

CAPABILITY GUIDE · AUGUST 2026

AI you can audit

Governed AI across the workplace-safety compliance lifecycle — capture through audit. A person stays the control of record, and the evidence behind every decision is checkable in the product.

SafeGora (safe + agora — the ancient Greek gathering place) is the OSHA compliance platform uniting employers, consultants, auditors, and carriers.

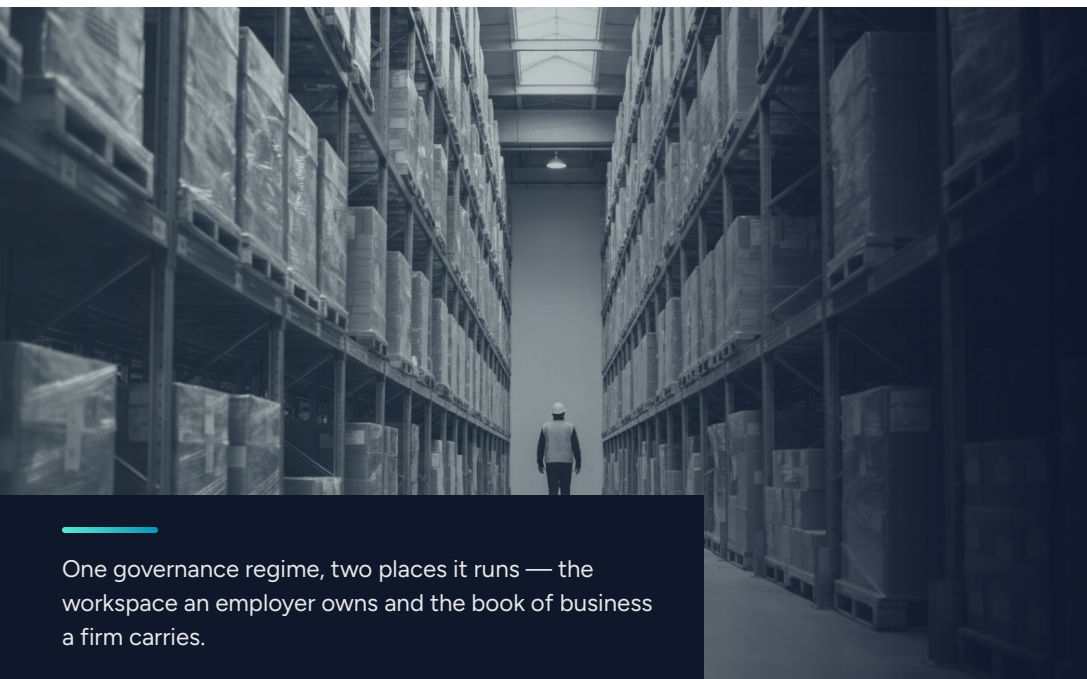
On the floor — the work every AI-assisted decision in this guide has to get right, on the record.

CONTENTS AND PATHS

Inside this guide

This guide inventories governed AI across both planes — the employer platform and the firm plane — and how each capability is checked. It is not an operator runbook; the four companion guides named on the last page carry day-to-day workflow.

SECTION	PAGE
01 Why governed AI now	3
02 How SafeGora governs AI	4
03 Ollie, under the same rules	5
04 Incidents and filings	6
05 Field capture and findings	7
06 Contractor and chemical files	8
07 Ergonomics and risk signals	9
08 In the product	10
09 Proof you can check	11
10 Two planes, one regime	12
11 Evaluate in 10 minutes	13
• Disclosures and related guides	14



One governance regime, two places it runs — the workspace an employer owns and the book of business a firm carries.

HOW TO READ THIS

Buyer

read 1, 2, 9, 10 and 11: moment, regime, proof, breadth, check

Security reviewer

read 2, 9 and 10: data boundary, decision ledger, class fences

EHS leader

read 3 to 7: the domains your program already runs on

Firm operator

read 3 to 9 as delivery tools, then 10 for the fences

1 THE MARKET MOMENT

Why governed AI now

Buyers stopped asking to see the demo and started asking to see the governance. A capability nobody outside the vendor can check is a risk the buyer carries — which is why governance now has to travel with the feature.

