



What Should Be There and What Is: The Two Ledgers a Warehouse Keeps, and the Drift Between Them

Architecture, Computational Methods, and Field Evidence from MileSoft Inventory Tracking

MileSoft Engineering Research Group

MileSoft Software Technologies

<https://milessoft.net/>

Corresponding authors: *Deepak Daryani* (deepakdaryani@milessoft.net), *Tanuj Daryani* (tanujdaryani@milessoft.net)

Working paper — Version 1.0 — 24 August 2026

Permanent link: <https://milessoft.net/research/products/inventory-tracking>

Abstract

The product's own material draws a distinction that this paper takes as its organising principle: a warehouse management system records what should be where, derived from the instructions it issued, while a tracking layer records what is actually where, derived from what sensors observed. They are different ledgers over the same building, and inventory accuracy is not a property of either one. It is the size of the set on which they disagree.

This paper presents the architecture and computational methods of MileSoft Inventory Tracking from that position. We define the drift set formally, show that a facility running only the prescriptive ledger has no instrument capable of measuring it, and characterise the three channels that find drift — fixed-reader boundary crossings, scans incidental to picking, and cycle counts — by when each fires and what each is blind to.

Three results follow. Passive capture changes the economics of detection rather than merely its speed: a fixed reader fires on the next boundary crossing at no operator cost, so its detection rate is not bounded by pick frequency the way an operator-scan channel is. Detecting divergence at task generation rather than at the pick face converts a failed pick into a re-routed one, and the cost ratio between those two outcomes is where the product's accuracy result comes from. And record accuracy becomes a continuously measurable quantity rather than an annual one, which is the precondition for ever replacing an exhaustive count with a sample.

A multi-client third-party logistics deployment reported inventory accuracy rising from 91% to 99.4%, and that figure is attributed in its own source to this module — handheld and radio-frequency capture with counts running as exceptions alongside picking. The same deployment's dock, dwell and onboarding results are attributed to other modules and are not claimed here. The product page carries no screenshots, so every figure is an authored schematic.

Keywords: Real-time inventory tracking; Inventory accuracy; RFID warehouse tracking; Location drift detection; Cycle counting exceptions; SKU to bin location map; Movement history audit trail; Mispick prevention; Passive RFID gate reads; Multi-location inventory visibility; Warehouse stock visibility; Slotting velocity data

1 Introduction

The product's material opens with a distinction worth quoting in substance: there are warehouses where inventory accuracy means the system says we have it and we do, and warehouses where it means the system says we have it and we can find it. The gap between the two is where labour and margin disappear.

That gap has a precise structure, and naming it is what this paper does. A warehouse keeps two ledgers over the same physical building. One is derived from instructions — the transactions that were issued and confirmed. The other is derived from observations — what a reader or a scanner actually saw. They agree until something moves without an instruction, and then they do not.

The consequence that organises the rest of the paper is this: a facility keeping only the first ledger owns no instrument capable of measuring its own error. It will find out how wrong it was at the next count, which is why accuracy in this industry is so often an annual figure about a state that no longer exists.

1.1 Where this sits among the neighbouring systems

Four systems in this domain touch inventory position and they are easy to conflate, so the boundaries are drawn here explicitly.

Table 1. Four systems, four questions. Each is treated in its own paper in this library, and the distinctions below are what keep them from being the same paper.

System	The question it answers	Its record
Warehouse management	Where should this be, and may this movement proceed?	Prescriptive — guarded transitions
Inventory tracking	Where is this actually, and when was it last seen?	Observational — sensor events
Annual inventory tagging	What is here, exhaustively, as of a stated date?	Periodic census
Inventory control	How much should we hold, and what is not moving?	Analytical — over the enterprise record

The reason a facility needs more than one is that each is blind where another sees. The prescriptive ledger cannot detect what moved without an instruction; the observational ledger cannot see where no reader is; the census sees everything once a year.

1.2 Contributions

1. The drift set defined formally, with accuracy as its complement and the observation that one ledger cannot measure it (Section 4.1).
2. Three detection channels characterised by firing condition and blind spot (Section 4.2).
3. Passive capture as a change in detection economics rather than latency, with the coverage term that bounds it (Section 4.3).
4. Pre-emption at task generation, and the cost ratio against discovery at the pick face (Section 4.4).
5. An eight-dimension capability reference framework for tracking layers (Section 6.3).

2 Background and Related Work

Treating a location record as a measurement rather than as a fact is the move that makes this literature applicable, and it is not the usual framing in warehouse software.

2.1 A location record is a measurement

ISO 5725 separates trueness — closeness of a mean to a reference — from precision, the scatter of repeated measurements. Applied to location, the two failure modes are systematically different.

A tracking layer with poor trueness is wrong in a patterned way: a reader positioned so that it records units as having entered a zone they only passed, a scan step performed at the wrong point in a process. A layer with poor precision is wrong sporadically: an occasional missed tag, an intermittent read.

