

How BC Cycle Tourism Reads the Terrain

Character Scores: Methods and Limitations

BC Cycle Tourism Society

Published July 2026

Preliminary — working data, subject to revision

Researched and drafted with AI assistance (Claude models, Anthropic)

Licence: Creative Commons Attribution 4.0 International (CC BY 4.0). Share it and build on it, including commercially, as long as you credit BC Cycle Tourism Society.

<https://creativecommons.org/licenses/by/4.0/>

The character map data this document describes is a derivative of OpenStreetMap and is licensed separately, under the Open Database Licence (ODbL) 1.0. <https://opendatacommons.org/licenses/odbl/1-0/>

Contains information from OpenStreetMap, © OpenStreetMap contributors, available under the Open Database Licence (ODbL) 1.0. <https://opendatacommons.org/licenses/odbl/1-0/>

Suggested citation: BC Cycle Tourism Society (2026). How BC Cycle Tourism Reads the Terrain: Character Scores, Methods and Limitations [report]. Version July 2026. <https://bcccycletourism.ca/research/>

Version	Date	Notes
1.0	July 2026	First published version

Heather Piwowar · heather@bcccycletourism.ca
bcccycletourism.ca/research/

1. Introduction and scope

BC Cycle Tourism Society characterizes routes and bike touring options around British Columbia (BC). Part of that work is a province-wide **character map**: every rideable stretch of BC's cycling network, scored on seven dimensions of what a ride is actually like. Three of those scores describe qualities almost everyone wants more of: **scenic** (what you see), **peaceful** (what you hear and feel beside you), and **flow** (how far you can ride before something stops you). Four describe preferences that riders genuinely differ on: **comfort** (traffic stress), **ruggedness** (how rough the surface is), **grade** (how much climbing), and **remoteness** (how far from resupply). The map is browsable at bccycletourism.ca/research/character-map/.

This document explains how those scores are computed, what was checked, and what the scores cannot claim. It is one half of a pair, and the division of labour is simple: **this document covers the map** (how every rideable stretch of network earns its seven scores), and our companion document, *How BC Cycle Tourism Rates Routes: Methods and Limitations*, on the same research page, **covers the routes** (how a published route's GPS track is matched onto that scored network, how its labels are rolled up from the scores, and how its difficulty tier is assigned). Where both documents describe the same score, the method is the same method; if you read both you have the complete system.

Everything here is preliminary working material, shared because it may be useful to trail organizations, tourism planners, researchers, and mappers, not because it is finished. The honest limitations are in Section 10, and they are not small.

2. The seven scores, in rider language

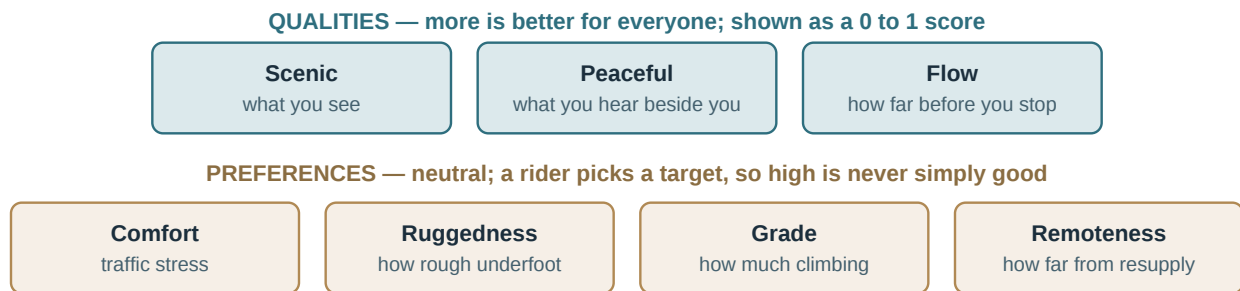


Figure 1. The seven character scores fall into two families that must never be mixed.

The family split matters more than any individual score. The three **qualities** are close to universal: a water view is a view for everyone, calm is calm, and an uninterrupted path is pleasant whether you are seven or seventy. They are scored 0 to 1, where higher is better.

The four **preferences** are different in kind. More climbing is not better or worse; it depends who is riding. A new rider wants low ruggedness and low grade; a loaded gravel tourer may want exactly the opposite. So the preference scores describe, they do not rank. When we show them, a high value means "this is a lot of climbing" or "this is far from a grocery store", never "this is good" or "this is bad". Treating a preference as a quality is the classic way terrain maps mislead, and the design guards against it throughout.

One consequence is worth stating plainly: some of the most scenic and flowing riding in BC is on roads with the smallest comfortable audience, wild coastline on a shoulderless rural highway, for example. Because the axes are kept separate, that riding stays visible on the map with its character intact, rather than being filtered out by a comfort judgement the rider never asked us to make.

Score	Family	What it answers	Computed on
Scenic	quality	Is there something worth looking at: water, forest, farmland, views?	100 m cells
Peaceful	quality	Is it calm here, or is traffic, rail, or industry roaring beside you?	100 m cells
Flow	quality	Can you keep moving, or is it stop sign, signal, crossing, repeat?	100 m cells
Comfort	preference	How much interaction with motor traffic does riding here involve?	whole ways
Ruggedness	preference	What surface is under your tires, pavement to rough gravel?	whole ways
Grade	preference	How hard does the climbing get, both the grind and the steep pitch?	whole ways
Remoteness	preference	How far are you from the next place to buy food?	whole ways

3. The unit of measurement

3.1 Edges and cells

All seven scores are computed on the OpenStreetMap (OSM) cycling network, one shared substrate, so the scores can always be compared for the same stretch of road. OSM's natural unit is the way, split at intersections into *edges*. Edges are topologically tidy but wildly uneven in length: a downtown block is 100 metres, a rural road can be 20 kilometres. An attribute that varies along the road, and scenery does, cannot honestly hang on a unit like that: a 3-kilometre road that begins at a beach and ends at an industrial park has no single scenic truth.

