

Certified Bounds on Achievable Legibility under a Path Cost Budget

Munawar Kazmi

Abstract—Legible motion is optimised rather than decided, and for good reason: the optimum is intractable. What an optimiser reports is the clarity it found, which does not establish that no clearer trajectory exists, so a question of the form “can any trajectory within this budget be read correctly here” has not been answerable. This paper answers it for a bounded planar problem without computing the optimum, by bracketing it instead. Given a world of convex polygonal obstacles, a finite goal set, a Boltzmann-rational observer and a ceiling on path cost, it returns a trajectory achieving a stated legibility together with a bound that no trajectory within the budget exceeds, and reports the gap between them rather than hiding it. Over 8 published scenarios at 4 cost ceilings the gap is at most 0.0567 and no bound is violated. The same argument bounds trajectories constrained to avoid keep-out zones, which turns the usual comparison of two searches into a certified lower bound on what such a constraint costs in legibility.

Index Terms—Legible motion, human-robot interaction, motion planning, formal guarantees.

I. INTRODUCTION

A ROBOT that moves legibly lets an observer infer where it is going before it arrives. The formulation of Dragan, Lee and Srinivasa [1] makes this precise: an observer holds a posterior over candidate goals, updates it from the motion so far under an assumption of efficiency, and legibility is that posterior’s mass on the true goal averaged along the trajectory with early motion weighted more heavily.

Work since has optimised this objective, and with reason: computing the optimum is intractable in general, being PSPACE-complete in the stochastic sequential formulation and NP-hard even for deterministic environments [2]. What a trajectory optimiser reports is therefore what a particular search found under a particular budget, which is a different statement from what is achievable. The distinction matters whenever the interesting claim is negative. A scenario designer who wants to assert that a world is genuinely hard, that no trajectory within a twenty five per cent cost budget can be read correctly there, cannot get that from a search: a search that failed to find one has not shown that none exists.

This paper closes that gap for a bounded problem. The contribution is an instrument that returns two numbers rather than one: a trajectory achieving a stated legibility, and a bound no trajectory within the budget can exceed. The interval between them is the answer, and its width is reported.

Three things follow that are worth stating separately. The bound is exact in the geometry rather than sampled: cost-to-go comes from a visibility graph and carries no discretisation. The lower end is constructed from the bound itself rather than searched for, which matters because a search is exactly what cannot be trusted here. And because the argument is about a

feasible set, it transfers unchanged to the safety-constrained problem, which yields a certified statement about the price of a constraint.

II. RELATED WORK

The formulation this paper bounds is that of Dragan, Lee and Srinivasa [1]. Legibility there is the observer’s posterior in the true goal, averaged along the trajectory under a weight that favours early motion, and divided by the integral of that weight. The objective used here is theirs and is not modified.

Two statements in that paper shape what a result about it is allowed to say. The first is that legibility cannot be driven to one: with more than one candidate goal a robot can make a trajectory more and more legible without ever reaching a score of one, at ever greater cost. A value near one is therefore not a proportion of an attainable ideal, and is not reported as one here. The second follows immediately from it. Because the objective alone runs away with cost, that paper adds a regulariser to keep the robot from departing too far from what the observer expects. The cost ceiling in this work is the same idea in hard form and is not an invention of ours.

Its companion [3] makes the bound on cost an explicit constraint, maximising legibility subject to a ceiling on trajectory cost, which it calls a trust region of predictability. That paper is direct about what lies outside it: the legibility model can only be trusted inside the trust region. This is why every number reported here is reported at a stated ceiling, and why an unbounded optimum is treated as a diagnostic rather than as a result.

A separate line formalises observer-aware behaviour in stochastic sequential settings. Miura and Zilberstein [2] unify legibility with explicability and predictability as observer-aware Markov decision processes, in which each property is a choice of reward over the observer’s belief, and establish the complexity of the resulting problem: it is PSPACE-complete for a Bayesian observer, and remains NP-hard when restricted to stationary policies or to deterministic environments. Approximation and partially observable extensions follow, for instance [4].

That result places this paper. It does not compute the optimum either; it brackets it. The setting is the continuous one of [1], with exact cost-to-go rather than a discretised state space, and the two ends of the bracket are a trajectory that exists and a bound no trajectory within the budget can pass. Optimality remains computable in discretised stochastic formulations, up to the discretisation, so the claim here is specifically about the continuous setting under a path cost budget.

