



The Light Is Not the Poka-Yoke: Interlocks, Sensor Confirmation, and Where the Error-Proofing Actually Lives

Architecture, Computational Methods, and Field Evidence from MileSoft Pick-to-Light

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/pick-to-light>

Abstract

A pick-to-light station is named after its least important component. The lamp tells an operator which bin to open, and telling is a warning function whose effectiveness depends on attention. What makes a wrong pick impossible is the interlock holding the other bins shut — a control function that is independent of attention, training and fatigue.

This paper presents the architecture and computational methods of MileSoft Pick-to-Light from that position. Every published outcome for this class of system rests on the interlock rather than the indicator, and a specification that includes lamps and omits locks is a specification for a different product.

Three results follow. Confirmation source decides whether the pick log is evidence: a button press records that an operator asserted a pick, while a bin sensor records that something left a bin, and only the second survives an audit that questions the operator. Variant discrimination moves from operator memory into the work order, which is the mechanism behind a ramp-up collapsing from weeks to days — the operator is no longer required to distinguish six close-cousin parts, so there is nothing left to learn. And the residual error budget after deployment does not merely shrink; it changes shape, because the interlock guarantees which bin opened and can say nothing about how many pieces came out or what somebody put in.

A Tier-1 automotive deployment reported zero wrong-part picks since rollout against a 1.4% baseline, and operator ramp-up falling from four to six weeks to five days. Both are attributed in their own source to this module. The same deployment's containment and throughput results are attributed to a companion station and are not claimed here. This product page currently displays screenshots belonging to another product, so every figure in this paper is an authored schematic.

Keywords: Pick-to-light system; Poka-yoke kitting; Bin interlock; Zero defect kitting; Wrong part picking; Light directed picking; Operator ramp-up; Close variant discrimination; Sensor confirmed picking; Work order kit list; Assembly error proofing; Pick audit trail

1 Introduction

The product's own material opens with a distinction between a wrong way and a right way to think about kitting errors. The wrong way says operators need more training. The right way, stated in the lean canon since the 1980s, says that if a wrong outcome is physically possible it will eventually happen, so the physics is what needs fixing.

This paper agrees and then presses the point one step further than the product's name does. If the argument is that physics rather than training should carry the guarantee, then it matters a great deal which part of the station is doing the physics.

A lamp is not physics. It is a very fast instruction, and an instruction is precisely the thing whose reliability depends on the operator — which is the dependency the whole argument was constructed to remove. What removes it is the lock.

1.1 Why the naming matters commercially

This would be pedantry if it did not have a purchasing consequence, and it does.

Systems in this category are quoted, compared and cost-reduced by component. Lamps and controllers are the visible, countable, obviously-necessary part; interlocks are extra hardware on every bin with a wiring and maintenance cost attached. A budget-constrained specification that keeps the lamps and drops the locks looks like the same product with a smaller number on it.

It is not. It is a warning system, and its error rate is the operator's error rate multiplied by the probability that the operator missed the lamp — a smaller number than before deployment, and not zero, and not a poka-yoke.

The published outcome for this class of system is zero wrong picks. That number is only available from a control function. Any specification that cannot make it should not be quoting it.

1.2 Contributions

1. The control-warning decomposition applied to the station's own components, with the residual error rate of each (Section 4.1).
2. Confirmation source — sensor against button — and its effect on the evidential value of the pick log (Section 4.2).
3. Discrimination load as the mechanism behind ramp-up, and why removal rather than acceleration explains weeks becoming days (Section 4.3).

4. The residual error budget after deployment, and why its largest remaining term is upstream of the station (Section 4.4).
5. An eight-dimension capability reference framework for light-directed picking (Section 6.3).

2 Background and Related Work

The theory this station implements is unusually settled, unusually old, and unusually often cited without its central distinction being applied.

2.1 Control and warning, as originally drawn

Shingo separates the control function of an error-proofing device — the process cannot continue when a deviation is detected — from the warning function, in which the operator is signalled and may proceed. He is explicit that the first is preferable wherever it is achievable.

The reason is not that warnings are useless but that their reliability is a property of the person receiving them, and it degrades under exactly the conditions that produce errors: fatigue, time pressure, an unfamiliar variant, the end of a shift.

Applying that distinction to a pick-to-light bay is immediate and, in the literature, oddly rare. The lamp signals; the lock controls. They are not two aspects of one device — they are Shingo's two categories, sitting side by side on the same bin.