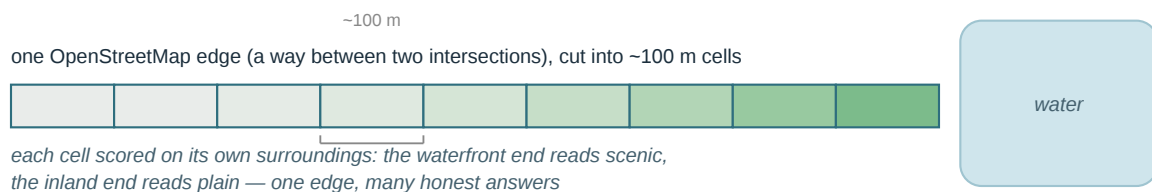


Figure 2. Cells let one road tell more than one truth about its surroundings.

So for the three context scores, every edge is cut into uniform cells of roughly 100 metres, and each cell is scored on its own surroundings. Province-wide that is about 2.85 million cells. The 100-metre size is not arbitrary: it matches the radius of the sideways look described below. Cut much finer and adjacent cells see the same neighbourhood and just repeat each other; cut much coarser and the waterfront end of a road is averaged against its inland end, which is the failure the cells exist to avoid.

The four preference scores stay on whole ways instead. Ruggedness and comfort are read largely from the way's own tags (its surface, its road class), which do not vary within a way; remoteness varies so slowly across the landscape that cells would add precision without information; grade is computed from an elevation profile sampled along the whole way. Long ways are still sampled internally, roughly every kilometre for remoteness and every 20 metres for grade, so a 20-kilometre backroad is not summarized by one point.

3.2 Two ways of looking: sideways, and along

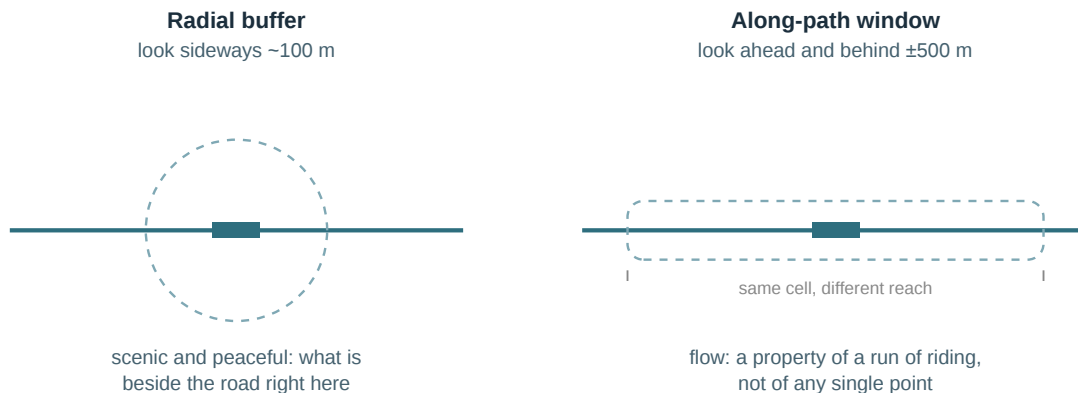


Figure 3. The same cell, two reaches: a radial buffer for what surrounds you, an along-path window for what riding is like.

Scenic and peaceful are properties of a *place*: what is beside the road right here. Each cell looks sideways into a buffer of roughly 100 metres and scores what falls inside: water, trees, farmland, an arterial road, a rail line. That local look is enough, because a view or a noise source 400 metres away is genuinely a different experience from one 40 metres away.

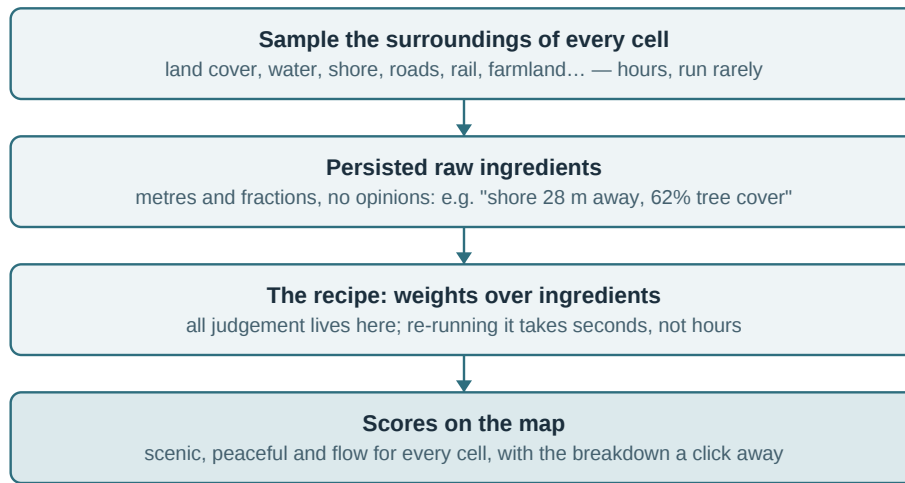
Flow is not a property of a place at all. You cannot read "can I keep moving?" off one 100-metre cell; it barely contains one intersection. Flow is a property of a *run* of riding, so each cell instead looks along its own way, 500 metres ahead and behind, and counts what would interrupt a rider in that window. The cell remains the storage and display unit for all three scores; only the reach differs. This distinction, sideways for context, along for movement, is the single most load-bearing design choice in the model.

3.3 What gets scored at all

Before any scoring, the network passes a shared rideability gate, the same gate for all seven scores: sidewalks, pedestrian-only footways, staircases, hiking-only trails, technical mountain-bike trails without bicycle designation, and ways tagged no-bicycles are excluded rather than scored. The alternative is worse than a gap: a rated hiking trail would read as a rideable gravel option, and an alley network would clutter every urban block with three parallel scored lines. Early versions of the map made exactly that mistake, scoring sidewalks and driveways brightly; the fix moved the exclusion to the front of the pipeline, where it now protects every score at once.

4. Scenic: what you can see

4.1 Ingredients, then a recipe



measure once, season to taste: calibration never re-measures the world

Figure 4. The architecture that made calibration possible: raw ingredients are measured once; opinions live in a recipe that re-runs in seconds.

