

From Manual to Machine: Leveraging Best-of-Breed Approaches in Clinical Data Management

Joseph Lebers, ThoughtSphere, Lodi NJ, USA
Christina Dinger, ThoughtSphere, Kansas City KS, USA

ABSTRACT

With the explosion of AI technology, CDM solutions have flooded the market, offering novel AI approaches for conducting data reviews and shortening delivery timelines. Many of these “black box” solutions are experimental at best and require massive change management to implement successfully. This puts drug developers in a tough position—deliver cutting-edge innovation, but with a high risk of failure.

AI-tech approaches that blend ML and AI capabilities into established and effective human-driven CDM processes offer great promise. By seamlessly marrying Human and AI capital, CDM teams gain efficiency, reduce the change management burden, and advance clinical research with minimal risk.

Whether you're evaluating a vendor solution, or devising a change management strategy, join us as we discuss technology adoption that takes a best-of-breed approach. We will explore AI/ML use cases that optimize traditional CDM tasks while eliminating risks of data hallucinations, bias, and the “black box” paradigm.

INTRODUCTION

Through an enhanced understanding of disease mechanisms and greater access to standardized clinical data, pharmaceutical sponsors are increasingly adopting machine learning (ML) to develop novel drugs for patients.¹⁻² With access to new open-source LLM's (large language models), breakthroughs in pure data science, and even coding assistance offered by LLM's, we only expect the availability of ML solutions to increase.

The reduced barriers to entry for implementing ML models offer a vision of unprecedented precision treatments and lower clinical trial costs. This high accessibility must be met with rigorous evaluation to ensure alignment with sponsors' research and business objectives.⁴ Traditionally, the healthcare field is risk-averse – requiring inventors and implementers to thoroughly test technology while managing the quality of findings and potential risks to patients.⁵ While the FDA has provided some guiding principles on ML in medical devices,⁶ a comprehensive regulatory standard like Good Manufacturing Practices (GMP) for ML models is not yet in place. This should urge sponsors and CRO's to take greater responsibility and oversight in evaluating ML solutions they seek to adopt.

This paper will explore the potential issues that arise when vendor software does not align with sponsor drug development objectives. This serves as guidance to pharma sponsors so they can evaluate ML solutions by vendors, and guidance to vendors so they can build solutions aligned with sponsor business objectives. Providing examples based on classification models – which are common in Clinical Data Management (CDM) functions – key indicators and strategies are explored to assist professionals in

delivering true value in data review. This paper aims to equip sponsors and software vendors with approaches to ensuring that ML applications are truly impactful to patient safety and business objectives.

MACHINE PROBLEMS, OR GAPS IN HUMAN COMMUNICATION?

With greater access to digital data sources and revolutionary breakthroughs in data science, ML is increasingly used in the healthcare sector and clinical trials to address issues of patient recruitment,⁷ personalized gene therapies,⁸ and the detection of data collection errors in trial data.

With a less rigorous validation process expected in comparison to traditional biophysical and pharmacological technologies, this introduces risk that the technology sector will pursue data science objectives that do not result in increased patient safety and efficiency of clinical trials. For example, non-experts in ML may view model accuracy as a sole measure of performance leading to ineffective models being deployed in practice.⁹ Just as following the GMP standard prevents “contamination, mix-ups, deviations, failures, and errors,” a regulatory standard for ML solutions can assure that models are reliable and impactful to clinical trials.¹⁰

It is in vendors’ best interests to ensure the outputs from applied models truly provide value to sponsors and CRO’s – and this can be supported by having a collaborative approach to evaluating model performance and impact.¹¹ Investing in ML development comes with monetary and resource costs, so vendors are motivated to build effective products to make a return on those investments. The reduced barriers to entry for implementing machine learning models sets the stage for vendors looking to make ROI on impressive ML feats that do not truly answer sponsor research questions. Technology vendors stand to benefit from working closely with sponsors, employing experienced domain experts, and evaluating the true impact of the model during the exploration process with those sponsors.¹¹

Investment in AI in the clinical trials industry is projected by some to reach \$20.16 billion by 2033.¹² Expecting that sponsors and CROs will invest more in machine learning solutions, it serves them well to critically understand the performance and impact of any model applied to clinical trials. Since a key feature of ML models is that they can improve in performance based on user feedback, it is optimal to adopt a ML solution for long-term use. Failed implementations not only cause a monetary loss, but a loss in knowledge for the organization.