The distinction matters for remedy. Systematic error is a configuration or process-design problem and is cheap to fix once identified; sporadic error is a coverage and hardware problem and is not. Reporting a single accuracy percentage merges them, which is why Section 4.1 defines the drift set rather than a percentage.

2.2 Why a failed pick is so expensive

De Koster, Le-Duc and Roodbergen establish that travel dominates order-picking time. That result is usually invoked to justify slotting, and it has a second consequence relevant here.

If travel dominates, then a picker who arrives at a face and finds nothing has spent the expensive part of the operation and obtained none of the value. The cost of a mispick exception is not the seconds of confusion at the face; it is the walk there, plus the walk to wherever the stock actually is, plus the disruption to a sequence that was optimised as a whole.

This is what makes pre-emption at task generation worth so much more than detection at the face, and Section 4.4 states the ratio.

2.3 Observation events and what they require

EPCIS, standardised as ISO/IEC 19987 with its vocabulary in ISO/IEC 19988, records events along four dimensions — what, when, where and under which business step. That is exactly the shape of an observational record, and adopting it means a tracking layer's output is a traceability record rather than a private log.

Instance identity is the precondition. A tag identifying only a part number tells a reader what passed; a serial or container code tells it which one, and only the second supports a statement about a particular unit's position. GS1 supplies the container code and ISO/IEC 15459-1 the equivalent outside that scheme.

3 System Overview

The system holds an observational position record per unit, a movement history behind it, a capture layer of handhelds and fixed readers, a comparison process against the prescriptive ledger, and a task generator that raises counts and re-routes as exceptions in existing work queues.

Every figure in this paper is an authored schematic. This product's page carries no screenshots of any kind, and no image belonging to another product is used to stand in for one.

3.1 The two ledgers

The prescriptive ledger updates when an instruction is issued or confirmed: a putaway task completed, a pick confirmed, a dispatch closed. It is complete with respect to intent and blind to everything else.

The observational ledger updates when a unit is seen: a handheld scan, a fixed reader at a zone boundary, a conveyor read. It is complete with respect to coverage and blind to intent — it does not know why a unit moved, only that it did.

TWO LEDGERS OVER ONE WAREHOUSE, AND THE SET WHERE THEY DISAGREE

Accuracy is not a property of either ledger. It is the size of the set on which they disagree, and neither ledger can measure it alone.

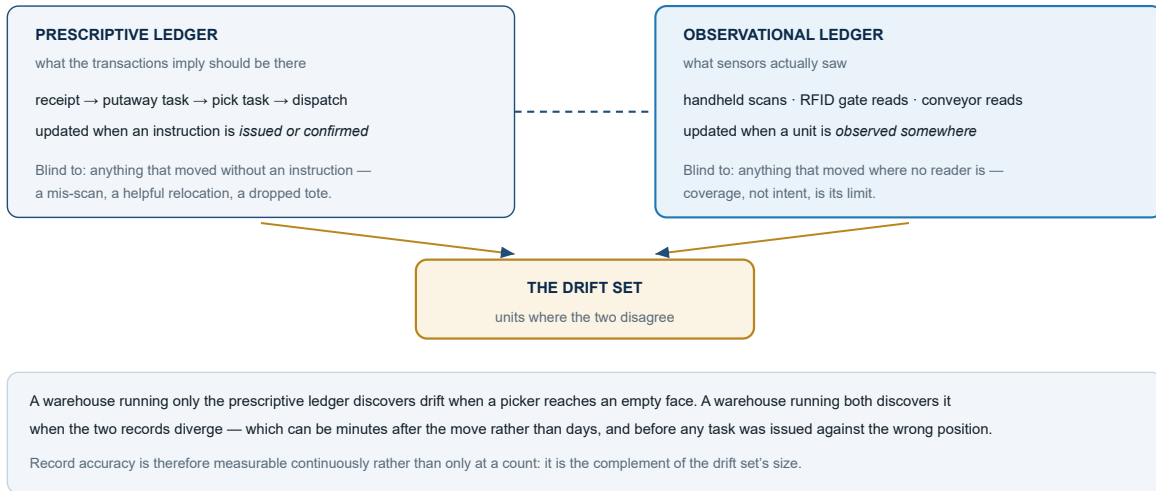


Figure 1. Schematic. Schematic (not product UI): the two ledgers and the set on which they disagree. The prescriptive record is blind to anything that moved without an instruction; the observational record is blind to anything that moved where no reader is. Accuracy is the complement of the drift set, and neither ledger alone can measure it.

The move that a facility should not make is to treat one as the correction of the other. The observational ledger is not simply right: a unit whose last read was at a zone entry three hours ago may have moved twice since. What it supplies is an independent estimate, and disagreement between two independent estimates is information neither produces alone.

