

Runtime Safety Filtering for Learned Small UAS Separation Policies under GNSS Degradation

Action filtering vs observation filtering

Alex Zongo · Peng Wei

Department of Mechanical and Aerospace Engineering
George Washington University, Washington, DC, USA.

01 MOTIVATION

A learned separation policy can resolve conflicts

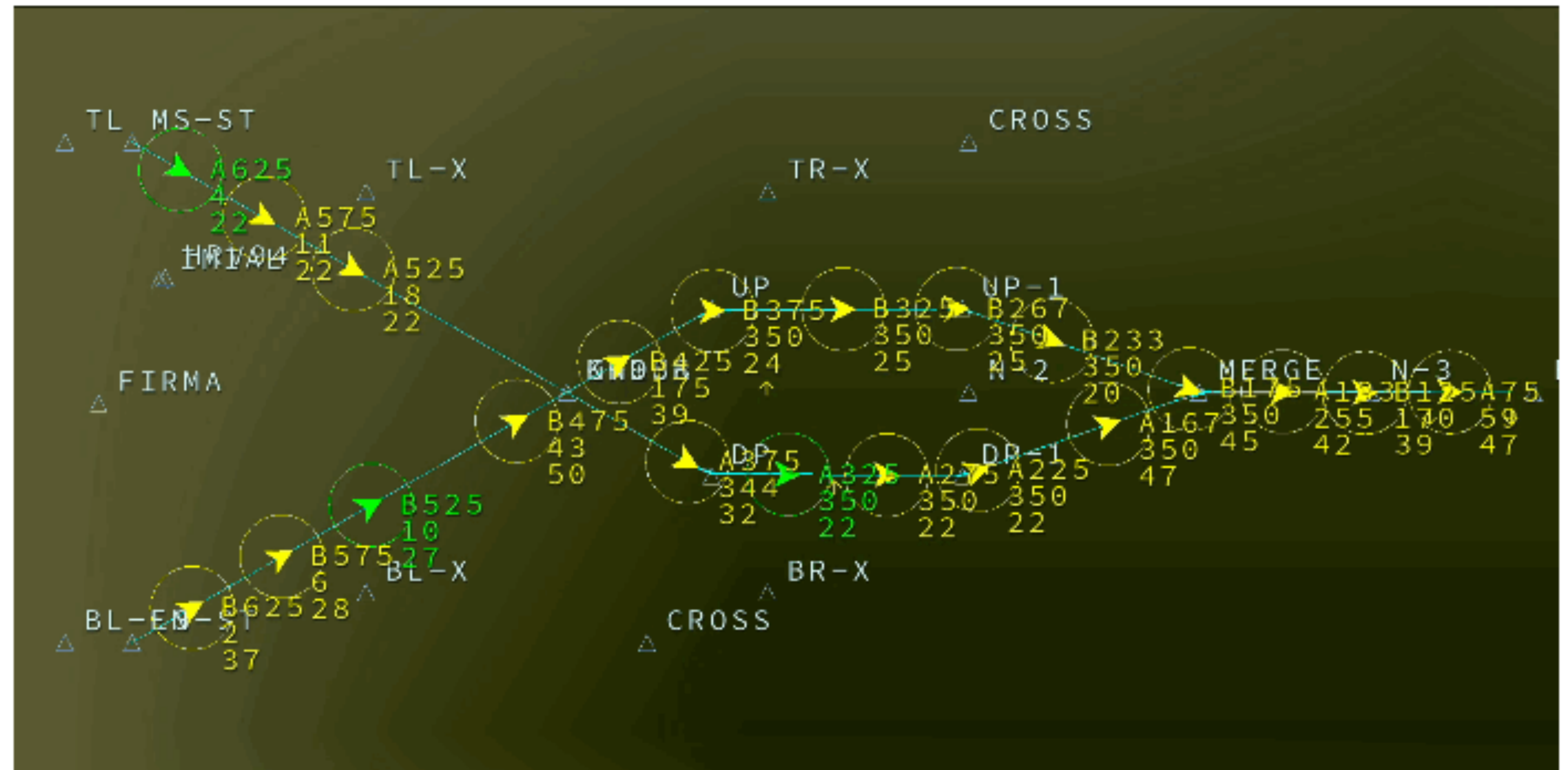
SIMULATION · NO GNSS DEGRADATION (NOMINAL GNSS)

With nominal GNSS, the policy produces a **separation-preserving action**.

The policy uses nearby aircraft information to adjust ownship speed in real time.