Disclosure and transparency expectations are arriving in regulated work faster than procurement cycles can absorb them. A capability claim now has to be checkable by the person who signs for it. Procurement tests the governance before it tests the feature.

Evidence has to form while the work happens, not be reconstructed afterwards. SafeGora runs AI in every layer of the lifecycle — capture, determination, verification, filing, corrective action and audit. AI that arrives at the end of a process can only comment on it; AI present throughout can be checked against it.

A person stays the control of record wherever an output could bind a compliance outcome. Your own auditors check the evidence behind it — what was cited, how it was tiered, who accepted or rejected it — continuously, inside the product.

WHAT BUYERS NOW CHECK

Inference

where it runs, and under whose cloud project

Approval

what never auto-approves, whatever confidence the model reports

Citations

whether each was verified against what was retrieved

The chain

whether your own auditor can check it end to end

51

US jurisdictions

50 states plus DC, one substrate

7

governance principles

applied to every capability

50+

governed capabilities

counted across eight domains, page 12 · August 2026

3

consequence tiers

informational, advisory, determinative-class

Asserted trust lasts as long as the sales cycle. Checkable evidence lasts as long as the record.

2 THE GOVERNANCE REGIME

How SafeGora governs AI



Seven principles govern every AI capability SafeGora ships, and no capability is exempt. They fix what resolves the jurisdiction, what has to be cited, what a person must approve, and what the ledger records.

Get the jurisdiction wrong and every citation downstream of it is wrong. Establishment-scoped AI resolves regulatory context from the physical jurisdiction of the work: the worksite first, then the establishment, then the employer's principal location, with the federal framework as the floor everywhere.

Divergence is applied by category — written-program mandates, stricter reporting triggers, heat and wildfire rules, distinct workers'-compensation fund structures — from structured, versioned regulatory data. A model never guesses an obligation the substrate does not carry.

THE DRAFT STATE

- AI-assisted draft – pending review. Not the record until a person approves it.

THE LEDGER ENTRY

- capability=recordability-triage
 - advisory · **human accepted**
 - ledger entry written

#	Principle	What it means
1	The right jurisdiction, always	Context resolves from the worksite, then establishment, then principal location; federal floor everywhere
2	Verified citations or honest abstention	Citations are checked against what was retrieved; thin grounding produces abstention
3	A human remains the control of record	Informational informs; advisory proposes with review; determinative-class enters a mandatory review queue
4	Tamper-evident accountability	Every governed decision is hash-chained: capability, model, confidence tier, human action
5	Bounded spend, operator control	Per-capability budgets, per-run ceilings, live cost visibility, audited kill-switches
6	Your data never trains a model	No training, no retention for model improvement, no brokering of your content
7	Released only when proven	Pre-release accuracy floors on golden evaluation corpora; failing releases never ship

AI drafts. People decide. Evidence holds.