The scenic pipeline separates measurement from judgement. A sampling pass visits every cell once and records raw, opinion-free **ingredients**: what fraction of the surrounding buffer is water, tree cover, natural grass, lawn, built-up land; how many metres to the nearest shoreline, river, cliff, waterfall; whether a viewpoint is within 300 metres; how much of the buffer is farmland or protected area; how close the nearest busy road is. That pass is expensive, hours over millions of cells, and runs rarely. The **recipe**, the weights that turn ingredients into a score, re-runs in seconds. Every calibration decision in this section was made by adjusting the recipe and looking at the map again, without re-measuring the world. The same architecture, and mostly the same ingredients, also feed the peaceful score.

4.2 What the land cover sees, and what it cannot

The base layer is satellite-derived land cover: the European Space Agency (ESA) WorldCover product, 10-metre resolution, resampled onto a 30-metre provincial grid. Land cover tells each cell how much of its buffer is water, forest, grassland, cropland, or built-up. The choice of product matters more than it sounds. Our first province-wide attempt used a 30-metre North American land-cover product that classifies entire residential neighbourhoods, front yards, boulevards, street trees and all, as "built". On a canonically leafy Vancouver street, all 232 scored cells read exactly zero tree cover: the signal was not weak, it was absent from the source, and no weight can boost a zero. WorldCover sees those same yards as tree cover, closely tracking the high-resolution regional dataset we had calibrated against, and it covers the whole province in one consistent scheme, with no classification seam to calibrate across.

Land cover is complemented by OpenStreetMap features: parks and protected areas, viewpoints, cliffs, waterfalls, rivers, beaches and wetlands, farmland polygons, and the province's Agricultural Land Reserve (ALR) boundaries, plus a provincial Freshwater Atlas overlay that corrects satellite misreads of narrow water bodies.

4.3 The shoreline problem

The single most instructive failure in the scenic build was the shoreline. A cell's 100-metre buffer physically cannot see water on the far side of a wide beach or marsh, and BC's best dyke and beach riding sits exactly there: Richmond's West Dyke Trail runs behind 300 to 600 metres of Sturgeon Bank marsh, and sandy beaches read as "barren" in land-cover terms. The first map scored some of the region's most loved waterfront riding as dull inland path.

The fix was a new ingredient rather than a new opinion: distance to the nearest OpenStreetMap shoreline feature (beach, coastline, sand, wetland). These features are semantically the coastal foreground, so they reach across the marsh that the buffer cannot. With shore distance in the recipe, the West Dyke rose from 0.30 to 0.69 and Spanish Banks from 0.33 to 0.61, while no cell anywhere dropped. Then came the second lesson: the first version of the shore term reached too far, quietly inflating 5,108 kilometres of network that had no visible water at all, downtown blocks two streets from the harbour, inland fens. The reach was tightened from 450 metres to 175 so that full credit goes to paths genuinely on the foreshore (the real dyke trails sit about 22 to 28 metres from the shore features) and credit fades quickly behind them. Both the failure and the diagnosis were only possible because raw distances, not pre-cooked scores, were persisted.

4.4 The recipe as shipped

The score is a weighted sum of ingredients, clamped to 0 to 1, plus a small set of deliberately non-linear boosts. The linear part rewards water in the buffer (weight 0.50, plus 0.30 more for water within 35 metres), tree cover (0.28), natural grass and marsh (0.32, valued above mowed lawn at 0.15), a nearby viewpoint (0.15), being on a bridge (0.20), and mildly penalizes built-up land (−0.10) and busy-road proximity (−0.20). Rivers (within 120 m), cliffs (150 m), waterfalls (300 m), and protected areas add smaller terms. Three boosts do the work the linear terms cannot:

- **The near-water boost** (up to 0.40): "being on the water" is the strongest single scenic experience, stronger than the sum of its parts. The boost fires on the maximum of close water, areal water, shore proximity, and river proximity, so seawalls, dyke paths, and riverside trails all light up through whichever signal reaches them. Where satellite land cover misreads a narrow inlet as built (False Creek is the canonical case), shore proximity substitutes for areal water at a one-sided discount, 0.40 of full strength, so a seawall reads "water on one side" rather than "surrounded by water". A bridge deck directly over water earns a related boost; a highway overpass over pavement earns nothing.
- **The urban-canopy boost** (up to 0.30): a leafy street is a real scenic experience that a modest linear tree weight cannot express. The boost uses a threshold curve, near zero below 30 percent canopy and full above 65 percent, and is gated by built-up fraction so it is specifically an *urban* tree signal: zero in a forest (which scores through the ordinary tree term), full on a tree-tunnel residential street.
- **The pastoral boost** (up to 0.28): open farmland reads as scenic, pasture, hayfield, orchard, vineyard, but ploughed fields are invisible to the positive land-cover classes, so the boost keys on OpenStreetMap farmland polygons and the ALR. It fades out as built-up fraction rises, which keeps greenhouse complexes and town edges muted. Glass is not the pretty kind of farm.

Each gate earns its keep: an ungated boost leaks. The canopy boost without its urban gate would double-count every forest; the pastoral boost without its built gate scored greenhouse strips as countryside; the first, ungated version of the shore term is the cautionary tale above.

5. Peaceful: what you hear and feel beside you

Peaceful is not scenic's sibling so much as its complement, and the model keeps them strictly apart because BC is full of places that have one without the other. The Sea-to-Sky highway shoulder is spectacular and deafening. A residential grid in Burnaby is calm and visually dull. A separated path beside a rail line can even be flowing and still not peaceful. Three axes, what you see, what roars beside you, what interrupts you, and each gets its own score.

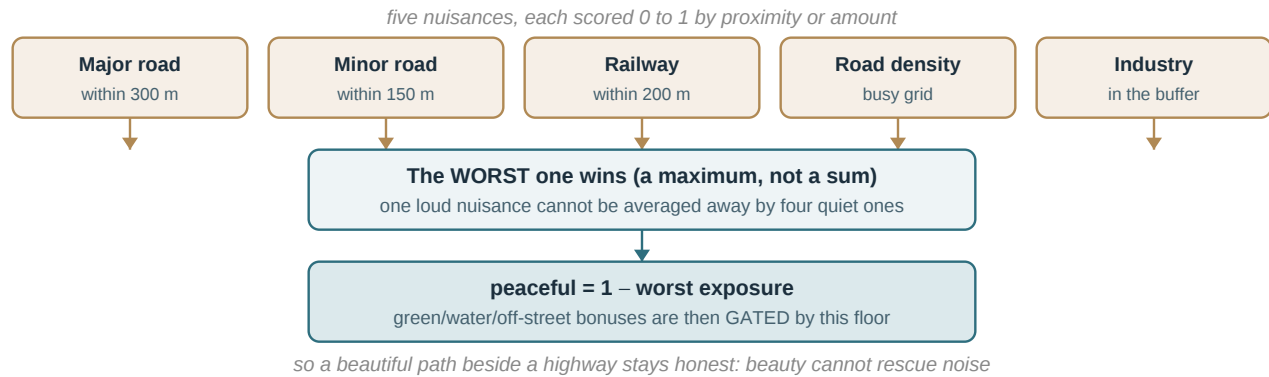


Figure 5. Peaceful is built on exposure, and exposure is a maximum: the worst nuisance sets the floor.