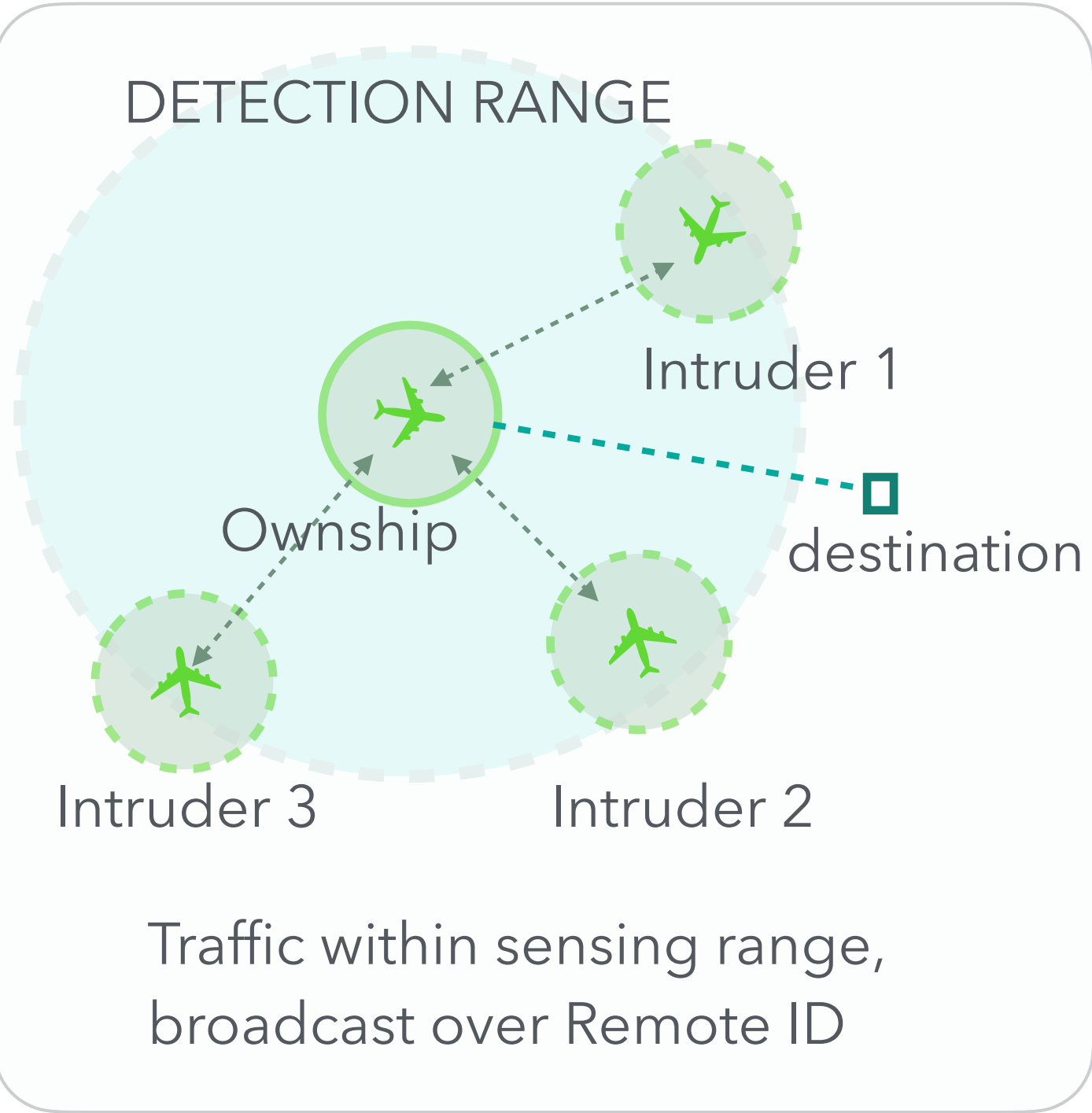
100 m separation standard
10 km routes

Next: What information does the policy receive, and how can GNSS degradation corrupt it?

