



Eurasia Specialized Veterinary Publication

International Journal of Veterinary Research and Allied Science

ISSN:3062-357X

2025, Volume 5, Issue 2, Page No: 139-150

Copyright CC BY-NC-SA 4.0

Available online at: www.esvpub.com/

Diagnostic Accuracy and Clinical Utility of Point-of-Care Biomarkers for Early Recognition of Sepsis and Systemic Inflammation in Dogs and Cats: A Systematic Review of Veterinary Evidence, 2017–2025

Wei Zhang^{1*}, Chen Hui², Michael Tan¹

¹Department of Veterinary Emergency and Critical Care Medicine, Faculty of Veterinary Science, University of Melbourne, Melbourne, Australia.

²Department of Veterinary Internal Medicine and Sepsis Research, Faculty of Veterinary Medicine, National University of Singapore, Singapore.

*E-mail ✉ wei.zhang@unimelb.edu.au

ABSTRACT

Early recognition of sepsis in dogs and cats is difficult because systemic inflammation, infection, organ dysfunction, and treatment effects overlap clinically. Point-of-care biomarkers could shorten the interval between suspicion and reassessment, but analytical validity, diagnostic discrimination, prognosis, and decision utility are distinct claims. To systematically evaluate veterinary evidence published from 2017 through 2025 on point-of-care or clinically rapid biomarkers for recognizing sepsis or systemic inflammation in dogs and cats, and to distinguish measurement performance from sepsis-specific diagnostic and clinical utility. Dogs and cats represented in eligible veterinary clinical and analytical studies, including emergency, critical-care, pyometra, pneumonia, septic peritonitis, pyothorax, and hospitalized populations. A systematic review was conducted using prespecified eligibility, structured searching, duplicate handling, staged screening, data extraction, diagnostic-accuracy fields, and risk-of-bias assessment. Evidence was synthesized narratively when populations, thresholds, reference standards, platforms, or clinical purposes were not sufficiently comparable for quantitative pooling. The evidence base was heterogeneous and predominantly canine. Acute-phase proteins, inflammatory mediators, endothelial and hemostatic markers, cell-death products, routine hematologic ratios, and local-fluid tests addressed different clinical questions. Several biomarkers detected inflammation or tracked severity without reliably establishing infection-specific discrimination. Point-of-care analytical validation was more mature for some assays than sepsis-specific clinical validation. Feline evidence was sparse. Point-of-care biomarkers should not be treated as interchangeable sepsis tests. Their value depends on the target condition, comparator spectrum, timing, reference standard, assay credibility, and intended use. A function-specific interpretation separating inflammation detection, infection attribution, prognosis or monitoring, and actionability is clinically more defensible than a single biomarker ranking.

Keywords: Sepsis, Systemic inflammation, Point-of-care testing, Biomarkers, Dogs, Cats

Received: 02 September 2025

Revised: 14 October 2025

Accepted: 16 October 2025

How to Cite This Article: Zhang W, Hui C, Tan M. Diagnostic Accuracy and Clinical Utility of Point-of-Care Biomarkers for Early Recognition of Sepsis and Systemic Inflammation in Dogs and Cats: A Systematic Review of Veterinary Evidence, 2017–2025. *Int J Vet Res Allied Sci.* 2025;5(2):139-50. <https://doi.org/10.51847/Ifv3CaWJ45>

Introduction

Sepsis recognition in small animals remains clinically difficult because infection, dysregulated inflammation, and organ dysfunction are incompletely aligned. Contemporary veterinary guidance emphasizes that no single validated definition or diagnostic test captures small-animal sepsis, while emergency-department data show that

systemic inflammatory response syndrome (SIRS) criteria are common in sick dogs and cats and therefore insufficiently specific when used alone [1, 2]. Organ dysfunction adds important severity information but does not independently establish infectious causation [3]. Biomarker evidence inherits this problem: even promising inflammatory signals can change apparent performance when the clinical phenotype or SIRS definition changes [4]. This review therefore evaluates point-of-care and clinically rapid biomarkers according to what they actually establish—measurement credibility, systemic-inflammation detection, infection discrimination, prognosis or serial monitoring, and decision consequence—rather than assuming that all positive biomarker signals answer the same sepsis question.

Review objective

The objective was to determine how accurately and usefully point-of-care or clinically rapid biomarkers identify sepsis and systemic inflammation in dogs and cats, with explicit separation of four functions: analytical measurement, early inflammatory recognition, discrimination of infection or sepsis from clinically relevant mimics, and prognostic or monitoring utility. The review also examined how species, infection source, illness spectrum, sampling time, assay platform, threshold selection, and reference-standard credibility constrain inference. This four-function lens is proposed as an organizing synthesis rather than a validated decision rule. It is intended to prevent an analytically valid assay from being misclassified as a sepsis diagnostic, or an outcome-associated biomarker from being interpreted as proof of infection. Clinical actionability was considered the most downstream claim and therefore required evidence beyond biomarker association alone.

Materials and Methods

A systematic-review protocol was specified before synthesis. Eligible records were identified with veterinary species terms, sepsis or systemic-inflammation concepts, biomarker and point-of-care terms, and infection-source terms relevant to small-animal emergency practice. Records were deduplicated before staged title/abstract and full-text assessment. Extraction separated population and setting, specimen, assay or platform, sampling time, threshold, comparator, reference standard, diagnostic measures, prognostic outcomes, serial behavior, and reported feasibility. Primary veterinary empirical evidence was distinguished from methodological and reporting sources used to support review conduct. Diagnostic and prognostic claims were coded separately. No unreported threshold, accuracy statistic, effect size, sample characteristic, or outcome was inferred. Because the eligible literature included different target conditions, platforms, comparators, and reference standards, synthesis was prespecified as narrative unless quantitative compatibility was demonstrable.