3.2 Capture, active and passive

The product's material describes a hybrid: radio-frequency identification for bulk pallet movements and high-value stock, barcode for serialised items and operator-confirmed steps, with fixed readers at zone entry and exit points and on conveyors.

The important division is not the technology but whether a person is involved. An operator scan is an active read: it costs a second of someone's time and it happens where a process step requires it. A fixed reader is a passive read: it costs nothing per event and it happens wherever stock crosses an instrumented boundary, including when nobody intended anything.

Section 4.3 shows why that division dominates the design. Passive reads detect movements that no process step was ever going to record, which is precisely the population that generates drift.

The failure modes differ as well as the economics. A barcode scan fails loudly — nothing is recorded and the operator sees it. A gate read fails quietly, reading fourteen of fifteen tags, and produces a record that looks complete.

3.3 Movement history as a first-class record

The system retains every move with its actor, time and reason rather than only the current position. The product's material frames this as the audit trail no spreadsheet can produce, and as history queries answered in seconds rather than by archaeology.

Keeping history rather than only state is what makes drift diagnosable rather than merely detectable. Knowing that a unit is in the wrong place is an exception; knowing that it was last read leaving zone C at 14:20 by a fixed reader, with no corresponding putaway instruction, is a cause.

It also supplies the velocity data the putaway recommendations depend on. Slotting suggestions that consider how fast a unit actually moves need a movement record to compute from, and the product's material describes re-slot candidates surfacing weekly rather than as an annual project.

3.4 Counts and re-routes as exceptions in existing queues

When the ledgers disagree, the system raises work rather than a report. A count task is inserted into a picker's queue — the product's material describes a picker finishing a zone pick and immediately receiving a count for an adjacent bin — and a task about to be issued against a diverged position is re-routed instead.

Placing the response inside an existing queue is what removes the operational freeze. It also places the work near the discrepancy in both space and time, which is why the count is cheap: the picker is already there.

4 Computational Methods

Four computations carry the paper: what drift is, how it is found, what passive capture changes, and what finding it early is worth.

4.1 The drift set and its complement

Let U be the tracked units, with $p(u)$ the position the prescriptive ledger asserts and $o(u)$ the position the observational ledger last recorded.

the drift set at time t is the units whose two ledgers disagree; record accuracy is one minus its share of the population (drift)

$p(u)$ is derived from confirmed instructions and $o(u)$ from the most recent observation. Neither is privileged; the drift set is a statement about disagreement rather than about which is wrong.

Two properties follow immediately and both matter operationally. Alpha is defined at every instant rather than only at a count, so accuracy becomes a monitored quantity. And a facility with only one ledger has an empty drift set by construction — not because it is accurate but because it cannot form the comparison.

This is the honest reading of a warehouse that reports no accuracy problem between counts. It has no drift set because it has one ledger, and the count is the first moment two independent estimates exist to compare.

The definition also admits a subtlety Section 3.1 raises: staleness is not drift. A unit last observed three hours ago may have moved since without the ledgers disagreeing yet, so $o(u)$ carries an age and a comparison against a very old observation is weak evidence. A usable implementation therefore weights disagreement by observation age rather than treating all mismatches alike.

4.2 Three channels, and what each is blind to

Drift is found by one of three mechanisms, and their hazard rates compose.

the detection rate for a unit is the sum of three channel rates; expected time to detection is (channels) its reciprocal

$\chi(u)$ is the rate at which unit u crosses an instrumented boundary, $f(u)$ its pick frequency, $\mu(c)$ the scheduled count rate, and $\Theta(u)$ the interval from a record becoming wrong to the disagreement being found.

THREE CHANNELS THAT FIND DRIFT, AND WHEN EACH ONE FIRES

Illustrative shape, not measured values. Time runs left to right from the moment a unit ends up somewhere the record does not expect.