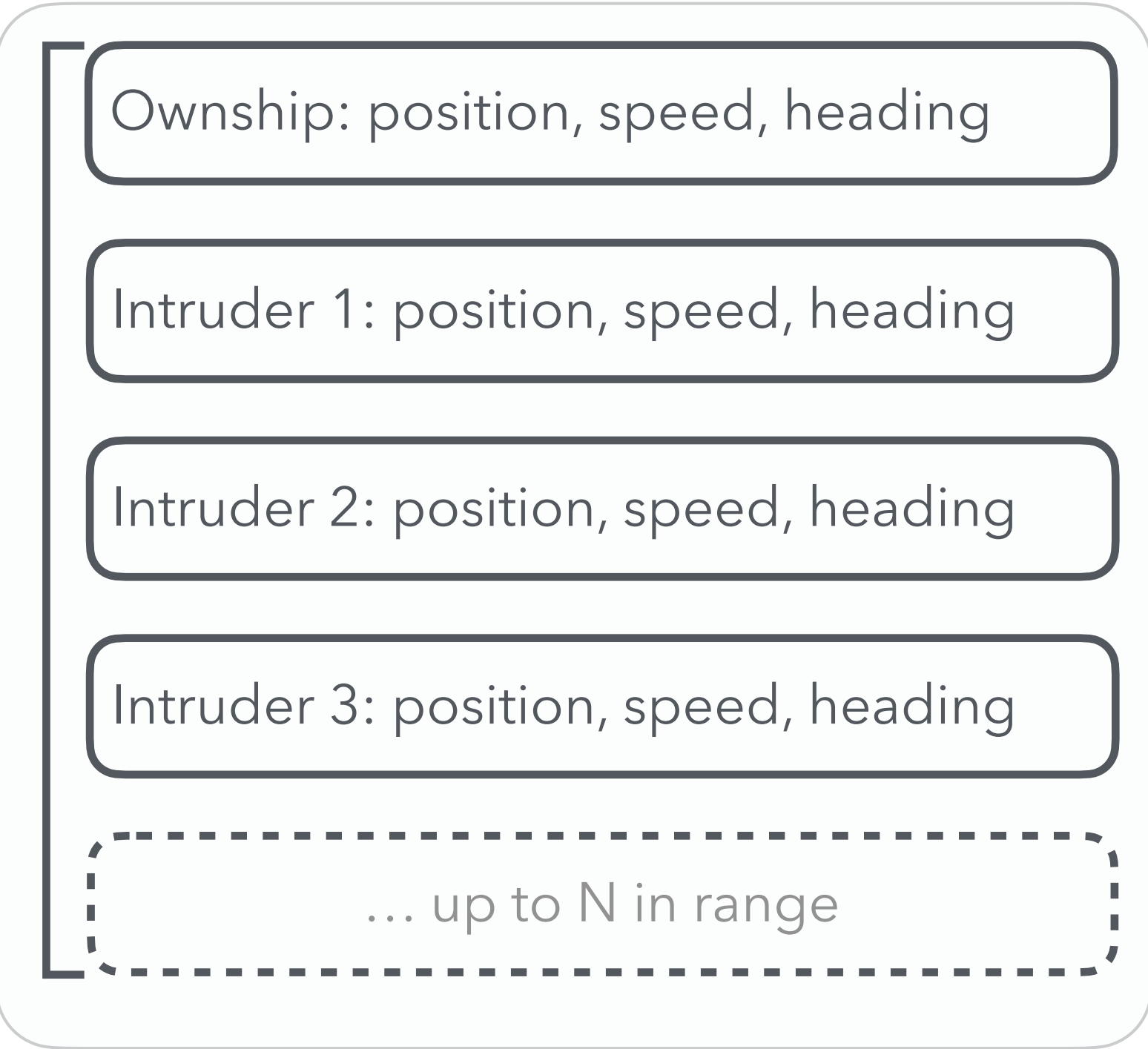


How a decentralized learned separation policy works

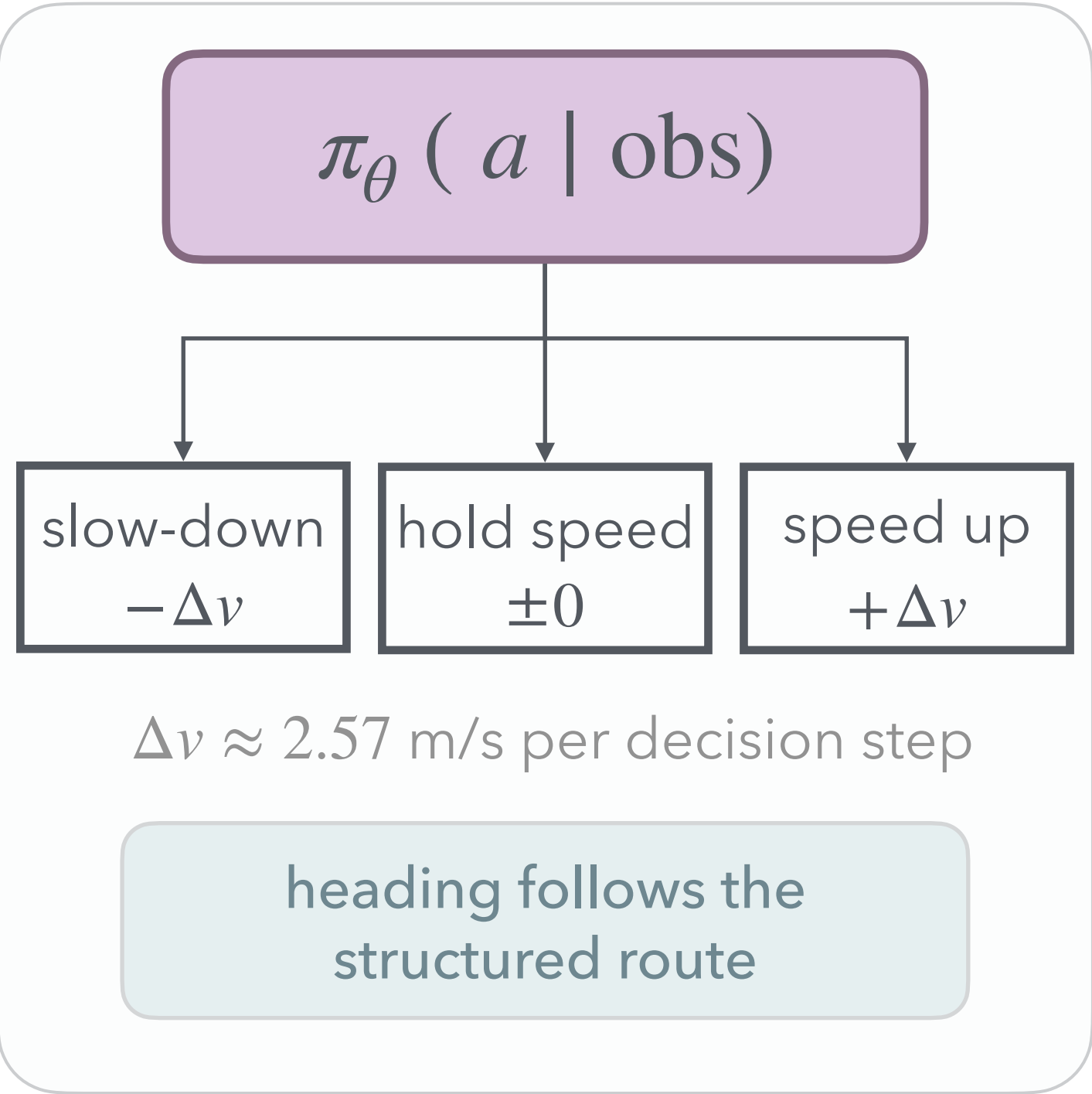
1. LOCAL TRAFFIC



2. LOCAL OBSERVATION



3. POLICY TO ACTION



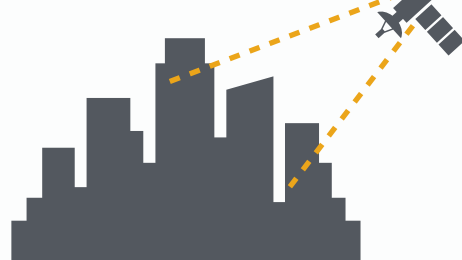
A speed change by one UAS alters what its neighbors observe.

02 THE PROBLEM

GNSS degradation distorts the state the policy relies on

1. WHY URBAN GNSS DEGRADES

Urban environments disrupt navigation integrity



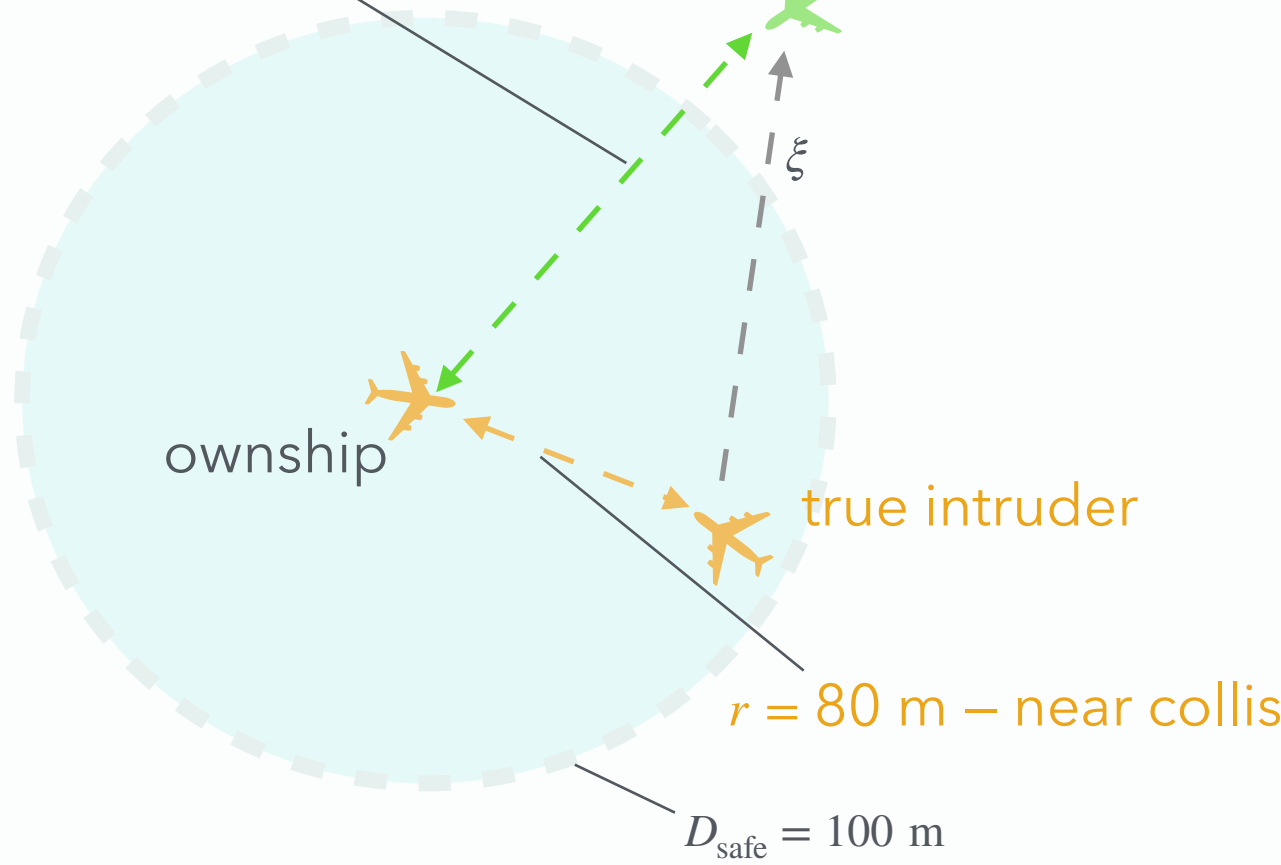
Multipath
Reflections from building faces bias position estimation

