

From: Justin Boge

To: Margit Kelley; Emily Hicks; Patrick Ward

Date: Aug. 20, 2026

The Badger Lens is essentially a seven-stage screen for thinking through health care AI policy before deciding whether additional regulation is actually needed.

The idea comes largely from how I think about risk as an anesthesiologist. We rarely eliminate risk. We define what we are trying to accomplish, determine what can go wrong and how consequential it could be, identify the protections already in place, add safeguards proportionate to the remaining risk, and then continue monitoring in case circumstances change.

Applied to AI policy, the Badger Lens asks us to:

1. Define the use and intended benefit.
2. Grade the consequence and exposure if it fails.
3. Determine what existing law or regulatory authority already applies.
4. Examine the evidence and whether it fits the real-world population and workflow.
5. Identify the incremental risk created or amplified by the technology.
6. Match any additional safeguards and accountability to the remaining risk.
7. Continue monitoring and reassess as the technology, data, workflow, or vendor changes.

The basic principle is that the regulatory response should follow the consequence and residual risk of the use rather than simply the fact that AI is involved.

That means the screen can lead to different answers. Existing law may already be sufficient. A narrow additional guardrail may be appropriate. A high-consequence use may warrant stronger evidence, qualified human review, auditability, appeal or override mechanisms, and clearer accountability. And, as you suggested, the same approach might also be useful for defining proportional guardrails around a regulatory sandbox.

I see the Badger Lens as a policy-triage tool rather than another comprehensive AI-governance framework. Its purpose is to give people with very different clinical, technical, legal, regulatory, and business backgrounds a common sequence of questions to ask when evaluating a proposal.

I have qualitatively pressure-tested the working version against state statutory examples, representative health care use cases, and several national and international frameworks, but I would not call it validated at this point. I'd welcome the committee pressure-testing it. If the framework holds up under scrutiny from the range of expertise around that table, then it may prove genuinely useful.

Respectfully,

Justin

From: Justin Boge

To: Margit Kelley; Emily Hicks

Date: Aug. 24, 2026

### **Executive Summary: FDA GenAI Medical-Device Discussion Paper and Implications for BADGER LENS**

In August 2026, the FDA Center for Devices and Radiological Health released [Considerations for the Regulation of Generative AI-Enabled Medical Devices: Discussion Paper and Request for Feedback](#). The document is expressly a discussion paper, not draft or final guidance, proposed policy, or a statement of FDA's future regulatory expectations. FDA is requesting stakeholder feedback by October 19, 2026.

The paper is relevant to the Wisconsin Legislative Council Study Committee because its proposed organizing logic closely parallels BADGER LENS. FDA is considering a two-axis risk heuristic based on:

1. the consequence of relying on an incorrect output; and
2. the degree to which the system merely provides information, directs a person toward an action, or independently takes action.

The FDA is also exploring a risk-proportionate, competency-based evaluation model. Under that approach, the final configured product, not simply its underlying foundation model, would undergo defined benchmarking and real-world or clinically representative confirmation. Possible confirmation methods range from retrospective evaluation and shadow deployment to independent clinician adjudication and prospective clinical studies. The necessary rigor would depend on intended use and risk.

BADGER LENS already addresses intended use and benefit, consequence and exposure, existing authority, evidence and real-world fit, incremental technology risk, safeguards and accountability, and lifecycle reassessment. The FDA paper therefore supports the framework's general architecture, but it does not validate BADGER LENS or establish that any particular state policy response is legally required.

Several refinements are advisable before the September committee meeting:

1. Decompose multi-function products into their distinct administrative, informational, action-directing, action-taking, and medical-device functions.
2. Add the system's degree of directiveness, independence, and human supervision to consequence grading.
3. Make evidence criteria more explicit, including pre-specified metrics, independent review, subgroup performance, and risk-proportionate real-world confirmation.
4. Identify responsibility among developers, foundation-model suppliers, deployers, institutions, clinicians, payers, and other affected actors.
5. Specify monitoring triggers for model, data, workflow, population, and vendor changes, along with re-benchmarking, incident response, rollback, and retirement criteria.

6. Add a federalism and jurisdiction check to determine whether a proposed Wisconsin action is grounded in traditional state health, insurance, professional-practice, consumer-protection, Medicaid, or procurement authority.
7. Recognize controlled pilots or regulatory sandboxes as implementation tools when potential benefit is meaningful but evidence remains incomplete.