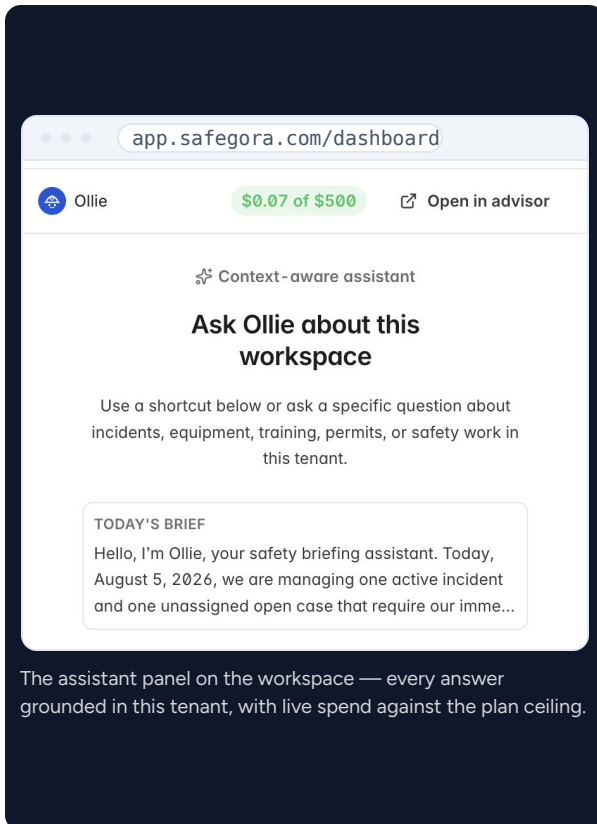
3 ACROSS THE WORKSPACE

Ollie, under the same rules



Ollie is present across the employer platform and the firm plane. It steps aside on the dedicated AI hub, the incident wizard mid-capture, and kiosk paths — those surfaces keep their own AI rails. Everywhere it runs, the regime is the same: cited, reviewed, logged, budgeted, and honest about its limits.

- ✓ capability=assistant-answer
 - informational
 - **citations verified**
 - ledger entry written



One persona, one contract. Ollie drafts, retrieves and explains, grounded in your workspace data and the jurisdiction corpus, with citations shown. It does not render legal determinations, it does not file anything without human approval, and it does not claim to be a person. Where it can help you act, it proposes and waits.

ALSO IN THE ASSISTANT

Voice

in the assistant, incident intake and seven field workflows

Offline

speech capture on disconnected sites, syncing when connectivity returns

Evidence panel

promote a thread to the full advisor workspace, intact

Budget

live usage against your plan; degraded modes are explicit

- 📖 Ollie is SafeGora's AI assistant. Answers are AI-generated and can be incomplete or wrong – verify anything you act on. Ollie shows the regulation text it relied on.

Ask, with receipts

Ollie answers regulatory and operational questions in the flow of work, not in a search box. Every answer shows its work. The grounding is checked against what was actually retrieved before the answer reaches you, and the status of that check travels with the answer rather than sitting in a help article.

- ✓ Read the citation and its verification state in the same view
- ✓ See the grounding status — grounded, partly grounded, or needs review
- ✓ Get honest abstention when the grounding is thin, not a guess
- ✓ Spot unverifiable material by its flag, never a silent assertion

What the answer relied on is always visible.

Act, with approval

When Ollie can help you do the work — drafting a corrective action, preparing a filing step — it does not act behind the scenes. Proposals arrive as reviewable cards. You accept or reject each one, and nothing executes silently. It is the same posture the rest of SafeGora runs on, expressed at the surface you use most.

- ✓ Get next actions that fit the page you are already on
- ✓ Review drafts for corrective actions and filing steps, as proposals
- ✓ Confirm before anything is sent or spent
- ✓ Find every governed decision in the decision ledger

Nothing executes until a person accepts it.

4 CAPTURE TO FILING

Incidents and filings



What matters about incident AI is where it ends. Every AI-assisted moment on this path terminates in a regulator-ready artifact your recordkeeper can stand behind, owned and certified by a person.



Stage	What the AI does	What a person controls	What comes out
Capture	Parses voice or narrative into structured wizard fields	Confirms or edits every field	A structured case in the incident wizard
Narrative	Drafts a complete narrative from the structured facts	Recordkeeper edits and owns the text	Case narrative ready for the record
Triage	Proposes severity and a cited recordability path	Accepts, modifies or rejects	Determination grounded in 29 CFR Part 1904
Precursor screening	Screens for precursors of serious-injury potential	Escalates investigation when warranted	Priority queue for deep investigation
Root cause	Proposes causes and corrective-action drafts	Confirms before anything applies	Items in the corrective-action register
Photo classification	Proposes occupational injury and illness codes	Accepts or rejects each code	Classification recorded on the case
Filing rail	Assembles forms and jurisdiction first-report packs	Certifies and files as required	OSHA Forms 300, 300A and 301, plus first-report packs

Accepted determinations feed OSHA Forms 300, 300A and 301, and the Injury Tracking Application (ITA) submission path. The same case renders workers'-compensation first reports against each jurisdiction's current official form across the 51-jurisdiction substrate — 50 states plus DC. Every pack clears a deterministic render-verification gate before it is offered: a certified render pack, not a free-form PDF.