In the continuous setting the tools remain local. Recent work incorporates legibility into cost-based local planning for mobile robots in real time [5], which is the right instrument for driving a robot and is not one that can decide a negative statement about a world.

Two findings in the founding papers bear directly on sections below, and one line of subsequent work bounds what may be claimed from the first of them.

Dragan et al. [1] observe that legibility is not obstacle avoidance, and that a legible trajectory exaggerates away from competing goals even when that carries it closer to a wall. They identify the phenomenon and do not measure it.

That is not to say the interaction goes unreported. Bastarache et al. [6] plan for legibility among moving agents and report minimum pairwise distance alongside it, together with minimum time-to-collision, which they use as a proxy for legibility and safety together. So the contribution here is not that safety appears beside legibility for the first time. It is narrower, and it is threefold: the safety quantity is satisfaction of a stated static constraint rather than proximity to a moving agent, it is bounded over all admissible trajectories rather than measured on the ones a planner produced, and the two together yield a price that is certified rather than compared.

Lastly, the companion paper [3] reports that observers may stop reasoning about the declared goals altogether once motion becomes strange enough, and may come to believe the robot is pursuing none of them, a belief its follow-up study measures directly. The observer used here cannot represent that belief, and Section VI says so.

Two scoping points follow from more recent work. Guarantees are not absent here: Liu et al. [7] give a convergence and completeness analysis for their multi-agent legibility algorithm. That guarantees an algorithm reaches what it converges to, which is a different object from a bound on what any trajectory could achieve, and it is the second sense the word certified carries in this paper. Nor is there one legibility: recent generative work produces motion across a controllable range of intent expressiveness rather than a single most legible trajectory, scored by a potential field of its own construction [8]. What is bounded here is the formulation of Dragan et al. and nothing else. Beyond these, the section is not a survey of everything since 2013.

III. PROBLEM

A world is a compact planar region containing finitely many disjoint convex polygonal obstacles, a start S , and a finite goal set $\mathcal{G}$ with a true goal $A \in \mathcal{G}$. Write $C^*(x, G)$ for the length of the shortest obstacle-avoiding path from x to G . This is computed exactly, by shortest path search over the visibility graph on the obstacle vertices, so nothing below rests on a discretisation of the geometry.

An observer of [1] assumes the robot is efficient and scores each goal by how much the motion so far has cost against the best it could have done. After the robot has travelled a distance ℓ and arrived at x , the belief in goal G is

$$b_G(x) \propto p_G \exp \beta (C^*(S, G) - \ell - C^*(x, G)), \quad (1)$$

normalised over $\mathcal{G}$, with p_G a prior and $\beta > 0$ a rationality coefficient. Write $b = b_A$ for the belief in the true goal.

A trajectory ξ of length L is sampled at $N + 1$ points $x_0 \dots x_N$ spaced equally in arc length, and scored by averaging b over them under the weight $f(t) = T - t$ of [1], divided by the sum of those weights:

$$\text{Leg}(\xi) = \frac{\sum_i b(x_i) (T - t_i)}{\sum_i (T - t_i)} = \frac{\sum_i b(x_i) (1 - i/N)}{\sum_i (1 - i/N)}. \quad (2)$$

Its cost is constrained by a ceiling c on the ratio of L to $C^*(S, A)$, which is the trust region of [3] expressed as a ratio, for the reasons given in Section II.

Two consequences of these definitions do the work below, and neither is apparent on first reading.

The belief is a field. The travelled distance ℓ enters every goal's score in (1) identically, so it cancels under normalisation. The observer's belief is therefore a fixed scalar function of position over the free space, not a quantity carried along the path, and (2) is a weighted average of a fixed field. Everything the bound needs to know about a world can be computed once, before any trajectory is considered.

The objective does not see duration. Sampling at equal arc length puts $t_i = iT/N$, so the weights are $T(1 - i/N)$ and the factor T cancels between the numerator and the denominator of (2). That is the second equality there. The weighting is over normalised arc length, and a bound does not have to range over the interval of durations a cost budget permits. What survives is the sample count N , which follows the length.