Blockage
Urban canyons create gaps and unstable fixes

Spoofing / jamming
Delayed signals, false or replayed signals, indistinguishable from the truth

2. CONSEQUENCE

A closing conflict can be reported as clear traffic



$r = 180 \text{ m}$ – appears clear

reported position

ownship

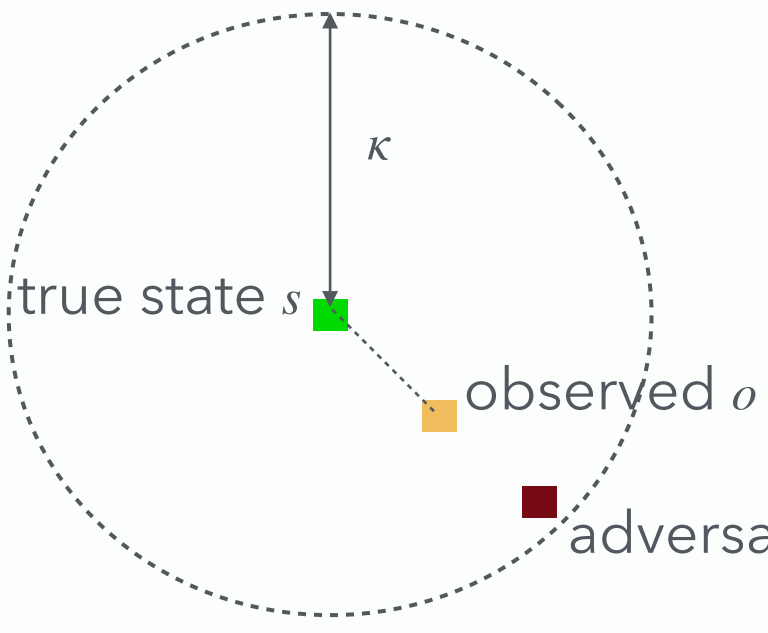
true intruder

$r = 80 \text{ m}$ – near collision

$D_{\text{safe}} = 100 \text{ m}$

3. DEGRADATION MODEL

A bounded state degradation model



true state s

observed o

adversarial ξ

κ

$o = (1 - R)s + R\xi$, where $\|\xi - s\|_{\infty} \leq \kappa$

Perturbation is bounded componentwise

$R = 0 \cdot \text{nominal}$ $\xrightarrow{\text{orange arrow}} R \uparrow \cdot \text{degraded}$

The separation policy then acts on observation o , while safety is measured on true state s .

Robustness can be built in – or handled at runtime

ADVERSARIAL / ROBUST TRAINING

- Shapes how the policy responds to corrupted or threatening observations
- Robustness depends on the assumed disturbances and training distribution
- A train-deployment gap can remain when real degradation exceeds those assumptions

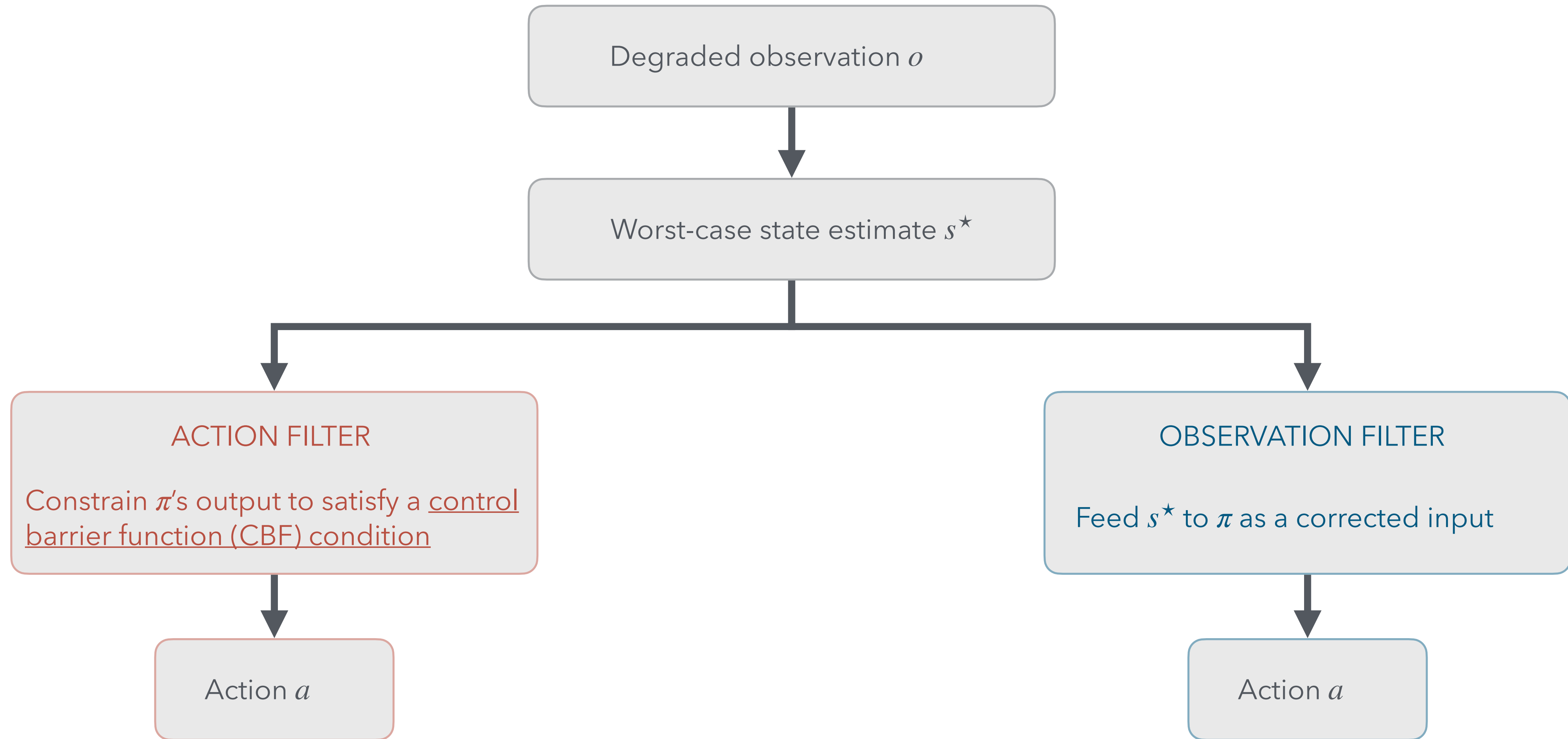
RUNTIME FILTERING

- Acts during execution, independent of how the policy was trained
- Requires no retraining of the deployed policy
- Uses the current degradation bound at each decision step
- Has no training-time knowledge and inherits whatever the policy learned

How should a runtime filter attach to a learned policy?