In Clinical Data Management (CDM), misaligned prediction targets and model performance can lead to missed safety signals, incomplete data reconciliation, and even lost patient lives. A deeper understanding of business problems and ML hypotheses ensures patient safety while generating ROI for both the sponsor and vendor.

DETERMINING KEY MEASURES IS SERIOUS BUSINESS

Like with any technology that impacts patient lives and pharmaceutical drugs, key decisionmakers at sponsors and CRO’s must be vigilant in deeply understanding their current processes, workforce, and the outcomes of the new technology, in this case, ML prediction models. Analyzing multiple model performance metrics and applying those metrics to the current business not only ensures that the safety of drugs is carefully evaluated, but also that implementing the ML model is an improvement over the current state and a worthy ROI for the organization. In the following section is an exploration leveraging a binary classification model as an example. A classification model will predict the status of data as *True* for a target condition (e.g., patients likely to drop out of the study, patients likely to experience a life-threatening adverse event, etc.) or *False* for a target condition.

MOVING BEYOND ACCURACY AS THE END-ALL MEASURE

While there are a plethora of methods to evaluate the performance of ML models, accuracy is often marketed to business users as an indicator of high performance. For “low-stakes” applications of ML such as recommending online products, movies, and potential dates, a judgement on accuracy may suffice for these use cases. However, in the healthcare space and especially in central data monitoring for clinical trials, detecting the target condition can be vital to patient health or proper analysis of drug safety and efficacy. Especially when applying models on data problems with class imbalances (when true signals are a “needle-in-the-haystack”), accuracy may not reflect the impact in practice.¹³ This is critical because missed safety signals can jeopardize the life of patients, the clinical trial itself, and subsequently the massive investment in the trial by the sponsor.

Understanding the prediction rate among *True Positives*, *True Negatives*, *False Positives*, and *False Negatives* for a classification model gives a more holistic view of what to expect from model insights compared to accuracy score, which is a summarized metric.

The Confusion Matrix is an oft-used evaluation of model performance. There are more robust frameworks for evaluating the performance of classification models such as Area Under Curve (AUC) and Precision-Recall Curves¹⁴ but for the scope of this paper we analyze the Confusion Matrix.

	Prediction is Correct	Prediction is Incorrect
Condition is Detected	True Positive (TP)	False Positive (FP)
Condition is Absent	True Negative (TN)	False Negative (FN)

Table 1: Standard Confusion Matrix for evaluating performance of a binary classification model.

Measure	Formula	Definition
Accuracy	$\frac{TP + TN}{TP + TN + FP + FN}$	<i>Out of all the predictions, how many were correctly labeled as True and False?</i>
Precision	$\frac{TP}{TP + FP}$	<i>When the model predicts a True case, what is the rate of being correct?</i>
Recall / Sensitivity	$\frac{TP}{TP + FN}$	<i>How proficient is the model at identifying all True cases?</i>
F-Measure (F1 Score)	$\frac{2TP}{2TP + FP + FN}$	<i>A harmonic mean which balances Precision and Recall.</i>

Table 2: Common model performance metrics with representative formula and English explanation.

	High Performance	Low Performance
Accuracy	Model is proficient at differentiating between True and False cases.	Deficient at discerning between True and False cases.
Precision	Positive predictions by the model are reliable.	Positive predictions by the model are unreliable.
Recall / Sensitivity	The model is proficient at identifying most True cases, with minimal False cases predicted as True.	The model is deficient at predicting most True cases. AND/OR The model is sufficient at predicting most True cases, but there are mass amounts of False cases predicted as True.

F1 Score	The model balances proficiency in predicting most True cases (high Recall), with proficiency in differentiating between True cases and False cases (high Precision).	The model is deficient at detecting the target condition (low Recall), AND/OR The model is deficient at differentiating between True Positives and False Positives (low Precision).
-----------------	--	--

Table 3: Interpretation of high and low performance across model performance metrics.