The peaceful score starts from **exposure**: proximity to major arterials (felt within about 300 metres), minor arterials (150 metres, weighted 0.6), railways (200 metres, weighted 0.9, including elevated rapid transit), the density of motor roads in the buffer, and industrial land use. The pivotal choice is that these combine by *maximum*, not by sum: the single worst nuisance sets the floor, because that is how nuisance works. A quiet side street beside a rail line is dominated by the rail line; the absence of four other nuisances does not average the trains away.

On top of the floor, small bonuses recognize genuinely calming context: being on an off-street path (0.15), tree presence, being on the water, a burbling creek within 60 metres (the one place streams enter any score). All bonuses are *multiplied by* the exposure floor rather than added alongside it, so beauty cannot rescue noise: a gorgeous path pinned to a highway stays honestly un-peaceful.

The design was accepted against named places chosen to be hard to get right, including one engineered trap: a popular separated path that runs directly beneath elevated rapid transit. Without the rail ingredient the model scored it serene; with it, the path reads 0.16, which anyone who has stood under the guideway will recognize. Rapid transit is tagged inconsistently in OpenStreetMap (Vancouver's SkyTrain is tagged as "subway" even where elevated), so the rail extract deliberately spans rail, light rail, subway, tram, and monorail, and drops tunnelled track, which is silent at street level.

Named calibration spot	Peaceful	Why it is a test
Seymour Valley Trailway	0.95	forest path far from roads: the top of the scale
Pitt polder dyke roads	0.85	rural calm without wilderness
Stanley Park seawall	0.55	calm, but a busy park causeway runs nearby
Stanley Park causeway shoulder	0.26	metres from the seawall, opposite experience
BC Parkway beside SkyTrain	0.16	the engineered trap: flowing, separated, loud

Table 2. Peaceful acceptance spots, Metro Vancouver. The causeway/seawall pair, metres apart with opposite scores, is the discrimination the model must make.

6. Flow: how far before something stops you

Flow measures uninterrupted movement: how far a rider can go without stopping or yielding. It is the computable core of what riders call "fun", and deliberately the *shared-floor* part: a gentle flowing path is pleasant for a beginner and an expert alike. (The thrill of a fast twisty descent is real but rider-relative, and it lives with the preference dials, not here.)

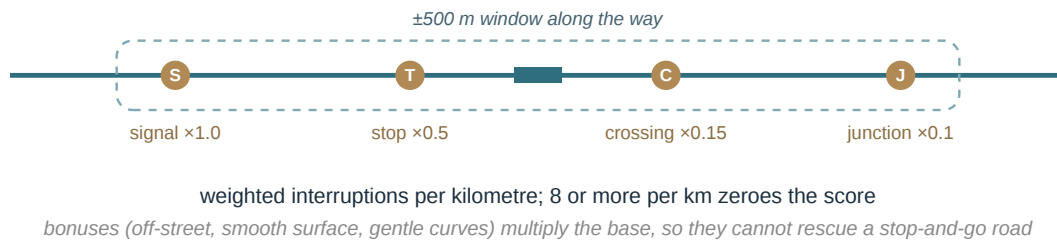


Figure 6. Flow reads the road itself over a ±500 m window: interruptions are the workhorse, everything else is seasoning.

Within each cell's along-path window, the model counts tagged interruptions and weights them by how badly they break a ride: a traffic signal costs a full point, a stop sign half, a pedestrian crossing 0.15, an ordinary road junction 0.1. The weighted count per kilometre is the score's backbone; at eight or more per kilometre, flow reaches zero. Bonuses for being off-street (0.20), smooth surface (0.10, rough surfaces –0.25), and gentle curviness (0.10) multiply the base, the same gating discipline as peaceful: an off-street path riddled with stop signs cannot buy its way back.

Two of those weights carry findings worth recording. **Crossings are weighted low on purpose:** OpenStreetMap tags one crossing node per intersection leg, up to four per corner, so their raw count over-states reality and a light weight corrects the double-counting rather than the riding. **Junctions nearly did not survive at all.** In the Metro Vancouver calibration they measured about 20 per kilometre *everywhere*, on flowing greenways and downtown grids alike, because the network is full of minor path spurs; a weight on a constant is not information, and the weight was set to zero. At province scale the picture inverted: rural rail trails have essentially zero junctions per kilometre while small-town grids have ten to twenty, so a small weight (0.1) now separates them. The same ingredient was noise at one scale and signal at another, which is exactly why raw counts are persisted even when their weight is zero.

Two honest structural notes. First, a short dead-end stub with no tagged interruptions would read as perfect flow, so scores fade out on fragments shorter than about 90 metres; the proxy is blunt, because OpenStreetMap splits ways at every intersection and way length alone cannot tell a block of a continuous greenway from an isolated stub. Second, this is *corridor* flow, the potential of the road itself, not the realized stop-count of any particular journey through it, which depends on every turn taken. Quiet suburban crescents therefore read flowy, technically truthfully. Both limits reappear in Section 10.

Named calibration spot	Flow	Reads as
Stanley Park seawall	1.00	uninterrupted off-street riding
Richmond West Dyke	0.87	long flowing dyke path
Central Valley Greenway	0.73	flowing, a few at-grade crossings
BC Parkway	0.58	dragged by ~12 at-grade crossings per km

Named calibration spot	Flow	Reads as
Knight Street arterial	0.38	moving traffic is not the same as flow for a rider
Downtown Vancouver grid	0.00	signal, crossing, signal: stop-and-go

Table 3. Flow acceptance spots, Metro Vancouver (median over cells). The interruption inventory behind them: 99,531 tagged nodes in the calibration region.