IV. THE BOUND

A. Reachability

Let ξ be admissible, meaning it runs from S to A with $L \leq cC^*(S, A)$, and let x be its point at normalised arc length s . The arc from S to x is at least the shortest path between them, and likewise from x onwards, so

$$C^*(S, x) \leq sL \quad \text{and} \quad C^*(x, A) \leq (1 - s)L. \quad (3)$$

Write $R(s)$ for the set these two conditions cut out. Every sample of every admissible trajectory lies in the $R(s)$ for its own s , so substituting the largest belief available there into (2) bounds every admissible trajectory at once:

$$\text{Leg}(\xi) \leq \frac{\sum_i (\max_{x \in R(i/N)} b(x)) (1 - i/N)}{\sum_i (1 - i/N)}. \quad (4)$$

Both radii in (3) grow with L , so $R(s)$ grows with it, and evaluating (4) at the largest admissible length bounds every shorter one. One evaluation therefore covers the whole budget rather than a single path length.

What (4) discards is the coupling between samples: it lets each one sit wherever its own $R(s)$ allows, independently of the rest. It bounds a trajectory that need not be a trajectory. Section VI reports how much that costs, which is less than it appears.

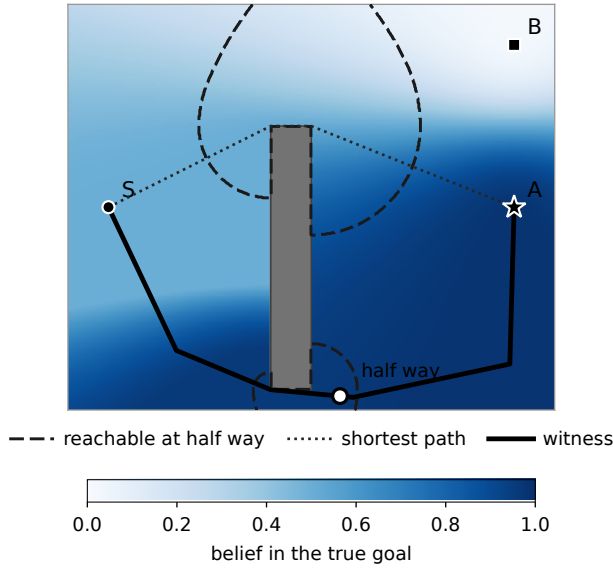


Fig. 1. The mechanism in one world, with a wall between the start and both goals. Shading is the observer's belief in the true goal A , which depends on position alone and is computed from exact geodesic cost-to-go. The outlined region is where (3) permits a trajectory to be when it is half way along, and it has two parts because the wall can be passed on either side. The shortest path takes the upper one and reads ambiguously for most of its length; the witness takes the lower one, which costs path and buys belief. The marked point is the witness at its own half way point, which lies inside the region as any admissible trajectory must.

B. Bounding the belief over a cell

Evaluating (3) on a lattice requires bounding the belief over a cell rather than at a point, and near an obstacle that is where the difficulty lies: two points a millimetre apart on opposite sides of a wall are a wall's length apart in the geodesic metric, so no small bound on the variation of cost-to-go is available inside such a cell.

Writing the belief as odds against the true goal removes the need for one. Let

$$\Delta_G(x) = (C^*(S, G) - C^*(x, G)) - (C^*(S, A) - C^*(x, A)), \quad (5)$$

so that the belief b in the true goal is

$$b(x) = \left(1 + \sum_{G \neq A} \frac{p_G}{p_A} e^{\beta \Delta_G(x)}\right)^{-1}. \quad (6)$$

Bounding b from above needs a lower bound on Δ_G , which is a *difference* of cost-to-go values, not an upper bound on either. The upper bound is what an obstacle denies; the difference is available. Each cost-to-go is 1-Lipschitz in the geodesic metric, so moving a distance D within a cell shifts Δ_G by at most $2D$, and for every x in a cell whose lattice point is p ,

$$b(x) \leq (1 + \text{odds}(p) e^{-2\beta D})^{-1}. \quad (7)$$

The obstacle problem reduces to one scalar per cell: an upper bound D on the geodesic distance from a point of the cell to its centre.

C. The detour bound and its precondition

Take a cell of radius r about a lattice point p . Away from any obstacle $D = r$, since the segment from p to any point of the cell is free and is therefore the geodesic between them.

Near an obstacle no universal bound exists. An obstacle thinner than a cell separates two of the cell's points by its own length rather than by the cell's, and that length is unrelated to r . What exists instead is a bound under a precondition, and the precondition is checkable. Suppose the obstacle does not pass clean through the cell and no second obstacle meets it. Then three facts compose:

- 1) from a point outside a convex body, the ray heading directly away from that body's nearest point never re-enters it, so any point of the cell reaches the cell's boundary along a free ray;
- 2) p is the centre, so its own such ray is of length exactly r , and any other point of the cell reaches the boundary in at most $2r$;
- 3) the free part of that boundary is a single arc, since the obstacle does not cross the cell, so the two arrival points are joined along it in at most $2\pi r$.

Hence $D \leq (3 + 2\pi)r$, which scales with the lattice as r does.

The precondition is decided for each cell separately, by counting where each obstacle's boundary crosses that cell, and a cell that fails it is bounded by one, which holds unconditionally. Deciding it per cell rather than per world is not fastidiousness. Minimum width is a global property of an obstacle, and a convex polygon with a sharp vertex can be many cells wide overall while its tip is thinner than one, so a test applied to the whole world would certify cells that the argument does not cover.

In practice almost no cell needs the wrapping bound at all. Consider a cell that no obstacle passes through and that contains no obstacle vertex. Two edges of a convex polygon meet only at a vertex, so that cell meets the obstacle's boundary within a single edge, which cuts the cell with a half plane and leaves the free part convex. Then $D = r$ there as well. Only cells holding a corner are left, and for a lattice of spacing h there are $O(1)$ of those per obstacle against $O(1/h)$ cells along its boundary.

D. The witness

The lower end of the interval is constructed rather than searched for, because a search is the thing this paper exists to avoid trusting. The bound already identifies, for each fraction of the way along, the reachable cells where the belief is highest. A witness takes the best of those as anchors and threads a trajectory through them, joining consecutive anchors by exact geodesics so the result cannot pass through an obstacle however the anchors fall, and pulling them back towards the shortest path until the cost ceiling holds. The result is scored by the same metric as everything else, so what it claims is what it measures. Against a local search at 500 evaluations it produces the better trajectory in 13 of the 32 cases below, by up to 0.1485; where the search wins, the search's value is used.

TABLE I

EACH WORLD AT THE COST CEILING WHERE ITS INTERVAL IS WIDEST. *ach* IS THE BETTER OF A LOCAL SEARCH AND A WITNESS BUILT FROM THE BOUND; *bnd* IS THE UPPER BOUND NO TRAJECTORY WITHIN THE BUDGET CAN EXCEED. THE LAST COLUMN IS THE SAME WORLD'S NARROWEST GAP OVER ALL FOUR CEILINGS. NO BOUND IS VIOLATED ANYWHERE IN THE GRID.

world	widest				narrowest
	<i>c</i>	<i>ach</i>	<i>bnd</i>	gap	gap
door_pair	1.50	0.8507	0.8776	0.0269	0.0066
fan_middle	1.25	0.4343	0.4909	0.0567	0.0144
fan_outer	1.10	0.6410	0.6692	0.0282	0.0162
keep_out_shortcut	1.05	0.7870	0.8068	0.0198	0.0067
narrow_gap	1.50	0.7584	0.8041	0.0456	0.0268
open_pair	1.05	0.7870	0.8068	0.0198	0.0067
pillar_aisle	1.05	0.8022	0.8205	0.0184	0.0071
wall_choice	1.10	0.5787	0.5915	0.0128	0.0064

V. RESULTS

The instrument is run on the 8 scenarios of a published legible motion benchmark, at 4 cost ceilings, on a lattice of 0.0125 world units, under the observer that can see the obstacles. Table I reports each world at the ceiling where its interval is widest, together with its narrowest gap over all four.

Across all 32 pairs there are 0 violations: no achieved trajectory ever exceeds the bound claimed against it. The gaps run from 0.0064 to 0.0567. The widest is in *fan_middle* at ceiling 1.25, a world with no obstacles at all, which is worth noting because the obstacle machinery is where the argument is loosest and is not where the widest interval sits.