Protocol and reporting framework

Reporting followed PRISMA 2020 principles, with diagnostic-test-specific fields drawn from PRISMA-DTA and search documentation aligned with PRISMA-S [5-7]. These frameworks were used to preserve the distinction among the index test, target condition, comparator or reference standard, threshold, patient spectrum, and review-selection process. Their use was reporting-oriented: adherence to a reporting checklist was not treated as evidence that an individual primary study had low risk of bias or that the review search was necessarily exhaustive. For this veterinary context, the reporting framework was additionally adapted to record species, infection source, emergency or intensive-care setting, local-fluid versus systemic testing, and whether the assay was genuinely point-of-care, rapid laboratory testing, or a research biomarker without bedside deployment. These distinctions were carried forward into extraction and synthesis.

Eligibility criteria

Peer-reviewed journal articles published from 2017 through 2025 were eligible when they involved dogs, cats, or directly applicable review methodology; had a verifiable DOI; met the prespecified journal-quality requirement; and supported at least one defined review domain. Primary review evidence required empirical evaluation of a biomarker, rapid test, clinically accessible analyte, relevant comparator, infection source, reference standard, prognosis, serial response, or analytical performance pertinent to sepsis or systemic inflammation. Studies were excluded for incompatible species, publication dates outside the window, duplicate versions, absence of a verifiable DOI, topic adjacency without claim-level relevance, or materially incompatible clinical targets. Methodological papers informed reporting, appraisal, or synthesis but were not treated as included veterinary

diagnostic evidence units. Local infection tests were retained only when their target condition was explicitly preserved rather than relabeled as systemic sepsis.

Population and clinical setting

The target population comprised dogs and cats in settings where early systemic inflammation or infection recognition could alter diagnostic or monitoring decisions. The evidence encompassed emergency-department patients, critically ill animals, hospitalized cats, and syndrome-specific populations such as pyometra, pneumonia, septic peritonitis, and pyothorax. These groups were not considered interchangeable. A biomarker that separates diseased animals from healthy controls addresses a different spectrum from a test evaluated against noninfectious critical illness, and a local cavity-effusion test addresses source infection rather than systemic sepsis. This distinction is central because clinically useful specificity is most challenged by inflammatory mimics, prior treatment, organ dysfunction, and mixed infection sources. Species was retained as an explicit boundary: canine findings were not transferred to cats unless feline evidence supported the same inference.

Index biomarkers

The index-biomarker domain included acute-phase proteins, procalcitonin, serum amyloid A (SAA), inflammatory cytokines and chemokines, endothelial or glycocalyx markers, coagulation and fibrinolysis variables, cell-free DNA and nucleosomes, CBC-derived ratios, routine biochemical variables, and rapid local-fluid infection tests. In a prospective canine comparison, C-reactive protein (CRP) behaved as an inflammation-responsive marker, whereas procalcitonin did not provide robust separation between sepsis and noninfectious SIRS [8]. A point-of-care luminometer showed clinically relevant accuracy for septic cavity effusions against aerobic culture, but its target was local septic effusion rather than systemic sepsis [9]. In hospitalized cats, point-of-care feline SAA measurement correlated with reference testing, yet SAA or the SAA-to-albumin ratio did not establish sepsis-specific diagnostic or prognostic validity in the small clinical cohort [10].

Reference standards

Reference standards were treated as clinical adjudication structures rather than assumed gold standards. Depending on the study question, they included microbiologic culture, cytologic or biochemical evidence from infected body-cavity fluid, syndrome definitions based on clinical infection plus systemic inflammation, and organ-dysfunction classifications. Each carries a different error structure. Culture can be affected by sampling site, prior antimicrobial exposure, organism burden, and specimen quality; SIRS-defined sepsis can incorporate the same physiologic abnormalities being evaluated; and healthy-control comparisons can exaggerate apparent discrimination. Organ dysfunction strengthens severity characterization without independently establishing infection. Consequently, diagnostic performance was interpreted only within the target condition actually verified by each study. The review therefore used a proposed reference-standard credibility distinction: direct source confirmation, clinically adjudicated infection, inflammation-based syndrome classification, and severity-only classification were kept analytically separate rather than combined.

Table 1 presents what each veterinary sepsis reference standard verifies—and what it cannot verify.

Table 1. What each veterinary sepsis reference standard verifies—and what it cannot verify

Reference basis	Material or clinical basis	Target condition verified	Factors altering verification	Use of a positive biomarker	Inference unavailable	Separate confirmation and evidentiary origin
Direct source confirmation	Culture or source-specific microbiologic evidence	Supports infection at sampled site	Sampling quality; prior antimicrobials	Strengthens infection attribution	Negative culture does not exclude infection	Confirmation: Reconcile result with treatment and clinical course Origin: Evidence-derived
Septic effusion adjudication	Local-fluid microbiology, cytology, or biochemical pattern	Supports an infected cavity	Effusion type; sampling protocol	Guides source-directed management	Local infection is not equivalent to systemic sepsis	Confirmation: Reassess systemic state separately Origin: Evidence-derived

Clinical infection plus systemic inflammation	Compatible source plus inflammatory syndrome	Supports clinically adjudicated sepsis	Case definition; comparator spectrum	Raises sepsis probability	Incorporation and spectrum bias remain possible	Confirmation: Seek independent source or trajectory evidence Origin: Evidence-derived
Organ-dysfunction classification	New or worsening organ dysfunction	Defines severity consequence	Comorbidity; treatment effects	Escalates monitoring and support	Does not establish infectious cause	Confirmation: Reassess cause and trajectory Origin: Evidence-derived
Integrated adjudication	Concordant source, inflammation, organ state, and assay credibility	Higher-confidence contextual interpretation	Species; time; platform; treatment	Supports contextual action, not a universal score	Requires prospective validation	Confirmation: Repeat discordant domains rather than average them Origin: Proposed synthesis