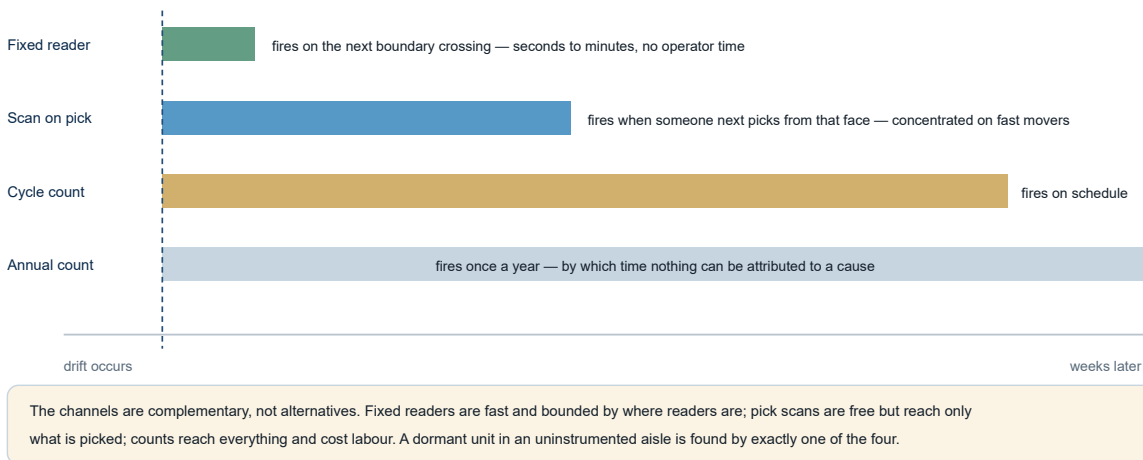


Figure 2. Schematic. Schematic (illustrative shape, not measured values): when each channel fires after a unit ends up somewhere unexpected. A fixed reader fires on the next boundary crossing and costs no operator time; a pick scan fires when someone next visits that face; a cycle count fires on schedule; an annual census fires once, by which point nothing can be attributed to a cause.

Table 2. The three channels, what triggers each, and the population each cannot reach.

Channel	Fires when	Blind to
Fixed reader	stock crosses an instrumented boundary	movement within an uninstrumented zone
Scan on pick	someone picks from that face	slow-moving and dormant stock
Cycle count	the schedule reaches that location	everything, between visits

The composition is what makes the design work and also what makes it uneven. A fast-moving unit crossing instrumented boundaries has a large eta and is found in minutes; a dormant unit in an uninstrumented aisle has eta equal to $\mu(c)$ alone and waits for the schedule.

That unevenness is defensible rather than a flaw, because it aligns detection effort with cost: an error on a fast mover will cause a failed pick soon, and an error on dormant stock will not cause anything until somebody wants it. But it should be stated rather than hidden inside an accuracy percentage, because the percentage is dominated by the population that is easiest to check.

4.3 What a fixed reader actually buys

The product's material describes fixed readers updating location without an operator scan step. The gain is usually described as labour saved; the more important gain is coverage of movements no process step would have recorded.

*boundary-crossing rate is the instrumented fraction times the unit's total movement rate; (passive)
detection improves linearly in coverage; and the cost per passive read falls as coverage and
movement rise*

$m(u)$ is how often unit u moves at all, β the fraction of movements that cross an instrumented boundary, and $C(\text{readers})$ the fixed cost of the reader estate. Unlike an operator scan, the marginal cost of a passive read is essentially zero.

Three consequences follow. Coverage is the design variable rather than reader count — a reader on a boundary nothing crosses contributes nothing, and β is a property of where readers sit relative to the flow. Detection improves linearly in β while cost is fixed, so the estate has a break-even movement volume below which handheld scanning is cheaper. And crucially, the population passive reads catch is disjoint from what pick scans catch: a unit relocated by a helpful operator with no instruction generates no scan and does generate a boundary crossing.

This is the specific reason a facility that already scans diligently still has drift. Diligent scanning covers the movements that have process steps; drift is generated by the movements that do not.

4.4 Finding it before the picker does

The product's material describes flagging drift before a picker hits an empty bin. That is a statement about when a check runs, and its value is the difference between two outcomes.

cost at the pick face = travel there + search + re-sequencing the round + the probability-weighted cost of shorting the order; cost at task generation = a lookup; the value of pre-emption is the difference

$w(\text{travel})$ is the walk to a face that turns out to be empty, $w(\text{search})$ the hunt for the stock, $w(\text{resequence})$ the disruption to a route optimised as a whole, π the probability the order ships short, and $C(\text{short})$ the cost when it does.

The first term is where Section 2.2 becomes operational. Because travel dominates picking, the walk to an empty face is the expensive part and it is spent before anything is learned. Detection at task generation avoids it entirely by issuing the task against the position the observational ledger actually reports.

The fourth term is the one that reaches the customer. A short-shipped order carries a cost far outside the warehouse — a service-level penalty, a split delivery, a lost sale — and its probability rises sharply when the discovery happens late in a wave with no time to recover.

This is the mechanism behind an accuracy improvement that does not require the picking process to change. Nobody scans differently and nobody walks differently; the tasks are simply issued against a better estimate of where things are.

5 Reported Outcomes and Field Evidence

This module carries the clearest single-product attribution in the warehouse suite, and that makes it worth reporting carefully rather than enthusiastically.

5.1 The result attributed to this module

A multi-client third-party logistics operator ran inventory accuracy at around 91% before deployment, meaning roughly one pick in eleven involved an exception, with service-level penalties from an anchor client eroding margin.

After deployment the operator reported 99.4%. The published account attributes it specifically: handheld and radio-frequency capture replaced manual bin checks, cycle counts ran continuously alongside picking as exceptions rather than scheduled freezes, and exception detection flagged location drift before a picker reached an empty bin — catching discrepancies in real time rather than at the next count.