2.2 Prevention against verification

The product's material draws a second distinction worth formalising: scan-to-verify catches a wrong pick after it happens; pick-to-light prevents it before.

The detection-cost argument in the pokayoke literature says why this is worth so much. A wrong part caught at the station costs a re-pick; the same part caught at end of line costs disassembly; caught at the customer it costs containment. Every step of that escalation is a multiple, and prevention removes the whole sequence rather than shortening it.

Verification also carries a compliance dependency that prevention does not. A scan step is something an operator can skip when the line is running hot, and the population of skipped verifications correlates with the population of errors, because both are produced by pressure.

This is the structural criticism of verification-based error-proofing: its coverage is worst exactly when its coverage matters most.

2.3 What the standards require of the device

IATF 16949 obliges automotive suppliers to verify error-proofing devices on a schedule and to have a reaction plan for a device that fails. That obligation makes sense only for a control function — there is no meaningful verification schedule for a lamp — and it is one more reason the interlock is the component the quality system actually depends on.

ISO 13849-1 governs the design of the interlock circuit itself, including the performance level a guard release must achieve. A bin lock released by a controller with no safety-rated path is a lock whose failure mode is to open, which is the wrong direction.

IEC 62264 supplies the interface by which a work order reaches the station, and IEC 62541 the transport. Both matter to Section 4.3, because the variant decision is removed from the operator only if the station knows the kit without being told by a person.

3 System Overview

A bay is a row of bins, each with an indicator, a lock and a sensor, driven by a controller that holds the current kit and streams every event to the systems above. The kit itself arrives from a work order rather than from an operator selection.

Every figure in this paper is an authored schematic. This product's page currently displays screenshots belonging to a different MileSoft product, and presenting those as this system's interface would be a fabricated figure.

3.1 What is on a bin

The product's material describes three components per bin, and separating them is the whole argument of this paper.

Table 1. The three per-bin components, the function each performs in Shingo's sense, and what fails if it is absent.

Component	Function	Absent
Indicator	Warning — says which bin and how many	Correct picks get slower; wrong picks stay impossible
Interlock	Control — the other bins do not open	Wrong picks become possible again; the guarantee is gone
Sensor	Observation — records that something was taken	The log becomes an assertion rather than a measurement

The right-hand column is the useful one. Removing the indicator degrades throughput; removing the interlock removes the product's reason to exist; removing the sensor leaves the physical guarantee intact and destroys the evidence of it.

THE LAMP TELLS; THE LOCK DECIDES

A kitting bay of eight bins. Only one component in this picture makes a wrong pick impossible, and it is not the one the product is named after.

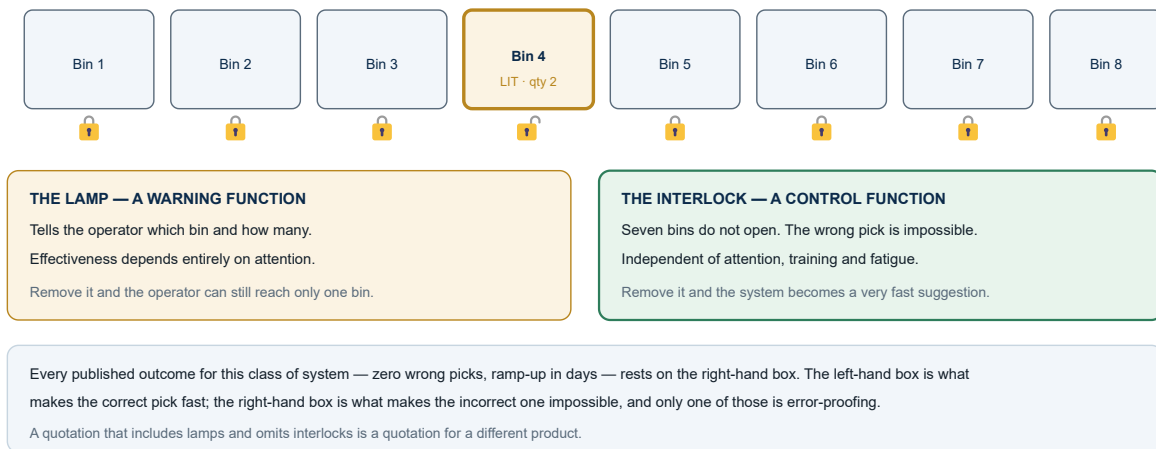


Figure 1. Schematic. Schematic (not product UI): a kitting bay mid-pick. One bin is lit and unlocked; seven are locked. The lamp is a warning function whose effectiveness depends on attention; the interlock is a control function that does not. Every published outcome for this class of system rests on the second.