Information sources

Information-source planning prioritized PubMed-indexed veterinary literature and supplementary targeted retrieval of directly relevant diagnostic, biomarker, methodological, and infection-source evidence. Search documentation was designed to retain the database or interface, search date, complete syntax, limits, and retrieved-result count because omission of these elements materially reduces reproducibility [11]. A result set exposed by an interface was not assumed to be equivalent to a complete database export. Search-derived numerical quantities were therefore eligible for reporting only when generated by the recorded retrieval and deduplication process and were never inferred from bibliography size or a desired inclusion total. Citation chasing and targeted known-item checking were used as sensitivity checks, but supplementary discovery was kept methodologically distinct from database-derived retrieval quantities.

Search strategy

Search strings combined controlled and free-text terms for dog, canine, cat, feline, sepsis, SIRS, systemic inflammation, biomarker, point-of-care testing, CRP, SAA, procalcitonin, lactate, cytokines, chemokines, endothelial injury, coagulation, cell-free DNA, nucleosomes, and major septic syndromes. Search-string sensitivity was checked against known relevant records rather than assumed from face validity [12]. Search filters can trade sensitivity against precision, so opaque filters were avoided when direct concept terms could be used [13]. Search quality assurance also recognized that information-specialist involvement is associated with stronger search reporting and reproducibility, although such involvement cannot guarantee exhaustive retrieval [14]. Searches were restricted to the prespecified publication window; species and clinical-target decisions were applied during screening when this reduced the risk of inadvertently excluding relevant mixed-species or methodological records.

Study selection

After deduplication, records were screened against the prespecified year, species, publication, DOI, journal-quality, clinical-target, and evidence-fit criteria. Full texts were sought for records that could plausibly inform diagnostic performance, analytical validity, prognosis, monitoring, reference-standard interpretation, or clinically relevant heterogeneity. Full-text exclusion required a single interpretable primary reason, such as incompatible species, ineligible publication type, nonqualifying target condition, insufficient empirical information, duplicate article version, or failure of the prespecified journal or DOI requirement. Reporting and diagnostic-accuracy guidance informed staged selection, but methodological papers were kept distinct from veterinary evidence units. Selection numbers were not inferred from the final bibliography or harmonized by assumption; any numerical study flow reported elsewhere in the review must originate from the same underlying screening record.

Data extraction

A structured extraction form separated study identity from the clinical claim the study could support. Fields included country, species, sample and setting, infection phenotype, study design, index biomarker or device, specimen, platform, timing, threshold, comparator, reference standard, analytical method, sensitivity or specificity when reported, prognostic endpoint, serial behavior, feasibility, and authors' limitations. Analytical validation and clinical validation were coded separately. For local-fluid tests, the extracted target condition remained septic effusion or septic peritonitis rather than being broadened to systemic sepsis. For prognostic studies, mortality or severity associations were not recoded as diagnostic discrimination. Missing values were recorded as not reported

rather than inferred. Extraction also preserved alternative explanations including treatment exposure, healthy-control spectrum, assay interference, and case-definition dependence for subsequent heterogeneity and bias interpretation.

Risk-of-bias assessment

Diagnostic-accuracy evidence was appraised by domain rather than converted into a numerical quality score. Patient selection, index-test conduct and interpretation, reference-standard credibility, and flow or timing were evaluated using diagnostic-accuracy risk-of-bias principles [15]. Indirectness was considered when the population, platform, target condition, or reference standard differed from the clinical question of early canine or feline sepsis recognition [16]. Inconsistency, imprecision, and possible publication bias were treated as body-of-evidence limitations when studies were small or methodologically heterogeneous [17]. Prognostic evidence was appraised separately because association with mortality, disease severity, or biomarker trajectory does not establish diagnostic accuracy [18]. Particular attention was directed to healthy-control designs, incorporation of SIRS elements into target definitions, absence of noninfectious critically ill comparators, small feline cohorts, and lack of independent external validation.

Diagnostic-accuracy measures

For studies with a defined target condition and reference standard, extraction captured reported sensitivity, specificity, likelihood ratios, receiver-operating-characteristic area under the curve, threshold, and uncertainty measures when available. Metrics were interpreted at the exact threshold, specimen, time point, platform, and comparator used by the investigators. Analytical agreement, precision, linearity, and assay interference were classified as measurement-performance outcomes and were not converted into sepsis sensitivity or specificity. Likewise, prognostic associations and serial biomarker decline were not treated as diagnostic-accuracy measures. Quantitative pooling was considered only if populations, target conditions, reference standards, assay platforms, and thresholds were sufficiently compatible; otherwise, preserving study-specific estimates was judged more defensible than producing a pooled value with uncertain clinical meaning. No unreported diagnostic metric was calculated from incomplete published data.

Evidence synthesis

Synthesis was organized by the clinical function a biomarker could defensibly perform rather than by analyte concentration alone. Evidence was first separated into analytical validity, inflammatory-state recognition, infection or sepsis discrimination, prognosis or serial monitoring, and potential decision utility. Studies using healthy animals as the principal comparator were interpreted differently from studies including noninfectious critically ill animals. Local-source tests were retained as infection-detection evidence but were not pooled conceptually with systemic biomarkers. Thresholds and diagnostic estimates remained study specific when assays, populations, reference standards, or sampling times differed. This approach avoided implying equivalence among markers whose biological pathways, intended uses, and validation stages were substantially different.

