)^\beta, \quad T < T_c, \quad (2)$$

for some exponent β . This is not a peculiarity of iron; the same functional form, with the same exponent, governs the density difference between liquid water and steam near their critical point, and the magnetization of dozens of unrelated ferromagnetic materials. That universality – the same number, regardless of the material – is itself a clue we will return to at the end of this note: it is a

hint that what actually determines β is something about large-scale behavior common to all these systems, not any microscopic detail specific to iron, water, or any other individual material.

Landau's 1937 theory is deliberately *phenomenological*, and that choice matters beyond just whether the theory is correct. Landau did not start from any microscopic model of atoms, spins, or molecules; he wrote down a free energy using only symmetry and smoothness assumptions, so that the same equation applies equally to Weiss's magnet, van der Waals's fluid, or an alloy, without committing to any of their microscopic details.¹ (Landau's theory is also chronologically *later* than Ising's specific microscopic spin model, which we meet in Section 2 – Ising published in 1925, Landau in 1937 – even though Landau's framework is the more general, model-independent one of the two. The precise link between them, made concrete rather than assumed, is the first thing we do in Section 2.)

Landau's argument, given this generality, is short enough to give in full. Write the free energy as an expansion in powers of the magnetization M , keeping only the terms allowed by symmetry (no odd powers, since flipping every spin should cost the same energy):

$$F(M) = R M^2 + U M^4, \quad U > 0, \quad (3)$$

where R and U are treated as smooth (analytic) functions of temperature – meaning, concretely, that they are given by an ordinary, convergent power series in T , with no dependence at all on the length scale over which you are looking at the system – an assumption Section 3 onward will directly challenge – with R changing sign at T_c : $R \propto (T - T_c)$. Minimizing F over M gives $M = 0$ for $R > 0$ (above T_c , no spontaneous magnetization) and, for $R < 0$ (below T_c),

$$0 = \frac{\partial F}{\partial M} = 2RM + 4UM^3 \implies M = \sqrt{\frac{-R}{2U}} \propto (T_c - T)^{1/2}. \quad (4)$$

Landau's theory, in other words, predicts $\beta = 1/2$. Flat out.

For nearly forty years, this is what physicists believed. Landau's theory was elegant, general, and derived from nothing more exotic than symmetry – there was no obvious reason to distrust it, and every qualitative feature it predicted (a sharp transition, a continuously vanishing order parameter) matched experiment perfectly. Then the measurements got precise enough to pin down the exponent itself, and they kept insisting on something else entirely.

Here is the number that started this whole story: real experiments give $\beta \approx 0.32$ to 0.36 depending on the material, with the most careful measurements (on fluids very close to their critical point, where the analogous exponent is easiest to isolate) pinning it at $\beta = 0.325 \pm 0.005$. That is not experimental noise sitting on top of $1/2$. It is a different number, reproducibly, across every material anyone has ever measured – and Landau's derivation above contains no error. Every step of (3)–(4) is correct, given its assumptions. The discrepancy was not one bad experiment, or one awkward material; it was everywhere, in every substance anyone tried, and it stayed put for decades. Something fundamental had to be missing, and the problem is buried in an assumption so natural nobody thought to question it for the better part of forty years: that R and U are ordinary,

¹Landau was as rigorous as he was self-assured. Landau required his own students to pass a legendarily difficult qualifying exam, the “Theoretical Minimum,” before he would work with them – only about forty-three candidates passed it between the 1930s and 1961 – and Leonard Susskind later borrowed the name outright for his own popular physics lecture series, crediting Landau explicitly. Landau also kept a running, semi-serious logarithmic scale ranking the physicists of his century by sheer genius, and rated himself on it too: a 2.5 for most of his career, a rank he later revised to a 2 after his own Nobel Prize – literally re-scoring himself as the superior scientist. Neither extraordinary mathematical ability nor extraordinary confidence was enough to reveal the assumption this section is about to expose. Wilson's breakthrough was not greater technical virtuosity, but seeing that the theory itself had been asking the wrong question.

smoothly-varying (analytic) functions of temperature at all length scales. Readers with a modeling background in other fields will recognize the shape of this failure: it is the same way a mean-field model of an epidemic or a social network can get the qualitative story right (there is a threshold, above which an outbreak or a cascade takes off) while getting the exponents governing how it takes off wrong, because the mean-field approximation ignores correlations between neighboring individuals in the same way Section 1.1 flagged Weiss ignoring them between neighboring atoms. Section 3 shows where that assumption fails, and Section 7 shows what has to replace it.

2 The Ising Model, and Its Cousin QUBO

To go further we need a concrete microscopic model, not just a free energy written down by appeal to symmetry. The simplest one that captures the physics above is the *Ising model*, introduced by Ising in 1924 (who, in one dimension, promptly proved it does *not* have a phase transition at all – a fact we will use directly in Section 3).

Place a variable $s_i \in \{-1, +1\}$ (“spin up” or “spin down”) at every site i of a lattice, and define an energy

$$\mathcal{H}(\{s_i\}) = -J \sum_{\langle i,j \rangle} s_i s_j, \quad (5)$$

where the sum runs over neighboring pairs of sites and $J > 0$ favors neighboring spins pointing the same way (ferromagnetism). At temperature T , the probability of a given spin configuration is given by the Boltzmann distribution, and it is conventional to absorb temperature into a single dimensionless coupling

$$K \equiv \frac{J}{k_B T}, \quad (6)$$

so that large K means low temperature (spins strongly prefer to align) and $K \rightarrow 0$ means high temperature (spins are essentially random). The partition function is

$$Z = \sum_{\{s_i\}} \exp\left(K \sum_{\langle i,j \rangle} s_i s_j\right), \quad (7)$$

summed over all 2^N configurations of N spins. Every thermodynamic quantity of interest – magnetization, free energy, the exponent β itself – is in principle extractable from Z . The entire difficulty of statistical mechanics is that this sum is intractable to do directly for any realistic N .

2.1 Where Landau’s Free Energy Actually Comes From

Section 1.3 introduced Landau’s free energy (3) as a phenomenological guess, justified only by symmetry. We can now make the connection to the Ising model concrete, rather than merely asserted, by doing the simplest possible approximation to (7): assume every spin sees the same average magnetization m as its neighbors (the mean-field, or Bragg-Williams, approximation), so that a lattice with N spins and q neighbors per site has, on average, $N(1+m)/2$ spins pointing up and $N(1-m)/2$ pointing down. This gives an (approximate) free energy per spin, in units of $k_B T$, combining the mean-field interaction energy with the exact entropy of $N(1 \pm m)/2$ independently-placed spins:

$$f(m) = -\frac{1}{2}qK m^2 + \frac{1+m}{2} \ln \frac{1+m}{2} + \frac{1-m}{2} \ln \frac{1-m}{2}. \quad (8)$$

Expanding the entropy term in powers of m (a standard Taylor series, verified directly by symbolic expansion) gives

$$\frac{1+m}{2} \ln \frac{1+m}{2} + \frac{1-m}{2} \ln \frac{1-m}{2} = -\ln 2 + \frac{1}{2}m^2 + \frac{1}{12}m^4 + O(m^6), \quad (9)$$

so that, dropping the m -independent constant,

$$f(m) = \frac{1}{2}(1-qK)m^2 + \frac{1}{12}m^4 + O(m^6). \quad (10)$$

This is Landau's form (3), with $R = \frac{1}{2}(1-qK)$ and $U = \frac{1}{12} > 0$ – not assumed, but derived from the microscopic Ising Hamiltonian under one explicit, clearly stated approximation: treat each spin's neighbors as if they were frozen at their average magnetization, rather than fluctuating individually. Landau's phenomenological framework and the Ising model are, after this approximation, the same object. R changes sign, as required, at $qK = 1$, giving a mean-field critical coupling

$$K_c^{\text{MF}} = \frac{1}{q}. \quad (11)$$

For a two-dimensional square lattice, $q = 4$ nearest neighbors, so $K_c^{\text{MF}} = 0.25$. The exact value, from Onsager's celebrated 1944 solution (Section 4), is $K_c = 0.4407$ – mean-field theory is not just wrong about the exponent β , it is wrong about the transition point itself, by more than forty percent. Put in terms that matter for building something rather than just computing a number: a 40% error in predicting the temperature at which a phase transition happens is not a rounding error you can absorb with a calibration constant; it means the mean-field model has fundamentally misjudged how much fluctuations change the physics, and no amount of re-fitting a single parameter will fix a model that is missing an entire physical effect. Both errors have the same root cause: the mean-field approximation above threw away fluctuations entirely (every spin sees only the *average* of its neighbors, never their actual, fluctuating values), which is the same flat, scale-independent assumption Section 1.3 flagged in Landau's original reasoning. Section 3 begins fixing this.