3.2 Where the kit comes from

Kits resolve from enterprise and manufacturing-execution work orders, so the operator never selects the variant. A kit barcode starts the sequence and the station loads the pick list for that unit.

This is the component that makes Section 4.3 work, and it is easy to underrate because it is software rather than hardware. An operator who chooses the kit from a list has not been relieved of the discrimination — the choice has simply moved from the bin to the screen, and a close-cousin variant selected wrongly there produces a perfectly executed wrong kit.

The discrimination is removed only when nothing a person does determines the variant. If there is a dropdown, the memorisation is still required, and the ramp-up result should not be expected.

3.3 The environment the hardware lives in

The product's material notes that the indicator and confirmation hardware are sized for real industrial conditions — oil mist, vibration and dust — and this is worth taking seriously rather than treating as boilerplate.

Every guarantee in Section 4 assumes the lock holds and the sensor fires. A sensor whose lens fogs in oil mist reports no pick from a bin that was emptied; a lock whose mechanism binds under vibration is a lock that either fails open or stops the line. The environmental specification is therefore part of the error-proofing argument rather than an installation detail.

3.4 What leaves the station

Every pick is captured with operator, bin, kit and timestamp, and the stream feeds cell-level effectiveness reporting and audit playback.

Operator authentication by badge or biometric is what makes the attribution reliable. An event stream attributing every pick to whoever last logged into the bay is a record that will not survive the first dispute about who did what.

4 Computational Methods

Four computations carry the paper: what each component contributes to the error rate, what a confirmation is worth as evidence, why ramp-up collapses, and what is left afterwards.

4.1 The residual error rate of a warning and of a control

Let $p(0)$ be the probability that an operator would select the wrong bin unaided at a bay of close-cousin variants.

with a warning only, the error rate is the unaided rate reduced by however often the operator actually reads the indicator; with a control, it is the probability the lock fails plus the probability the bin held the wrong part; the first is far larger (control)

$\epsilon(\text{attn})$ is the fraction of picks on which the operator both notices and acts on the indicator, $p(\text{lock})$ the mechanical and control failure rate of the interlock, and $p(\text{stock})$ the probability the correct bin was stocked with the wrong part.

Two properties separate the terms qualitatively rather than by size. Epsilon is a human quantity that varies with fatigue, pressure and familiarity — it is lowest at the moments when $p(0)$ is highest, so the two failure modes are positively correlated and the product is worse than either factor suggests. The lock's failure rate is a hardware quantity, stable and measurable, and IATF 16949 obliges it to be verified on a schedule.

The consequence is the one Section 1.1 draws commercially. A specification retaining the indicators and dropping the locks does not deliver a slightly worse version of the published outcome; it moves the error rate from a hardware term to a human one, which is the entire distinction the poka-yoke argument was constructed around.

Note that $p(\text{stock})$ appears in the control expression and not in the warning one. Removing the operator's selection error exposes an upstream error that was previously hidden underneath it — Section 4.4 develops this.

4.2 What a confirmation is worth as evidence

The product's material specifies that confirmation comes from a bin sensor rather than a button press, so the record reflects what was taken. That distinction has a precise consequence for what the log can support.

a button confirmation is defeated by an operator advancing without picking, or pressing ahead to prime the next step; a sensor confirmation is defeated only by a failed read (confirm)

$p(\text{skip})$ is the rate of advancing without taking, $p(\text{prime})$ the rate of pressing ahead of the action, and $p(\text{miss})$ the sensor's failure-to-detect rate. Only the last is a hardware property with a specification attached.

The two skip terms are not hypothetical. Advancing a station faster than the physical work is a well-documented operator adaptation to rate pressure, and it produces a log that is internally consistent, complete, and describes a sequence that did not happen.

This is what decides whether the record can be used as evidence. A per-pick log built from button presses is a record of assertions and an auditor is entitled to say so; a log built from sensor observations is a measurement with a stated error rate, and the difference is not visible on any screen showing either one.