Table 2 presents clinical claims permitted at successive levels of point-of-care biomarker evidence.

Table 2. Clinical claims permitted at successive levels of point-of-care biomarker evidence

Evidentiary question	Analytical validity	Inflammatory signal	Sepsis discrimination	Prognostic/serial signal	Actionable POC result
What is actually observed	Acceptable assay agreement or precision	Marker differs from health or changes with inflammation	Difference against relevant noninfectious comparators	Association with outcome or temporal change	Rapid result plus credible target-condition evidence
Biomarker claim permitted	Marker can be measured credibly	Biological response detected	Possible diagnostic information	Severity or trajectory information	Potential bedside decision support
Clinical role	Measurement may enter clinical evaluation	Supports contextual suspicion	May refine infection probability	May support monitoring	May influence testing or treatment sequence

Evidence still absent	Disease discrimination unproved	Infection specificity uncertain	Threshold and reference standard matter	Does not establish diagnosis	Outcome benefit unestablished
Validation step	Validate in target population	Compare with clinical mimics	External validation required	Repeat within clinical course	Test decision impact prospectively

Results and Discussion

The retained evidence demonstrated substantial biological and methodological diversity. Recent canine studies described acute-phase and lipoprotein abnormalities, endothelial activation, and hemostatic or fibrinolytic perturbations in sepsis [19-21]. These findings supported multiple pathobiological domains rather than a single dominant biomarker pathway. Importantly, elevation in septic animals did not automatically establish discrimination from noninfectious critical illness, particularly when healthy animals supplied the comparator. Temporal context also mattered. Serial assessment of hyaluronic acid and related endothelial variables illustrated that biomarker concentrations can change during critical illness and alongside therapeutic exposures, including fluid administration [22]. Across the evidence base, the strongest recurring distinction was therefore between detecting a disturbed biological state and identifying its infectious cause. Analytical feasibility, prognostic association, and rapid turnaround were additional dimensions, but none could substitute automatically for clinically relevant diagnostic specificity.

Study-selection flow

The available selection record documented an initial targeted discovery subset of 60 surfaced records across five prespecified query families. Eight duplicate records within that subset were removed, leaving 52 unique surfaced records. Title/abstract screening excluded 16 records on prespecified eligibility grounds, leaving 36 potentially eligible records within the verified subset. The available selection records did not contain a complete uncapped database-export denominator or a reconciled full-text-to-final-inclusion chain. Consequently, no additional selection number was reconstructed, estimated, or inferred.

Because a conventional completed flow would falsely imply verification beyond the available selection record, **Figure 1** makes the boundary between verified selection accounting and unavailable downstream counts explicit.

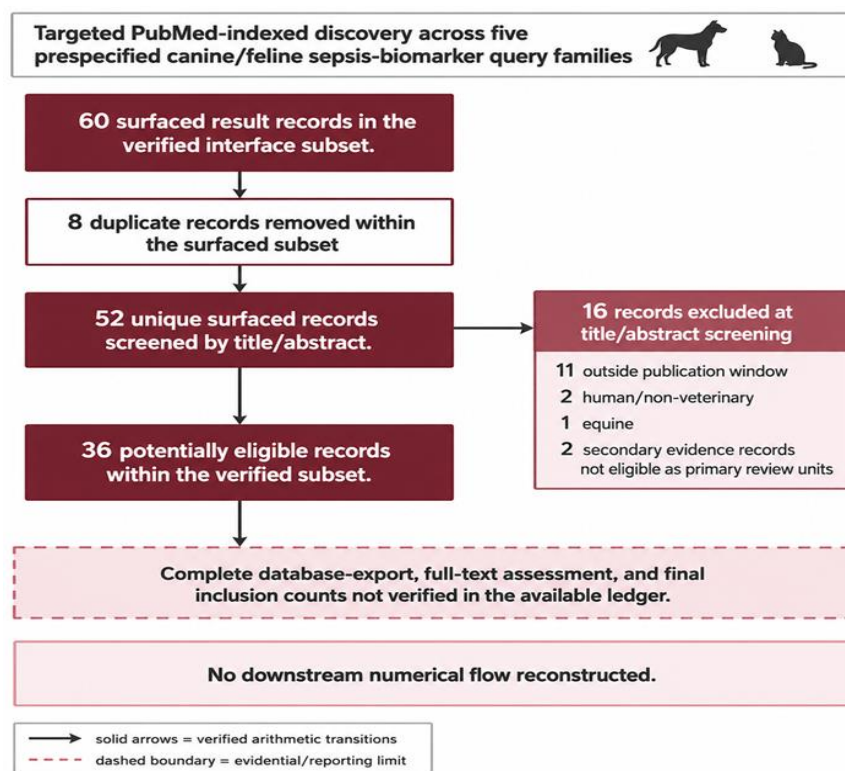


Figure 1. Verified boundary of the canine and feline sepsis-biomarker selection record

The flow reports only selection quantities documented in the available selection record: 60 surfaced records, 8 duplicates removed, 52 unique records screened, 16 records excluded, and 36 potentially eligible records. A terminal verification boundary prevents unverified full-text and inclusion quantities from being visually inferred. Solid connectors denote verified arithmetic transitions; the dashed terminal boundary denotes unavailable downstream accounting. Top-to-bottom review-selection flow showing 60 surfaced records, removal of 8 duplicates, screening of 52 unique records, exclusion of 16 records, and 36 potentially eligible records. The flow ends at a dashed verification boundary indicating that complete database-export, full-text exclusion, and final inclusion counts are not available and are not extrapolated.