Attribution, stated plainly. This result is this module's. The same deployment's dock idle, truck dwell and client onboarding figures are attributed to other modules and none of them is claimed here.

The published mechanism maps onto Section 4 exactly, which is unusual and worth noting. Handheld and radio-frequency capture is the observational ledger of Section 4.1; counts as exceptions alongside picking is the composed detection rate of Section 4.2; and flagging drift before the picker arrives is the pre-emption of Section 4.4.

Three caveats belong with the number. It is one operator, and the 91% baseline describes a facility doing manual bin checks — a low starting point. No denominator is published: an accuracy figure over locations, over units and over picks are three different measurements and the account does not say which. And Section 4.2 shows detection is uneven across the population, so a figure weighted toward frequently visited faces will read higher than one weighted by stock value.

5.2 Design statements published for this module

Table 3. Statements published about this module, separated by kind so a reader can tell a design claim from a measurement.

Statement	What it asserts	Kind
99% or better accuracy	a product-page benefit with no baseline or denominator	Claim
Live location and movement history	current position plus every move, by whom, when and why	Design
Fixed readers at zone and conveyor points	location updated without an operator scan step	Design
Counts as exception tasks	counting interleaved with picking, no floor freeze	Design
Drift flagged before the empty bin	divergence detected at task issue rather than at the face	Design
Velocity-aware putaway suggestions	re-slot candidates surfaced weekly rather than annually	Design

The first row is weaker than the attributed deployment figure above it and should not be read as confirming it. A benefit claim of 99% or better with no baseline, no denominator and no population is a marketing statement, and this paper treats it as one.

5.3 Modelled search and handling time

The published return model for this module prices the saving as search and stock-handling time. Its assumptions are printed here so a reader can substitute their own.

Table 4. Modelled scenario — not a deployment result. Assumptions are the published defaults of the return-on-investment model for this product.

Assumption	Value
Staff on stock handling	8
Share of their time recovered	12%
Where the saving comes from	searching for stock the record placed elsewhere
Modelled recovery	the equivalent of just under one full-time role

Twelve per cent is an assumption, and Section 4.4 says what it depends on: the frequency of failed picks, which is the drift rate times the pick rate against drifted positions. A facility at 99% accuracy has little to recover; one at 91%, like the deployment above, has a great deal.

The model is therefore self-limiting in a way worth understanding before buying. Its value is proportional to how bad the current drift is — and a facility that cannot measure its drift, per Section 4.1, does not know which case it is in.

6 Discussion

6.1 A facility that cannot measure its drift cannot manage it

The conclusion of Section 4.1 is uncomfortable and worth stating without hedging. A warehouse keeping only a prescriptive ledger has no drift set, because a drift set requires two independent estimates and it has one.

This explains a pattern familiar to anyone who has worked in distribution: a facility that believes its records are good, an annual count that finds otherwise, and a year of small unexplained exceptions in between that nobody connected to the eventual number. The exceptions were the drift being discovered one pick at a time, at the most expensive moment available.

The corrective is not more diligence. Diligence covers movements that have process steps, and drift is generated by movements that do not — so the second ledger is not a redundancy over the first but a different observer with different blind spots.

6.2 Continuous accuracy is what makes sampling defensible

The companion paper on annual tagging argues that a facility at high, measured record accuracy can eventually replace an exhaustive count with a stratified sample, and that the transition is gated on the accuracy being measured rather than asserted.

Section 4.1 is what supplies the measurement. Alpha is defined at every instant, so a facility can present an auditor with a continuously monitored accuracy series rather than a single annual observation — which is a materially stronger basis for a sampling argument than a count result on its own.

One honest limit belongs with this. Section 4.2 shows detection is concentrated on what moves, so a continuously monitored alpha over-represents active stock. An auditor is entitled to ask for the figure stratified by movement class, and a facility should be able to produce it.

The blind spots of the two regimes are why a facility needs both. Continuous monitoring watches what moves; a periodic census reaches the stock that does not, and no amount of the first substitutes for the second on dormant inventory.

6.3 A capability reference framework for tracking layers

Table 5. Capability reference framework. Each dimension separates a tracking layer from a second copy of the warehouse system's records, and each is answerable by demonstration on a live floor rather than by reading a specification.

Dimension	Question the system must answer by demonstration
D1 Two independent records	Is the observed position derived independently, or written by the same transaction?
D2 Drift reported	How many units are currently in the drift set, and what is that as a percentage?
D3 Observation age	Does a position carry the time it was last observed, or only the position?
D4 Passive coverage	What fraction of movements cross an instrumented boundary?
D5 Pre-emption	Move a pallet without an instruction. Is the next task issued against the old position?
D6 Quiet-failure detection	Does the system notice a gate reading fourteen of fifteen tags?
D7 History with cause	For one drifted unit, can the system show the last observation and the missing instruction?
D8 Stratified accuracy	Report accuracy separately for fast, slow and dormant stock. Do the three agree?