The same distinction runs through this library at other scales — a warehouse's prescriptive and observational ledgers, a torque controller's generated result against a typed one. It is the same argument each time: a record written by the actor is weaker than a record written by an instrument.

4.3 Why ramp-up collapses rather than shortens

The deployment reports operator ramp-up falling from four to six weeks to five days. That is not a training improvement of degree, and treating it as one obscures the mechanism.

training load grows with the number of variant pairs and inversely with how visually distinguishable each pair is; with the variant resolved from the work order, there are no pairs to distinguish and the load goes to zero (discrimination)

$d(i,j)$ is the visual distance between variants i and j . Six close-cousin variants give fifteen pairs, and the deployment's own description — small visual differences, very different downstream calibration — says every d is small.

The sum over pairs is what makes close-cousin families so expensive to train on. Adding a seventh variant to six adds six new pairs rather than one new part, and every one of them is a discrimination an operator must hold under time pressure.

Removing the discrimination does not make that learning faster. It makes it unnecessary, and what remains to be learned is the physical routine — reach, take, place — which is a matter of days in any assembly operation. Five days is not a compressed training programme; it is the training programme that is left once the hard part is deleted.

This also predicts where the result will not reproduce. A bay with two visually distinct parts has almost no discrimination load to remove, and its ramp-up will not fall from weeks to days because it was never weeks.

4.4 What is left, and where it moved

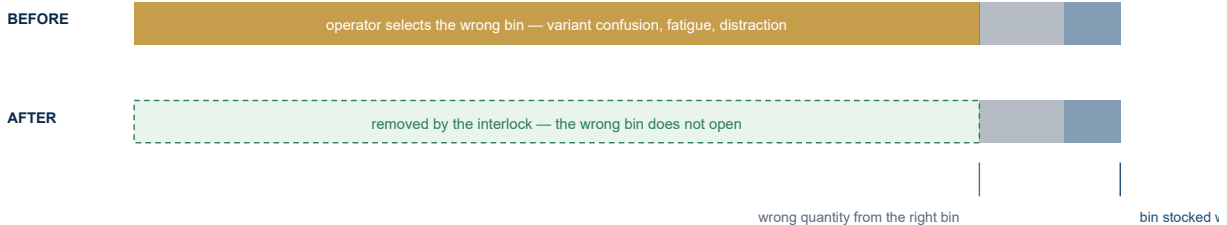
An interlock guarantees which bin opened. Two things it cannot guarantee remain, and after deployment they are the whole error budget.

the residual kit-error rate is the lock's failure rate plus the wrong-quantity rate plus the upstream stocking-error rate, and the last term dominates (residual)

$p(\text{qty})$ is the probability of taking the wrong number of pieces from a correctly opened bin, and $p(\text{stock})$ the probability that bin was replenished with a close-cousin variant.

WHAT THE INTERLOCK COVERS, AND WHAT IS LEFT AFTERWARDS

Illustrative proportions, not measured values. The bar is the wrong-part error budget of a kitting bay before deployment.



THE RESIDUAL IS NOT ZERO, AND ITS SHAPE CHANGED

The interlock guarantees which bin was opened. It cannot guarantee how many pieces came out, and it cannot know what was put in.

After deployment the dominant remaining failure is upstream: a bin correctly opened, correctly picked, and correctly wrong — because somebody restocked it with a close-cousin variant, and the station has no way to see that.

A zero-wrong-pick record is a statement about picks. It is not a statement about kits, and the difference lives in the two narrow bars.

Figure 2. *Schematic.* Schematic (illustrative proportions, not measured values): the wrong-part error budget before and after. The interlock removes the operator's bin-selection error, which was almost all of it. What remains is wrong quantity from a correctly opened bin, and a bin correctly opened but stocked upstream with the wrong part — which the station has no way to see.

The stocking term is the important one and it is genuinely awkward. The station sees a lock open on the correct bin, a sensor confirm a take, and a pick logged against the right part number — a complete and consistent record of a wrong pick, produced by a system working exactly as designed.

This has a direct design implication that the product's own architecture supports without the material drawing it: the replenishment of a bin should itself be an error-proofed operation, resolved against the same work-order system, and it is the natural next station in a programme of this kind.

A zero-wrong-pick record is a statement about picks. It is not a statement about kits, and the gap between the two is exactly $p(\text{qty})$ plus $p(\text{stock})$.