Characteristics of included studies

The empirical evidence mapped in the review was predominantly canine and concentrated in referral, emergency, critical-care, and syndrome-specific populations. Feline evidence was comparatively sparse and more often involved small hospitalized cohorts, pyometra, or inflammatory profiling. Study purposes also differed: some evaluated rapid or point-of-care analytical performance, whereas others examined diagnosis, prognosis, disease severity, or longitudinal biomarker behavior.

Comparator spectrum varied markedly. Healthy animals were common controls, but clinically more demanding comparisons involved noninfectious SIRS or other critically ill animals. Infection phenotypes included pneumonia, pyometra, septic peritonitis, pyothorax, and broader sepsis cohorts. These differences limited direct cross-study ranking because the apparent value of a biomarker depended on whether it was being asked to separate disease from health, infection from sterile inflammation, survivors from nonsurvivors, or changing disease states over time.

Biomarker categories

Feline cytokine and chemokine profiling supports a biologically complex inflammatory response in sepsis and septic shock, but these mediators remain research-oriented rather than validated bedside diagnostics [23]. Local-fluid biomarkers represent a different category: paired biochemical measurements in canine septic peritonitis can assist source-infection recognition, but their target is the infected cavity rather than systemic sepsis itself [24].

Serial acute-phase measurements constitute a monitoring category. In dogs with pneumonia, septic peritonitis, or pyometra, longitudinal CRP, SAA, and other biomarker trajectories changed during treatment, suggesting potential value for assessing resolution [25]. Conversely, acute-phase and lipoprotein changes also occur across systemic inflammatory phenotypes, reinforcing that inflammatory responsiveness is not synonymous with infection specificity [26]. Thus, biomarker categories were most informative when classified by biological pathway and intended clinical function, rather than analyte family alone.

Diagnostic performance

Diagnostic performance was strongly dependent on the clinical question encoded by the study design. Local-fluid testing could provide useful discrimination for septic effusions, while systemic inflammatory biomarkers frequently performed a broader role as indicators of inflammation or illness severity. A marker that differs substantially between sick and healthy animals can still perform poorly when the clinically relevant comparator is noninfectious critical illness.

Thresholds were therefore not treated as transportable constants. Assay platform, specimen type, infection source, case definition, treatment exposure, and sampling time could each change apparent performance. Similarly, receiver-operating-characteristic results from narrowly selected disease phenotypes were interpreted as study-specific evidence rather than universal canine or feline sepsis thresholds. The evidence supports clinically bounded diagnostic interpretation; it does not support a single cross-platform biomarker cutoff or a universal ranking of point-of-care tests.

Prognostic and clinical utility

Prognostic information was distinguished from diagnostic discrimination. Biomarkers associated with organ dysfunction, septic shock, mortality, or longitudinal improvement may contribute clinically after sepsis is suspected or established, yet outcome association does not prove that the same marker can identify infection among undifferentiated patients. Serial trajectories can similarly aid reassessment without validating an initial diagnostic threshold.

Clinical utility was interpreted even more conservatively. Rapid turnaround, assay feasibility, and biological responsiveness are necessary attributes for some point-of-care applications, but they do not establish that biomarker-guided management improves survival, antimicrobial stewardship, length of hospitalization, or welfare. Evidence of actual decision impact was limited. Consequently, the review distinguishes availability of a result, clinical interpretability of that result, and demonstrated benefit from acting on it as progressively stronger and currently noninterchangeable claims.

Sources of heterogeneity

Clinical phenotype was a major source of heterogeneity. Evaluation of neutrophil-to-lymphocyte ratio in canine bacterial pneumonia illustrates that performance in one infection phenotype cannot be assumed to transport to all septic syndromes [27]. In dogs and cats with septic peritonitis or pyothorax, bilirubin abnormalities additionally illustrate how species, infection site, organ dysfunction, and sampling interval can influence prognostic interpretation [28].

Comparator selection also altered apparent accuracy. Pyometra-versus-healthy designs can yield strong separation while providing limited evidence about discrimination from noninfectious inflammatory mimics [29]. Reference-standard heterogeneity compounds this problem because microbiologic findings differ by source, sampling procedure, and organism distribution even within canine pyometra [30]. **Table 3** presents how five sources of heterogeneity alter comparability and transfer of veterinary sepsis-biomarker performance.

Table 3. How five sources of heterogeneity alter comparability and transfer of veterinary sepsis-biomarker performance

Heterogeneity question	Infection phenotype	Comparator spectrum	Species	Sampling time/treatment	Reference standard
Study contrast	Pneumonia, pyometra, cavitary sepsis	Healthy versus critically ill mimics	Canine-dominant versus sparse feline evidence	Baseline versus serial/post-treatment	Culture, local adjudication, SIRS-based classification
Inference distorted	Disease-specific performance generalized	Inflated apparent discrimination	Cross-species transfer	State change attributed to infection alone	Different target conditions treated as equivalent
Clinical consequence	Misapplied threshold	False confidence in specificity	Unsupported feline inference	Misread trajectory	Inconsistent diagnostic labels
Transfer limit	Source-specific spectrum	Spectrum bias	Species biology/assay behavior	Temporal and therapeutic context	Verification/incorporation bias
How comparability can be tested	Revalidate by phenotype	Test against clinical mimics	Obtain feline validation	Interpret sequentially	Reconcile source evidence

Risk-of-bias findings

Recurring risk-of-bias concerns arose from selected referral populations, small cohorts, disease-versus-healthy comparisons, inconsistent target-condition definitions, and reference standards that could share components with the index classification. Studies comparing septic animals only with healthy controls were particularly vulnerable to spectrum effects because the bedside diagnostic challenge is usually separation of infection from other inflammatory critical illnesses.