WHY ACCURACY DOES NOT GIVE THE FULL PICTURE

Sponsors and CRO's ideally want to leverage ML models that are 100% accurate – meaning they perfectly predict True and False cases in the input clinical data. However, due to data availability and the complexity of the data learned by the model, this may be impractical. Extremely high accuracy in classification models can even indicate the model is overfitted to (e.g. the model is “memorizing” rather than learning) the training dataset.¹⁵ Accuracy can also be misleading when there is an imbalance of classes in the data – meaning different expected counts of Positive vs Negative cases (e.g., a model can consistently predict which patients do not have a rare disease, but fails to predict the few patients who do).¹⁶ Considering ML models are often developed to detect serious risks to patient health and drug safety, the ability for the model to predict all the True conditions (measured as Recall or Sensitivity) is paramount to healthcare professionals.

DETECTING THE TARGET CONDITION IS KEY TO THE IMPLEMENTATION

Sensitivity scores can be maximized by a vendor's data scientists to catch all True cases of the prediction target – such as discrepant datapoints or patients at risk of developing cancer. Clinical data scientists are driven to maximize sensitivity scores to eliminate these points of risk in trials and healthcare research, because the industry is dedicated to patient care. If the data management team were to follow predictions by a low-sensitivity model without supplemental traditional review processes, they will likely miss the high-risk observations they seek to address.

ML engineers can leverage strategies to tweak model weightages and architecture to influence the distribution of True, False, Positive, and Negative predictions generated. In pursuit of maximizing the sensitivity of a model to ensure the prediction of True Positives, one method data scientists leverage is to “cast a wider net” to ensure that most, if not all, True Positives are predicted. But by making this adjustment the model may become more likely to generate False Positives, that are not true risks. If you have ever expanded a threshold window in an edit check, for example when detecting lab values that change from the patient's baseline, there would be more *potentially anomalous* lab values for a Medical Monitor to review, but the proportion of *actual non-anomalous* data that is reviewed, increases (Sensitivity increases; Precision decreases). In the case of CDM, the data management team can feel confident in their review since they covered more of the data, but there is a cost in efficiency when spending time and resources on non-anomalous signals.

FALSE POSITIVES CAN IMPACT EFFICIENCY

Even as part of traditional data management and review processes, clinical observations are reviewed multiple times and by multiple study team personnel to ensure the safety of interventions. While predicting True Positives is paramount for the model, data scientists must be mindful of introducing False Positives in the predictions. In a case like central data review where the data management team is directed to review *all data* and not just those that appear anomalous, reviewing False Positive predictions can be a strategy to supplement the data review process – but it may not be the most efficient strategy.

Due to the importance of detecting Positive cases in healthcare and clinical trials, there is a *penalty difference* in classification errors. And depending on the data management team itself, there may be a penalty difference in the time, effort, and cost it takes to identify and review True Positives versus False Positives.¹⁷ Some factors that can influence the time and effort it takes to review anomalous data are:

- Whether or not the ML recommendation system provides all contextual info needed to review data quickly.
- In what scenarios there is a hand-off from data manager to medical monitor, or to a data management lead.
- The traditional data management and review processes put in place (listing reviews, edit checks, clean patient tracker, patient profile reviews, etc.) to supplement the ML recommendation system.

A close relationship between sponsors and technology vendors during the exploration phase can ensure that True Positives are detected by the model, and that the review process for study staff offers improved efficiency over the current traditional process.

PROJECTING ADDED VALUE OF THE MODEL

Implementing a ML recommendation system can, and is expected to, alter the data management and review process for clinical study teams. For any business decision maker, understanding the efficiency of the data management function and its processes is foundational when evaluating ML solutions. During implementation, study team members will not know that a Positive prediction is a False Positive until they review it. So, based on model performance, there is a risk study team members may review hundreds or thousands of predicted False Positives before reaching an important True Positive. This is critical because based on the cases being predicted, patients can experience harm and the sponsor risks failure of their clinical study.