5 Reported Outcomes and Field Evidence

This module carries one of the clearest single-product attributions in the MileSoft portfolio, which makes careful separation of what belongs to it more important rather than less.

5.1 The results attributed to this module

A Tier-1 automotive supplier building radar modules ships six close-cousin variants with small visual differences and very different downstream calibration. Its operator pick-error rate ran at approximately 1.4%, and training a new operator to hold cycle time on more than two variants took four to six weeks.

After installing this module on every kitting bay — only the correct bin lighting, other bins physically locked until the right pick is confirmed — the deployment reports zero wrong-part picks since rollout, and operator ramp-up of five days.

Attribution, stated plainly. Both figures are attributed in their own source to this module's installation. The same deployment's 92% reduction in customer containment incidents is attributed to the radar alignment station, and its 18% throughput gain to the programme; neither is claimed here.

The mechanism in the published account maps onto Section 4 precisely, which is worth noting because it is rare. The lock is named as the reason the wrong pick is prevented rather than flagged, and the removal of variant memorisation is named as the reason ramp-up falls — the two mechanisms of Sections 4.1 and 4.3 respectively.

Three caveats belong with the numbers. Zero is a count over an unstated period and an unstated volume, and a rate cannot be computed from it. The 1.4% baseline is a single site's, on an unusually confusable variant family — Section 4.3 predicts the ramp-up result specifically will not reproduce where variants are visually distinct. And zero wrong-part picks is not zero wrong kits, for the reasons Section 4.4 states.

5.2 Figures published for the module generally

Table 2. Published figures for this module, with their provenance.

Figure	Value	Provenance
Wrong-part picks removed	over 90%	Product page — a general claim, distinct from the deployment's zero
Operator ramp-up	3-5 days from 4-6 weeks	Product page, consistent with the deployment's 5 days
Confirmation source	bin sensor rather than button press	Product page — a design statement
Interlock behaviour	adjacent bins locked until the correct pick is confirmed	Product page — a design statement

The first row is the general claim and the deployment's zero is a specific observation; they are not the same statement and the paper keeps them apart. Over ninety per cent is the honest figure to carry into a new site, because Equation (residual) says the remaining terms do not go away.

5.3 Modelled wrong-part cost avoided

The published return model prices the saving as wrong-part cost avoided. Its assumptions are printed here so a reader can substitute their own.

Table 3. Modelled scenario — not a deployment result. Assumptions are the published defaults of the return-on-investment model for this product.

Assumption	Value
Picks per month	50,000 (600,000 per year)
Baseline wrong-part rate	1.5% (9,000 per year)
Share of wrong-part picks removed	90%
Modelled errors avoided	8,100 per year
Cost per picking error	site-specific; the published default illustrates only

The 90% is the assumption to interrogate and, unusually, this paper can say what it should be replaced by. It is the operator-selection share of a site's own wrong-part budget — the wide bar in Figure (residual) — which a site can estimate from its own defect reason codes rather than adopting a vendor figure.

The cost term is where the model is most conservative. A picking error found at the station costs a re-pick; the same error reaching a customer costs containment, and the deployment above sits in an industry where that distinction is measured in orders of magnitude.

6 Discussion

6.1 How this product gets bought wrongly

Section 4.1 has a procurement consequence sharp enough to state as a warning rather than an observation.

Interlocks are the expensive, unglamorous, per-bin part of a quotation. They add hardware to every bin, wiring, a safety-rated release path under ISO 13849-1, and a verification schedule under IATF 16949. Indicators are cheaper, more visible, and demonstrate well.

A value-engineering exercise that keeps the indicators and drops the locks produces a system that looks identical in a demonstration, costs materially less, and cannot make the claim the business case was written on. Equation (control) says the error rate moves from a hardware term to a human one, and the human one is correlated with exactly the conditions that produce errors.

If a proposal for this class of system does not have a line item per bin for the interlock and its verification, the proposal is for a warning system. That may be a reasonable purchase; it is not the purchase that produces zero.

6.2 Solving one error exposes the next

Section 4.4 shows the residual budget after deployment is dominated by upstream stocking error. This is a general pattern in error-proofing programmes and it is worth naming, because it is routinely mistaken for a failure of the system just installed.