🕒 capability=narrative-draft
 • advisory • **recordkeeper modified**
 • ledger entry written

WHAT THE PATH PRODUCES

Case

a structured, reviewable record from voice or narrative

Narrative

recordkeeper-owned text with a cited path to recordability

Corrective actions

opened with owners, due dates and evidence attached

Filing packs

certified first reports for the jurisdiction where the work happened

5 WALK, CAPTURE, REVIEW

Field capture and findings



Walk-throughs, notes, voice, photos and documents become citation-backed, human-reviewed findings — findings that open corrective actions and shareable reports in the same register the incident path uses.

- ✓ capture • cited, confidence-tiered analysis • human review
- finding or corrective action
- ledger entry



app.safegora.com/inspections/new

FINDINGS

Record hazards or use AI vision / dictation. Review all AI output before completing.

Category and severity use fixed OSHA-aligned options for reporting.

Hazard inspection findings — the surface's own instruction to record hazards or use AI vision and dictation, and to review every AI output before completing.

Site scan

Walk the site and capture photos as you go. Analysis proposes hazards as citation-backed findings, checked against the regulatory context that applies where the work is happening. A dismissal is written to the ledger exactly like an acceptance. Recurring conditions are recognized across scans, so a repeat condition surfaces as a pattern rather than as five unrelated findings.

- ✓ Findings cite the jurisdiction of the site, not headquarters
- ✓ Findings carry verified citations and a stated confidence tier
- ✓ Accept or dismiss each finding before it enters the record
- ✓ Accepted findings open corrective actions in one step
- ✓ Export a shareable report for stakeholders and clients

Nothing enters the record until a person confirms it.

Job hazard analysis

Draft a job hazard analysis directly from a photographed finding. The hazards come from what was actually observed on your site, not from a library template that happens to match the trade. A person reviews and accepts before it becomes the operational record.

- ✓ Drafted from the condition your photo actually shows
- ✓ Task-based, tied to the finding that produced it
- ✓ Reviewed and accepted before it stands as the record
- ✓ Opens the same corrective actions and shareable reports as a scan finding

The hazards come from your site.

WHAT PEOPLE STILL DO

Confirm the structured records before they stand

Accept or dismiss each proposed finding

Own the routing and the owners

Authorize what enters the corrective-action register

ALSO IN THE FIELD

Inspection intake

raw notes and voice become structured hazard records

Inspection prep

assembles what a compliance officer would ask, before arrival

Voice reports

dictate, transcribe, extract, route; offline capture syncs later

Document analysis

the same citation and review discipline as desk-side AI

6 VERIFICATION WITH ENFORCEMENT

Contractor and chemical files

Two document domains, one discipline. Contractor files are verified against your requirement profiles by rules rather than judgment calls; chemical ingredients are indexed against versioned public lists that carry their own provenance.



Contractor documents

Certificates, injury summaries and credentials are read by multimodal AI, and every extracted field arrives with its provenance. A deterministic rules engine — not free-form model judgment — decides pass or fail against your requirement profiles. Expiring documents run a staged renewal loop that re-verifies on resubmission.

- ✓ Every field carries a page anchor, verbatim source snippet and confidence score
- ✓ Deficiency flags cite the requirement item and the exact certificate field that failed
- ✓ Certificates never auto-approve, regardless of model confidence
- ✓ Permits warn or hard-block on lapsed verification, with an audited variance override

Extraction may propose; rules evaluate; humans approve.

Chemical register

Safety data sheets are read by multimodal AI with page-anchored verbatim quotes and a confidence score on every field, and nothing commits until a person approves it. Ingredients are then cross-referenced against versioned, provenance-stamped public-domain federal lists. Identical inputs produce identical indexing, every time, against a stated dictionary version. Concern rollups run at ingredient, chemical and establishment grain with the factors disclosed.