European governance provides a useful comparison. The EU AI Act emphasizes intended purpose, risk management, data governance, logging, human oversight, provider and deployer responsibilities, post-market monitoring, and incident reporting. The European Health Data Space adds interoperability, patient control, and secure health-data reuse. These European mechanisms should be treated as a menu of potential safeguards rather than imported wholesale into Wisconsin law.

The recommended approach is therefore a BADGER LENS v0.3 refinement, not a redesign. It should remain a neutral, consequence-based policy-triage screen that helps determine whether existing law is sufficient, whether a targeted safeguard is needed, whether enhanced controls are justified, or whether a narrow restriction or hold is warranted. It should not be presented as a substitute for FDA requirements, NIST, the EU AI Act, professional judgment, or legal analysis.

**Principal sources:** FDA August 2026 GenAI-enabled medical-device discussion paper; OMB Memoranda M-25-21 and M-25-22; HHS AI Strategy; Regulation (EU) 2024/1689, as amended; and Regulation (EU) 2025/327 on the European Health Data Space.

Justin

<https://www.fda.gov/media/194242/download>

# THE BADGER LENS

## v0.2.1 | Seven-stage consequence-based screen for health-technology policy

*Working discussion tool - internally stress-tested, not externally validated*

**Starting question:** What is the technology being asked to do, who could be affected, and what happens if it is wrong?

**Core rule:** Use an AI-specific regulatory response only after identifying a material gap in existing law or oversight. The response should rise with consequence, likelihood/exposure, and residual risk - not merely with the presence of AI.

### The Seven-Stage Analyzer

|   |                                             |                                                                                                                                                                              |
|---|---------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1 | <b>Define use &amp; intended benefit</b>    | What function is being performed, for whom, in what workflow, and what benefit is expected? Is AI necessary, or would a simpler tool suffice?                                |
| 2 | <b>Grade consequence &amp; exposure</b>     | If it is wrong, how severe, likely, scalable, reversible, or inequitable could the harm be? How often and how broadly are people exposed?                                    |
| 3 | <b>Check existing law &amp; authority</b>   | What federal/state law, professional standard, regulator, contract, or institutional control already governs the conduct? What material gap remains?                         |
| 4 | <b>Test evidence &amp; real-world fit</b>   | Has the system been validated for the intended purpose, population, setting, and human workflow? Are subgroup, comparator, and operational data adequate to the consequence? |
| 5 | <b>Identify incremental technology risk</b> | What risk is newly introduced or materially amplified by AI/automation - e.g., scale, opacity, automation bias, drift, cyber risk, data use, vendor changes, or dependency?  |
| 6 | <b>Match controls &amp; accountability</b>  | What additional safeguards are proportionate to residual risk, and who is accountable for decisions, monitoring, escalation, correction, and response?                       |
| 7 | <b>Monitor, revalidate &amp; exit</b>       | What is monitored after deployment? What triggers revalidation, restriction, rollback, or retirement when the model, data, population, workflow, or vendor changes?          |

### Legislative Workflow

USE -> CONSEQUENCE/EXPOSURE -> EXISTING AUTHORITY -> EVIDENCE/FIT -> INCREMENTAL RISK -> CONTROLS/ACCOUNTABILITY -> LIFECYCLE

### Possible Policy Outputs

**A. Existing framework sufficient:** No new AI-specific rule identified; ordinary legal, professional, privacy, security, and quality controls still apply. **B. Targeted safeguard:** Address a defined gap with a narrow requirement such as disclosure, documentation, auditability, or specified human review. **C. Enhanced high-consequence controls:** Stronger evidence, qualified review, appeal/override, security, incident response, and lifecycle monitoring. **D. Restrict or hold:** Use a narrow restriction, pause, or prohibition when evidence is inadequate or residual risk remains unacceptable.

### What the v0.2.1 Stress Test Supports

This version was qualitatively applied to 20 enacted statutory examples from 17 states that are relevant to health-AI governance (not all are AI-specific) and to 10 representative healthcare use cases. It was also cross-walked to NIST AI RMF, FDA/IMDRF device-focused principles and guidance, WHO health-AI materials, FUTURE-AI, and the EU AI Act. The exercise did not validate the tool; it identified wording, scope, and evidence gaps and supported keeping an explicit Evidence & Fit stage.