A. What a row licenses

Take *wall_choice* at a cost ceiling of 1.50, the world drawn in Fig. 1. This is that world's narrowest interval, the 0.0064 in the last column of its row in Table I. The bound is 0.8068 and a trajectory attaining 0.8004 is exhibited, the one drawn. In words, and this is the form a scenario designer needs:

No trajectory from *S* to *A* in this world, spending at most 50 per cent more path than the shortest one, attains a legibility of 0.81 or above under this observer. One attaining 0.8004 exists.

Three properties of that statement are worth separating.

It is negative, and it is about every trajectory rather than about the ones someone looked at. No amount of further searching can overturn it, and a better optimiser would not weaken it. This is what a search cannot supply at any budget: a search that returns 0.8004 leaves open whether 0.81 was merely missed.

It is not vacuous. The two ends are 0.0064 apart, so any threshold outside that band is decided: above it, unreachable; at or below the achieved value, reached. Only a threshold falling inside the band is left open. For contrast the shortest path in this world scores 0.5457, so the budget does buy clarity here, and the question of how much is answered rather than bracketed by ignorance.

It is conditional in ways that must travel with it. It holds for this observer, at this rationality coefficient, under this cost ceiling, and it says nothing about a person watching the robot.

TABLE II

WHAT RESPECTING A KEEP-OUT ZONE CERTIFIABLY COSTS. *free* IS A TRAJECTORY THAT EXISTS WITHIN THE BUDGET AND MAY CROSS A ZONE; *safe bnd* IS A BOUND NO TRAJECTORY THAT AVOIDS ONE CAN EXCEED. WHERE THE FIRST IS LARGER, THE DIFFERENCE IS A CERTIFIED LOWER BOUND ON THE PRICE OF THE CONSTRAINT.

world	<i>c</i>	<i>free</i>	<i>safe bnd</i>	price
keep_out_shortcut	1.05	0.7870	0.7911	none
keep_out_shortcut	1.10	0.8079	0.8009	0.0070
keep_out_shortcut	1.25	0.8353	0.8225	0.0128
keep_out_shortcut	1.50	0.8594	0.8661	none
pillar_aisle	1.05	0.8022	0.7558	0.0464
pillar_aisle	1.10	0.8208	0.8370	none
pillar_aisle	1.25	0.8456	0.8572	none
pillar_aisle	1.50	0.8682	0.8772	none

Section VI is not decoration; the statement above is only as good as the model it quantifies over. That applies with particular force here, since 1.50 is the loosest ceiling reported and model error grows with the cost ratio. The rows at tighter ceilings carry less of that risk and wider intervals, and Table I gives both ends of that trade.

The same reading applies to every row of Table I, which is the point of reporting the suite rather than one world. A scenario suite whose worlds carry properties of this kind can assert that a world is hard, rather than that a particular planner found it so.

B. What safety costs

Keep-out zones are regions a trajectory may enter and is penalised for entering, rather than obstacles that block motion. The usual way to measure what avoiding them costs is to run two searches and compare, which cannot establish that the constraint costs anything: a search that did worse under a constraint may simply have been a worse search.

The bound transfers to the constrained problem with one change. A keep-out zone constrains the robot and not the observer, so it enters the reachability conditions of (3) and leaves the belief field of (6) untouched. The two are therefore computed against different blocking sets over the same lattice.

Putting the constrained bound beside an unconstrained trajectory that exists gives a certified lower bound on the price of the constraint: something is achievable within the budget, and nothing that respects the zone can match it, so the constraint costs at least the difference. Of the 8 world and ceiling pairs carrying keep-out zones, 3 certify a positive price, the largest 0.0464 in *pillar_aisle* at ceiling 1.05. The remainder certify nothing and are reported as nothing rather than as a negative price.

The largest of the three arises in *pillar_aisle* at ceiling 1.05, where the cheapest route is not safe to begin with, so the constraint binds from the start rather than only once the robot tries to communicate.

Worked in full, in the world built to ask this question. In *keep_out_shortcut* the zone sits exactly where a trajectory would deviate to signal the true goal, and the cheapest route to that goal does not enter it. At a ceiling of 1.25:

The shortest route scores 0.6968 and is already safe, so the constraint costs nothing to a robot that does

not try to be understood. A trajectory scoring 0.8353 exists within the budget, and it crosses the zone. No trajectory within the same budget that stays out of the zone attains 0.8225. So being understood here costs at least 0.0128 of legibility once the zone is respected.