D1 is the one that decides whether the product is real. A tracking layer written by the same confirmation that updates the warehouse record is not a second observer — it is the same ledger with a second name, and its drift set is empty by construction.

6.4 Generalisability

The two-ledger structure generalises to any custody domain where intent and observation are recorded separately: asset registers against RFID sweeps, clinical stock against dispensing records, tool cribs against sign-out logs. In every case the same result holds — one ledger cannot report its own error.

The channel composition of Section 4.2 generalises to any inspection regime mixing scheduled and use-triggered discovery, and the observation that use-triggered detection concentrates on what is used generalises with it.

What does not generalise is the accuracy improvement. Ninety-one to 99.4 per cent is a distance from a specific starting point, and the recoverable fraction at another site depends entirely on where its drift currently sits — a quantity Section 4.1 says most facilities cannot yet measure.

7 Threats to Validity and Limitations

1. The accuracy result has no stated denominator. Accuracy over locations, over units and over picks are three different measurements, and the published account does not say which was used.

2. The baseline was manual bin checks. Ninety-one per cent describes a facility without any observational layer, so the improvement measures the distance from that rather than the value against a competent alternative.
3. Detection is uneven across the population. Section 4.2 shows fast-moving stock is checked constantly and dormant stock is not, so any single accuracy figure is weighted toward the population easiest to observe.
4. Staleness is not drift, and the paper's own model only partly handles it. A position last observed hours ago is weak evidence, and weighting disagreement by observation age is stated as a requirement rather than demonstrated.
5. Passive reads fail quietly. A gate reading fourteen of fifteen tags produces a plausible observational record, and no general method proves that an observation which never occurred should have.
6. Coverage is a capital decision outside the software. Beta in Equation (passive) is set by where readers are installed relative to the flow, and a facility can buy the system without buying the coverage that makes it work.
7. The pre-emption model assumes the observational position is better. Where the observational ledger is stale and the prescriptive one is correct, re-routing on divergence sends a picker to the wrong place instead.
8. No figure in this paper is a product capture. This product's page carries no screenshots, so nothing here demonstrates that the described interface exists in the form modelled.

The seventh limitation is the one an implementation must resolve explicitly. Divergence says the two records disagree; it does not say which is right, and a system that always prefers one of them will be confidently wrong in a predictable direction.

8 Future Work

- Publishing alpha continuously and stratified. Reporting the drift set as a live percentage, split by movement class, would turn Section 4.1 from a definition into an instrument and would give an auditor the stratification Section 6.2 says they are entitled to.
- Age-weighted divergence. Treating a mismatch against a three-hour-old observation differently from one against a three-minute-old observation is stated here as a requirement and deserves a published rule.
- Resolving which ledger is right. Divergence detection should carry a policy — prefer the fresher observation, prefer the confirmed instruction, or raise a count — rather than an implicit preference.
- Coverage as a reported design metric. Beta is measurable from movement history against reader positions, and publishing it would let a facility see whether its reader estate sits on the flow or beside it.
- Quiet-failure detection on gate reads, by reconciling expected against observed tag counts on a known pallet, so that a partial read is an exception rather than a plausible record.

9 Conclusion

A warehouse keeps two ledgers over the same building: one derived from the instructions it issued, one from what its sensors saw. Inventory accuracy is not a property of either. It is the size of the set on which they disagree, and a facility keeping only the first owns no instrument capable of reporting its own error.

Three results follow. Drift is found by three channels with different blind spots — a fixed reader bounded by coverage, a pick scan bounded by velocity, a count bounded by cost — so detection is fast on what moves and slow on what does not, and any single accuracy percentage is weighted toward the population easiest to check. Passive capture changes the economics rather than the latency: the movements a fixed reader catches are precisely the ones no process step was ever going to record, which is why a facility that scans diligently still drifts. And detecting divergence when a task is issued rather than when a picker reaches the face converts an expensive walk into a lookup, which is how fulfilment accuracy improves without the picking process changing at all.

The published deployment moved from 91% to 99.4% by exactly this route, and the account's own explanation maps onto those three mechanisms. The caveats are real — one operator, an unstated denominator, a low baseline — and they are stated in Section 5.1 rather than buried.

The framework of Section 6.3 is offered as the durable contribution, and its first dimension decides whether the rest applies: is the tracked position derived independently, or written by the same confirmation that updates the warehouse record?

Appendix A Nomenclature

Table 6. Symbols and abbreviations used in this paper.