As ML implementations stand to affect business processes and human capital in the CDM space, business stakeholders should deeply understand the efficiency and workload of the central data review team. It is essential to benchmark how quickly, accurately, and costly it is for the specific data management team to review data, so the sponsor can determine if the ML implementation makes a positive impact to efficiency.¹⁷

COMPARE ROCII TO ROI

In determining if a new ML implementation offers true efficiency improvements over the traditional process, the business must first determine the critical success factors they are aiming to improve – such as cost, time-to-query, and resources needed to complete data review. An *ROCII* (Return on competitive intelligence investment) can be applied to business processes to estimate the value of insights, including ML insights from business intelligence tools and clinical decision-making systems.¹⁸ Measuring the ROCII and ROI in planning the adoption of a new system is flexible and depends on the business priorities of the organization and specific use case.

Examples:

- Positive Cases Detected, per Dollar Spent
- Time to Raise a Query, per Dollar Spent
- Data Manager Assigned to a Study, per Dollar Spent
- Critical Safety Issue Detected, per Record Reviewed
- At-Risk Patients Discontinued Before a Life-Threatening Event, per Dollar Spent

When the ROCII of the new business intelligence tool surpasses the ROI of the original process, the sponsor can consider the business intelligence tool an improvement in reaching their business objectives.

To reliably estimate ROI, sponsors must deeply understand the capabilities and limitations of the review team, and costs associated with their manageable workload. Model performance metrics can be leveraged to estimate the ROCII. Leveraging realistic benchmarks of value provides more reliable estimates of ROI and ROCII. By benchmarking the data management team in their current process, and running a pilot of the model on real study data with the same team, components of ROCII and ROI can be estimated.

ESTIMATING RISKS AND EXTRAPOLATING COSTS

There are many ways to estimate the risk in manual processes and business intelligence systems including risk matrix, decision tree, and expected value analyses. In the following hypothetical case study, *expected value* is calculated for simplicity and demonstration. Decision theory can be applied to the before-and-after of implementing a new ML recommendation system to understand the risk to business objectives through expected value.¹⁹ There are three basic elements that play into estimating the risk to business objectives, when dealing with a decision problem: *decisions to be made*, *uncertain events*, and *value of outcomes*.²⁰

Business Objective: Reduce the <i>time</i> taken to review all clinical data and raise necessary EDC queries to the site for a clinical study, by fully migrating data review activity to a ML recommendation system.							
Model Prediction Target: (Binary Classification) Data records needing a query marked as TRUE Data records not needing a query marked as FALSE							
Risk: Reviewers may spend time reviewing records which do not result in a query.							
Estimated Workload: 3,600,000 clinical data records							
Current Traditional Process				New ML Process			
Activity	Outcome	Uncertainty	Value (time)	Activity	Outcome	Uncertainty	Value (time)
Find issue in Listing Review	Query	-0.20	8	Find issue in Positive Predictions	Query	-0.70 (Precision)	4
	No Query	0.80	10		No Query	0.30	5
Find issue in risk review	Query	-0.50	5	Find issue in Negative Predictions	Query	-0.20	6
	No Query	0.50	7		No Query	0.80	7
Find issue in CPT	Query	-0.20	9				
	No Query	0.80	9				
Find issue in data check output	Query	-0.90	4				
	No Query	0.10	6				

Table 4: A hypothetical case study to demonstrate how sponsors and vendors can evaluate value added by ML implementation.

This hypothetical business problem, along with uncertainty and value were synthesized for the sake of explanation and simplicity. *Uncertainty* in the Current Traditional Process reflects the likelihood that a data record being reviewed will fit either a Positive or Negative outcome. *Uncertainty* was given a negative weightage ratio because the hypothetical DM team considers time spent on Positive cases to be worth the time reviewing. The *values* in the Current Traditional Process reflect the time needed to review data records found through the specific Activity, and finding the respective Outcome. The business needs to define the uncertainty and value in the context of their team and business objectives.

As there are different penalties (time taken to review a record and raise query) for classes, outcomes of Query (*Positive*) and No Query (*Negative*) are broken out separately. The outcomes of Query and No Query represent the *decision to be made* by the data management team. Since the Current Traditional Process deploys various review strategies, the differences in value and uncertainty (likeliness to occur) are broken out separately as well. Since the data management team seeks to reduce time taken to fully clean their data, the *value* is measured in the minutes it takes to identify a data issue, review it, and raise a query on it (if deemed Positive).