**Scope note:** Badger Lens is a policy-triage layer. It does not replace applicable law, FDA requirements, NIST risk management, professional standards, institutional governance, or other domain-specific frameworks.

**Reference base (abbrev.):** NIST AI RMF 1.0 (2023, revision underway); FDA AI/SaMD materials, Transparency principles (2024), PCCP final guidance (Dec. 2024), AI-enabled device lifecycle draft guidance (Jan. 2025), and FDA/Health Canada/MHRA + IMDRF GMLP work; WHO Ethics & Governance of AI for Health (2021) and Regulatory Considerations on AI for Health (2023); Lekadir et al., FUTURE-AI, BMJ 2025;388:e081554; Regulation (EU) 2024/1689; Wisconsin Legislative Council Staff Brief SB-2026-02 (Aug. 5, 2026).

# BADGER LENS v0.2.1

## Verification & Qualitative Stress-Test Appendix

### 1. Purpose and Method

This appendix documents an internal, qualitative stress test of the Badger Lens. It is not a validation study and should not be described as one. The analyzer was applied to a purposive sample of 20 enacted statutory examples from 17 states relevant to health-AI governance, 10 representative healthcare use cases, and five major external governance/regulatory reference families. The sample is not exhaustive or statistically representative of state legislation.

Legal-source anchor: Wisconsin Legislative Council Staff Brief SB-2026-02 (Aug. 5, 2026), which cites the underlying state statutes and enactments. Selected primary state sources were rechecked during this verification review; where a primary enacted text conflicted with earlier shorthand, the enacted text controls.

### 2. Skeptical Review: Corrections Made in v0.2.1

**Sample description:** Changed “20 enacted health-AI laws” to “20 enacted statutory examples relevant to health-AI governance.” At least one example (Utah preauthorization review) is an AI-relevant human-judgment baseline rather than an AI-specific statute.

**Louisiana:** Corrected the AI-transcription row. Enacted R.S. 37:22.1 requires verbal disclosure before recording for AI transcription; the enacted text does not require patient consent.

**FDA dates/status:** Corrected the PCCP final guidance to December 2024. The January 2025 AI-enabled device lifecycle document is draft guidance. GMLP attribution is described as FDA/Health Canada/MHRA work that was subsequently internationalized through IMDRF.

**WHO terminology:** The 2023 WHO document is described as “Regulatory Considerations on Artificial Intelligence for Health”; WHO states that it is not itself a regulatory framework or policy.

**Validation language:** Removed “the seven stages survived the test.” The present exercise is qualitative and internal; outputs are illustrative analyst judgments, not validated scores or legal conclusions.

**Stage 2:** Added likelihood/exposure explicitly, so consequence is not treated as severity alone.

**Stage 5:** Changed “what AI uniquely changes” to “incremental technology risk.” Many risks are amplified by AI rather than unique to it.