- ✓ Nothing enters the register until a person approves the extraction
- ✓ Every match carries its statutory citation and the dictionary version behind it
- ✓ Federal list coverage spans toxics release, hazardous air pollutants and extremely hazardous substances with planning quantities
- ✓ Exposure limits come from OSHA and NIOSH tables; PFAS matches disclose their basis

List membership is a cited fact, not regulatory advice.

Document or control	What the AI reads	What the rules engine enforces
Certificate of insurance	Limits, dates and endorsements, with page anchors and snippets	Requirement-profile match; a human verifier's acceptance is required
Workers'-compensation coverage	Document type and coverage dates, with page anchors and snippets	Monopolistic-fund categories reject a private-carrier certificate; elective categories route to attestation-or-certificate evidence
Annual injury summary	Totals and rates read from the face of the document	Arithmetic check, cross-validation and baseline comparison, both values cited
Permit binding	The bound contractor's documents, re-read on every resubmission	Warn or hard-block on lapse, with an audited variance override

THE LEDGER ENTRY

✓ capability=coi-verification · determinative-class
 · verifier accepted · ledger entry written

7 SCREENING AND SIGNALS

Ergonomics and risk signals

Screening-grade assessment you can re-check by hand, and forward-looking signals that stay advisory until a person disposes of them. Both run under the same tiering, citation and review discipline as everything else.

Ergonomics screening

Capture a short task video on any phone. The AI performs observational postural coding — the categorical judgments a trained rater makes by eye — and published instruments score those codes deterministically. Scope is screening: prioritization and engineering-control justification, with a qualified ergonomist retained where one is required.

- ✓ Deterministic scoring of RULA, REBA and the Revised NIOSH Lifting Equation
- ✓ A per-factor audit trail: worksheet factor, coded value, evidence frame
- ✓ Observational coding only — no skeletal reconstruction, no joint-angle measurement
- ✓ Results correlated against your own coded injury records
- ✓ Reviewed findings reach the corrective-action register with owners and due dates

Bring your own ergonomist and check the math.



Signals and prediction

Continuous scanning across your operational data raises categorized, severity-ranked signals with acknowledgment and aging visibility, so nothing quietly gets old. Forward-looking hazard categorization with timeframes is reviewed like any other advisory output: proposed, then accepted, modified or rejected by a person, then written to the ledger.

⊘ capability=hazard-prediction • advisory
• reviewer rejected • ledger entry written

Nothing files itself.

- ✓ Predictions never silently create hazards, incidents or obligations
- ✓ Disposition stays with the operator: acknowledge, escalate, or open work
- ✓ An establishment compliance score is recomputed continuously for prioritization
- ✓ Inspection and audit preparation is assembled before anyone arrives
- ✓ Multi-step filing work advances as proposals under approval gates

THE PROACTIVE LAYER

Daily brief

assembled when you ask, from open risk and priority context

Committee agendas

drafted from current operational context, for committee review

Toolbox talks

drafted from recent field findings, reviewed before crew delivery

Knowledge checks

generated quizzes, reviewed before use as training material

Asset advisories

cited to a structured obligation dictionary, never free-form invention

8 FIELD AND DESK In the product



The same regime, on the surfaces people actually use: capture in the field on a phone, review and filing on a larger screen. Nothing about the governance changes when the screen does.

Field and desk are the same system. A finding captured on a phone at a jobsite and a determination certified at a desk share one record, one evidence trail and one ledger. Whoever reviews the work later sees exactly what the person in the field saw.

- ⊗ capability=site-scan-finding
 - advisory
 - reviewer dismissed
 - ledger entry written

WHAT A LEDGER ENTRY RECORDS

DECISION	LEDGER	CAPABILITY · MODEL · CONFIDENCE TIER · HUMAN ACTION
capability		recordability-triage · narrative-draft · coi-verification · hazard-prediction
model		recorded with every governed decision
confidence tier		informational · advisory · determinative-class
human action		✓ accepted · modified · rejected

Field Work

Start field tasks, resume local drafts, and repair offline sync from this device.