The QUBO connection. Readers who have encountered Quadratic Unconstrained Binary Optimization (QUBO) problems – the standard input format for quantum annealers and many classical heuristics – already know this object under a different name. A QUBO problem asks for binary variables $x_i \in \{0, 1\}$ minimizing

$$E(\{x_i\}) = \sum_{i \leq j} Q_{ij} x_i x_j, \quad (12)$$

for some given matrix of coefficients Q_{ij} . The substitution

$$x_i = \frac{1 + s_i}{2} \quad \Longleftrightarrow \quad s_i = 2x_i - 1 \quad (13)$$

converts (12) into an Ising Hamiltonian of the form (5) plus a (physically irrelevant) constant and a set of single-site “field” terms, and vice versa: any Ising model can be rewritten as a QUBO. This is not an analogy; it is a change of variables, exact term by term. Finding the ground state of an Ising Hamiltonian – the configuration that minimizes (5), i.e. the $T \rightarrow 0$ (large- K) limit of the physics above – is, after (13), literally the same computational problem as solving the corresponding QUBO.

What is actually quantum here? A quantum annealer builds a lattice of coupled qubits whose interaction strengths are programmed to match your Q_{ij} matrix, starts every qubit in a known

ground state, and slowly deforms the Hamiltonian toward the one encoding your problem. Classical simulated annealing, faced with a rugged energy landscape, can only climb over a barrier between two local minima; a quantum annealer can, under the right conditions, tunnel through it directly. Gate-model quantum computers take a different route to the same objective: the Quantum Approximate Optimization Algorithm (QAOA) encodes the Ising energy as the cost function of a parameterized circuit, tuned by a classical optimizer. Different hardware, different mechanism, same underlying object – which is why quantum annealers, coherent Ising machines, and QAOA are all, at bottom, aimed at the same NP-hard problem.

One caveat matters for what follows: a general QUBO matrix Q_{ij} typically produces both single-site fields *and* non-nearest-neighbor couplings once translated via (13). Our field-free, nearest-neighbor ferromagnet (5) is the simplest special case – chosen deliberately because it is the one this note can solve by hand. Section 4.4 returns to why that simplicity is not merely cosmetic.

	Ising language	QUBO language
Variable	$s_i \in \{-1, +1\}$	$x_i \in \{0, 1\}$
Objective	minimize $\mathcal{H}(\{s_i\})$	minimize $E(\{x_i\})$
Dictionary	$s_i = 2x_i - 1$	

We will stay in Ising language for the rest of this note, both because it is Wilson’s own language and because the ± 1 variables make the symmetry argument behind (3) more transparent. Every equation below has a QUBO-language counterpart obtained mechanically from (13).

3 An Exact, Hand-Computable Renormalization Group

Wilson’s insight, at its core, is a strategy for computing Z in (7) without ever summing over all 2^N configurations at once: sum over *some* of the spins exactly, producing a new, smaller problem of the same form but with a renormalized coupling K' ; then repeat. We will do this exactly – no approximation of any kind – for the one-dimensional Ising chain, where it happens to be possible to carry every step out by hand. Our goal in this section is to see, in the simplest setting available, how coarse-graining turns a fixed coupling into a scale-dependent one, and how the fixed points of that flow encode whether a phase transition exists at all. Everything in this note, through Section 4, is *real-space* RG – coarse-graining directly on the lattice, spin by spin. Wilson’s own momentum-space formulation, working in Fourier space with a continuum field and a momentum cutoff, is the same idea in different clothes, more powerful for the ϵ -expansion of Section 5; the two pictures are related by a Fourier transform, not by any difference in physical content.

Operations researchers do a version of this instinctively: we aggregate warehouses into regions, individual customers into demand classes, and continuous time into planning buckets, because the finest-grained description is too unwieldy to reason about directly. Wilson’s insight was to do the same thing to a physical system – provided you also let the model’s own parameters change as you coarsen it, rather than pretending the aggregated model has the same fixed constants as the original. That single proviso, which sounds almost too simple to matter, is the entire content of what follows.

3.1 Setting up the decimation

Consider a chain of spins $s_1, s_2, s_3, \dots$ with Hamiltonian $K \sum_i s_i s_{i+1}$ (we have absorbed the sign of J into K ; nothing below depends on it). Group the spins into odd-indexed $(s_1, s_3, s_5, \dots)$ and

even-indexed ($s_2, s_4, \dots$), and sum out every even-indexed spin exactly, holding the odd-indexed spins fixed. Figure 1 shows the picture: each even spin sits between two odd neighbors, and it only ever appears in the Hamiltonian through its coupling to those two neighbors, so it can be summed out one at a time.

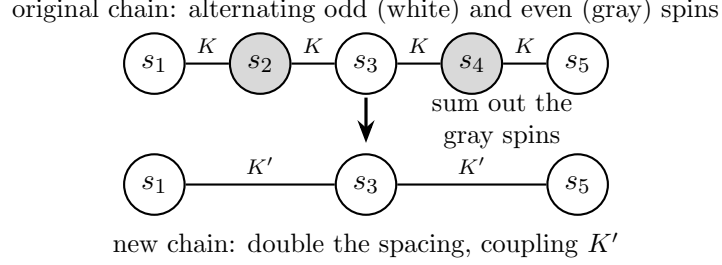


Figure 1: One decimation step. Summing the even-indexed (gray) spins out of the partition function exactly reproduces the original model on the odd-indexed spins alone, at a new coupling K' and double the lattice spacing.

Focus on one even spin, s_2 , coupled to its two odd neighbors s_1 and s_3 .² Its contribution to the partition sum, summed over its two possible values, is

$$\sum_{s_2=\pm 1} \exp [K s_1 s_2 + K s_2 s_3] = \exp [K(s_1 + s_3)] + \exp [-K(s_1 + s_3)] = 2 \cosh [K(s_1 + s_3)]. \quad (14)$$

This expression depends on s_1 and s_3 only through their sum, and by the symmetry $s_1, s_3 \in \{-1, +1\}$ there are really only two distinct cases:

$$s_1 = s_3 \text{ (aligned)} : \quad s_1 + s_3 = \pm 2 \implies 2 \cosh(2K), \quad (15)$$

$$s_1 = -s_3 \text{ (anti-aligned)} : \quad s_1 + s_3 = 0 \implies 2 \cosh(0) = 2. \quad (16)$$

We now ask: is there a new coupling K' and an overall constant A (a multiplicative renormalization of the free energy, physically irrelevant to K' 's flow) such that

$$A e^{K' s_1 s_3} = 2 \cosh [K(s_1 + s_3)] \quad (17)$$

reproduces (15)–(16)? Since $s_1 s_3 = +1$ in the aligned case and -1 in the anti-aligned case, (17) demands

$$A e^{K'} = 2 \cosh(2K), \quad A e^{-K'} = 2. \quad (18)$$

Two equations, two unknowns (A and K') – and this system solves exactly, with no approximation:

$$\boxed{e^{2K'} = \cosh(2K) \iff K' = \frac{1}{2} \ln [\cosh(2K)]}. \quad (19)$$

(One can check, by dividing the two equations in (18) the other way and using standard hyperbolic identities, that (19) is exactly equivalent to the perhaps more familiar-looking form $\tanh K' = (\tanh K)^2$; both give identical numbers, verified below to six decimal places.)

²This calculation is the standard textbook entry point to real-space renormalization – essentially the same derivation appears in Kadanoff's own writing, in Nelson and Fisher's classic treatment [Nelson and Fisher, 1975], and worked in comparable detail in Walker's book [Walker, 2023]. It is used here for the same reason introductory electromagnetism always derives Gauss's law for a sphere before anything harder: it is the simplest instance where every step is exact, not because it is a novel calculation.

Equation (19) *is* a renormalization group transformation: an exact rule for how the coupling changes when you integrate out short-distance degrees of freedom (here, every other spin) and rescale. Nothing about its derivation involved approximation, truncation, or an expansion parameter. This is the real thing, in miniature. On one page, we have carried out the same conceptual operation for which Wilson won the Nobel Prize: we replaced a microscopic description with a coarser one while preserving the observable physics exactly, and in doing so, the coupling itself changed, from K to K' . That is renormalization, in full, not a sketch of it. Wilson’s great leap was to realize that this simple-looking recursion was not a trick peculiar to the Ising model – it was a new language for organizing physical theories themselves, one in which a coupling’s “true” value is not a fact about the world but a question that only makes sense once you’ve said at what scale you’re asking it. In plain terms: Equation (19) answers the question “if I only care about the physics at twice the original spacing – a coarser resolution, one level up – what effective coupling K' reproduces the same probabilities for the spins that remain?” Readers with a background in multiscale or hierarchical modeling will recognize the shape of this question even before seeing the answer.

3.2 Running the numbers

Equation (19) is a single line any reader can put into a calculator, and I’d genuinely encourage you to: there is a particular satisfaction in watching a Nobel-Prize-winning idea reduce to six taps on a \$10 calculator, and it is not an experience popular accounts of RG ever let you have. Starting from $K_0 = 1$ (an arbitrary but not small coupling – comfortably in the “low temperature, spins want to align” regime) and iterating $K_{n+1} = \frac{1}{2} \ln[\cosh(2K_n)]$ gives Table 1.

Iteration n	K_n
0	1.000000
1	0.662501
2	0.350061
3	0.113672
4	0.012812
5	0.000164
6	≈ 0

Table 1: The exact decimation recursion (19), iterated from $K_0 = 1$. Every entry is reproducible on an ordinary calculator: $K_1 = \frac{1}{2} \ln[\cosh(2)] = \frac{1}{2} \ln(3.7622) = 0.662501$, and so on.

Every finite starting coupling K_0 , however large, flows to the same place: $K^* = 0$. Picture it physically before the RG language: imagine looking at the chain from farther and farther away, each step in Table 1 corresponding to doubling your viewing distance. Every time you double it, the effective interaction between what you can now resolve gets weaker. Zoom out enough, and the chain looks like a random, uncorrelated string of spins, no matter how strongly aligned its near neighbors actually were up close – Table 1’s numbers just put that intuition on a calculator. In RG language, $K^* = 0$ is a *fixed point* of the transformation (19) – in plain terms, a parameter value that does not change under coarse-graining: plug it in and it maps to itself, $\frac{1}{2} \ln[\cosh(0)] = 0$. It is also an *attractive* fixed point, meaning that if you start from almost any finite K and repeatedly coarse-grain, you get pulled into that fixed point rather than pushed away from it – so the fixed point, not the specific starting value of K , is what controls the long-distance physics. If you find it more natural to think in terms of code than equations: if you wrote a short program that iteratively

coarsens a 1D Ising chain and updates K by (19) at each step, you would see the same flow shown in Table 1, for any finite starting K you tried.

This is not a disappointing answer; it is the *correct* one, and here is why. Ising proved, by direct (non-RG) methods, back in 1924, that the one-dimensional chain has no spontaneous magnetization at any temperature $T > 0$ – ordering survives only exactly at $T = 0$, i.e. $K = \infty$. The flow in Table 1 is telling us this fact, extracted purely from the structure of the recursion, with no separate calculation of the magnetization required: because every finite K flows to the disordered fixed point $K^* = 0$ under coarse-graining, the physics looks disordered at long distances for every $T > 0$. The RG machinery has correctly reproduced a fact independently known to be true.

Why is this true microscopically? The RG flow confirms it but doesn’t by itself explain it. Take an ordered chain – every spin aligned – and flip one entire contiguous stretch of it, however long. That costs a fixed, finite amount of energy no matter how long the flipped stretch is: you only pay at the two boundaries, the two “domain walls” where aligned meets flipped, and nothing in between. But the *number of places* you could put those two domain walls grows with the length of the chain – there are far more ways to flip a long stretch somewhere in a long chain than a short one. Energy cost stays fixed; entropy (the number of ways to do it) grows without bound. At any temperature above absolute zero, however low, that growing entropy eventually overwhelms the fixed energy cost, and the chain is better off (lower free energy) full of domain walls than uniformly ordered. This is the microscopic reason order cannot survive in one dimension, and it is also a preview of what changes in two: flipping a whole *patch* of a 2D lattice costs energy proportional to the patch’s *perimeter*, not a fixed two-wall cost, and that changes the energy-versus-entropy competition enough to allow a real phase transition to survive. Section 4 makes that difference precise.

4 A Real Phase Transition: The 2D Ising Model

Section 3’s calculation was exact, but in one sense it was too easy: the one-dimensional chain has exactly one fixed point, $K^* = 0$, and it governs the physics at every temperature. There is no phase transition to find in one dimension – only a confirmation that none exists. Real ferromagnets live in two or three dimensions, and there the story changes: a *second*, nontrivial fixed point appears, corresponding to the critical temperature T_c of Section 1, and it separates two genuinely different behaviors – spins that end up aligned (an *ordered* phase, nonzero average magnetization) from spins that end up random (a *disordered* phase, zero average magnetization). A phase transition is the sharp switch between the two as temperature crosses T_c , and its signature in RG language is exactly this second fixed point. This section stays in two dimensions, where Onsager’s exact 1944 solution gives us something to check our work against, and builds the simplest hand-computable extension of Section 3’s decimation trick that can actually produce one.

4.1 Why 1D’s trick doesn’t carry over directly

Here is the obstacle. In one dimension, summing out a spin only ever generates a new coupling between the two spins on either side of it – decimation stays inside the world of nearest-neighbor chains forever, which is why (19) was a single clean equation. In two dimensions, summing out one spin couples *all* of its remaining neighbors to each other: next-nearest-neighbor terms appear, then four-spin terms, and the list keeps growing at every step. Chasing that proliferation exactly is a real undertaking – it is what forced Onsager into a famously difficult exact solution, and later forced Wilson to build the systematic ϵ -expansion machinery referenced at the end of this section.

There is, however, a much cheaper route to a genuine phase transition – two fixed points instead of one – using nothing but a calculator, due to Migdal and Kadanoff.

4.2 The Migdal-Kadanoff approximation

The trick is to turn the 2D lattice back into something Section 3 already knows how to handle – a 1D chain – by physically relocating bonds before decimating. Figure 2 shows the move: take a vertical bond and slide it sideways until it lies exactly on top of a horizontal bond. Two bonds occupying the same edge, each pulling with strength K , combine exactly like two identical springs wired in parallel: their couplings simply add, $K + K = 2K$. Do this once for every vertical bond in the lattice, and what is left looks, edge for edge, like the one-dimensional chain we already solved – except every bond now carries strength $2K$ instead of K .

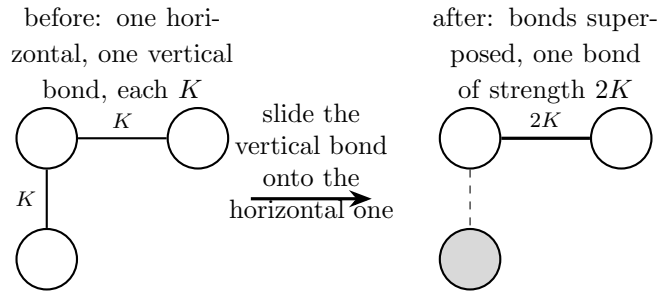


Figure 2: The bond-moving step. A vertical bond is slid sideways until it lies exactly on top of a horizontal bond; since both then connect the same two sites, their couplings simply add, $K + K = 2K$. Repeating this across the lattice turns a two-dimensional problem into a chain of effective bonds $2K$, which Equation (19) already knows how to decimate.

This move is the approximation, and it is worth being plain about that: sliding bonds around a lattice is not something you get to do for free in general, and no error bound comes attached to it. What you get in exchange is the one-dimensional recursion (19) from Section 3, now run at $2K$ in place of K :

$$K' = \frac{1}{2} \ln[\cosh(2 \cdot 2K)] = \frac{1}{2} \ln[\cosh(4K)]. \quad (20)$$

That one line – still nothing but a calculator – is the entire Migdal-Kadanoff recursion for the square lattice. The rest of this section finds out exactly what it costs you.

4.3 Finding the fixed point, and checking it against the exact answer

Unlike (19), Equation (20) has *two* fixed points where $K^* = \frac{1}{2} \ln[\cosh(4K^*)]$: the trivial $K^* = 0$ (disordered, as in 1D), and a second, nontrivial solution found numerically to be

$$K_{\text{MK}}^* = 0.3047. \quad (21)$$

Table 2 shows the flow on both sides of this new fixed point, starting from couplings just below and just above it.

This is the first genuinely new phenomenon the renormalization group discovers in this note, and here is why it matters. One dimension, in Section 3, could only ever tell us why no phase transition exists there – a confirmation, not a discovery. Two dimensions reveal something the one-dimensional case was structurally incapable of producing: the appearance of a second, nontrivial fixed point,

starting below K_{MK}^*		starting above K_{MK}^*	
Iteration	K_n	Iteration	K_n
0	0.2547	0	0.3547
1	0.2241	1	0.3913
2	0.1786	2	0.4573
3	0.1180	3	0.5808
4	0.0537	4	0.8199
5	0.0115	5	1.2939
6	0.0005	6	2.2412

Table 2: The Migdal-Kadanoff recursion (20), started 0.05 below and 0.05 above the nontrivial fixed point $K_{\text{MK}}^* = 0.3047$. Below it, the flow runs to $K = 0$ (disorder); above it, the flow runs away toward $K = \infty$ (complete order). This is the qualitative structure of a real phase transition, and it is something the one-dimensional case of Section 3 could never produce, no matter how it was iterated.

separating two genuinely distinct phases of matter. That mathematical object – a fixed point that is neither the trivial disordered one nor a runaway to infinity, but a balance point the flow can approach from either side – is what Landau’s theory, built entirely without reference to scale, could never produce, no matter how carefully its free energy was written down. Wilson himself offered a useful physical picture of this structure in a 1979 piece written for a general audience [Wilson, 1979]: think of every possible combination of couplings as a point on an undulating landscape, with coarse-graining playing the role of gravity, always pulling the point downhill. Almost every starting point rolls into one of two valleys – complete order, or complete disorder – but one very specific location, a mountain pass connecting two peaks, lets a point linger indefinitely without committing to either side. That pass is the nontrivial fixed point; a system held at its critical temperature is a marble balanced right on the ridge.

It is just as important to be clear about what this rough approximation gets *right* as about what it gets wrong: it correctly identifies that a phase transition exists at all, correctly places it as separating an ordered low-temperature phase from a disordered high-temperature one, and correctly reproduces the qualitative shape of the flow – all of that structure is not in doubt. What it cannot be trusted for is the precise *location* of the transition, which is the same number Section 2.1’s cruder mean-field theory also got wrong, if by a larger margin. A real-space calculation can in fact do much better with more care: Wilson’s own square-lattice spin-decimation calculation, described in the same 1979 article, tracked on the order of two hundred separate couplings rather than the single K we followed above, and came out within roughly 0.2% of Onsager’s exact answer [Wilson, 1979] – a striking demonstration that the gap between our 31% and the exact result is a matter of how much of the proliferating coupling structure you are willing to track, not a fundamental limitation of the real-space approach itself.

Where does a number like $K_c = 0.4407$ actually come from? Onsager’s method builds a *transfer matrix* T : a matrix that advances the whole lattice by one row of spins at a time, so the partition function becomes a matrix product, $Z = \text{Tr}(T^N)$, and the free energy per spin is controlled entirely by T ’s largest eigenvalue. Diagonalizing T exactly for an infinite 2D lattice is a substantial undertaking in its own right – Onsager did it via a clever mapping onto free fermions, and the calculation is famous for being difficult even by the standards of exact solutions – but the strategy itself is simple to state: find how T ’s largest eigenvalue depends on K , and the free energy, the critical point, and the exponents all fall out of it. We will not diagonalize T here; we will only use

what falls out of having done so.

The comparison is this: Onsager’s exact 1944 solution gives $K_c = 0.4407$ – the value satisfying $\sinh(2K_c) = 1$, a remarkably clean condition that follows from an exact symmetry (Kramers-Wannier duality) relating the square-lattice Ising model at coupling K to the same model at a different coupling. The intuition is closer to a mirror than to an equation: duality says the physics at strong coupling (very ordered, low temperature) and the physics at some corresponding weak coupling (very disordered, high temperature) are secretly the same calculation wearing two different disguises. Most values of K get mapped to some *other* value under this mirror – but the critical point is the one special coupling that gets mapped to itself, the one temperature at which the mirror image is indistinguishable from the original. That self-consistency is what pins K_c down to a single clean number rather than leaving it to be found by brute-force search. With K_c in hand, the exact solution also gives a magnetization exponent $\beta = 1/8 = 0.125$, for the two-dimensional square-lattice Ising model. Our Migdal-Kadanoff estimate, $K_{MK}^* = 0.3047$, is about 31% low. That gap is real, and here is why: the bond-moving step is an uncontrolled approximation – there is no small parameter that guarantees it gets better and better, the way, say, a well-chosen perturbative expansion would. Wilson’s own Nobel lecture cites Migdal-Kadanoff bond-moving explicitly as one of several early real-space approaches that were useful for getting the qualitative picture right without being quantitatively reliable. What it buys you, cheaply, is what Section 2.1’s mean-field theory could not: a genuine two-fixed-point flow, obtained from a real (if approximate) coarse-graining argument rather than from throwing fluctuations away entirely.

4.4 Why Planarity Matters: A Bridge Back to QUBO

Section 2 noted that solving for an Ising ground state and solving the corresponding QUBO are, after the substitution (13), the same computational problem. How hard is that problem, actually? Two-dimensional Ising models sit on a genuinely special, sharp boundary of difficulty – one that has nothing to do with the approximate real-space RG of this section, and everything to do with the shape of the underlying graph.

The reason Onsager could solve the 2D square-lattice Ising model exactly at all, with no magnetic field, isn’t a numerical accident: it’s a deep structural fact about *planar* graphs – graphs, like the square lattice, that can be drawn on a flat page with no two edges crossing. Kasteleyn [Kasteleyn, 1961] and, independently, Fisher [Fisher, 1966] showed that the full partition function (7) of *any* Ising model on *any* planar graph, with *any* pattern of coupling strengths J_{ij} on its edges (not just a uniform K on a regular lattice), can be computed exactly, in an amount of computation that grows only polynomially with the number of spins N – by an ingenious mapping onto the problem of counting perfect matchings (dimer coverings) of an auxiliary planar graph, solved via a determinant (a Pfaffian). Onsager’s original solution predates this general result and used a different method entirely, but the Kasteleyn-Fisher machinery is what tells us *why* an exact solution had to exist for a whole family of two-dimensional problems, not merely for the one particular lattice Onsager happened to solve.

This tractability, however, is fragile in two specific ways, both made precise by Barahona [Barahona, 1982]: turn on an external magnetic field, or make the underlying graph non-planar (allow even a single pair of crossing edges – as happens, generically, the moment you go to three dimensions, or add even one long-range or next-nearest-neighbor coupling to a 2D lattice), and the same problem – computing the ground state, or the partition function – becomes NP-hard. This is a genuinely sharp dichotomy, not a fuzzy statement about difficulty in degree: planar and field-free is polynomial; almost anything else is (believed to be) exponentially hard in the worst case.

This is why it’s worth building quantum annealers, coherent Ising machines, and specialized

classical solvers around QUBO in the first place. A general QUBO instance – an arbitrary Q_{ij} matrix, which after (13) becomes an Ising model with both a field *and* a generically non-planar coupling graph – sits squarely in Barahona’s hard regime. Our discussion in this section, the clean, hand-checkable 2D square lattice with no field, is the tractable exception rather than the typical case: it’s solvable because it’s planar and field-free, and every real optimization problem QUBO gets used for in practice lives outside that narrow island. Wilson’s renormalization group and Kasteleyn’s exact planar solution are two different responses to the same underlying difficulty (an intractably large configuration space): Wilson’s works approximately, for any dimension and any coupling structure, by coarse-graining; Kasteleyn’s works exactly, but only on the special planar, field-free island where the graph’s geometry does the work instead.

5 Why This Matters Beyond Magnets

Why should any of this – a made-up ferromagnet, an old chemistry-class gas law, some hand-iterated hyperbolic cosines – matter to anyone who is not a condensed-matter theorist? Two answers, one about the physics we have not yet done, and one that connects directly back to a very different and very current field.

5.1 What full accuracy actually requires

Getting from Migdal-Kadanoff’s rough 31% to genuinely accurate three-dimensional numbers ($\nu \approx 0.630$, $\beta \approx 0.325$) requires Wilson’s actual machinery: the continuum Landau-Ginzburg field theory, coarse-grained in momentum space rather than by moving lattice bonds, organized as a systematic expansion in $\epsilon = 4 - d$ (the distance of the real dimension $d = 3$ from the value $d = 4$ at which Landau’s mean-field theory becomes exact) [Wilson, 1971, Wilson and Fisher, 1972], later carried to high order by Brézin, Le Guillou, and Zinn-Justin [Brézin et al., 1976]. Unlike Migdal-Kadanoff, this expansion improves systematically order by order with a quantifiable error at each step – substantially longer than anything else in this note, and a subject for a dedicated companion note rather than a detour here, but its leading term is still one line you can check by hand, and it is worth seeing the five steps that get you there:

- Step 1. **Start with Landau.** Section 1.3 wrote the free energy as $F(M) = RM^2 + UM^4$, with R and U treated as fixed numbers.
- Step 2. **Make it a field.** Promote M to a field $\phi(x)$ that can vary in space, and add the simplest term that penalizes rapid variation, $(\nabla\phi)^2$. This is the Landau-Ginzburg-Wilson action – Landau’s free energy, restored to having fluctuations at every length scale.
- Step 3. **Solve it for a general number of field components, n .** Wilson’s RG, run on this field theory near $d = 4$ (where the ϕ^4 term is exactly marginal), finds a nontrivial fixed point and a formula for ν as a function of $\epsilon = 4 - d$ and n .
- Step 4. **Plug in $n = 1$.** The Ising model has one component per spin, so setting $n = 1$ in Wilson’s general $O(n)$ result gives the specific numbers below.
- Step 5. **Watch the coupling flow.** Iterate the one-loop equation for the coupling numerically, starting above and below the fixed point, and see it converge – the continuum analogue of Tables 1 and 2.

For the Ising universality class, the first-order result is

$$\nu \approx \frac{1}{2} + \frac{\epsilon}{12}. \quad (22)$$

Plug in $\epsilon = 1$ (that is, $d = 3$) and you get $\nu \approx 0.583$ – not exact, but already a real step up from mean-field’s flat $\nu_{\text{MF}} = 1/2$, purely from the first term of a systematic series. And because it *is* systematic, you can watch it improve: the next term in the series, computed the same way, is

$$\nu \approx \frac{1}{2} + \frac{\epsilon}{12} + \frac{7\epsilon^2}{162}, \quad (23)$$

which at $\epsilon = 1$ gives $\nu \approx 0.627$ – one more term, and the gap to the exact 0.630 has shrunk from 0.047 to 0.003. That is the difference [Brézin et al., 1976] bought over [Wilson, 1971, Wilson and Fisher, 1972]: not a new idea, just the next, honestly quantifiable term in the same expansion.

The exponent formulas above are one payoff of the fixed point Step 3 promised; another is to watch the fixed point itself get reached, the same way Tables 1 and 2 watched the lattice couplings flow. At one loop, the dimensionless quartic coupling λ – in the action $S = \int d^d x \left[\frac{1}{2}(\partial\phi)^2 + \frac{\lambda}{4}(\phi^i\phi^i)^2 \right]$ – obeys

$$\frac{d\lambda}{dt} = \epsilon\lambda - \frac{n+8}{8\pi^2}\lambda^2, \quad t \equiv -\ln\mu, \quad (24)$$

so that t increasing means flowing toward the infrared (lower μ , longer distance), exactly as $b > 1$ meant coarsening in Sections 3–4; with this sign convention, λ^* is infrared-attractive, which is exactly the convergence Table 3 shows from both sides. Setting the right-hand side to zero gives the nontrivial Wilson-Fisher fixed point,

$$\lambda^* = \frac{8\pi^2\epsilon}{n+8}, \quad (25)$$

which for $n = 1$, $\epsilon = 1$ evaluates to $\lambda^* = 8\pi^2/9 \approx 8.773$. Iterating (24) with a simple Euler step ($\Delta t = 0.05$), starting once below and once above λ^* , gives Table 3 and Figure 3.

Step	starting below ($\lambda_0 = 6.000$)	starting above ($\lambda_0 = 11.500$)
0	6.0000	11.5000
1	6.0948	11.3213
2	6.1879	11.1568
3	6.2790	11.0053
4	6.3683	10.8652
5	6.4556	10.7357
6	6.5408	10.6156

Table 3: The one-loop coupling flow (24), Euler-iterated from two starting values, one below and one above the Wilson-Fisher fixed point $\lambda^* \approx 8.773$. Both sequences move monotonically toward λ^* , though here it is a continuous flow taken in small steps, so six rows show clear convergence underway rather than arrival.

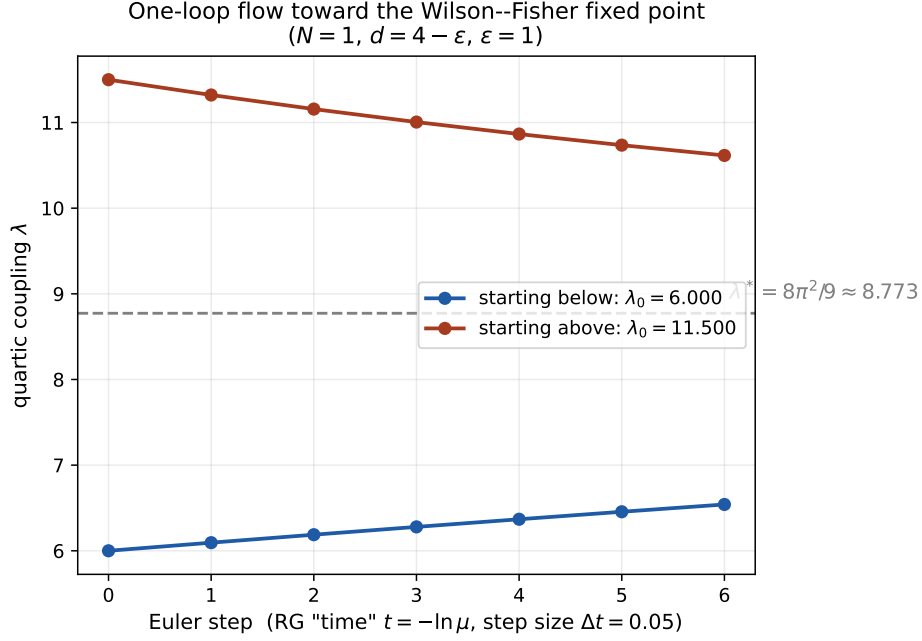


Figure 3: The same flow as Table 3, shown graphically: the coupling λ moves toward the Wilson-Fisher fixed point $\lambda^* \approx 8.773$ (dashed line) from both above and below, at one loop, for the Ising universality class ($n = 1$, $\epsilon = 1$).

5.2 Criticality as a resource, not just a curiosity

Physicists haven't cared about the divergence of the correlation length for a century out of pure curiosity; it's fast becoming an engineering resource in its own right. Modern quantum sensors are increasingly designed to operate deliberately close to a critical point, because nature herself amplifies tiny perturbations there – which is the physical fact everything in Sections 3–4 has been building toward: the divergence of the correlation length ξ as $T \rightarrow T_c$. In plain terms, the correlation length is the typical distance over which a local change – flipping one spin, slightly perturbing one atom's environment – is still felt by the rest of the system; near an ordinary temperature, that influence dies off within a few lattice spacings, but as $T \rightarrow T_c$, ξ grows without bound, so a tiny, local perturbation ends up influencing the system over macroscopic distances. A diverging correlation length means a diverging susceptibility – the system's response to an infinitesimally small external nudge grows without bound near the critical point. That is not merely a mathematical curiosity; it is the property that a growing body of work in quantum metrology has identified as a genuine sensing resource. Operating a quantum sensor close to a phase transition – exploiting the same divergent sensitivity this note has spent several sections computing exponents for – can, under the right conditions, enhance a sensor's ability to detect a weak signal, an idea now studied under the name *critical quantum metrology*, following Zanardi, Paris, and Campos Venuti's founding 2008 analysis [Zanardi et al., 2008] and a substantial subsequent literature spanning trapped ions, superconducting circuits, and cold-atom platforms [Garbe et al., 2020]. Concretely: a sensor built from atoms, spins, or superconducting circuits, tuned to sit close to its own critical point, can in principle respond more sharply to a small change in an external magnetic field, or a small shift in temperature, than an equivalent sensor operated far from criticality – the same divergent susceptibility that makes a real magnet's response to an external field blow up near T_c is, in a

deliberately engineered quantum system, the resource being exploited. Zanardi, Paris, and Campos Venuti make this concrete: for a critical system of size L , the achievable precision in estimating the parameter driving the transition improves by an extra factor of $1/L$ right at the critical point, compared to the same system away from criticality – a genuine, quantifiable scaling advantage, not merely a qualitative sense that things become more sensitive near T_c . The exponents ν and β computed by hand in this note are not idle numbers: they are the quantities that determine how fast that sensing advantage grows as a real device is tuned toward its own critical point.

6 The Relativistic Case: Couplings That Run With Energy

At first glance this may feel like a completely different subject. We have left magnets behind and entered particle physics. But mathematically we are about to do exactly the same thing: ask how a coupling changes when we change the scale at which nature is probed.

That this works at all is not an accident of notation. Wilson’s own path to the renormalization group ran through magnets and the Kondo problem, both resolutely non-relativistic, years before its implications for particle physics were understood – the connection to QED and QCD came later, once physicists recognized their own running couplings as instances of the same mathematical object. Nothing that follows involves particle creation, antiparticles, or any of the machinery Dirac and Feynman needed to survive relativity; it requires only a coupling that depends on the scale at which you probe it, exactly as before, now measured in energy instead of length.

6.1 The one-loop QED formula, stated and used

In quantum electrodynamics (QED), the electric charge measured in a low-energy experiment – a capacitor on a laboratory bench, say – is not the same number as the charge that governs a collision at much higher energy. Before writing down why, it helps to have a picture in mind, the same way Section 3 gave you the picture of zooming out on a spin chain before handing you the recursion. Imagine trying to measure an electron’s charge from far away. The electron is never truly alone: quantum mechanics allows short-lived electron-positron pairs to flicker in and out of existence in the surrounding vacuum, and these virtual pairs orient themselves to partially cancel the electron’s field, screening it a little, the way a cloud of oppositely-charged dust would. This isn’t hand-waving; a useful way to picture it is that the same uncertainty principle everyone meets in an introductory quantum course, applied to energy and time rather than position and momentum ($\Delta E \cdot \Delta t \gtrsim \hbar/2$), permits extremely short-lived particle-antiparticle fluctuations: a fleeting pair can appear and annihilate again before anything could catch it in the act. The vacuum, in this picture, is not empty at all – it is a restless, boiling soup of such pairs, winking in and out of being everywhere, all the time. Here is the twist that separates Wilson’s picture from the one it replaced – and the intuitive-sounding alternative is wrong: it is tempting to think that if you could only probe closely enough – with enough energy – you would eventually punch through the screening cloud and find a single, “bare,” unscreened electron at the center, its one true charge finally exposed. Wilson’s renormalization group says no such core exists: however closely you look, you find only a new description of the electron valid at that particular magnification, with no solid particle waiting underneath. The scale-dependent charge in (26) below is not a mask hiding some truer number; it *is* the true number, at that scale, and demanding its value at infinite energy is exactly as ill-posed as demanding the last digit of the irrational number from Section 1.1. A low-energy probe, approaching from far away, sees the charge after most of that screening cloud has had its effect; a high-energy probe, capable of resolving much shorter distances, punches through more of the cloud and sees a slightly larger, less-screened charge underneath. In effect, the measured coupling

depends on how closely you look – exactly the conceptual move Sections 3 and 4 made for magnets, now expressed in the language of quantum fields: coarse-graining a spin chain hides microscopic fluctuations, while probing QED at higher energy peels back a layer of vacuum screening. The mathematics is different; the organizing question – how do the effective parameters change as the resolution changes? – is identical.

The one-loop result for how the fine-structure constant $\alpha = e^2/4\pi\epsilon_0\hbar c$ depends on the energy scale Q at which it is probed is a standard result of quantum field theory, and we simply state it here rather than re-derive it from Feynman diagrams (that derivation is a genuinely separate undertaking, requiring the machinery of loop integrals that is outside this note’s scope):³

$$\frac{1}{\alpha(Q)} = \frac{1}{\alpha(\mu)} - \frac{2}{3\pi} \sum_f N_{c,f} Q_f^2 \ln \frac{Q}{m_f}, \quad (26)$$

summed over every charged fermion f with mass $m_f < Q$, with electric charge Q_f (in units of e) and color multiplicity $N_{c,f}$ ($N_{c,f} = 1$ for leptons, 3 for quarks). This is the non-relativistic analogue, in QED’s own language, of the coupling flowing under coarse-graining in Sections 3–4: $\ln(Q/m_f)$ plays the role our $\ln L$ played, with energy scale replacing length scale. There, L was the length scale at which you examined the system, and coarser L meant a blunter probe; here, Q is the momentum or energy scale of a collision, and larger Q means a finer probe that resolves more of the screening cloud just described.

We can carry this out by hand for the three charged leptons, whose masses are well known, running from the laboratory value $\alpha(m_e)^{-1} = 137.036$ up to the energy of the Z boson, $M_Z = 91.19 \text{ GeV}$, a benchmark scale measured at CERN’s LEP collider:

Fermion	m_f	$\ln(M_Z/m_f)$	contribution to $1/\alpha$	running total $1/\alpha$
– (below thresholds)	–	–	–	137.036
electron	0.511 MeV	12.09	–2.566	134.470
muon	105.66 MeV	6.76	–1.435	133.035
tau	1776.86 MeV	3.94	–0.836	132.199

Read down the last column and you are watching the coupling run: start at the laboratory value $1/\alpha = 137.036$, cross the electron’s mass threshold and subtract its contribution, cross the muon’s and subtract again, cross the tau’s and subtract once more, and you land on $1/\alpha(M_Z) \approx 132.20$. This is the same kind of iteration as Tables 1 and 2 earlier – a running value updated step by step – just three steps instead of six, because here each step is a whole particle threshold rather than a fixed factor of two in length.

Here is why anyone outside particle theory would care about this specific number: this strengthening of α with Q is the correction that has to be built into predicting collision rates (cross sections) at experiments operating near the Z -boson energy – using the low-energy value $1/137$ instead of the correct, run value at M_Z would get real, measurable collider predictions wrong.

Two things stand out about this number. First, the direction and rough size are right: the coupling has strengthened, from $1/137$ toward $1/128$ -ish, purely from three known lepton masses plugged into a formula stated (not derived) above – this is the actual mechanism behind the blog-post claim that the fine-structure constant runs from $1/137$ to roughly $1/128$ at high energy.

³For readers who want the fine print: Equation (26) is the leading-log, one-loop result, equivalent up to scheme-dependent finite pieces to running in the $\overline{\text{MS}}$ scheme; it is exact in the leading-logarithm approximation but omits higher-loop corrections, which are known and included in precision electroweak fits but do not change the qualitative story below.

Second, and just as important: 132.2 is not 127.9, the actual measured value of $1/\alpha(M_Z)$, and the gap is not a mistake in the arithmetic above. The remaining contribution comes from quarks – and for the light quarks (up, down, strange), (26) cannot honestly be applied with a bare quark mass, because quarks are confined and never appear as free particles with a sharp, unambiguous mass scale in the way electrons do. The real calculation replaces the light-quark terms with experimentally measured hadronic scattering data via a dispersion relation, not a perturbative loop integral – a genuinely different (and non-perturbative) piece of physics, not papered over here. For scale: the dispersion-relation evaluations of this hadronic piece give $\Delta\alpha_{\text{had}}^{(5)}(M_Z^2) \approx 0.0276$ [Jegerlehner, 2003] – comparable in size to the entire leptonic contribution above, and enough, once combined properly, to close essentially all of the remaining gap to 127.9. The missing physics is real and quantitatively accounted for, not a loose end.

6.2 Why QCD runs the other way: asymptotic freedom

The strong coupling constant does the opposite of α : it grows at low energy and shrinks at high energy, the phenomenon called asymptotic freedom, discovered by Politzer, Gross, and Wilczek and cited explicitly in Wilson’s own Nobel lecture as one of the renormalization group’s major post-1971 payoffs. The reason is a single sign – the cleanest illustration in this whole note of how much physics can hinge on the sign of a coefficient. The strong coupling’s flow with energy scale Q is governed, at one loop, by the differential equation

$$Q \frac{d\alpha_s}{dQ} = -\frac{b_0}{2\pi} \alpha_s^2 + O(\alpha_s^3), \quad (27)$$

which makes the sign dependence transparent by inspection: since $\alpha_s^2 \geq 0$ always, the sign of $d\alpha_s/dQ$ is the *opposite* of the sign of b_0 . A positive b_0 forces α_s to shrink as Q grows; a negative b_0 would force it to grow instead. The one-loop coefficient itself is

$$b_0 = 11 - \frac{2}{3}n_f, \quad (28)$$

where n_f is the number of quark flavors light enough to be active at the energy in question (at most 6 in our universe, and the top quark is too heavy to matter below the energies where b_0 is normally evaluated, so effectively $n_f \leq 5$ in every regime of practical interest). The “11” comes from gluons interacting with themselves – a genuinely non-abelian effect with no analogue in QED, where photons carry no electric charge and cannot interact with each other this way – while the “ $-\frac{2}{3}n_f$ ” comes from quark loops, structurally the same sign and origin as the lepton loops in (26) above. For any physically realized $n_f \leq 5$: $b_0 = 11 - \frac{2}{3}(5) = 11 - 3.33 = 7.67 > 0$, and a positive b_0 is what makes the coupling *shrink* as energy increases (the opposite sign convention from (26), where the fermion loops alone – QED has no self-interacting-boson term to compete with them – always make $1/\alpha$ decrease with energy, i.e. α itself increase). Gluon self-interaction is strong enough to overwhelm the quark contribution for any realistic number of flavors; had nature had 17 or more quark flavors, b_0 would flip sign and QCD would run the same direction as QED instead.

Equation (27) can be solved exactly at one loop, the same way (19) was solved for the Ising chain: $1/\alpha_s(Q) = 1/\alpha_s(Q_0) + \frac{b_0}{2\pi} \ln(Q/Q_0)$, a straight line in $\ln Q$. Starting from the measured value $\alpha_s(M_Z) \approx 0.1179$ [Particle Data Group, 2022] and stepping up in energy by factors of ten, with $b_0/2\pi \approx 1.221$:

Q	Q/M_Z	$1/\alpha_s(Q)$	$\alpha_s(Q)$
91.19 GeV ($= M_Z$)	1	8.482	0.1179
1 TeV	10.97	11.41	0.0877
10 TeV	109.7	14.22	0.0703
100 TeV	1097	17.03	0.0587

Each row is one more decade of energy plugged into the same straight line; α_s shrinks steadily as Q climbs, the opposite direction from α 's table above. This is asymptotic freedom in numbers rather than in name: by the time a collision probes energies ten times the Z -boson mass, the strong coupling has already dropped by about a quarter.