Other concerns were assay-specific. Analytical interference, variable sampling time, prior treatment, and lack of prespecified or externally validated thresholds could influence observed performance. Feline evidence additionally carried substantial imprecision because of small samples and limited replication. These limitations did not nullify the reported biological signals; rather, they narrowed the claims those signals could support. Risk-of-bias judgments were consequently used to qualify interpretation rather than to create a composite quality score or exclude potentially informative negative findings.

The central finding is that “point-of-care biomarker for sepsis” describes several scientifically different propositions. Some assays demonstrate that an analyte can be measured rapidly and reproducibly. Others detect

systemic inflammation, differentiate selected infection phenotypes, track biological severity, or change during recovery. These achievements are clinically relevant, but none automatically proves the others.

The evidence therefore does not support a single best biomarker across dogs and cats. Instead, it supports a function-specific interpretation in which measurement credibility precedes disease discrimination, and disease discrimination precedes claims of actionable clinical utility. This proposed hierarchy is intentionally nonnumeric. It also preserves discordance: a strong inflammatory signal with weak infection confirmation should not be averaged into a synthetic “sepsis score.” Such discordance is itself clinically informative and should prompt reassessment of infection source, comparator diagnoses, sampling time, and assay credibility.

Interpretation of diagnostic performance

Simple and readily available markers require the same interpretive discipline as novel assays. Neutrophil-to-lymphocyte and platelet-to-lymphocyte ratios in critically ill dogs showed limited or inconsistent prognostic relationships and should not be treated as sepsis-specific indicators [31]. A three-group study of thiamine status further demonstrates that a biomarker can distinguish critical illness from health without meaningfully distinguishing septic from nonseptic critical illness [32].

Analytical validation represents another boundary. A whole-blood point-of-care coagulation analyzer can be assessed for agreement, precision, and interference without thereby demonstrating sepsis discrimination [33]. Likewise, a quantitative fluorescent canine CRP assay may reliably measure CRP while leaving infection-specific clinical validity unanswered [34]. Diagnostic performance should therefore be interpreted as a chain of linked but separate questions: Was the analyte measured credibly? Did it distinguish the relevant target condition? Would that distinction change a clinically defensible decision?

Clinical implications

For bedside practice, a positive point-of-care biomarker should be interpreted alongside the suspected infection source, physiological state, organ dysfunction, assay characteristics, and temporal trajectory. A high inflammatory marker in an unstable animal may appropriately escalate attention while still leaving infectious causation uncertain. Conversely, a modest biomarker value should not override strong source evidence when assay timing or sensitivity is uncertain.

A proposed clinical interpretation architecture follows from the synthesis: first establish analytical credibility; next define what biological state the marker reflects; then ask whether the comparator evidence supports infection attribution; finally determine whether the result changes an existing diagnostic, monitoring, or treatment decision. This architecture is not a validated scoring system. Its purpose is to prevent analytical performance, inflammatory sensitivity, prognostic association, and clinical effectiveness from being collapsed into a single claim of “sepsis accuracy.”

Evidence gaps

Prospective studies should preferentially enroll diagnostically difficult animals rather than rely predominantly on healthy controls. Clinically useful validation requires noninfectious inflammatory comparators, predefined target conditions, independently interpretable reference standards, prespecified thresholds, transparent handling of culture-negative infection, and adequate representation of major infection phenotypes.

Feline evidence requires particular expansion. Future work should also distinguish analytical validation from diagnostic validation and diagnostic validation from decision-impact evaluation. Serial paired testing could clarify whether biomarker trajectories add information beyond routine clinical reassessment. External validation across hospitals and assay platforms is necessary before thresholds are transported between populations. Ultimately, pragmatic studies should test whether access to a rapid biomarker changes decisions and patient-centered outcomes relative to otherwise comparable care; such studies are needed because current associations and assay-performance data cannot establish clinical benefit by themselves.

Limitations

Early bedside studies of circulating cell-free DNA and nucleosomes illustrate both the potential of rapid prognostic biomarkers and the limitations imposed by small cohorts and restricted external validation [35]. Canine point-of-care CRP validation similarly demonstrates that establishing analytical performance does not establish

sepsis-specific diagnostic validity [36]. Feline point-of-care SAA assays have undergone analytical evaluation, but this evidence cannot be converted into clinical sepsis accuracy without target-condition studies [37]. Emerging feline evidence remains particularly fragile. Testican-1 findings in pyometra were derived from a very small SIRS-defined cohort and therefore require independent replication before broad interpretation [38]. The review is additionally constrained by marked heterogeneity in phenotype, comparator spectrum, reference standards, sampling time, and assay platforms. These limitations preclude defensible universal thresholds and restrict direct quantitative comparison across the evidence base.

Conclusion

Point-of-care biomarkers can contribute to early veterinary recognition of inflammation, infection-related states, prognosis, and serial reassessment, but these functions are not interchangeable. Current evidence is stronger for selected analytical and biological signals than for universal sepsis-specific discrimination or demonstrated decision benefit. Clinicians should interpret results within species, infection source, comparator spectrum, timing, reference-standard credibility, and assay context. The proposed function-specific interpretation separates measurement, inflammation detection, infection attribution, prognosis, and actionability, while emphasizing that discordant findings require reassessment rather than automatic numerical integration.