Before the policy sees the state – or after it proposes an action?

Two ways to add a runtime safety filter



Both filters start from the same worst-case state estimation

1. SAFETY MARGIN

Barrier function as safety margin

Unsafe | Safe

$h = 0$

$$h(s) = (\|r\|^2 - D_{\text{safe}}^2) + \alpha \cdot r^T v_{\text{rel}}$$

$\|r\|^2 - D_{\text{safe}}^2$
How much room there is between two UAS

$\alpha \cdot r^T v_{\text{rel}}$
Whether two UAS are closing in on each other
 α sets how far ahead the margin looks

2. WORST CASE

What looked clear now reads as a conflict

Current observation Worst-case view

$D_{\text{safe}} = 100 \text{ m}$

ownship

180 m – appears clear

96 m – WORST CASE

$s^* = \arg \min_{s \in \mathcal{U}(o, R_{\text{max}})} h(s)$

3. IN CLOSED FORM

A closed-form displacement

$\mathcal{U}(o, R) = \{s : \|s - o\|_{\infty} \leq R\kappa\}, \quad \delta \triangleq R\kappa$

observed o true state s

Recall $o = (1 - R)s + R\xi$, where $\|\xi - s\|_{\infty} \leq \kappa$

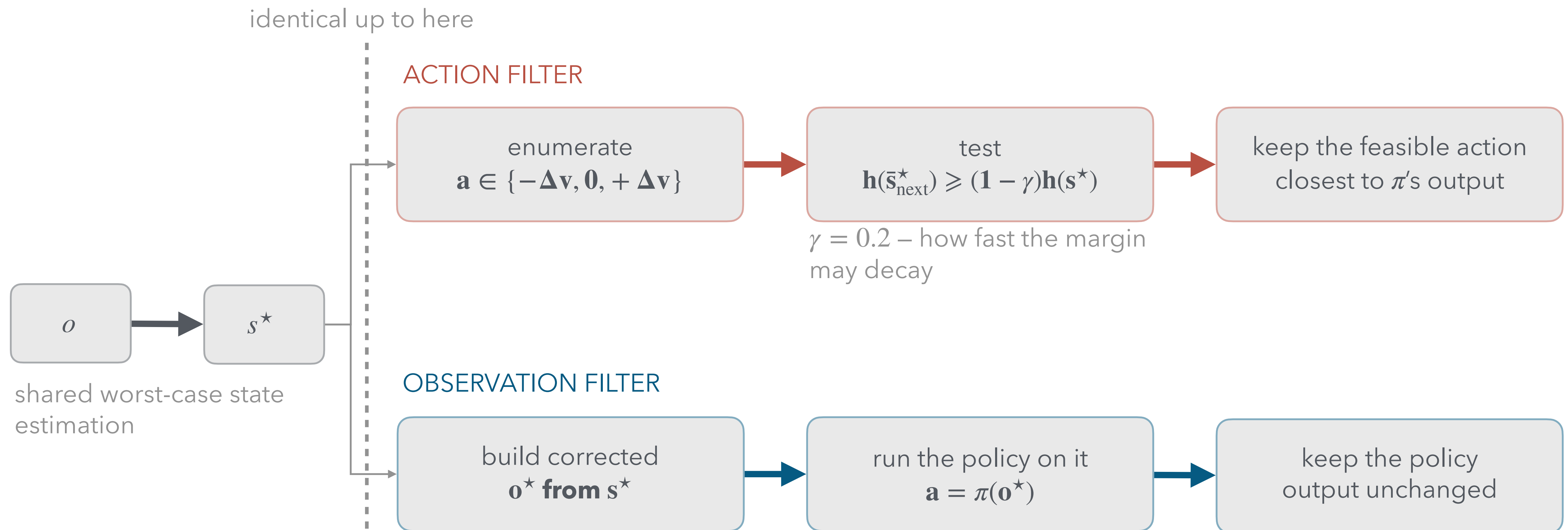
$R = 0 \cdot \text{nominal}$ $R \uparrow \cdot \text{degraded}$

With R_{max} as a design choice, $s^* \leftarrow o \pm R_{\text{max}}\kappa$

Both filters compute s^* from the same safety margin h . However, only the action filter lets h veto the policy.

03 APPROACH

The filtering approaches differ in what happens next



Same estimate, different use. The discrete-time CBF condition applies only to the action filtering.