7 What This Teaches

Let me pull the threads together, because by this point it would be easy to lose track of why any of this mattered in the first place. Four calculations, and underneath all of them, one mechanism. Section 3 gave us a fully exact, hand-computable decimation, and it did something almost embarrassingly modest: it confirmed a fact – no order in one dimension – that Ising had already nailed down by other means back in 1924. Section 4 did something braver, and got it only partly right: an approximate real-space recursion produced, for the first time in this note, an actual nontrivial fixed point and the real shape of a phase transition, off by about thirty percent from Onsager's exact answer – an error I'd rather state plainly than paper over. Section 5 took that same divergent correlation length and asked why anyone outside condensed-matter theory should care, landing on current work treating criticality itself as a quantum-sensing resource. And Section 6 took the identical idea and ran it in energy instead of length, correctly explaining both the direction and, for three known leptons, a good chunk of the size of the fine-structure constant's climb with energy – plus the one sign that makes the strong coupling do the opposite.

Here is the shift in perspective, laid out side by side, the way you would lay an old policy next to a new one:

Feature	Old View (Bethe, Schwinger, Feynman)	Modern View (Wilsonian RG)
The infinities	A flaw in the mathematics, to be cleverly subtracted away	A symptom of using a low-energy theory at infinitely high energy
Bare parameters	Real, unmeasurable numbers hidden behind the physical ones	Cutoff-dependent bookkeeping constants specifying an effective theory, not directly measurable numbers
The cutoff	A regulator to be sent to infinity as soon as possible	A physical boundary marking where new, unknown physics takes over
Validity of the theory	Expected to hold at every energy scale, without exception	An effective field theory, honest about the scale up to which it applies

Nothing on the right-hand side of that table required throwing out the left. Physicists did not discard Feynman diagrams or Schwinger's perturbative loops – they are still, today, the workhorse

for computing an observable like the electron’s magnetic moment to spectacular precision, and Wilson’s RG gets you the exact same number if you push it far enough. What changed is not the arithmetic but its interpretation: Feynman himself called the old subtract-the-infinity procedure a “dippy process,” hocus-pocus that happened to work; Wilson supplied the missing logical floor underneath it, by showing that low-energy physics is naturally, and for a good reason, insulated from whatever is happening at energies far above where anyone has yet to look.

In one sentence, then: what Wilson did was show that a coupling – a charge, a mass, a critical temperature – is not a fixed number waiting to be measured, but a function of the scale at which you ask the question, and that this function has a structure (a flow, with fixed points) organizing everything else a theory can tell you. Everything else in this note – the decimation of Section 3, the bond-moving of Section 4, the running of α and α_s in Section 6 – is that one idea, applied to a different system.

None of the first three needed relativity, particle creation, or path integrals – none of the machinery Dirac and Feynman built to survive Einstein. Even the fourth, QCD’s asymptotic freedom, traces its decisive sign to gluon self-interaction, which is a fact about non-abelian gauge theories, not about relativity per se. What ties all four together is Landau’s original mistake, corrected: he treated the coefficients of a free energy as fixed numbers instead of functions of the scale you’re looking from. That correction is also the real payoff of the universality this note opened with – the fact that the same β shows up for magnets and for fluids stops being a coincidence and becomes something explained: systems that look nothing alike up close, but share the same symmetry and dimensionality, flow under coarse-graining to the same fixed point. Physicists call that shared destination a *universality class*, and the exponents we computed by hand belong to the class, not to any one material. Fix the scale-dependence, the way Wilson did, and one wrong universal number becomes several separately checkable right ones – in magnetism, in particle physics, and now in quantum sensing.

Some care is needed about the size of this claim: “Wilson found a better approximation than Landau’s” undersells it, and “Wilson invented scale-dependence out of nothing” oversells it. The individual pieces existed before him: Kadanoff’s block-spin picture [Kadanoff, 1966], Gell-Mann and Low’s running charge in QED [Gell-Mann and Low, 1954], Stueckelberg and Petermann’s original renormalization group [Stueckelberg and Petermann, 1953]. What Wilson did was weld those separate hints into one systematic, computable framework, carried all the way through to an actual fixed-point structure – fundamentally changing what most physicists meant by a theory. A coupling’s “true” value, in this language, isn’t a fact about the world; it’s a question that only makes sense once you’ve said at what scale you’re asking it. This is, at bottom, the same sense in which an irrational number isn’t a hidden final digit waiting behind a curtain, but is genuinely constituted by the never-terminating sequence of rational approximations that converge to it: you don’t compute more digits of π and get closer to some pre-existing true π sitting somewhere beyond them; the convergent sequence itself, in full, is what π *is*. This isn’t just a loose metaphor: mathematical physicists working to put the renormalization group on a fully rigorous footing construct it as exactly this kind of limit – a sequence of successive coarse-grainings shown, with actual proof, to converge [Kupiainen, 2023]. A running coupling works the same way. Its value at one loop, at two loops, at the next energy scale up, is not a rough draft of some perfect number waiting at the end of the calculation – there is no end of the calculation. The flow is the object.

Later, with Kogut, Wilson gave this idea its lasting name [Wilson and Kogut, 1974]: every theory here is really an *effective field theory*, valid only up to some cutoff, with its measured couplings standing in for whatever more complete physics operates beyond it. Landau’s error, in this language, was simply not knowing his own theory this way. *Nature Physics* marked the 50th anniversary occasion with a Focus issue collecting short commentaries from working physicists on

how far the idea has since traveled [Phillips, 2023]. This note has stayed narrowly within magnets, QED, and QCD; the collection is a good place to see the same machinery doing very different work elsewhere: probing candidate theories of quantum gravity [Eichhorn, 2023], organizing strongly-coupled supersymmetric field theories [Song, 2023], explaining long-range order in flocks of birds and swarms of bacteria [Tu, 2023], and sharpening precision QCD calculations tested against collider data [Boito, 2023].

If you have spent a career building models of complex systems rather than studying physics, none of this should feel like a foreign language. A parameter that looks constant at the scale you happened to write it down at is usually hiding real dependence on scale, and finding that dependence is often the whole difference between a model that is qualitatively suggestive and one you would actually bet on.

A The Mathematics Behind Section 5.1

Section 5.1 walked through five steps from Landau’s free energy to a genuine numerical flow, but kept the equations to a minimum so the main argument stayed readable. This appendix writes out the mathematics of each step in full, including the one derivation the main text only stated the result of: where $\nu \approx \frac{1}{2} + \frac{\epsilon}{12}$ actually comes from.

Step 1: Landau’s free energy

Section 1.3 wrote

$$F(M) = RM^2 + UM^4, \quad R \propto (T - T_c), \quad (29)$$

with R and U treated as ordinary numbers: smooth functions of temperature, with no dependence on length scale. Minimizing over M gave the mean-field result $\beta = 1/2$ – the wrong number this whole note has been correcting.

Step 2: Promoting M to a field

Real magnets fluctuate from point to point, so replace the single number M with a field $\phi(x)$ that can vary in space, and add the simplest term that penalizes rapid variation:

$$S[\phi] = \int d^d x \left[\frac{1}{2} (\nabla \phi)^2 + r \phi^2 + u \phi^4 \right], \quad (30)$$

identifying $r \leftrightarrow R$ and $u \leftrightarrow U$. This is the Landau-Ginzburg-Wilson action: Landau’s free energy, with fluctuations at every length scale restored via the gradient term.

Steps 3 and 4: The $O(n)$ generalization, and $n = 1$

Let ϕ become an n -component field, $\phi = (\phi_1, \dots, \phi_n)$, with $\phi^2 \equiv \sum_i \phi_i^2$:

$$S[\phi] = \int d^d x \left[\frac{1}{2} (\nabla \phi)^2 + r \phi^2 + \frac{\lambda}{4} (\phi^2)^2 \right]. \quad (31)$$

$n = 1$ recovers the single-component Ising theory of Step 2; $n = 2, 3$ describe the XY and Heisenberg universality classes, not pursued in this note. Wilson’s momentum-space RG, run near $d = 4$ where

λ is exactly marginal, produces two coupled one-loop flow equations (stated here, not re-derived from Feynman diagrams – that derivation genuinely is a separate undertaking):

$$\frac{d\lambda}{dt} = \epsilon\lambda - \frac{n+8}{8\pi^2}\lambda^2, \quad (32)$$

$$\frac{dr}{dt} = 2r - \frac{n+2}{8\pi^2}\lambda r, \quad (33)$$

with $t = -\ln \mu$ as in Section 5.1. Equation (32) is Equation (24), repeated; Equation (33) is new here – it is the one-loop correction to the “mass” r , and it is what actually controls ν , not λ itself.