7. The four preference scores

The four preference scores are computed per way rather than per cell (Section 3.1), and each is a classifier in its own right. This section gives each its full treatment as it appears on the map; the companion document covers how these same scores become route labels, including the full validation crosstabs for comfort.

7.1 Comfort: how much traffic you share, and what stands between you and it

Comfort estimates traffic stress on a six-rung ladder, C1 to C6, and each rung is deliberately phrased as a condition, never a verdict: it names the traffic you would share and what stands between you and it, from separated paths and slow streets at C1 to fast highway traffic with a painted line at most at C6. The touring cyclists this map serves ride the whole ladder, so a high rung is a fact about the road, not advice to stay away. Above C3 the second half of each label stops changing, and that is the point: from there on, built infrastructure amounts to paint, and only the speed climbs.

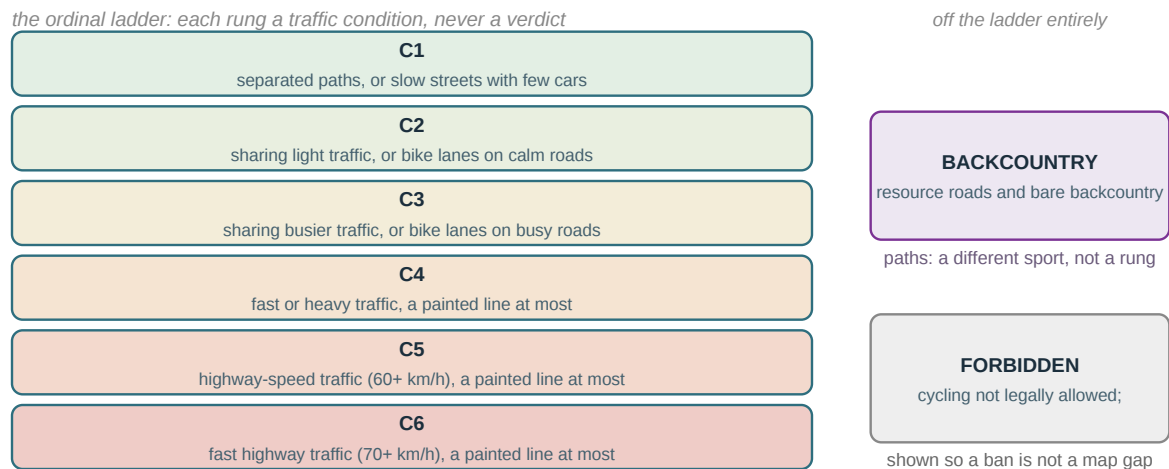


Figure 7. The comfort scale: six ordinal rungs, and two classes that are deliberately off the ladder.

The classifier is adapted from HUB Cycling's bikeway comfort framework, and we thank HUB for publishing it. The class labels, C1 to C6, are our own: HUB's classification of Metro Vancouver is field-reviewed in a way an automated province-wide model cannot match, so our classes carry our names, and HUB's names appear in this document only where we compare against HUB's own labels. The classifier reads, for every way: its road class, speed limit, lane count, parking, surface, and any cycling infrastructure. **Physical separation wins first:** a separated path or protected lane reads C1 regardless of the road beside it, because HUB's own field labels show separated paths rated most-comfortable even along arterials. Below that, painted lanes and shared roads are classified by speed and traffic volume together, a painted lane on a quiet street is very different from the same paint at 70 km/h.

Two of those inputs do not exist in OpenStreetMap and are synthesized, which is worth stating bluntly. Where no speed limit is tagged, the classifier assumes one from road class. Traffic volume is not measured anywhere in the pipeline: it is estimated from road class and adjusted upward by proxies, four or more lanes, a centre turn lane, scheduled transit service (a bus stop within 25 metres), a signed truck route. These

proxies were validated against the gold standard below, but they are proxies, and they are the weakest input feeding the most rider-facing score.

The bottom of the ladder is our own extension below HUB's scale: on true highway classes with no physical separation, high speed pushes a way to C5 and then C6. Each rung of built infrastructure buys exactly one 10 km/h step of grace: a bare road flags at 60 km/h, a painted lane at 70, a signed bicycle shoulder at 80, and paint never exempts, because paint is not protection at highway speed. Where the speed is guessed rather than tagged, the guess deliberately errs scary: an untagged rural secondary highway is assumed fast and flags C5, on the principle that a road we are guessing about should look cautionary rather than quietly fine. Two classes sit off the ladder entirely: BACKCOUNTRY (resource roads and bare backcountry paths, prime gravel-touring terrain but a different sport from everyday traffic comfort) and FORBIDDEN (cycling legally excluded, shown explicitly so a banned freeway renders as a ban rather than a gap in the map).

Comfort is the one score with a true gold standard: HUB Cycling's 17,622 labelled segments in Metro Vancouver. The shipped rule set, re-measured 26 July 2026, agrees exactly on 59.5 percent of length-weighted samples and within one class on 82.1 percent, with the strongest agreement at the comfortable end (85 percent of HUB-MOST kilometres exact). Part of the gap is deliberate: the C5 and C6 rungs sit below HUB's scale entirely, so every kilometre placed there scores as a near-miss or worse against a scale that cannot express them. The full crosstabs, including where we still disagree with HUB and why, are published in the companion document. The validation boundary matters more than the score: every gold label is urban Metro Vancouver, so province-wide comfort is extrapolation, tuned to err scary when guessing.

7.2 Ruggedness: what is under your tires

Ruggedness grades every rideable way R0 to R4, a tire-choice scale: the question it answers is what bike you would want, not whether the way is good.

Level	Surface	What it is like	Typical tire
R0	Paved	Asphalt or concrete pavement.	road bike
R1	Champagne gravel	Smooth, hard-packed, well-maintained dirt or fine gravel; rides almost like pavement.	25–32 mm
R2	Standard gravel	Maintained dirt and gravel roads with minor potholes, washboard, or loose corners.	28–35 mm
R3	Chunky gravel	Poorly maintained roads with larger stones, ruts, or thick sand.	33–40 mm
R4	Extreme gravel	Unmaintained and rugged: deep jagged stone, roots, erosion; borders mountain-bike terrain.	40 mm +

Table 4. The five ruggedness levels. R4 evidence is sparse in OpenStreetMap today; it is reserved for explicitly tagged technical ways.