Symbol / term	Meaning
$p(u)$	Position the prescriptive ledger asserts for unit u , from confirmed instructions
$o(u)$	Position the observational ledger last recorded for unit u , from a sensor read
$\Delta(t)$	The drift set — units whose two ledgers disagree at time t
$\alpha(t)$	Record accuracy — one minus the drift set's share of the population
$\Theta(u)$	Time from a record becoming wrong to the disagreement being found
$\eta(u)$	Composed detection rate across the three channels
$\mu(r), \mu(p), \mu(c)$	Per-event detection rates for fixed reader, pick scan and cycle count
$\chi(u)$	Rate at which unit u crosses an instrumented boundary
$m(u)$	Rate at which unit u moves at all, instrumented or not
β	Instrumented coverage — the fraction of movements crossing a reader
$f(u)$	Pick frequency of unit u
$C(\text{face}), C(\text{gen})$	Cost of discovering drift at the pick face, and at task generation
$w(\text{travel}), w(\text{search})$	Walk to an empty face, and the hunt for the stock
$\pi, C(\text{short})$	Probability the order ships short, and the cost when it does
Active / passive read	A read requiring operator time, against one that does not

Appendix B Worked Numerical Examples

Appendix B.1 How long a wrong record survives, by what kind of stock it is

A facility runs cycle counts at 4 per location per year. Consider three units. A fast mover crosses an instrumented boundary 90 times a year and is picked 260 times. A slow mover crosses 6 times and is picked 9. A dormant unit does neither.

Applying Equation (channels) with a fixed-reader detection rate of 1 per crossing and a pick-scan rate of 1 per pick: the fast mover has $\eta = 90 + 260 + 4 = 354$ per year, so expected time to detection is $1/354$ of a year — about one day.

The slow mover has $\eta = 6 + 9 + 4 = 19$, giving 19 days. The dormant unit has $\eta = 4$, giving 91 days.

A single accuracy percentage over this population is dominated by the first case, because errors there are corrected within a day and rarely appear in any snapshot. The stock that is actually wrong at any given moment is disproportionately the stock nobody visits — which is also the stock a periodic census exists to reach.

A facility reporting 99.4% accuracy and holding 15% of its value in dormant stock should be asked for the figure on that 15% alone. It will not be 99.4%.

Appendix B.2 What pre-emption is worth per exception

A picker walks a mean 55 metres to a face at 1.1 metres per second, so 50 seconds of travel. On finding it empty: 3 minutes of search, 90 seconds of re-sequencing disruption, and a 12% chance the order ships short at a cost of 240.

Applying Equation (preempt) with a loaded labour cost of 0.09 per second: $C(\text{face}) = (50 + 180 + 90) \times 0.09 + 0.12 \times 240 = 28.8 + 28.8 = 57.6$. The lookup at task generation costs essentially nothing, so V is about 57.6 per pre-empted exception.

Note the split. Half the cost is warehouse labour and half is the short-ship risk, and the second half falls outside the building entirely — on a customer, an anchor client's service level, a penalty clause. A cost model counting only picker time undervalues pre-emption by a factor of two.

At 91% accuracy across 4,000 picks a day, roughly 360 picks a day meet a drifted position. Pre-empting even 80% of them is worth $360 \times 0.8 \times 57.6 = 16,600$ a day in this model — which is the order of magnitude that makes an accuracy programme fund itself.

Appendix B.3 Where a reader should go

A facility has movement history showing 240,000 unit-movements a year. Of those, 96,000 cross the four zone boundaries, 62,000 cross the conveyor line, and 82,000 are within-zone relocations crossing nothing.

Instrumenting the zone boundaries alone gives $\text{beta} = 96,000 / 240,000 = 0.40$. Adding the conveyor gives 0.66. The remaining 82,000 within-zone movements cannot be reached by boundary instrumentation at all.

Applying Equation (passive), detection improves linearly in beta — so the conveyor readers, at 26 percentage points of coverage, are worth roughly two thirds of what the zone readers are worth, and should be priced against that rather than against their unit cost.

The within-zone third is the honest limit. Those movements will be found by pick scans and counts, on the timescales Appendix B.1 computes, and no reader placement changes that. A facility whose drift is dominated by within-zone relocation is buying the wrong instrument.

Data provenance

Every quantitative claim in this paper resolves to one of the entries below. Figures marked as modelled estimates are not deployment measurements; their assumptions are stated.

Metric	Reported value	Provenance
Inventory accuracy	99.4% (from 91%)	Operator-reported — Tier 1 Manufacturers (3PL, Madhya Pradesh, India) /case-studies/warehouse — Continuous exception-triggered cycle counting replaced scheduled freeze counts.
What distinguishes tracking from a warehouse management system	records what is actually where, against what should be there	Product specification /products/inventory-tracking — The product's own FAQ formulation: a WMS records what should be where from directed transactions; tracking records observed position.
Location update without an operator scan step	fixed RFID readers at zone entry, exit and on conveyors	Product specification /products/inventory-tracking
The gap the product defines itself against	the system says we have it, versus we can find it in 60 seconds	Product specification /products/inventory-tracking
Inventory accuracy claimed on the product page	99% or better	Product specification /products/inventory-tracking — A product-page benefit claim without a baseline; the attributed deployment figure of 99.4% from 91% is a separate measurement.
Search and stock-handling time saved	12% across 8 staff (modelled)	Modelled estimate 8 staff on stock handling in scope; 12% of their time recovered through real-time location; Model and defaults published in src/data/roiModels.ts