The risk contribution for each of these combinations of review activity and outcome, can be represented by the formula:

$$\text{Risk} = (\text{Uncertainty}) \times (\text{Value of Cases})$$

The risk contributions of each approach in the case study must be summed to realize the Total Risk.

$$\text{Risk}_{\text{Total}} = \sum P(X_i) \times X_i$$

For the Current Traditional Process:

$$\begin{aligned} \text{Risk}_{\text{Manual}} &= -0.20(8) + 0.80(10) + (-0.50)(5) + (0.50)(7) + (-0.20)(9) + (0.80)(9) + (-0.90)(4) \\ &\quad + (0.10)(6) \\ \text{Risk}_{\text{Manual}} &= 3.1 \end{aligned}$$

For the New ML Process:

$$\begin{aligned} \text{Risk}_{\text{ML}} &= (-0.70)(4) + (0.30)(5) + (-0.20)(6) + (0.80)(7) \\ \text{Risk}_{\text{ML}} &= 9.8 \end{aligned}$$

Final Comparison:

$$\text{Risk}_{\text{Manual}} < \text{Risk}_{\text{ML}}$$

In this hypothetical case study, there is more risk associated with the New ML Process. Since the measures are synthetic, this is not meant to show that ML implementations are less efficient – but demonstrates that determining gains in efficiency and ROI can be very complex and highly-specific to an organization. The sponsors can create granular models (e.g. a breakdown by team member experience) to realistically estimate the risk in these processes. As next steps, the sponsor may multiply these risk scores by the count of clinical records being reviewed in the scenario, then multiplied by the cost of employing data managers per minute to estimate total cost of reviewing all data for the study. In the New

ML Process, the sponsor will also need to factor in the costs of the yearly license of the solution (if applicable).

Leveraging Traditional Approaches with ML Approaches

The selection process for every ML solution will look different based on the use case addressed, aspects of the data management team, complexity of clinical data, and the software vendor's implementation. Sponsors and CROs may want to leverage the power of models to predict Positive cases while avoiding potential misdirection and inefficiencies caused by those models. Traditional review approaches can be leveraged to more efficiently review potential False Positives and Negative cases within the data. In the following subsections we touch on a few strategies common in the industry.

DATA CHECKS

Data checks are leveraged to identify data discrepancies that fit under consistent logical conditions, employing operators like '=', '>', '<', 'in', etc. between values in and across clinical datasets. Data checks often address clear discrepancies such as start dates occurring after end dates, and concomitant medications marked as ongoing when they have a specified end date. Within a clinical decision-making tool, such as ThoughtSphere's ClinDM, data checks can even be enabled to automatically raise queries in the EDC system when a discrepancy is detected. To aid study team members, vendors can design their business intelligence tool to display relevant contextual data about the patient so they can verify true discrepancies quickly and holistically. ML models use complex, layered statistical computations and transformations which can be "overkill" for these types of discrepancies.

RISK-BASED CLINICAL DATA MANAGEMENT

Employing a Risk-Based Clinical Data Management strategy can guide data managers to reviewing truly impactful signals in their clinical data. Calling on ICH E6 & ICH E8 guidelines, teams can devise a review strategy based on the most impactful risks to the study data quality. Critical-to-Quality (CtQ) factors should be identified and thoroughly analyzed during the planning phase of the study, involving data managers and biostatisticians.²¹ Not only does it increase the quality of the Statistical Analysis Plan, but it also allows the study team to conduct holistic risk assessments and track risks to quality throughout the trial. For example, in ThoughtSphere's platform CtQ's can be created and linked to review objects (data checks, listing reviews, risk reviews, clean patient tracker) to enable tracking of critical data and processes at every step of the review process. By combining a ML approach with risk-based assessments, the data most critical to success of the study and safety of patients can be reviewed on priority.