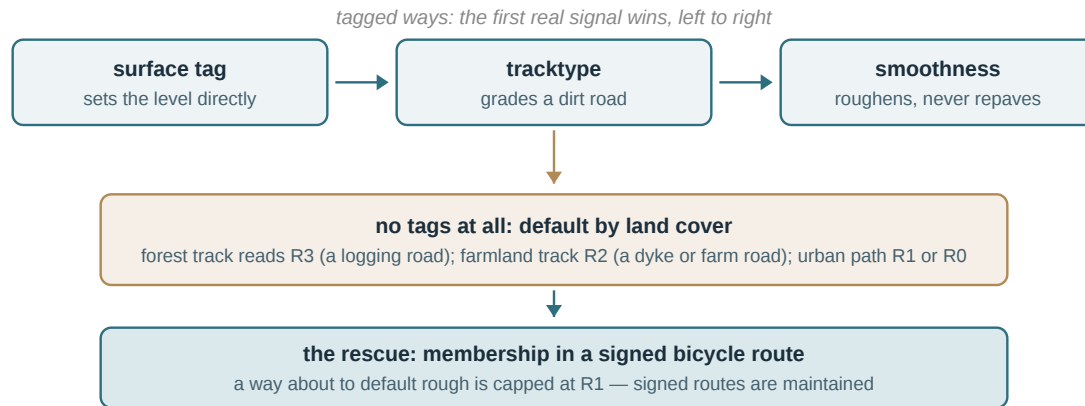


Figure 8. How a way gets its level: real tags first, land-cover-gated defaults when there are none, and a rescue for signed routes.

For tagged ways, the level fuses three OpenStreetMap tags in strict precedence: the surface type sets the category, a track-grade tag refines dirt roads, and a smoothness tag is consulted only to *roughen*, never to repave. That last asymmetry earns its place: a cracked-but-concrete path stays R0, because smoothness complaints on pavement describe maintenance, not surface type, and letting them demote pavement was mis-colouring downtown paths as gravel. Mountain-bike and hiking difficulty tags bump toward R4.

Untagged ways, and about half of BC's unpaved kilometres carry no surface tag at all, get defaults gated by satellite land cover at the way itself: an untagged track through forest reads R3, because it is almost certainly a logging road; the same bare tag on farmland reads R2, a dyke or farm access road (this single distinction fixed an entire dyke-trail network that had been reading chunky); urban paths default smooth. One rescue runs before any unpaved default lands rough: membership in a signed, non-mountain-bike bicycle route (about 32,000 member ways province-wide) caps the level at R1, on the grounds that agencies do not sign touring routes down unmaintained tracks. Every way carries a confidence flag, tagged, rescued, or defaulted, roughly half the network is defaulted, so downstream display can hedge instead of asserting.

7.3 Grade: the grind and the wall

Grade measures pure climbing, and it measures two different fears. The **grind** is ascent per kilometre: the long slog that empties your legs. The **wall** is the steepest sustained 100-metre pitch: the short ramp that stops a loaded touring bike outright. They are different kinds of hard, a steady 3 percent for 20 km and a 15 percent pitch for 100 metres do not trade off against each other, so the published score is the *worse* of the two, never a blend: a flat ride with one brutal wall must not average out to "moderate". Both facets are computed in both directions (a climb one way is a descent the other), and the map colours the harder direction. The scale anchors at plain numbers: around 2 percent reads flat, 5 percent moderate, 8 percent hard, and a sustained 10 percent or more reads brutal.

Elevation comes from the Shuttle Radar Topography Mission (SRTM) terrain model, sampled every 20 metres along each way, and most of the method is defence against that model's 30-metre resolution. Bridges and tunnels are not scored at all: SRTM records the terrain and water under a deck, not the deck, so a bridge's true grade is unknowable from it and we do not assert one. The harder problem is the **phantom wall**. On a flat rail trail benched into a canyon side, a ten-metre horizontal error lands the elevation sample on the cliff face, manufacturing a steep "climb" on a railway grade; small terrain dips and spurs do the same on milder ground.

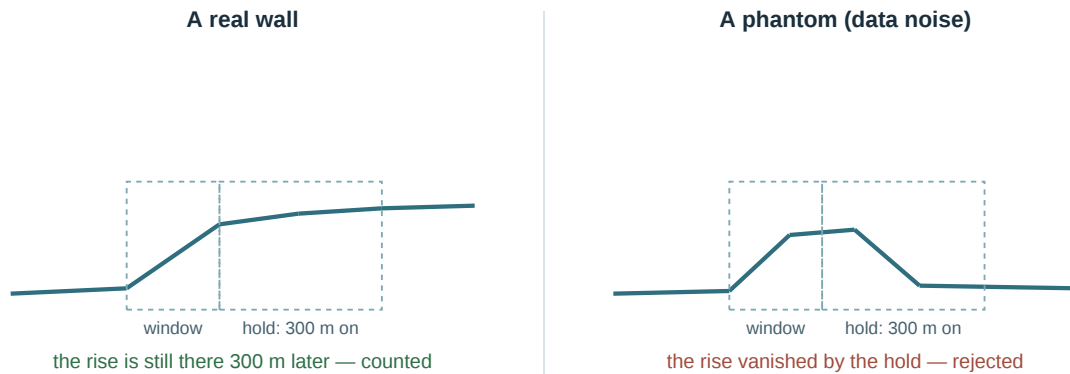


Figure 9. The held-rise test that separates real walls from phantom ones: a real climb is still there 300 m later; a phantom has already recovered.

The defence exploits the one signature the real thing has and the phantom lacks: **a real climb persists**. A candidate wall only counts if its rise is still present 300 metres past the measurement window; a noise hump has dropped back by then and is rejected. On benched terrain, detected from the cross-slope beside the way, the measurement window also stretches from 100 metres up to 500, which dilutes a slip artifact to the trail's real, nearly flat grade while a genuine sustained climb reads steep at any window length. The acceptance evidence: a famously steep Vancouver residential street reads brutal (score 0.85, against a surveyed 12.3 percent), Richmond's dykes read 0.00, the Cypress and Seymour mountain climbs stay brutal, and on the benched Kettle Valley Rail (KVR) corridor the share of kilometres misread as brutal fell from 33 percent to 4 with the machinery in place. What the machinery cannot fix, residual inflation on steep-sidehill backcountry tracks, is stated in Section 10.

7.4 Remoteness: how far to the next groceries

Remoteness asks one concrete question: how far are you from the ability to buy food? The signal is the distance from each way to grocery-type places, grocery and general stores, convenience stores, farm stands and markets, in our BC Bike Stops directory, about 2,600 of them province-wide when the remoteness layer was computed. Restaurants deliberately do not count: they are not resupply, and when they were tried, a single roadside inn made an entire remote plateau read as served.

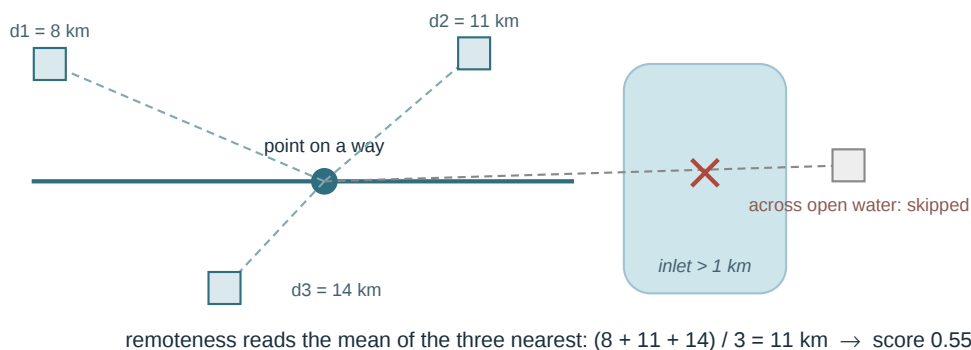


Figure 10. Mean of the three nearest, and the across-water rule: a store on the far shore of an inlet is not resupply.

The score reads the *mean of the three nearest*, not the single nearest, for three reasons: one mislisted or closed store cannot flip a valley; the mean encodes local density (one store in a sparse valley reads different from one store among many); and the gap between nearest and mean is a free confidence signal, when one

store is far ahead of the next two, the way is flagged single-source and the display can hedge ("one resupply point, confirm it is open"). A candidate store is skipped when the straight line to it crosses more than a kilometre of open water, killing the inlet-and-big-lake bleed where a store on the far shore "counted". Distances map onto a continuous 0-to-1 scale, logistic on the logarithm of distance, where a 10 km mean scores 0.5; the scale is continuous, with no named levels, deliberately unlike ruggedness, because resupply distance has no natural rungs. Long ways are sampled about every kilometre so a 20 km backroad is not one number; 527,353 ways are scored province-wide.

Calibration probe	Mean-of-3	Score	Reads as
Downtown Vancouver	0.1 km	0.00	town
Stanley Park seawall	2.1 km	0.11	urban park
Seymour Valley Trailway	7.2 km	0.39	near-city forest
Gray Creek, Kootenay Lake east shore	10.3 km	0.51	one store, otherwise sparse (single-source)
Myra Canyon, KVR	11.6 km	0.55	day-trip remote
KVR at Brookmere	33 km	0.83	remote
Highway 20 at Kleena Kleene, Chilcotin	84 km	0.94	deep remote

Table 5. Selected rungs of the remoteness calibration ladder (straight-line calibration values). The full ladder has fourteen probes and must stay in this order under any tuning.

8. Every score carries its reasons

A design rule runs through the whole system: **a score with no reason is not publishable**. Every scenic, peaceful, and flow score decomposes into the named ingredients that produced it, and the interactive map exposes that decomposition: click any stretch and the popup shows the breakdown, for example that a cell's scenic score rests on shore proximity and the near-water boost, or that a peaceful score is floored at 0.16 by a rail term of 0.82. The per-way scores do the same: comfort reports the road class and assumed speed behind each call, ruggedness reports whether its level came from a real tag or a default, remoteness reports the three distances themselves.

This is partly for us, because a score we cannot explain is a score we cannot debug, and every fix described in this document started from a breakdown that looked wrong. But it is mostly for you: it converts "the map says 0.7" into a checkable claim about the world, and when the map is wrong, the breakdown usually says exactly which input to blame.

9. Validation: what was checked, and how hard

Honesty requires a clear statement up front: **scenic, peaceful, and flow have no gold standard**. There is no labelled dataset of "how scenic this road is" to test against, and we did not manufacture false rigour by inventing one. What these scores have instead is structured, adversarial calibration: named acceptance spots chosen *before* tuning, including places designed to break the model, reviewed on the map by a person who has ridden them.

- **Acceptance ladders.** Each score keeps a ladder of named places whose ordering is not negotiable (Tables 2, 3, and 5). Tuning that fixed one rung was rejected if it broke another.
- **Engineered traps.** The BC Parkway beside the SkyTrain must read loud; the Stanley Park causeway and seawall, metres apart, must diverge; beach-fronted dykes must read as waterfront; greenhouse strips must not read as countryside; a flat benched rail trail must not grow a phantom climbing wall. Each trap exists because a version of the model once failed it.

- **Measured deltas, not impressions.** Changes shipped with their distributional effect measured: the shore fix raised 17.8 percent of cells and lowered none; its overreach was quantified (5,108 km of unearned credit) before being corrected.

The remoteness ladder (Table 5) shows the pattern for a per-way score: fourteen probes from downtown Vancouver to the deep Chilcotin, required to stay in that order under any tuning. Grade's acceptance set runs the same way: a famously steep 12 percent Vancouver street must read brutal, Richmond's dykes must read flat, the benched KVR must read like the gentle rail grade it is, and the Cypress and Seymour mountain climbs must stay brutal; the last two constraints pull against each other, and Section 10 records the residue.

Comfort is the exception that shows what the others lack: a true external gold standard (the HUB Cycling comparison summarized in Section 7, published as full crosstabs in the companion document). For the cell scores, the equivalent does not exist yet. Calibration against independently curated published routes, the judgement of experienced route authors used as a held-out check rather than a tuning input, is planned and not yet done.

10. Limitations

These are the known limits of the character map, stated without softening. Anything we publish that hedges should draw from this list, not shrink it.

1. **Scenic, peaceful, and flow have no external validation.** They were calibrated by structured eyeball against named places, almost entirely by one person. The acceptance ladders make the calibration repeatable and falsifiable, but they are not a gold standard, and no inter-rater check has been run. Someone else's eye would tune a different map.
2. **Calibration is Metro Vancouver; application is all of BC.** Nearly every named acceptance spot for the cell scores is in or near Metro Vancouver. The Interior, the North, and the coast outside the Lower Mainland are scored by a model tuned elsewhere, and their distinctive scenery (sagebrush benchland, alpine drama) is exactly where the model's ingredients are thinnest.
3. **Scenic measures proximity, not visibility.** There is no viewshed computation. A road beside a fjord it cannot see, behind a berm or a wall of condos, scores like one with the full view. Mountain backdrop, the defining scenery of much of BC, contributes nothing at all: elevation drama is not yet an ingredient. This is the largest known gap in the scenic score.
4. **Everything inherits OpenStreetMap's unevenness.** Flow's interruption counts exist only where signals and stop signs have been mapped; thinly mapped towns read flowier than they ride. Peaceful infers noise from road classification, not measured traffic: a classified-but-quiet rural highway reads louder than it is, an unusually busy minor road quieter. No traffic count enters any cell score.
5. **The satellite sees categories, not beauty.** Land cover cannot tell a majestic cathedral forest from a scrubby second-growth plantation, or a wildflower meadow from a mowed verge. The protected-area ingredient is the only, weak, proxy for forest quality.
6. **Flow's continuity is a blunt proxy, and cul-de-sacs read flowy.** Way length cannot tell a block of a continuous greenway from a stub, so the fade-out only catches fragments under about 90 metres, and quiet crescents with no tagged interruptions score high. Corridor flow also is not journey flow: it cannot count the stops a real route through the network would make.
7. **Crossings are structurally over-counted and patched, not fixed.** OpenStreetMap tags up to one crossing node per intersection leg; the low weight compensates on average but not per-place.
8. **Elevation misreads engineered grades, and some backcountry grade is inflated.** The 30-metre terrain model cannot see rail beds, benched trails, or bridge decks. Anti-phantom machinery contains the damage on the marquee rail trails, and bridges are unscored by design, but on steep-sidehill

backcountry tracks the residual noise still inflates grade, and no tuning can cool it without also erasing real climbs. The fix is better elevation (light detection and ranging (LiDAR) or recorded rides), not better arithmetic.

9. **Remoteness is straight-line distance to a directory.** No network distance (a store 2 km away across an unbridgeable river narrower than a kilometre still counts), no store hours, no seasonal closures, and no water sources, though water is the resupply axis that matters most exactly where remoteness is highest. Known directory gaps (two Kootenay-area stores at last audit) overstate remoteness in affected valleys; the error direction is at least the safe one.
10. **Half the ruggedness map is educated defaulting.** Roughly half of BC's unpaved ways carry no surface tag; their level comes from land-cover-gated defaults, flagged by a confidence field but displayed identically.
11. **Scores drift with OpenStreetMap, and no extract is stamped.** Every number is a function of one OpenStreetMap extract; edits, including fixes we prompt ourselves, change scores on rebuild, and we do not yet record which extract produced a published score, so figures cannot be exactly reproduced from a dated source. Recording it is planned work, not done work.
12. **What these scores do not claim.** They are not safety ratings, not suitability verdicts, and not recommendations. A comfort level is an audience size, not advice; a high grade is a fact about a hill, not a warning; a scenic 0.9 is a statement about surroundings, not a promise about your ride. Conditions, weather, construction, and traffic on the day are all outside the model.

11. What would improve this

The limitations above are also a work plan. In rough order of payoff:

- **Viewshed.** Computing what is actually visible from the road, water and mountain sightlines from elevation surfaces, is the single largest scenic upgrade: it converts proximity into visibility and finally lets mountain backdrop count. The elevation data exists; the computation is the work.
- **Interior calibration passes.** A structured acceptance ladder for the Okanagan, the Kootenays, and the North, built the same way as the Metro Vancouver one, ideally with local riders' eyes rather than ours.
- **A held-out expert check.** Scoring independently curated published day rides and testing whether they rank where expert curation says they should, used as validation, never as tuning.
- **Real traffic counts** for the peaceful score (and comfort): the BC Ministry of Transportation and Infrastructure publishes counts covering numbered highways best, exactly where the noise question matters most for touring.
- **LiDAR or recorded-ride elevation** for the marquee rail trails, where terrain elevation is most wrong and coverage is only a few hundred kilometres of named trail.
- **Network-distance remoteness with water sources**, replacing straight-line distance and grocery-only resupply.
- **Better OpenStreetMap, induced rather than awaited.** The pipeline already knows which missing tags cost the most: the untagged surfaces, unmapped stop signs, and missing shoreline features where one edit most improves the map. Published as a fix-list, that turns our weakest inputs into a concrete community mapping project.

12. Data and code availability

The character map is browsable at bccycletourism.ca/research/character-map/. The underlying data is derived from OpenStreetMap and carries OpenStreetMap records, which makes it a derivative database under the Open Database Licence (ODbL) 1.0, share-alike, with the attribution © OpenStreetMap contributors. This document itself is prose rather than data and is released under Creative Commons Attribution 4.0 (CC BY 4.0): share it and build on it, including commercially, with credit to BC Cycle Tourism Society.

The code behind the scores is not yet packaged for published release. If you are trying to reproduce a number, check a method, or build on a piece of it, tell us what you need and we will gladly send the relevant piece: heather@bccycletourism.ca.

For the route side of the system, how published routes are matched onto this scored network, how their labels and difficulty tiers are rolled up, and the full comfort validation crosstabs, see the companion document, *How BC Cycle Tourism Rates Routes: Methods and Limitations*, at bccycletourism.ca/research/.

The research here uses AI assistance (Claude models, Anthropic) for data collection, analysis, and writing. Feel free to ask for more details if you are interested.