References

- [1] de Koster, R., Le-Duc, T., and Roodbergen, K. J. (2007). Design and control of warehouse order picking: A literature review. *European Journal of Operational Research*, 182(2), 481-501.
- [2] Roodbergen, K. J., and de Koster, R. (2001). Routing methods for warehouses with multiple cross aisles. *International Journal of Production Research*, 39(9), 1865-1883.
- [3] Petersen, C. G. (1997). An evaluation of order picking routeing policies. *International Journal of Operations & Production Management*, 17(11), 1098-1111.
- [4] Gu, J., Goetschalckx, M., and McGinnis, L. F. (2007). Research on warehouse operation: A comprehensive review. *European Journal of Operational Research*, 177(1), 1-21.

- [5] International Organization for Standardization (2023). ISO 5725-1: Accuracy (trueness and precision) of measurement methods and results - Part 1: General principles and definitions. ISO, Geneva (superseding ISO 5725-1:1994).
- [6] International Organization for Standardization (2019). ISO 5725-2: Accuracy (trueness and precision) of measurement methods and results - Part 2: Basic method for the determination of repeatability and reproducibility. ISO, Geneva.
- [7] International Organization for Standardization / International Electrotechnical Commission (2024). ISO/IEC 19987: Information technology - EPC Information Services (EPCIS) Standard, version 2.0. ISO/IEC, Geneva. <https://www.iso.org/standard/85557.html>
- [8] International Organization for Standardization / International Electrotechnical Commission (2024). ISO/IEC 19988:2024 - Information technology - GS1 Core Business Vocabulary (CBV). ISO/IEC, Geneva, edition 3, March 2024; publishes GS1 CBV 2.0, ratified June 2022, which fixes the vocabularies for the EPCIS business step, disposition and business transaction type fields. <https://www.iso.org/standard/85558.html>
- [9] GS1 (2024). GS1 General Specifications. GS1 AISBL; defines the GTIN, SSCC and GLN identification keys. <https://ref.gs1.org/standards/genspecs/>
- [10] GS1 (2024). Serial Shipping Container Code (SSCC). GS1 identification key series; an 18-digit key for a logistic unit, carried under application identifier 00. <https://www.gs1.org/standards/id-keys/sscc>
- [11] GS1 (2023). GS1 Logistic Label Guideline. GS1 AISBL; the SSCC is the only mandatory field on a GS1 logistic label. <https://www.gs1.org/standards/gs1-logistic-label-guideline/1-3>
- [12] International Organization for Standardization / International Electrotechnical Commission (2014). ISO/IEC 15459-1: Information technology - Automatic identification and data capture techniques - Unique identification - Part 1: Individual transport units. ISO/IEC, Geneva. <https://www.iso.org/standard/54779.html>
- [13] International Organization for Standardization (2014). ISO 22400-1: Automation systems and integration - Key performance indicators (KPIs) for manufacturing operations management - Part 1: Overview, concepts and terminology. ISO, Geneva. <https://www.iso.org/standard/56847.html>
- [14] International Organization for Standardization (2014). ISO 22400-2: Automation systems and integration - Key performance indicators (KPIs) for manufacturing operations management - Part 2: Definitions and descriptions. ISO, Geneva; defines OEE as availability x effectiveness x quality ratio. <https://www.iso.org/standard/54497.html>
- [15] International Electrotechnical Commission / International Society of Automation (2025). Enterprise-control system integration - Part 1: Models and terminology. IEC 62264-1; ANSI/ISA-95.00.01-2025 (IEC 62264-1 Mod).
- [16] International Society of Automation (2013). Enterprise-control system integration - Part 3: Activity models of manufacturing operations management. ANSI/ISA-95.00.03-2013 (IEC 62264-3 Modified).
- [17] International Organization for Standardization (2022). ISO 28000: Security and resilience - Security management systems - Requirements. ISO, Geneva; published 15 March 2022. <https://www.iso.org/standard/79612.html>
- [18] Little, J. D. C. (1961). A proof for the queuing formula: $L = \lambda W$. Operations Research, 9(3), 383-387.
- [19] H. F. Dickie (1951). ABC Inventory Analysis Shoots for Dollars, Not Pennies. Factory Management and Maintenance, vol. 109, no. 7, July 1951, pp. 92-94; the origin of ABC classification, applying Pareto and Lorenz to stock value.
- [20] Shingo, S. (translated by A. P. Dillon) (1986). Zero Quality Control: Source Inspection and the Poka-Yoke System. Productivity Press, Cambridge, MA. Originally published in Japanese, 1985.