04 RESULTS

Action filtering barely moves the safety needle

NMAC

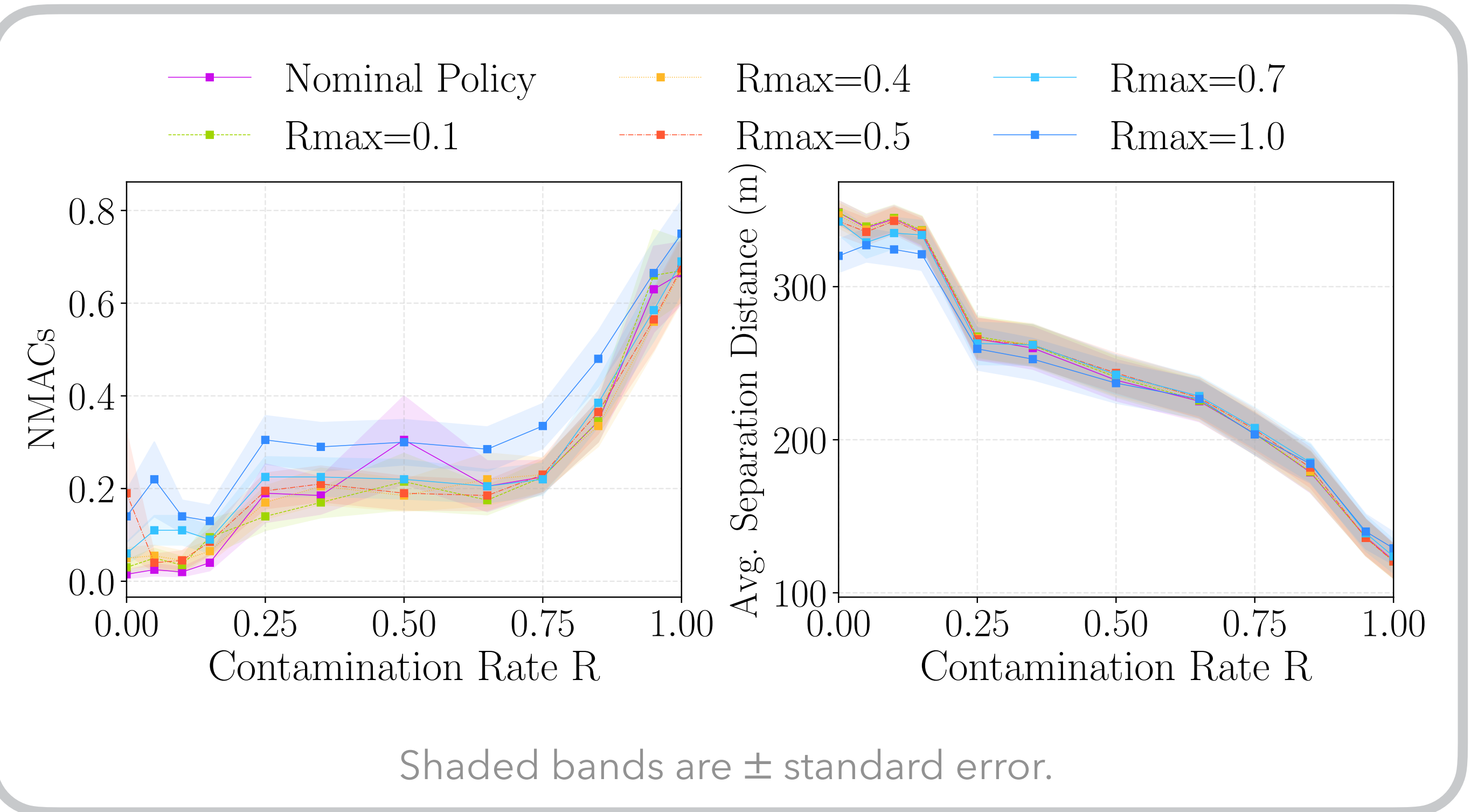
Near mid-air collision when two aircraft pass within $D_{\text{safe}} = 100$ m; counted per episode.

R

Contamination rate: how degraded GNSS actually is. 0 is nominal; 1 is fully adversarial.

R_{max}

What the filter assumes: a design choice fixed before deployment.



Every setting tracks the nominal policy. Despite changing R_{max} , action filtering does not separate from the unfiltered baseline.

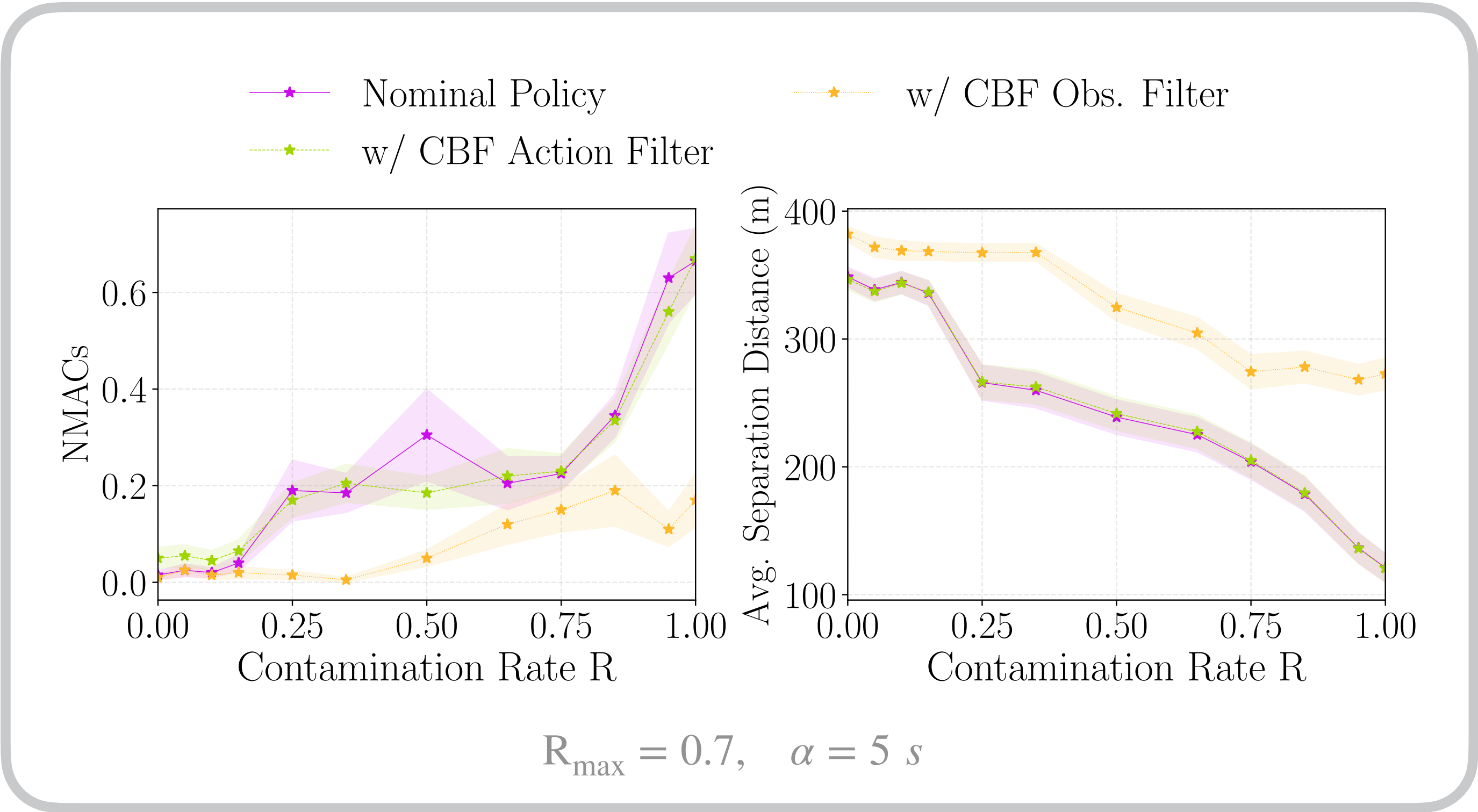
Observation filtering removes nine near collisions in ten

90%

fewer near mid-air collisions

0.24 → 0.02
NMACs per episode

250 m → 375 m
average separation



With the same worst-case estimate, the observation filter performs better on its own.