Before deployment, a bin stocked with the wrong part was largely invisible: it produced wrong kits indistinguishable from the operator's own selection errors, and it was absorbed into a 1.4% rate nobody could decompose. Removing the larger term does not create the smaller one — it reveals it.

The correct response is to treat replenishment as the next station rather than to conclude the picking system underperformed. A bin refill resolved against the same work-order system, with the same interlock discipline, closes the remaining term by the same mechanism.

This also explains a pattern in the deployment record. Zero wrong-part picks and a non-zero rate of downstream quality escapes are entirely consistent, and a site seeing both has not been misled — it has been shown where to build next.

6.3 A capability reference framework for light-directed picking

Table 4. Capability reference framework. Each dimension separates a poka-yoke from a guided picking display, and each is answerable by demonstration at a live bay rather than by reading a specification.

Dimension	Question the system must answer by demonstration
D1 Control not warning	Reach into an unlit bin. Does it open?
D2 Interlock integrity	What is the lock's failure mode — fail closed, or fail open?
D3 Sensor confirmation	Press the confirm without taking anything. Does the station advance?
D4 Variant from the order	Does any person choose which kit is being built, at any point?
D5 Verification schedule	How is each interlock proven still working, and how often?
D6 Attributable events	Is the operator on a pick record authenticated, or inherited from a login?
D7 Quantity coverage	Take three when the display says two. Does anything notice?
D8 Replenishment coverage	How does the system know the bin holds what it thinks it holds?

D1 takes three seconds and settles the category. D8 is the one nobody asks and the one Equation (residual) says will dominate the error budget after deployment.

6.4 Generalisability

The control-warning decomposition generalises to every error-proofing decision anywhere, and the observation that a warning's reliability is worst when it matters most generalises with it. The confirmation-source argument generalises to any system whose log is written by the actor rather than by an instrument.

The discrimination-load result generalises to any task where training time is spent learning to tell similar things apart — medication picking, specimen handling, fastener selection — and it predicts the same collapse rather than compression wherever the discrimination can be removed entirely.

What does not generalise is the ramp-up figure. Five days from four to six weeks is a result about a six-variant close-cousin family; Equation (discrimination) says a bay with visually distinct parts had little to remove and will see little change.

7 Threats to Validity and Limitations

1. Zero is a count, not a rate. The deployment reports zero wrong-part picks since rollout without stating the period or the volume, so no error rate can be computed and no confidence interval attached.
2. One site, one variant family. The 1.4% baseline and the six close-cousin variants are one supplier's, and Section 4.3 predicts the ramp-up result specifically will not reproduce on visually distinct parts.
3. Zero wrong picks is not zero wrong kits. Equation (residual) shows quantity error and upstream stocking error survive the interlock, and the deployment reports the first quantity rather than the second.
4. The interlock's own failure rate is unpublished. $p(\text{lock})$ is the term the whole guarantee rests on and no figure for it appears anywhere in the product's material.
5. Sensor miss rate is unpublished. Equation (confirm) makes $p(\text{miss})$ the only defeat mode for a sensor-confirmed log, and its value is not stated.
6. The 90% general claim has no population. It is a product-page figure with no site count, no baseline distribution and no definition of which errors were counted.
7. The modelled 1.5% wrong-part rate is a default, not a measurement, and Section 5.3 argues a site should substitute the operator-selection share of its own defect reason codes.
8. No figure in this paper is a product capture. This product's page currently displays another product's screenshots, so nothing here demonstrates the described interface exists in the form modelled.

The fourth limitation is the one to put to a vendor. A control function's value is bounded by its own reliability, and a supplier who cannot state the interlock's failure rate and failure direction is asking for the guarantee to be taken on trust.

8 Future Work

- Publishing the interlock failure rate and its direction. The whole guarantee rests on $p(\text{lock})$, and a stated figure with a verification method would let a buyer size the residual risk rather than assume it away.
- Error-proofed replenishment. Section 4.4 identifies upstream stocking as the dominant remaining term; a bin refill resolved against the same work order and the same interlock discipline closes it by the same mechanism.
- Quantity coverage. A weight or count sensor on the bin would address $p(\text{qty})$, the second surviving term, and would convert a take confirmation into a quantity measurement.
- Zero reported as a rate. Publishing wrong picks per million picks over a stated period would turn a count into a figure that can be compared, trended and given a confidence interval.
- Discrimination load measured per bay. Estimating the pairwise visual distance across a variant family would let a site predict, before purchase, whether its ramp-up will collapse or barely move.