CONCLUSION

In conclusion, both clinical research organizations and the software vendors can achieve greater success by taking a collaborative and transparent approach to implementing ML models in business processes. Sponsors should come to the table having a deep understanding of their team's strengths and weaknesses, and a critical eye on model performance. Software vendors must be aware of the criticality in reviewing valuable model predictions, so they develop their models effectively from the start. Sponsors can leverage ML models to predict Positive cases and avoid performance issues, by deploying traditional review strategies in combination with an advanced RBCDM strategy.

REFERENCES

1. Shah P, Kendall F, Khozin S, et al. Artificial intelligence and machine learning in clinical development: a translational perspective. *NPJ Digit Med*. 2019;2:69. doi:10.1038/s41746-019-0148-3.
2. Weissler EH, Naumann T, Andersson T, et al. The role of machine learning in clinical research: transforming the future of evidence generation. *Trials*. 2021;22:537. doi:10.1186/s13063-021-05489-x
3. Hutson M. How AI is being used to accelerate clinical trials. *Nature*. 2024;627:S2-S5. doi:10.1038/d41586-024-00753-x.
4. Jullien M, Bogatu A, Unsworth H, Freitas A. Controlled LLM-based reasoning for clinical trial retrieval. *arXiv Preprint*. 2024. arXiv:2409.18998.
5. Clark D, Dean G, Bolton S, et al. Bench to bedside: The technology adoption pathway in healthcare. *Health Technol*. 2020;10:537-545. doi:10.1007/s12553-019-00370-z.
6. U.S. Food and Drug Administration. Artificial intelligence and machine learning in software as a medical device. *U.S. Food and Drug Administration*. May 11, 2021. Accessed February 08, 2025.
7. Cai T, Cai F, Dahal KP, et al. Improving the efficiency of clinical trial recruitment using an ensemble machine learning to assist with eligibility screening. *ACR Open Rheumatol*. 2021;3:593-600. doi:10.1002/acr2.11289.
8. Schwarzer A, et al. Predicting genotoxicity of viral vectors for stem cell gene therapy using gene expression-based machine learning. *Molecular Therapy*. 2021;29(12):3383-3397
9. Yang Q, Suh J, Chen N-C, Ramos G. Grounding interactive machine learning tool design in how non-experts actually build models. In: *Proceedings of the 2018 Designing Interactive Systems Conference (DIS '18)*. Association for Computing Machinery; 2018:573-584. doi:10.1145/3196709.3196729
10. U.S. Food and Drug Administration. Facts about current good manufacturing practice (cGMP). *U.S. Food and Drug Administration*. Published August 1, 2021. Accessed February 08, 2025. <https://www.fda.gov/drugs/pharmaceutical-quality-resources/facts-about-current-good-manufacturing-practice-cgmp>.
11. Takeuchi H, Yamamoto S. Business analysis method for constructing business–AI alignment model. *Procedia Comput Sci*. 2020;176:1312-1321. doi:10.1016/j.procs.2020.09.140.
12. The Pharma Letter. Industry spending on AI set to reach \$3 billion by 2025. *The Pharma Letter*. Published February 12, 2025. Accessed February 12, 2025. <https://www.thepharmaletter.com/pharmaceutical/industry-spending-on-ai-set-to-reach-3-billion-by-2025>.
13. Balde B. Why is accuracy misleading? *Medium*. Published December 15, 2024. Accessed February 12, 2025. <https://medium.com/@becaye-balde/why-is-accuracy-misleading-9465975fa429>.
14. GeeksforGeeks. Precision-Recall Curve in ML. *GeeksforGeeks*. Published September 1, 2024. Accessed February 12, 2025. <https://www.geeksforgeeks.org/precision-recall-curve-ml/>.
15. Valverde-Albacete FJ, Peláez-Moreno C. 100% classification accuracy considered harmful: the normalized information transfer factor explains the accuracy paradox. *PLoS One*. 2014;9(1):e84217. doi:10.1371/journal.pone.0084217.
16. Scaringi JA, McTaggart RA, Alvin MD, et al. Implementing an AI algorithm in the clinical setting: a case study for the accuracy paradox. *Eur Radiol*. 2024. doi:10.1007/s00330-024-11332-z.

17. