

Article

Not peer-reviewed version

Universal Suffering Units (USU): A Calibrated Additive Unit of Experienced Suffering

[Denis Mikhaylov](#) *

Posted Date: 18 November 2025

doi: 10.20944/preprints202511.1315.v1

Keywords: Universal Suffering Units (USU); welfare metrics; pain intensity; calibration; uncertainty; DALY; QALY; Monte Carlo



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a [Creative Commons CC BY 4.0 license](#), which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Article

Universal Suffering Units (USU): A Calibrated Additive Unit of Experienced Suffering

Denis Mikhaylov

World Suffering Index (Independent Researcher), Sacramento, CA; denis@worldsufferingindex.org

Abstract

We present the Universal Suffering Unit (USU), a calibrated, additive unit of experienced suffering that scales with intensity, duration, and affected population, enabling aggregation across people, regions or countries, and time. We define USU as $k \cdot \sum I^p \cdot \Delta t$, where I is a 0–10 intensity rating, $p \geq 1$ is an exponent that modestly up-weights high intensities, and k is a calibration constant. We calibrate the unit so that a reference trajectory of renal colic (kidney-stone pain) equals 1.0 USU, propose a simple rule for co-occurring harms, and recommend reporting medians with 90% uncertainty intervals from Monte Carlo simulations. Using publicly available data, we illustrate the framework with two examples: dengue in Brazil (epidemiological weeks 1–23 of 2024) and flood-related displacement in Rio Grande do Sul, plus a year-over-year dengue comparison and a sensitivity analysis over p in {1.0, 1.25, 1.5}. These illustrations show how large numbers of moderate episodes and smaller numbers of longer, disruptive episodes can be expressed on a common experiential scale, while remaining interpretable via an anchor ladder. We discuss validation strategies, highlight ethical guardrails and misuse risks, and argue that USU is best used alongside DALY/QALY and routine operational indicators as a decision-support tool for comparing heterogeneous harms, rather than as a stand-alone welfare metric.

Keywords: Universal Suffering Units (USU); welfare metrics; pain intensity; calibration; uncertainty; DALY; QALY; Monte Carlo

1. Introduction

Allocating finite resources across qualitatively different harms requires a unit that permits transparent comparison. Time-based welfare metrics such as DALYs and QALYs summarize health lost over time, but they do not directly quantify the experienced burden itself. USU is designed to complement these frameworks by offering a calibrated unit of experience that aggregates over people and time while preserving interpretability [1–4].

2. Related Work

Pain intensity is routinely measured using numeric rating scales (NRS) and visual analog scales (VAS), validated in clinical and experimental contexts; these tools support interval/ratio interpretations when implemented carefully [5–7]. The revised IASP definition emphasizes that pain is an unpleasant sensory and emotional experience [8]. In population health, DALYs and QALYs summarize time-based health loss via disability weights and preference-based utilities, indispensable for long-run prioritization [1–4]. A well-developed literature shows peak-end weighting and duration neglect in remembered burden (cold-pressor, colonoscopy, lithotripsy studies) [12–14]. USU pragmatically incorporates these insights: it aggregates intensity-by-time, allows a convex intensity transform I^p ($p \geq 1$), and maintains interpretability by anchoring a familiar trajectory (renal colic = 1.0 USU). Reporting USU alongside DALY/QALY clarifies distinct dimensions: experienced burden vs time-based health loss [15].

3. Methods

3.1. Definition

$USU = k \cdot \sum_i \sum_t I_i(t)^p \cdot \Delta t$, where $I \in [0,10]$, $p \geq 1$, Δt in days by default, and k is a calibration constant. Additivity follows from linearity; monotonicity holds for $p \geq 1$.

3.2. Calibration

Fix a renal-colic reference trajectory $I_{ref}(t)$; compute $S_{ref} = \sum I_{ref}(t)^p \Delta t$ and set $k = 1/S_{ref}$ so the reference equals 1.0 USU. Anchor ladder: needle (~0.001), paper cut/stubbed toe (~0.01), fingertip burn (~0.1), renal colic (1.0).

3.3. Co-Occurring Harms

Within a person-day, rank by intensity and apply diminishing weights (1.0, 0.5, 0.25, ...). Alternatives (max, L_p) are in the supplement.

3.4. Uncertainty and Reporting

Treat counts/durations/intensity bands as uncertain, use simple distributions (e.g., triangular for intensity), compute totals via Monte Carlo, and report medians with 90% uncertainty intervals (5th–95th) [9–11].

3.5. Recipe

Specify events and population → assign intensity/duration bands → choose p (default 1.25) → fix renal-colic reference and compute k → apply overlap rule → run N draws → aggregate → report median + 90% UI → repeat for $p \in \{1.0, 1.25, 1.5\}$ with recalibrated k .

3.6. Ethics

This is a methodological study using only public, aggregate data; no research involving human or animal participants was conducted, and no ethics approval was required.

4. Results: Worked Examples (Illustrative, Sourced)

A) Dengue in Brazil, EW1–23 of 2024 (PAHO counts; CDC duration guidance). B) Rio Grande do Sul displacement, end-June 2024 (IDMC). Demonstrations, not national totals.

Table 1. Dengue in Brazil, EW1–23 2024 — medians with 90% UIs. Calibration $p=1.25$; renal-colic anchor = 1.0 USU. Sources: PAHO [16], CDC [11].

Component	Median USU	p5	p95
Non-severe cases	75,313,058	46,843,581	109,800,717
Severe cases	180,404	115,653	271,866
TOTAL	75,490,915	47,063,080	109,957,174

Table 2. Rio Grande do Sul floods (Brazil), 2024 — displaced stock end-June; medians with 90% UIs. Source: IDMC [17].

Component	Median USU	p5	p95
Displacement stock (end-June)	24,673,072	12,762,842	32,528,611

5. Additional Results: Year-Over-Year Comparison. Brazil Dengue, EW1–23 2023 vs 2024

PAHO reports 7,866,769 dengue cases in Brazil for EW1–23 2024, a 230% increase versus the same period in 2023 [16]. With identical mappings and fixed p and k, USU totals scale approximately with case counts. Medians: 2023 ≈ 22.9M; 2024 ≈ 75.5M; change ≈ +230%.

6. Sensitivity to the Intensity Exponent p

Recalibrating k per p keeps the renal-colic anchor at 1.0 USU. Dengue 2024 median (p = 1.00 / 1.25 / 1.50): ~63e6 / ~75e6 / ~90e6. Displacement medians rise similarly with p. See the CSV in the supplement.

7. Discussion

This paper has two main findings. First, a calibrated intensity-time unit such as USU is technically simple to construct and can be anchored to familiar experiences in a way that many readers find intuitive. Second, when applied even to rough, publicly available data, USU produces magnitudes that are easy to interpret and to compare across qualitatively different harms, such as infectious disease outbreaks and disaster-related displacement. The Brazil examples show that, with transparent assumptions, weeks of non-severe but widespread dengue and weeks of housing disruption for hundreds of thousands of people both generate tens of millions of USU, highlighting that “ordinary-seeming” events can represent a very large aggregate burden of experienced suffering.

In the dengue illustration, the vast majority of USU comes from non-severe cases rather than from the small fraction of severe disease. This is consistent with standard public-health intuition: large numbers of moderate episodes can dominate the total burden, even when each episode is far less intense than a small number of life-threatening cases. USU makes this trade-off explicit in experienced-burden terms, rather than only in clinical severity or mortality. The year-over-year comparison reinforces this point: under unchanged mappings, the ~230% increase in reported cases between EW1–23 of 2023 and 2024 is mirrored by a roughly proportional increase in total USU, making it straightforward to communicate the scale of deterioration to decision-makers.

The displacement example illustrates a complementary pattern. Here, the population is smaller but each episode is longer and sits at an intensity band associated with loss of housing, disruption of work and schooling, and reduced privacy and control. The resulting USU totals are comparable in order of magnitude to the dengue illustration, despite the very different mechanism of harm. In this sense, USU functions as a “bridge” between domains that are often analyzed with separate metrics (epidemiological indicators vs. disaster impact statistics). The unit is not intended to replace those domain-specific indicators, but to provide a common experiential denominator when choices must be made across them.

Conceptually, USU sits between time-based welfare measures such as DALY/QALY and short-term operational indicators. DALYs and QALYs summarize health loss over years of life and are designed for long-run priority-setting; they are less natural for very short-horizon trade-offs (for example, whether to shift surge staff from dengue wards to flood-displacement shelters this month). Conversely, operational dashboards often display counts (cases, occupied beds, displaced persons) without an explicit mapping to how bad these states are to live through. USU aims to fill this gap by

translating counts and durations into an anchored scale of experienced burden. In practice, we expect USU to be most useful when used alongside DALY/QALY and standard operational indicators, not in isolation.

From a methodological perspective, the examples also illustrate the role of the exponent p and the anchor choice. A convex transform I^p with $p > 1$ modestly up-weights the contribution of high-intensity experiences relative to low-intensity ones, which many people find normatively attractive. However, this convexity could quickly make USU totals dominated by rare extremes if not re-anchored. By recalibrating k for each p so that the renal-colic reference remains 1.0 USU, we keep the unit interpretable while still allowing moderate sensitivity to the tail of the distribution. The sensitivity table in the supplement shows that, for the examples considered, aggregate rankings and orders of magnitude are stable across $p \in \{1.0, 1.25, 1.5\}$; changing p changes the emphasis on the highest-intensity contributions but does not overturn basic qualitative conclusions.

A further question is how USU should be validated. The examples in this paper are illustrative rather than definitive, but they suggest several avenues. One is convergent validity: do USU time-series correlate as expected with external indicators such as emergency-department visits, hospital occupancy, sales of strong analgesics, or absenteeism? Another is face validity and expert judgment: do domain experts (clinicians, humanitarian practitioners, people with lived experience) find the relative magnitudes and rankings produced by a given set of mappings and overlap rules acceptable, after they are made transparent? Over time, empirical studies could refine default mappings from clinical categories into USU parameters, in the same way that disability-weight elicitation has refined DALY weights.

Finally, the examples highlight practical decision contexts where a calibrated experiential unit may be most useful. These include: ranking short-term crises for triage when “everything looks urgent”; communicating to the public or policymakers that a rise in cases or displacement corresponds, roughly, to “this many kidney-stone-equivalent episodes of suffering”; and comparing portfolios of interventions in terms of USU averted per unit of resource. None of these applications require that USU capture every aspect of welfare; they require only that the unit be transparent, reasonably calibrated, and used with appropriate guardrails.

8. Limitations and Ethical Guardrails

The main limitations of USU come from the fact that the intensity mappings and overlap rule are based on explicit choices, not fixed laws of nature.

A central ethical risk is over-interpreting USU as a complete welfare score or as a sufficient basis for high-stakes decisions. Because USU is explicitly designed to be additive, it could be misused to justify policies that impose very intense suffering on a small number of people in order to reduce many mild burdens elsewhere, or to treat “USU averted per dollar” as a single overriding objective. This risk is amplified by measurement error, by the fact that some harms (for example, social exclusion or loss of meaning) are only partially captured by current mappings, and by the temptation to “game” the metric through optimistic assumptions. For these reasons, USU estimates should always be accompanied by clear disclosures of assumptions and uncertainty, and interpreted within a broader ethical and democratic process rather than as a stand-alone decision rule.

USU is intended to inform comparison, not to replace moral judgment; in particular, it should not be used to justify severe harm to a few to offset mild harms to many.

9. Conclusions and Future Directions

This article has proposed the Universal Suffering Unit (USU), a calibrated additive unit of experienced suffering, anchored to a familiar reference pattern (renal colic) and illustrated with two real-world examples. The core contribution is not a particular set of numbers for Brazil in 2024, but a general recipe: map events into intensity–time profiles, calibrate a unit to an anchor trajectory, and aggregate across people and time in a way that is transparent and reproducible. Within this

framework, different analysts can disagree about mappings, the exponent p , or the overlap rule, but those disagreements can be expressed and debated explicitly.

The worked examples suggest that USU is capable of capturing both high-incidence, moderate-severity harms (non-severe dengue) and lower-incidence, longer-duration harms (displacement) on a common experiential scale, while remaining interpretable via the anchor ladder. They also show that simple year-over-year comparisons and sensitivity analyses are straightforward to compute and to communicate. At the same time, these illustrations are deliberately conservative: they use public aggregate data and stylized assumptions rather than detailed micro-data. A natural next step is to apply USU to richer datasets at national or sub-national level and to compare the resulting rankings with those from DALY/QALY and other burden-of-disease metrics.

Future work falls into at least four directions. First, empirical calibration: eliciting more precise mappings from clinical or experiential categories into USU parameters (intensity bands, durations, overlap rules), using both expert judgment and data from patient-reported outcomes. Second, domain-specific adaptations: for example, tuning default mappings for mental-health conditions, chronic pain, or conflict-related harms, and examining whether different domains warrant different default values of p . Third, integration with existing metrics and systems: embedding USU calculations into routine surveillance and humanitarian dashboards, so that decision-makers can see experienced-burden estimates alongside counts, DALYs, or cost indicators. Fourth, extensions to other welfare domains: for example, adapting the framework to certain animal-welfare contexts involving human-caused suffering, or exploring symmetric constructs for positive experiences.

More broadly, USU should be understood as a decision-support tool rather than a moral calculus. It does not answer normative questions about whose suffering “matters more,” and it does not settle debates about distributive ethics. What it can do is make assumptions explicit, reduce some dimensions of welfare to a calibrated and interpretable scale, and support clearer reasoning about trade-offs when choices must be made. The examples here are intended as a starting point, and the accompanying templates and supplementary files are provided in the hope that other researchers and practitioners will test, critique, and refine both the unit and its applications.

Data Availability Statement: The supplement includes: anchors, calibration summary, minimal template, dengue inputs/results (2024), displacement inputs/results, p -sensitivity CSV, and YoY inputs/results.

Acknowledgements: Drafting/editing assistance was provided by OpenAI GPT-5 Pro under author direction; the author verified all content and assumes responsibility for errors.

Conflicts of Interest: Independent researcher; founder of the World Suffering Index. Role did not influence methods.

References

1. Vos T, Lim SS, Abbafati C, et al. *Lancet*. 2020;396:1204–1222.
2. Salomon JA, Haagsma JA, Davis A, et al. *Lancet Glob Health*. 2015;3(11):e712–e723.
3. Murray CJL, Lopez AD, eds. *The Global Burden of Disease*. 1996.
4. Cieza A, Causey K, Kamenov K, et al. *Lancet*. 2020;396:2006–2017.
5. Huskisson EC. *Lancet*. 1974;2:1127–1131.
6. Price DD, McGrath PA, Rafii A, Buckingham B. *Pain*. 1983;17:45–56.
7. Jensen MP, Karoly P, Braver S. *Pain*. 1986;27:117–126.
8. Raja SN, Carr DB, Cohen M, et al. *Pain*. 2020;161(9):1976–1982.
9. Rubinstein RY, Kroese DP. *Simulation and the Monte Carlo Method*. Wiley; 2016.
10. O’Hagan A, Buck CE, Daneshkhah A, et al. *Uncertain Judgements*. Wiley; 2006.
11. CDC. Clinical Features of Dengue. <https://www.cdc.gov/dengue/hcp/clinical-signs/index.html>
12. Kahneman D, Fredrickson BL, Schreiber CA, Redelmeier DA. *Psychol Sci*. 1993;4:401–405.
13. Redelmeier DA, Kahneman D. *Pain*. 1996;66:3–8.
14. Fredrickson BL, Kahneman D. *J Pers Soc Psychol*. 1993;65:45–55.

15. Kahneman D, Wakker PP, Sarin R. QJE. 1997;112:375–405.
16. PAHO/WHO. Epidemiological Update, 18 Jun 2024. <https://www.paho.org/>
17. IDMC. 2024 Mid-year Update. <https://story.internal-displacement.org/2024-mid-year-update/>

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.