[Refresh](#)
[AI assist](#)

The field hub — field tasks, local drafts and offline-sync repair, on the surface crews open in the field.

app.safegora.com/operator/review-desk

Operator review desk

Cross-type contractor review queue — programs, COI deficiencies, and stats submissions. Human verdict is the record of truth.

Type

SLA

All types

All open

Refresh

The operator review desk — the cross-type contractor review queue, its type and SLA filters, and the surface's own line that a human verdict is the record of truth.

WHAT TO LOOK FOR

Review affordance

every proposal carries a way to accept or dismiss it

Citation state

the citation and whether it was verified, shown together

Ledger entry

the decision behind the record, written and checkable

Confidence tier

the tier recorded with the decision, informational through determinative-class

Queue position

determinative-class work waiting in the human review queue

9 FOR THE REVIEWER Proof you can check

Governance here is a working surface, not a policy page. Leadership sees trends over real usage, operators run the Approvals inbox against service levels, and your own auditors verify the hash chain — all checkable in the product.

Every governed AI decision is written to a hash-chained decision ledger: capability, model, confidence tier, and the human action taken — accepted, modified or rejected. Your auditors run the integrity check in the product, continuously — not once a year in a letter. The queue carries claiming, service-level due times and disposition history, and determinative-class work never auto-approves.

THE INTEGRITY CHECK

- ✓ chain=decision-ledger
 - **integrity check passed**
 - verified in product
 - run by the tenant auditor

Dashboard view	What it shows	Why a reviewer asks
Adoption and outcomes	Accepted, modified, and rejected decisions over real usage	Shows whether AI is actually helping
Trust	Verified citations, review service levels, confidence tiers	Proves review discipline against your own data
Accuracy	Evaluation score against the release threshold	Shows the release bar still holds after launch
Risk signals	Risk-signal trends over the same usage	Separates noise from consequential work
Finding to corrective action	The finding-to-action funnel	Tests whether drafts become closed work
Cost	Usage against per-capability budgets	Confirms the ceilings and kill-switches are real

Inference runs in SafeGora's own enterprise project on Google Cloud Vertex AI — no consumer AI services, no third-party model vendors, no long-lived model keys in the serving path. Customer data never trains any model; stateless inference retains nothing for model improvement. Personal information is minimized before any model call, workplace video is face-anonymized before frames persist, capture carries workforce-consent notices, and the database forces tenant isolation.

Security posture is published as it stands, including where independent audit engagements currently sit, and a monthly governance brief packages the same substance for recurring delivery from the reports editor. The same honesty runs through the product: screening-grade work is labeled screening-grade, unverifiable citations are flagged rather than presented, thin grounding produces abstention, and an empty dashboard reads "no data yet" instead of inventing a number.

app.safegora.com/ai/advisor


Compliance gap review
 Compare controls against obligations


Incident recordability
 Classify facts and required records


Training matrix
 Check role-based assignments


Inspection readiness
 Verify logs, evidence, and open actions

e.g., What are the fall protection requirements for my construction site?

AI responses are informational. Verify critical decisions with a qualified professional.

The AI advisor workspace — compliance-task cards and the prompt box, with the informational-AI disclosure line carried on the surface itself.

10 THE REFERENCE GRID Two planes, one regime

Employer workspaces run the full stack. The SafeGora practice consoles apply the same stack across a firm's book of business — under class-based access fencing and consent-governed data flows, and without relaxing a single part of the regime.



Concern	Employer plane	Firm plane
Who it is for	Employer tenant owners running their own workspaces	Consulting practices, public consultation programs, carrier loss-control teams
Access and data movement	Tenant isolation forced at the database layer	Class-based access fencing and consent-governed data flows
Consulting and public consultation	Employer side of the grant	Grant, then seat, then switch — revocable, audited, named people
Public consultation boundaries	The employer workspace is unchanged	Template authoring disabled for academic hosts; no money surfaces
Carrier loss control — SafeGora Carrier Loss-Control Practice OS	Participation is explicit and revocable; sponsorship is not access	Consent plane only — no grants, seats or switch; consent-scoped service indicators