Acknowledgments: None

Conflict of Interest: None

Financial Support: None

Ethics Statement: None

References

1. Cortellini S, DeClue AE, Giunti M, Goggs R, Hopper K, Menard JM, et al. Defining sepsis in small animals. *J Vet Emerg Crit Care (San Antonio)*. 2024;34(2):97-109. doi:10.1111/vec.13359
2. Spillane AM, Haraschak JL, Gephard SE, Nerderman BE, Fick ME, Reinhart JM. Evaluating the clinical utility of the systemic inflammatory response syndrome criteria in dogs and cats presenting to an emergency department. *J Vet Emerg Crit Care (San Antonio)*. 2023;33(3):315-26. doi:10.1111/vec.13293
3. Ciuffoli E, Troia R, Bulgarelli C, Pontiero A, Buzzurra F, Giunti M. New-onset organ dysfunction as a screening tool for the identification of sepsis and outcome prediction in dogs with systemic inflammation. *Front Vet Sci*. 2024;11:1369533. doi:10.3389/fvets.2024.1369533
4. Hagman R, Klemming C, Bengtsdotter E, Södersten F, Wang L, Wernersson S. KC-like chemokine as a biomarker of sepsis in dogs with pyometra. *BMC Vet Res*. 2024;20(1):411. doi:10.1186/s12917-024-04271-w
5. Page MJ, McKenzie JE, Bossuyt PM, Boutron I, Hoffmann TC, Mulrow CD, et al. The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*. 2021;372:n71. doi:10.1136/bmj.n71
6. McInnes MDF, Moher D, Thombs BD, McGrath TA, Bossuyt PM, PRISMA-DTA Group. Preferred reporting items for a systematic review and meta-analysis of diagnostic test accuracy studies: The PRISMA-DTA Statement. *JAMA*. 2018;319(4):388-96. doi:10.1001/jama.2017.19163
7. Rethlefsen ML, Kirtley S, Waffenschmidt S, Ayala AP, Moher D, Page MJ, et al. PRISMA-S: An extension to the PRISMA Statement for reporting literature searches in systematic reviews. *Syst Rev*. 2021;10(1):39. doi:10.1186/s13643-020-01542-z
8. Rompf J, Lutz B, Adamik KN, Marti E, Mirkovitch J, Peters LM, et al. Plasma procalcitonin and C-reactive protein concentrations in dogs with bacterial sepsis and non-infectious systemic inflammatory response syndrome. *Front Vet Sci*. 2025;12:1609020. doi:10.3389/fvets.2025.1609020
9. Vezzi NK, Lane SL, Thevelein BA, Brainard BM. Performance evaluation of a Test&Treat rapid detection kit for the diagnosis of septic effusions in dogs and cats. *J Vet Emerg Crit Care (San Antonio)*. 2024;34(2):110-5. doi:10.1111/vec.13372

10. Allen AA, DeStefano IM, Rozanski EA. Utility of serum amyloid A-to-albumin ratio in hospitalized cats. *J Vet Emerg Crit Care (San Antonio)*. 2025;35(5):495-501. doi:10.1111/vec.70047
11. Rethlefsen ML, Brigham TJ, Price C, Moher D, Bouter LM, Kirkham JJ, et al. Systematic review search strategies are poorly reported and not reproducible: A cross-sectional metaresearch study. *J Clin Epidemiol*. 2024;166:111229. doi:10.1016/j.jclinepi.2023.111229
12. Lagisz M, Yang Y, Young S, Nakagawa S. A practical guide to evaluating sensitivity of literature search strings for systematic reviews using relative recall. *Res Synth Methods*. 2025;16(1):1-14. doi:10.1017/rsm.2024.6
13. Escobar Liquitay CM, Garegnani L, Garrote V, Solà I, Franco JVA. Search strategies (filters) to identify systematic reviews in MEDLINE and Embase. *Cochrane Database Syst Rev*. 2023;9(9):MR000054. doi:10.1002/14651858.MR000054.pub2
14. Pawliuk C, Cheng S, Zheng A, Siden HH. Librarian involvement in systematic reviews was associated with higher quality of reported search methods: A cross-sectional survey. *J Clin Epidemiol*. 2024;166:111237. doi:10.1016/j.jclinepi.2023.111237
15. Yang B, Mallett S, Takwoingi Y, Davenport CF, Hyde CJ, Whiting PF, et al. QUADAS-C: A tool for assessing risk of bias in comparative diagnostic accuracy studies. *Ann Intern Med*. 2021;174(11):1592-9. doi:10.7326/M21-2234
16. Schünemann HJ, Mustafa RA, Brozek J, Steingart KR, Leeftang M, Murad MH, et al. GRADE guidelines: 21 part 1. Study design, risk of bias, and indirectness in rating the certainty across a body of evidence for test accuracy. *J Clin Epidemiol*. 2020;122:129-41. doi:10.1016/j.jclinepi.2019.12.020
17. Schünemann HJ, Mustafa RA, Brozek J, Steingart KR, Leeftang M, Murad MH, et al. GRADE guidelines: 21 part 2. Test accuracy: Inconsistency, imprecision, publication bias, and other domains for rating the certainty of evidence and presenting it in evidence profiles and summary of findings tables. *J Clin Epidemiol*. 2020;122:142-52. doi:10.1016/j.jclinepi.2019.12.021
18. Foroutan F, Guyatt G, Zuk V, Vandvik PO, Alba AC, Mustafa R, et al. GRADE guidelines 28: Use of GRADE for the assessment of evidence about prognostic factors: Rating certainty in identification of groups of patients with different absolute risks. *J Clin Epidemiol*. 2020;121:62-70. doi:10.1016/j.jclinepi.2019.12.023
19. Bulgarelli C, Ciuffoli E, Troia R, Goggs R, Dondi F, Giunti M. Apolipoprotein A1 and serum amyloid A in dogs with sepsis and septic shock. *Front Vet Sci*. 2023;10:1098322. doi:10.3389/fvets.2023.1098322
20. Gaudette S, Smart L, Woodward AP, Sharp CR, Hughes D, Bailey SR, et al. Biomarkers of endothelial activation and inflammation in dogs with organ dysfunction secondary to sepsis. *Front Vet Sci*. 2023;10:1127099. doi:10.3389/fvets.2023.1127099
21. Hill E, Zhu Y, Brooks MB, Goggs R. Dogs with sepsis are more hypercoagulable and have higher fibrinolysis inhibitor activities than dogs with non-septic systemic inflammation. *Front Vet Sci*. 2025;12:1559994. doi:10.3389/fvets.2025.1559994
22. Rigot M, Bersenas AM, Bateman SW, Blois SL, Monteith G, Wood RD. Serum hyaluronic acid in critically ill dogs and influence of intravenous fluid therapy. *PLoS One*. 2025;20(6):e0325809. doi:10.1371/journal.pone.0325809
23. Troia R, Mascalcioni G, Agnoli C, Lalonde-Paul D, Giunti M, Goggs R. Cytokine and chemokine profiling in cats with sepsis and septic shock. *Front Vet Sci*. 2020;7:305. doi:10.3389/fvets.2020.00305
24. Martiny P, Goggs R. Biomarker guided diagnosis of septic peritonitis in dogs. *Front Vet Sci*. 2019;6:208. doi:10.3389/fvets.2019.00208
25. Goggs R, Robbins SN, LaLonde-Paul DM, Menard JM. Serial analysis of blood biomarker concentrations in dogs with pneumonia, septic peritonitis, and pyometra. *J Vet Intern Med*. 2022;36(2):549-64. doi:10.1111/jvim.16374
26. Behling-Kelly E, Haak CE, Carney P, Waffle J, Eaton K, Goggs R. Acute phase protein response and changes in lipoprotein particle size in dogs with systemic inflammatory response syndrome. *J Vet Intern Med*. 2022;36(3):993-1004. doi:10.1111/jvim.16420
27. Stark MA, Simons PM, Lyons BM. Evaluation of neutrophil-to-lymphocyte ratio as a diagnostic marker in dogs with bacterial pneumonia. *J Vet Emerg Crit Care (San Antonio)*. 2025;35(3):269-73. doi:10.1111/vec.13475