Where Equations (32)–(33) themselves come from: a sketch

Wilson’s own method for deriving flow equations like these is momentum-shell integration, the continuum cousin of the real-space coarse-graining in Sections 3–4: split the field’s Fourier modes into a low-momentum piece ($k < \Lambda/b$) and a thin high-momentum shell ($\Lambda/b < k < \Lambda$), integrate the shell out of the path integral, and track how doing so shifts the couplings r and λ that describe the remaining low-momentum theory. Two diagrams do the work, at one loop:

- **The mass correction (a “tadpole”).** A single quartic vertex has one pair of its four legs closed into a loop – one propagator folded back on itself, the field-theory cousin of the vacuum-screening picture from Section 6.1 – while the other two legs stay external, shifting r .
- **The coupling correction (a “bubble” or “fish”).** Two quartic vertices are joined by two internal propagators, forming a closed loop with four external legs; because the interaction is $(\phi^2)^2$ rather than a single-field ϕ^4 , this bubble can be drawn three ways – the s -, t -, and u -channels – and all three contribute.

Both diagrams are evaluated by doing the loop integral only over the thin momentum shell, then re-expressing that shell’s contribution as a shift in r and λ per unit of $d\ell = d(\ln b)$ – the differential version of the discrete decimation used on the lattice in Sections 3–4.

Where the specific numbers $n+8$ and $n+2$ come from is combinatorics, not new physics: an $O(n)$ -symmetric field has n components, so a loop can either run through the *same* component twice (a fixed combinatorial weight, independent of n) or trace over *all* n components (a term proportional to n). The tadpole diagram, with one vertex and one independent contraction pattern, produces $n+2$; the bubble diagram, with three channels each contributing their own fixed piece plus an n -dependent trace, produces $n+8$. Both formulas correctly reduce to the standard single-scalar ϕ^4 result at $n=1$, which is the usual check that the combinatorics have been counted correctly.

What this sketch omits is the loop integral itself – the explicit d -dimensional momentum integral over the shell that actually produces the $1/(8\pi^2)$ prefactor, once a regularization scheme is fixed. Carrying it out would turn this appendix into the field-theory chapter this note’s abstract already declined to write. Equations (32)–(33) are given at this point, not derived further; everything after is algebra alone.

Where $\nu \approx \frac{1}{2} + \frac{\epsilon}{12}$ comes from

Setting the right-hand side of (32) to zero and solving for the nonzero root reproduces the Wilson-Fisher fixed point already stated as Equation (25):

$$\lambda^* = \frac{8\pi^2\epsilon}{n+8}. \quad (34)$$

Now linearize (33) around this fixed point. At $\lambda = \lambda^*$,

$$\frac{dr}{dt} = \left[2 - \frac{n+2}{8\pi^2} \lambda^* \right] r = \left[2 - \frac{(n+2)\epsilon}{n+8} \right] r \equiv y_t r, \quad (35)$$

so r grows or shrinks exponentially near the fixed point at rate y_t . This y_t is the *thermal eigenvalue*, and the correlation length exponent is its reciprocal:

$$\nu = \frac{1}{y_t} = \frac{1}{2 - \frac{(n+2)\epsilon}{n+8}}. \quad (36)$$

Setting $n = 1$ (Step 4) gives $y_t = 2 - \epsilon/3$, and expanding (36) for small ϵ ,

$$\nu = \frac{1}{2 - \epsilon/3} = \frac{1}{2} \left(1 - \frac{\epsilon}{6} \right)^{-1} \approx \frac{1}{2} \left(1 + \frac{\epsilon}{6} \right) = \frac{1}{2} + \frac{\epsilon}{12} + O(\epsilon^2), \quad (37)$$

which is exactly Equation (22). This is genuine one-loop mathematics, built entirely from Equations (32)–(33), which is why it is derived here rather than merely cited. The next term, $7\epsilon^2/162$, requires the *two-loop* correction to (33) – a real calculation, but a substantially harder one, which is why Equation (23) is credited to Brézin et al. [1976] rather than derived in this appendix.

Step 5: The numerical flow

With λ^* in hand, Section 5.1 iterated (32) numerically by Euler’s method (Table 3, Figure 3); nothing further needs adding here beyond what is already shown there.

References

- P. Zanardi, M. G. A. Paris, and L. Campos Venuti. Quantum criticality as a resource for quantum estimation. *Physical Review A*, 78:042105, 2008.
- L. Garbe, M. Bina, A. Keller, M. G. A. Paris, and S. Felicetti. Critical quantum metrology with a finite-component quantum phase transition. *Physical Review Letters*, 124:120504, 2020.
- P. W. Kasteleyn. The statistics of dimers on a lattice: I. The number of dimer arrangements on a quadratic lattice. *Physica*, 27:1209–1225, 1961.
- M. E. Fisher. On the dimer solution of planar Ising models. *Journal of Mathematical Physics*, 7:1776–1781, 1966.
- F. Barahona. On the computational complexity of Ising spin glass models. *Journal of Physics A: Mathematical and General*, 15:3241–3253, 1982.
- F. Jegerlehner. Hadronic vacuum polarization effects in $\alpha_{\text{em}}(M_Z)$. *Proceedings, hep-ph/0308117*, 2003.
- K. G. Wilson and J. Kogut. The renormalization group and the ϵ expansion. *Physics Reports*, 12:75–199, 1974.
- K. G. Wilson. Problems in physics with many scales of length. *Scientific American*, 241(2):158–179, August 1979.

- E. Ising. Beitrag zur Theorie des Ferromagnetismus. *Zeitschrift für Physik*, 31:253–258, 1925.
- L. D. Landau. On the theory of phase transitions. *Zh. Eksp. Teor. Fiz.*, 11:19, 1937.
- K. G. Wilson. The renormalization group and critical phenomena. *Nobel Lecture*, 8 December 1982.
- K. G. Wilson. Renormalization group and critical phenomena. I. Renormalization group and the Kadanoff scaling picture. *Physical Review B*, 4:3174, 1971.
- K. G. Wilson and M. E. Fisher. Critical exponents in 3.99 dimensions. *Physical Review Letters*, 28:240, 1972.
- E. Brézin, J. C. Le Guillou, and J. Zinn-Justin. Field Theoretical Approach to Critical Phenomena, in *Phase Transitions and Critical Phenomena*, Vol. 6, C. Domb and M. S. Green, eds. Academic Press, 1976.
- H. D. Politzer. Reliable perturbative results for strong interactions? *Physical Review Letters*, 30:1346, 1973.
- D. J. Gross and F. Wilczek. Ultraviolet behavior of non-abelian gauge theories. *Physical Review Letters*, 30:1343, 1973.
- M. Gell-Mann and F. E. Low. Quantum electrodynamics at small distances. *Physical Review*, 95:1300–1312, 1954.
- E. C. G. Stueckelberg and A. Petermann. La normalisation des constantes dans la théorie des quanta. *Helvetica Physica Acta*, 26:499–520, 1953.
- L. P. Kadanoff. Scaling laws for Ising models near T_c . *Physics*, 2:263, 1966.
- D. R. Nelson and M. E. Fisher. Soluble renormalization groups and scaling fields for low-dimensional Ising systems. *Annals of Physics*, 91:226–274, 1975.
- P. Weiss. L’hypothèse du champ moléculaire et la propriété ferromagnétique. *J. Phys. Theor. Appl.*, 6:661–690, 1907.
- J. D. van der Waals. *Over de Continuïteit van den Gas- en Vloeistofoestand*. PhD thesis, Leiden University, 1873.
- Particle Data Group et al. Review of Particle Physics. *Prog. Theor. Exp. Phys.*, 2022:083C01, 2022.
- A. Kupiainen. Rigorous renormalization group. *Nature Physics*, 19:1539–1541, 2023.
- P. W. Phillips. Fifty years of Wilsonian renormalization and counting. *Nature Physics*, 19:1521–1524, 2023.
- A. Eichhorn. The microscopic structure of quantum space-time and matter from a renormalization group perspective. *Nature Physics*, 19:1527–1529, 2023.
- J. Song. Supersymmetric renormalization group flow. *Nature Physics*, 19:1530–1532, 2023.
- Y. Tu. The renormalization group for non-equilibrium systems. *Nature Physics*, 19:1536–1538, 2023.
- D. Boito. Consequences of the renormalization group for perturbative quantum chromodynamics. *Nature Physics*, 19:1533–1535, 2023.
- J. S. Walker. *A Student’s Guide to the Ising Model*. Cambridge University Press, 2023.