Domain	What the AI does	What a person decides
Assistant	Cited answers, proposals as cards; firm-plane book chat	Accepts or rejects every proposal
Incidents	Voice to structured case, narrative, cited recordability triage	Owens the narrative and the determination
Fillings	OSHA Forms 300, 300A and 301, ITA support, certified first-report packs	Certifies and files
Field	Site scan findings, job hazard analysis drafts, structured intake	Accepts or dismisses each finding
Contractors	Extraction with provenance, deterministic requirement checks	Verifier acceptance; certificates never auto-approve
Chemicals	Safety data sheet extraction, versioned federal list indexing	Approves before anything enters the register
Ergonomics	Observational postural coding, deterministic instrument scoring	Accepts findings and opens corrective actions
Predictive and proactive	Risk signals, forward-looking categories, drafted routine content	Disposes of every signal; nothing files itself

AI outputs support the firm's professional judgment and are labeled as AI-assisted wherever they appear. SafeGora does not manufacture regulatory endorsements of any kind.

The work a practice issues to its own clients goes out under the practice's name and carries the firm's own professional authority.

11 RUN IT YOURSELF Evaluate in 10 minutes

A sequence you can run inside the product, in order, without a demo script and without a vendor in the room. Ten minutes is enough to decide whether the governance claims on the preceding pages hold on your own data.

1

Read the regime

Start with the seven principles on page 4. Confirm for yourself that no capability is exempt and that the tiering is stated rather than implied.

2

Scan the index

Take the capability grid on page 12 in one pass. The breadth is one column; what a person still decides is the column beside it.

3

Verify the chain

Open the decision ledger and run the integrity check yourself. The check is continuous, it lives in the product, and your auditors run it — not ours.

4

Open the queue

Look at claiming, due times and disposition history. Then find the determinative-class work that never auto-approves, and confirm on your own screen that it behaves that way.

5

Test the edges

Ask something outside the corpus, then something inside it. Watch for honest abstention, a flagged citation, an empty state that stays empty, and the ledger entry each one leaves.



LEAVE ABLE TO

Explain the governance regime in one sentence

the seven principles on page 4, in your own words

Name what never auto-approves, and why

determinative-class work, whatever the model reports

Name what the AI never decides on its own

recordability, filing and certificate acceptance

Show your committee the ledger entry behind a decision

capability, model, tier and the human action taken

⊗ capability=assistant-answer • grounding insufficient • answer abstained • ledger entry written

Put the regime in front of your reviewers

Bring your security, legal and EHS reviewers to the same ten minutes. Open the ledger, run the integrity check, and take the answer back to your committee with the evidence attached rather than described.

safegora.com

app.safegora.com


SafeGora is independent software — not affiliated with, approved by, or endorsed by OSHA or any government agency.

Nothing in this guide is legal advice; filing and compliance obligations remain with the employer or other legally responsible party.

AI outputs are decision support; a competent human remains the professional of record at every consequential step.

This guide contains no customer endorsements, named customers or attributed quotations.

Comparative vendor verification-and-validation material is available under NDA; it is not a third-party certification.

Disclosure lines carried in the product and on generated artifacts, quoted as they appear:

This is an AI-assisted assessment, not a legal determination. Final recordability should be determined by a qualified safety professional.

** Citation not verified against the jurisdiction corpus — human review required before corrective action.*

Note: where authored state-citation vocabulary is not yet available, citations reference federal OSHA (29 CFR). Verify against current state-plan rules before acting.



RELATED GUIDES

Employer workspace

SafeGora Employer Workspace Guide, on the workspace an employer owns

Consulting practices

SafeGora Consultant Practice OS, on the grant, seat and switch class

Public consultation

SafeGora for public consultation programs, on the publicly funded class

Carrier loss control

SafeGora Carrier Loss-Control Practice OS, on sponsorship without workspace access

Edition August 2026.

© 2026 OllieLabs LLC. All rights reserved. SafeGora and Ollie are marks of OllieLabs LLC. · safegora.com