28. Benham-Crosswell FJ, Payne NOA, Hjalmarsson LJ, Humm KR. Retrospective evaluation of hyperbilirubinemia in cats and dogs with septic peritonitis or pyothorax. *J Vet Emerg Crit Care (San Antonio)*. 2025;35(4):408-14. doi:10.1111/vec.70007
29. Çolakoğlu HE, Karaca E. Evaluation of SIRS and qSOFA in the diagnosis of sepsis in dogs with pyometra. *Theriogenology*. 2025;241:117420. doi:10.1016/j.theriogenology.2025.117420
30. Ylhäinen A, Mölsä S, Thomson K, Laitinen-Vapaavuori O, Rantala M, Grönthal T. Bacteria associated with canine pyometra and concurrent bacteriuria: A prospective study. *Vet Microbiol*. 2025;301:110362. doi:10.1016/j.vetmic.2024.110362
31. Dourmashkin LH, Lyons BM, Hess RS, Walsh K, Silverstein DC. Evaluation of the neutrophil-to-lymphocyte and platelet-to-lymphocyte ratios in critically ill dogs. *J Vet Emerg Crit Care (San Antonio)*. 2023;33(1):52-8. doi:10.1111/vec.13269
32. Berlin N, Pfaff A, Rozanski EA, Chalifoux NV, Hess RS, Donnino MW, et al. Establishment of a reference interval for thiamine concentrations in healthy dogs and evaluation of the prevalence of absolute thiamine deficiency in critically ill dogs with and without sepsis using high-performance liquid chromatography. *J Vet Emerg Crit Care (San Antonio)*. 2024;34(1):49-56. doi:10.1111/vec.13341
33. Conroy EM, Lyons BM, Koenig A. Evaluation of a whole blood point-of-care coagulation analyzer in dogs. *J Vet Emerg Crit Care (San Antonio)*. 2024;34(5):446-54. doi:10.1111/vec.13416
34. Choi J, Yoo MJ, Jang YJ, Na B, Seo SK, Moon J, et al. Development and clinical evaluation of a quantitative fluorescent immunoassay for detecting canine CRP. *Int J Vet Sci Med*. 2023;11(1):87-93. doi:10.1080/23144599.2023.2247250
35. Letendre JA, Goggs R. Determining prognosis in canine sepsis by bedside measurement of cell-free DNA and nucleosomes. *J Vet Emerg Crit Care (San Antonio)*. 2018;28(6):503-11. doi:10.1111/vec.12773
36. Hindenberg S, Keßler M, Zielinsky S, Langenstein J, Moritz A, Bauer N. Evaluation of a novel quantitative canine species-specific point-of-care assay for C-reactive protein. *BMC Vet Res*. 2018;14(1):99. doi:10.1186/s12917-018-1415-2
37. Escribano D, Ortín Bustillo A, Pardo Marín L, Navarro Rabasco A, Ruiz Herrera P, Cerón JJ, et al. Analytical validation of two point-of-care assays for serum amyloid A measurements in cats. *Animals (Basel)*. 2021;11(9):2518. doi:10.3390/ani11092518
38. Yaprakci Ö, Akkuş T. Evaluation of Testican-1 as a sepsis biomarker in pyometra cats. *Vet Med Sci*. 2025;11(5):e70514. doi:10.1002/vms3.70514