9 Conclusion

The product is named after its lamp and guaranteed by its lock. That is not a quibble: the two components sit in Shingo's two categories, and only one of them can produce the outcome this class of system is sold on.

Three results follow. A warning's residual error rate is the operator's, degraded by fatigue and pressure at exactly the moments errors occur, while a control's is a hardware failure rate with a verification schedule attached — so a specification that keeps the indicators and drops the interlocks has moved the guarantee from a measurable term to an unmeasurable one. Confirmation source decides whether the pick log is evidence, because a button press records an assertion and a sensor records an observation, and the two are indistinguishable on screen. And ramp-up collapses rather than shortens: resolving the variant from the work order does not make discrimination faster to learn, it removes the discrimination, which is why weeks become days rather than fewer weeks.

The fourth result is the one to carry into the next project. Removing the operator's selection error does not leave a smaller version of the same budget — it leaves a different one, dominated by a bin somebody stocked wrongly, which the station cannot see and which is now the largest thing between the plant and a wrong kit.

The framework of Section 6.3 is offered as the durable contribution, and its first dimension takes three seconds: reach into a bin that is not lit, and see whether it opens.

Appendix A Nomenclature

Table 5. Symbols and abbreviations used in this paper.

Symbol / term	Meaning
$p(0)$	Unaided probability that an operator selects the wrong bin
$\epsilon(\text{attn})$	Fraction of picks on which the operator notices and acts on the indicator
$p(\text{warn})$	Residual error rate under a warning function only
$p(\text{ctrl})$	Residual error rate under a control function
$p(\text{lock})$	Interlock failure rate — mechanical and control path
$p(\text{stock})$	Probability the correct bin was replenished with the wrong part
$p(\text{qty})$	Probability of taking the wrong count from a correctly opened bin
$p(\text{skip}), p(\text{prime})$	Advancing without picking, and confirming ahead of the action
$p(\text{miss})$	Sensor failure-to-detect rate
$p(\text{kit})$	Residual kit-error rate after deployment
$L(\text{train})$	Training load — the discrimination an operator must acquire
$d(i,j)$	Visual distance between variants i and j
Control function	The wrong action is made impossible (Shingo)
Warning function	The operator is signalled and may proceed (Shingo)

Appendix B Worked Numerical Examples

Appendix B.1 The lamps-only specification, costed

A bay runs an unaided wrong-bin rate of 1.4% across 600,000 picks a year. An indicator is noticed and acted on for 92% of picks. An interlock fails at 40 parts per million, and upstream stocking error runs at 30 parts per million.

Applying Equation (control): with indicators only, $p(\text{warn}) = 0.014 \times 0.08 = 0.00112$, or 672 wrong picks a year. With interlocks, $p(\text{ctrl}) = 0.00004 + 0.00003 = 0.00007$, or 42 a year.

The lamps-only specification delivers a 92% reduction, which sounds close to the published figure and is a different product. It leaves 672 wrong picks a year where the control function leaves 42 — a factor of sixteen — and 630 of those 672 are attributable to a moment of inattention rather than to any component.

Note also which residual is measurable. Forty-two errors from a stated hardware failure rate can be verified on a schedule and trended. Six hundred and seventy-two from operator attention cannot be, which is why the quality system cannot rely on them.

Appendix B.2 Two logs, one screen

A station records 40,000 picks a month. Under button confirmation, operators advance without taking on 0.3% of picks under rate pressure and confirm ahead of the action on a further 0.8%. Under sensor confirmation the failure-to-detect rate is 0.05%.

Applying Equation (confirm): the button log's probability that a confirmed pick actually happened is $1 - 0.003 - 0.008 = 0.989$. The sensor log's is 0.9995.

In a month that is 440 button-confirmed events describing picks that did not occur, against 20 sensor events missing picks that did. Both logs look complete; both contain a per-pick record with operator, bin, kit and timestamp; and one of them contains 440 fabrications generated by the system's own design rather than by anybody's dishonesty.

Note the direction of each error. A button log over-reports picks and therefore under-reports problems; a sensor log under-reports picks and therefore over-reports them. For an audit trail, erring toward reporting a problem that did not happen is the safer failure.