04 RESULTS

The advantage grows with what the filter assumes

$$R_{\max} = 0.1 - 0.5$$

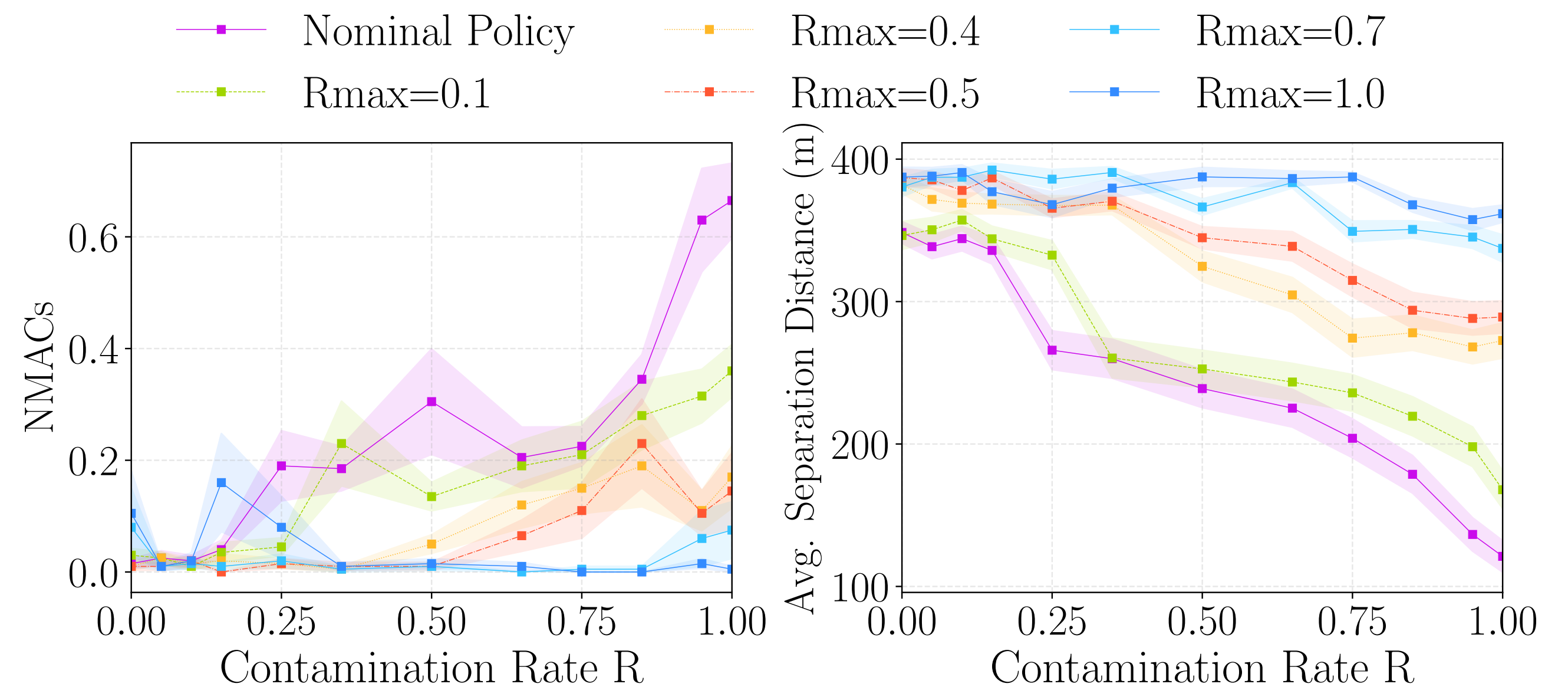
Little benefit. The assumed bound is too small to surface the danger.

$$R_{\max} = 0.7 - 1.0$$

NMACs are near zero; separation holds around 380 m as nominal policy collapses to 130 m.

No tuning

The filter never observes the true R ; $R_{\max}$ is fixed before deployment.

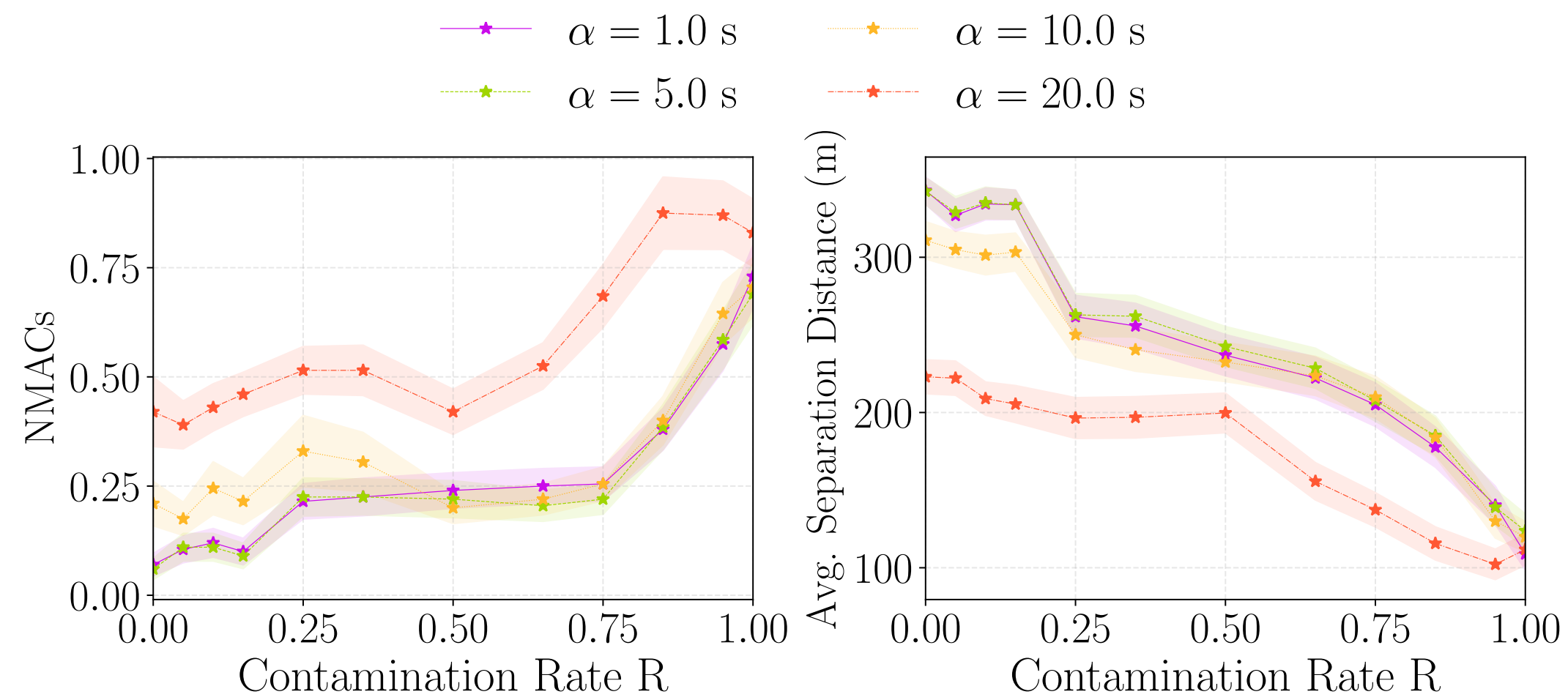


The filter is never told how bad the degradation really is.

04 RESULTS

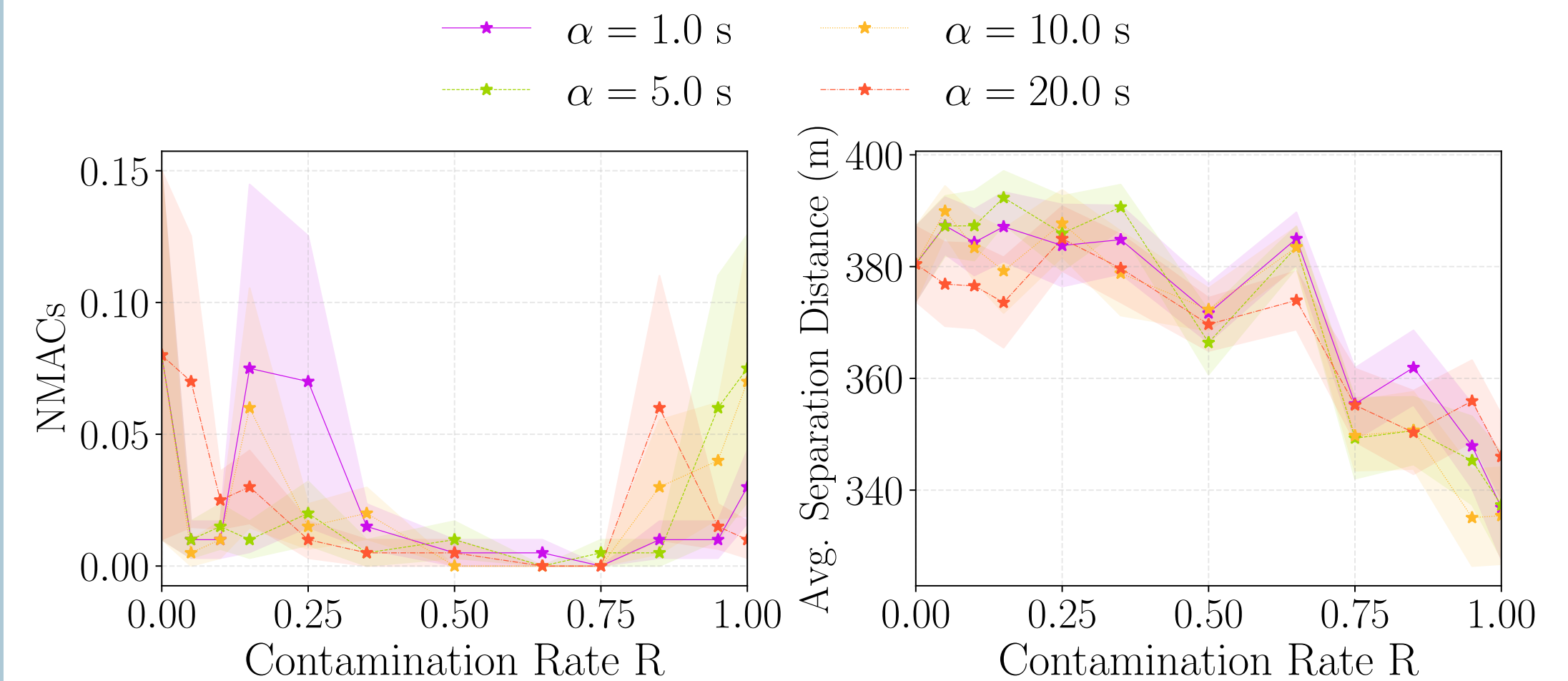
Action filtering breaks as α grows; observation filtering does not

ACTION FILTER



$\alpha = 20$ s more than doubles the NMACs.

OBSERVATION FILTER

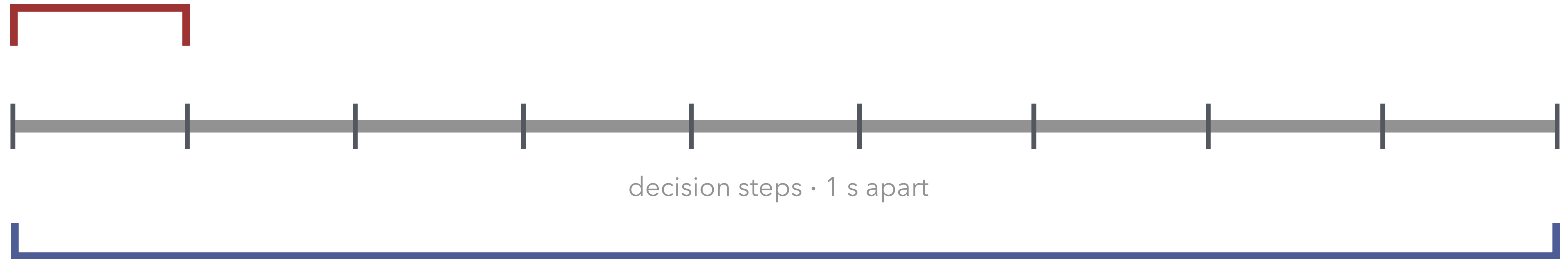


All α values' effects seem to coincide.

α is a barrier-function parameter. It matters only where the barrier can veto the policy's action.

The barrier looks one step ahead; the policy looks further

The CBF condition asks about one step



The trained policy learned to resolve the encounter over many steps

Action filtering replaces multi-step reasoning with a single-step test.

Preserve the trained policy's decision authority

When a learned controller has internalized safety behavior,
informing it outweighs overriding its decision.

ACTION FILTERING

No measurable safety gain; highly sensitive to the filtering assumptions.

OBSERVATION FILTERING

Could be more effective depending on how the policy was trained.

THE COST

Conservatism, e.g., earlier and more frequent slowing down.

Thank you!

Questions?

Runtime Safety Filtering for Learned Small UAS
Separation Policies under GNSS Degradation

Alex Zongo · Peng Wei

Department of Mechanical and Aerospace Engineering
George Washington University, Washington, DC, USA.

Contact: a.zongo@gwu.edu

Preprint: arXiv:2607.10014