### 3. External Framework Crosswalk

| Badger Lens stage           | NIST AI RMF                                                        | FDA / IMDRF (device-focused)                                                      | WHO health-AI materials                                   | FUTURE-AI                                                          | EU AI Act                                                                           |
|-----------------------------|--------------------------------------------------------------------|-----------------------------------------------------------------------------------|-----------------------------------------------------------|--------------------------------------------------------------------|-------------------------------------------------------------------------------------|
| 1 Use & benefit             | MAP: context, intended purpose, affected parties; go/no-go context | Intended use, user, clinical context; benefit-risk framing                        | Human autonomy; public benefit; intended health use       | Usability, universality, clinical context                          | Intended purpose; deployer/provider context                                         |
| 2 Consequence & exposure    | MAP/MEASURE: likelihood, magnitude, affected groups                | Benefit-risk; severity and patient impact                                         | Safety, well-being, equity; risk-benefit considerations   | Traceability risk management; fairness; robustness                 | Risk classification; Art. 9 risk management                                         |
| 3 Existing authority        | GOVERN 1.1: legal/regulatory requirements                          | Device jurisdiction and applicable statutory/regulatory pathways                  | Legal/regulatory context and governance responsibilities  | Cross-cutting legal/regulatory and governance considerations       | Horizontal AI duties plus sector-specific law                                       |
| 4 Evidence & fit            | MEASURE: validity, reliability, testing, context                   | GMLP; clinical/technical validation; safety/effectiveness                         | Safety, effectiveness, inclusiveness and intended setting | Universality, robustness, fairness, usability                      | Data governance; accuracy/robustness; conformity requirements for high-risk systems |
| 5 Incremental tech risk     | AI-specific/emergent, third-party, opaque and dynamic risks        | Adaptive change, transparency, cybersecurity, data/model change                   | Bias, privacy, opacity, security and automation concerns  | Traceability, explainability, robustness                           | AI-specific prohibited/high-risk duties; systemic/emergent risks                    |
| 6 Controls & accountability | GOVERN/MANAGE: controls, roles, escalation, residual risk          | Risk controls, labeling/transparency, human factors and regulatory accountability | Human autonomy; responsibility and accountability         | Traceability/governance; usability; explainability                 | Human oversight, logging, provider/deployer obligations                             |
| 7 Lifecycle                 | Monitoring, feedback, change management, decommissioning           | Total-product-lifecycle approach; PCCP; postmarket/change control                 | Responsiveness/sustainability; ongoing governance         | Lifecycle best practices from design through deployment/monitoring | Post-market monitoring; serious incident reporting; substantial change              |

**Interpretation:** The crosswalk shows substantial conceptual compatibility but not equivalence. Badger Lens is intended to ask a prior policy-triage question - whether an additional policy response is needed and how narrow it should be - rather than replace any of these frameworks.

4. State Statutory Sample (20 Examples / 17 States)

These examples are purposively selected from the Wisconsin Legislative Council survey of enacted state health-AI provisions. “Enacted” includes signed/chaptered provisions that may have a future effective date. The table is a screening sample, not a 50-state legal survey.