The structure of that claim is what distinguishes it from a comparison of two planners. The first figure is a trajectory that exists, the third is a bound over all trajectories, and only their combination is a statement about the constraint rather than about a search. A safety-constrained search reached 0.7890 here, which is consistent with the bound and, on its own, would have established nothing: a search that does worse under a constraint may only have been a worse search.

The price is a lower bound and not an estimate. The true cost of the constraint is at least 0.0128 and may be more, since the safe optimum lies somewhere at or below 0.8225. Reported the other way round, as a difference between two searches, the same world would have appeared to cost 0.8353 minus 0.7890, a larger number with nothing behind it.

VI. LIMITATIONS

The bound is exactly reproducible and is not validated against people. It certifies a statement about a stated observer model under a stated budget, and says nothing about whether a person watching the robot would agree with that model. Every result is reported at a ceiling for the reason [3] gives, that the model can only be trusted inside one.

The observer has no option to believe in neither goal. Its posterior sums to one over the declared goals however strange the motion becomes. This is not a hypothetical failure: Dragan and Srinivasa [3] report that observers do form exactly that belief once motion departs far enough from expectation, and measure it. It is a known direction of error here, and it grows with the cost ratio, so a result at a loose ceiling rests on the model further from where it was checked than one at a tight ceiling. Both are reported, and the ceiling is stated with every number for that reason.

The results are one lattice, one observer condition and one search budget over 8 worlds. Nothing here is a trend.

A. Where the residual width lives

Two candidates account for the intervals reported above, and they can be told apart by measurement.

The first is the relaxation of (4), which discards every constraint linking one sample to the next. It is natural to suppose most of the width lives there. It does not: restoring one such constraint, by requiring every admissible trajectory to pass through a common point at a stated fraction of the way along and maximising over that point, tightens the widest bound in the suite by less than a hundredth, and costs more than refining the lattice would.

The second is the lattice, and that is where most of it is. Table III recomputes the same intervals at four resolutions against a fixed achieved value. Refining from 0.05 to 0.00625 closes between 55 and 79 per cent of the gap across the 3 worlds tested. The bound is therefore not converged at the

TABLE III
GAP AGAINST LATTICE SPACING, AT COST CEILING 1.25. THE ACHIEVED VALUE DOES NOT DEPEND ON THE LATTICE, SO EACH ROW IS THE SAME INTERVAL MEASURED AT FOUR RESOLUTIONS. MOST OF THE RESIDUAL WIDTH IS THE LATTICE.

world	0.05	0.025	0.0125	0.00625
narrow_gap	0.0958	0.0634	0.0373	0.0331
pillar_aisle	0.0134	0.0093	0.0071	0.0060
wall_choice	0.0456	0.0216	0.0134	0.0096

resolution reported, and the numbers in Table I are conservative by roughly that margin.

What remains after refining is concentrated in the worlds with obstacles, and part of the reason is the constant of Section IV. The argument gives $D \leq (3 + \pi)r \approx 6.14r$, a worst case no configuration has been exhibited for. Sampling real points in real cells beside obstacle corners, measured to each point's own lattice point, finds a worst detour of $0.99r$, so the constant is loose by a factor of about 6.2. A bound must hold in the worst case and the worst case is rarely met, so this is not an error.

It is, however, worth less than it appears, for two separate reasons.

The first is that the constant barely moves the intervals. Halving it from $(3 + 2\pi)$ to $(3 + \pi)$ closed 1.6 per cent of the suite's total width, because twenty of the thirty two pairs draw no weight from the band and cannot move however tight it becomes, while refinement closes 55 to 79 per cent. Measured to each point's own lattice point rather than between two arbitrary points of a cell, and over the cell rather than over a box of the cell's diagonal, no sampled point in any world tested here is separated from its lattice point at all, so the wrapping argument is never binding in this suite.

The second is that most of the remaining looseness is not available to anyone. A measurement of slack says how much room there is in these worlds; it does not say how much of it a better argument could take. That is bounded from the other side, by exhibiting a configuration the constant must survive. Place a cell of unit radius with its lattice point free at the origin, and the obstacle

$$(-1.9001, -1.0448), \quad (0.8577, -0.5141), \\ (0.8763, 0.4817),$$

which is a single convex polygon that does not pass clean through the cell and so certifies the precondition of Section IV. The point $(-0.6055, -0.7958)$ lies in the cell and outside the obstacle, and its exact geodesic distance to the lattice point is 3.49 cell radii. No constant below that can hold. The whole of the improvement available to any sharper argument is therefore a factor of 1.76, and the remaining looseness against 0.99 is a property of the scenarios rather than of the proof.