Appendix B.3 Where the budget went

Before deployment a bay produces 8,400 wrong kits a year across 600,000 picks — a 1.4% rate. Decomposing by reason code afterwards: 8,190 were operator bin selection, 150 were wrong quantity from the right bin, and 60 were bins stocked with a close-cousin variant.

Applying Equation (residual) after the interlock removes the first term entirely: the residual is 42 lock and stocking events by the rates of Appendix B.1, plus the 150 quantity errors, for roughly 190 a year against 8,400 — a 97.7% reduction.

But the composition has inverted. Before, 97.5% of wrong kits were operator selection and 2.5% were everything else. After, essentially 100% is everything else, and quantity error alone is 79% of the remainder.

A plant reading only the headline would conclude the problem is solved. A plant reading the composition would install quantity sensing next, because that is now four fifths of what is left — and it would not have been visible as a priority before the larger term was removed.

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
Wrong-part picks since 0 rollout				Operator-reported — Tier-1 Manufacturers (Pune & Nasik plant, India) /case-studies/pokayoke — Baseline operator pick-error rate was approximately 1.4%.
Pre-deployment pick-error rate	operator	approx. 1.4%		Operator-reported — Tier-1 Manufacturers (Pune & Nasik plant, India) /case-studies/pokayoke — Six close-cousin radar-module variants with small visual differences.
Operator ramp-up to full cycle-time competence		5 days (from 4-6 weeks)		Operator-reported — Tier-1 Manufacturers (Pune & Nasik plant, India) /case-studies/pokayoke
Wrong-part reduction attributed to Pick-to-Light		over 90%		Product specification /products/pick-to-light
How a wrong pick is prevented rather than flagged		adjacent bins stay locked until the correct pick is confirmed		Product specification /products/pick-to-light
Source of pick confirmation		a bin sensor, not a button press		Product specification /products/pick-to-light — Stated on the product page as the reason the record reflects what was taken rather than what was asserted.
Close-cousin variants an operator previously had to discriminate by memory		6 radar-module variants with small visual differences		Operator-reported — Tier-1 Manufacturers (Pune & Nasik plant, India) /case-studies/pokayoke
Wrong-part picking errors avoided per year		8,100 of 9,000 (modelled)		Modelled estimate 50,000 picks per month, 600,000 per year; Baseline wrong-part rate of 1.5%; 90% of wrong-part picks removed; Model and defaults published in src/data/roiModels.ts

References

- [1] 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.
- [2] Nakajima, S. (1988). Introduction to TPM: Total Productive Maintenance. Productivity Press, Cambridge, MA.
- [3] International Organization for Standardization (2015). ISO 13849-1: Safety of machinery - Safety-related parts of control systems - Part 1: General principles for design. ISO, Geneva; a 2023 edition has since been published. <https://www.iso.org/standard/69883.html>

- [4] International Automotive Task Force (2016). Quality management system requirements for automotive production and relevant service parts organizations. IATF 16949:2016, first edition, 1 October 2016 (superseding ISO/TS 16949).
- [5] Verein Deutscher Ingenieure / Verband der Elektrotechnik (2025). VDI/VDE 2862 Blatt 1: Minimum restrictions for application of fastening systems and tools - Applications in the automotive industry. VDI/VDE, Dusseldorf.
- [6] 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>
- [7] 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>
- [8] 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).
- [9] 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).
- [10] International Electrotechnical Commission (2025). IEC 62541-1: OPC unified architecture - Part 1: Overview and concepts. IEC, Geneva. <https://webstore.iec.ch/en/publication/81513>
- [11] International Electrotechnical Commission (2020). IEC 62541-2: OPC unified architecture - Part 2: Security model. IEC, Geneva.
- [12] 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>
- [13] 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>
- [14] GS1 (2024). GS1 General Specifications. GS1 AISBL; defines the GTIN, SSCC and GLN identification keys. <https://ref.gs1.org/standards/genspecs/>
- [15] 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>
- [16] 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).
- [17] 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.
- [18] Automotive Industry Action Group (Chrysler, Ford, General Motors) (2010). Measurement Systems Analysis (MSA) Reference Manual. AIAG, 4th edition, MSA-4.
- [19] 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.
- [20] Lee, J., Bagheri, B., and Kao, H.-A. (2015). A cyber-physical systems architecture for Industry 4.0-based manufacturing systems. *Manufacturing Letters*, 3, 18-23.