| State / authority                                                     | Policy function                        | Verified mechanism / relevance                                                                                                                                                                                    | Badger stages |
|-----------------------------------------------------------------------|----------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------|
| AL - Ala. Code §27-1-17.2                                             | Insurer utilization review             | AI determinations must be patient-specific; disclosure, fairness/non-discrimination, privacy, and periodic accuracy/reliability review.                                                                           | 2,4,6,7       |
| CO - HB 26-1139 (enacted; eff. 1/1/2027)                              | Utilization review / coverage          | AI may assist and expedite approvals; adverse medical-necessity denial cannot issue solely from AI output without competent human review; individualized data, audit, fairness and periodic review requirements.  | 1,2,4,6,7     |
| MD - Md. Code, Ins. §15-10B-05.1                                      | Utilization review                     | Patient-specific data; provider role preserved in adverse coverage decisions; auditability, policies, quarterly performance review, privacy; AI may not deny, delay, or modify services.                          | 2,4,6,7       |
| WA - RCW 48.43.830; 41.05.845                                         | Prior authorization                    | AI not sole means to deny/delay/modify; licensed human review for medical-necessity denials; patient-specific inputs, auditability, fairness and periodic performance review.                                     | 2,4,6,7       |
| NE - Neb. Rev. Stat. §44-5443                                         | Prior authorization                    | AI-based algorithm may not be sole basis to deny/delay/modify; AI use disclosed; regulator may audit compliance.                                                                                                  | 2,6           |
| TX - Tex. Ins. Code §4201.156                                         | Adverse utilization review             | Automated decision system, including AI, may not make an adverse determination.                                                                                                                                   | 1,2,6         |
| UT - Utah Code §31A-22-650                                            | Preauthorization human-review baseline | AI-relevant but not AI-specific: adverse medical-necessity decisions require an appropriately qualified individual exercising independent medical judgment rather than relying solely on another recommendation.  | 2,3,6         |
| IN - Ind. Code §27-1-52-9                                             | Claims / downcoding                    | Automated process/system/tool, including AI, cannot be sole basis for medical-necessity downcoding without human medical-record review; provider-side automated claim submission also requires human review.      | 1,2,6         |
| LA - La. R.S. 37:22.1 (HB 475 / Act 649, 2026)                        | Clinical recording / AI transcription  | Healthcare professional must verbally disclose recording device/software/service before recording an encounter for AI transcription. Enacted text does not require consent.                                       | 1,2,6         |
| TX - Tex. Bus. & Com. Code §552.051(f); Health & Safety Code §183.005 | Provider / diagnostic AI disclosure    | Requires patient disclosure when AI is used in relation to health service/treatment and for diagnostic recommendations as specified by statute.                                                                   | 1,2,6         |
| CA - Cal. Bus. & Prof. Code ch. 15.5                                  | Professional identity                  | AI/nonhuman system may not imply that licensed-human health-care advice, reports or assessments are being provided when they are not.                                                                             | 1,2,6         |
| DE - 24 Del. C. §1920(h)-(i)                                          | Professional licensure / titles        | Nonhuman entity, including AI agent, may not be licensed for nursing roles or use protected nursing/doctor titles.                                                                                                | 1,2,6         |
| OR - Or. Rev. Stat. §678.027                                          | Professional titles                    | Restricts nonhuman entities from using protected nursing-related titles.                                                                                                                                          | 1,2,6         |
| CO - HB 26-1195 (enacted 2026)                                        | Psychotherapy / mental health          | Separates administrative/supplementary AI support from therapeutic communication; preserves clinician responsibility and imposes disclosure/consent requirements for specified recorded/transcribed therapy uses. | 1,2,6         |
| IL - 2025 HB 1806 / Public Act 104-0054                               | Mental health                          | Permits defined support uses with clinician responsibility; bars independent therapeutic decisions/direct therapeutic communication and specified emotion-detection uses.                                         | 1,2,5,6       |
| ME - 2026 Public Law ch. 687 (LD 2082)                                | Mental health                          | Permits administrative/supplementary support with provider review; imposes notice/consent for specified recording/transcription; bars independent therapeutic decisions.                                          | 1,2,6         |
| NV - NRS §629.610                                                     | Mental/behavioral health               | Bars AI from directly providing mental/behavioral health care; allows administrative support; requires independent provider review of generated/compiled information.                                             | 1,2,6         |
| TN - Tenn. Code §33-1-205                                             | Mental health identity                 | AI system may not be advertised/represented as a mental health provider or as able to act as one.                                                                                                                 | 1,2,6         |
| WA - RCW ch. 19.440                                                   | Companion chatbots                     | Disclosure that chatbot is AI/nonhuman; suicide/self-harm protocols; additional protections for minors and manipulative/dependency-related behaviors.                                                             | 1,2,5,6,7     |
| GA - Ga. Code §39-5-6                                                 | Companion chatbots                     | Disclosure, self-harm/minor safeguards, and restrictions on specified manipulative/dependency-inducing behavior and claims of humanness/sentence.                                                                 | 1,2,5,6,7     |

5. What Can and Cannot Be Inferred From the State Sample

- In this selected sample, legislatures commonly attach stronger requirements to consequential functions such as adverse coverage determinations, therapeutic communication, protected professional identity, and vulnerable-user chatbot interactions.
- Human review, disclosure, individualized information, auditability, non-discrimination, privacy, and periodic review recur, but not uniformly.
- State statutes often specify process safeguards more readily than clinical/scientific validation standards. That observation supports retaining a distinct Evidence & Fit stage, but it does not prove that state law generally omits validation.
- The sample cannot establish which regulatory model is most effective, whether a law improved outcomes, or whether a particular bill passed because it was risk-proportionate. Most enactments are too new for outcome evaluation.
- Badger Lens should therefore be presented as an organizing screen that can be empirically studied - not as a best-practice standard.