Two things about that configuration are worth stating, because both are easy to read past. Its obstacle has minimum width 0.86 cell radii and is thinner than the cell, so it would fail the global width test discussed in Section IV and passes the per-cell test that replaced it. It bounds the constant that is actually in force rather than the one an earlier version assumed.

And it is close to the boundary of admissibility: the supremum is approached as the obstacle nears the lattice point and is not attained, so 3.49 is a lower bound on the sharp constant and not the sharp constant itself, which is not determined here.

Where the residual width in obstacle worlds lives is therefore still open, and a sharper constant is not the answer to it.

VII. CONCLUSION

Legibility has been reported as a search outcome because its optimum is out of reach. This paper shows that the optimum does not have to be computed for the useful statements to be made. Bracketing it is enough, and the bracket is narrow: over 8 published worlds at 4 cost ceilings, every interval is two-sided, none is violated, and the widest is 0.0567.

What that buys is a negative statement quantified over every trajectory rather than over the ones a planner returned, and that changes what a benchmark can assert about itself. A scenario suite can currently record that its own search found nothing better in a world; it cannot record that a world is hard. With a bound it can carry the property outright, machine-checked at the same standing as its geometry: no trajectory within this budget is read correctly here. A designer building such a world can then be told whether it does what it was built to do, rather than whether one planner struggled with it, and a planner evaluated on it can be scored against what was achievable rather than against whatever the baseline happened to reach. The same argument, applied to a smaller feasible set, prices a safety constraint in the same terms.

The obvious extensions are a second observer condition, non-convex or non-polygonal obstacles, and moving goals. The one that would change the standing of all of it is different: the objective bounded here is exactly reproducible and has never been validated against people, and no amount of tightening addresses that.

CODE AND DATA AVAILABILITY

The implementation, the committed records behind every figure and table in this paper, the tools that regenerate them, and a working log recording what was tried and what was withdrawn are available at <https://github.com/munawarkazmi/legibility-bounds>. Every number stated here is produced by a tool in that repository rather than typed, and the suite reproduces from a clean clone.

REFERENCES

- [1] A. D. Dragan, K. C. T. Lee, and S. S. Srinivasa, “Legibility and predictability of robot motion,” in *Proceedings of the 8th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*, 2013, pp. 301–308.
- [2] S. Miura and S. Zilberstein, “A unifying framework for observer-aware planning and its complexity,” in *Proceedings of the Thirty-Seventh Conference on Uncertainty in Artificial Intelligence*, ser. Proceedings of Machine Learning Research, vol. 161. PMLR, 2021, pp. 610–620.
- [3] A. Dragan and S. Srinivasa, “Generating legible motion,” in *Proceedings of Robotics: Science and Systems*, Berlin, Germany, Jun. 2013.
- [4] S. Lepers, V. Thomas, and O. Buffet, “Observer-aware probabilistic planning under partial observability,” in *Proceedings of the 24th International Conference on Autonomous Agents and Multiagent Systems (AAMAS)*, 2025, extended abstract; complete version arXiv:2502.10568.

- [5] J. Amirian, M. Abrini, and M. Chetouani, “Legibot: Generating legible motions for service robots using cost-based local planners,” in *Proceedings of the IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)*, Pasadena, CA, USA, 2024, arXiv:2404.05100.
- [6] J.-L. Bastarache, C. Nielsen, and S. L. Smith, “On legible and predictable robot navigation in multi-agent environments,” in *Proceedings of the IEEE International Conference on Robotics and Automation (ICRA)*, London, United Kingdom, 2023, pp. 5508–5514.
- [7] Y. Liu, Y. Pan, Y. Zeng, B. Ma, and P. Doshi, “Active legibility in multiagent reinforcement learning,” 2024, arXiv:2410.20954.
- [8] W. Shi, C. Grislain, O. Sigaud, and M. Chetouani, “Controlling intent expressiveness in robot motion with diffusion models,” 2025, arXiv:2510.12370.