6. Ten Representative Use-Case Applications

| Use case                                     | Consequence / exposure | Key Badger Lens questions                                                                                                            | Illustrative output* |
|----------------------------------------------|------------------------|--------------------------------------------------------------------------------------------------------------------------------------|----------------------|
| 1. Appointment scheduling / clerical routing | Low                    | Privacy/security; error recovery; accessibility; existing workflow controls.                                                         | A/B                  |
| 2. Ambient clinical scribe                   | Moderate               | Recording/privacy; note accuracy; clinician verification; vendor data use; workflow fit; model updates.                              | B                    |
| 3. Draft patient-portal responses            | Moderate               | Hallucination/tone; urgency detection; clinician review; escalation; vulnerable patients.                                            | B                    |
| 4. Translation / consent support             | Moderate-high          | Language accuracy; informed-consent implications; population validation; fallback interpreter; accountability.                       | B/C                  |
| 5. Sepsis or deterioration prediction        | High                   | Local/population validation; false positives/negatives; alarm burden; override; drift; downtime.                                     | C                    |
| 6. Diagnostic imaging assistance             | High                   | FDA/device status where applicable; clinical validation; subgroup performance; workflow/human factors; incident monitoring.          | C                    |
| 7. Treatment / dosing recommendation         | High-critical          | Evidence quality; contraindications; clinician competence/override; audit trail; liability; model change.                            | C/D                  |
| 8. Therapeutic mental-health chatbot         | High-critical          | Direct therapeutic role; minors/self-harm; manipulation/dependency; professional practice law; escalation; evidence.                 | C/D                  |
| 9. AI-assisted adverse prior authorization   | High                   | Patient-specific facts; qualified independent review; appeal; auditability; discrimination; performance monitoring.                  | C                    |
| 10. Closed-loop procedural/device control    | Critical               | Device jurisdiction; fail-safe design; validation; human takeover; cybersecurity; incident response; rollback/PCCP where applicable. | C/D                  |

**\*Illustrative output only:** A = existing framework sufficient/no new AI-specific rule identified; B = targeted safeguard; C = enhanced high-consequence controls; D = narrow restriction/hold. These are analyst judgments, not legal determinations or validated scores.

7. Reference Verification Notes

1. NIST. Artificial Intelligence Risk Management Framework (AI RMF 1.0), NIST AI 100-1 (Jan. 2023); NIST has announced/referenced an ongoing revision process as of 2026. The Badger Lens crosswalk uses the GOVERN, MAP, MEASURE, and MANAGE functions and associated legal/compliance, risk, monitoring, and decommissioning concepts.
2. FDA. Artificial Intelligence in Software as a Medical Device. FDA identifies careful management across the medical-product lifecycle; its page lists Transparency principles (June 2024), final PCCP guidance (Dec. 2024), and draft AI-enabled device lifecycle/marketing-submission guidance (Jan. 2025).
3. Good Machine Learning Practice. FDA, Health Canada, and MHRA published 10 GMLP principles in 2021; IMDRF published a final international GMLP principles document in 2025 building on that work.
4. WHO. Ethics and Governance of Artificial Intelligence for Health (2021). WHO. Regulatory Considerations on Artificial Intelligence for Health (2023). The latter is a set of regulatory considerations and expressly is not itself a regulatory framework or policy.
5. Lekadir K, et al. FUTURE-AI: international consensus guideline for trustworthy and deployable artificial intelligence in healthcare. BMJ. 2025;388:e081554. Six principles: fairness, universality, traceability, usability, robustness, explainability; 30 lifecycle best practices.
6. European Union. Regulation (EU) 2024/1689 (AI Act). Risk-based obligations include high-risk risk-management requirements and post-market monitoring; application is phased and should not be described as though every provision was already fully applicable in August 2026.
7. Wisconsin Legislative Council. Staff Brief SB-2026-02, Study Committee on the Use of Artificial Intelligence in Health Care (Aug. 5, 2026). Primary source for the committee-specific state-law survey and statutory citations used in this appendix.
8. Louisiana. Enrolled HB 475 (2026), enacting R.S. 37:22.1: verbal disclosure is required before recording an appointment/treatment for AI transcription; consent is not in the enacted requirement.
9. Colorado. HB 26-1139 became law in 2026; its utilization-review provisions are effective Jan. 1, 2027. It permits AI assistance/expedited approvals while requiring patient-specific criteria, competent human review for adverse medical-necessity decisions, auditability, and periodic review.
10. Maryland §15-10B-05.1, Nebraska §44-5443, Washington RCW 48.43.830/41.05.845, Delaware 24 Del. C. §1920, Illinois Public Act 104-0054, and Maine Public Law ch. 687 were additionally checked against current official state legislative/code sources during the v0.2.1 verification review.

8. Proposed Next Validation Step

For academic or external policy claims, the next step should be prospective validation rather than more branding: structured expert content review; standardized case vignettes; independent scoring by clinicians, health-law attorneys, regulators, payers, technologists, and patient representatives; inter-rater reliability; sensitivity testing; and revision of ambiguous criteria. Until then, describe Badger Lens as an internally stress-tested working policy screen.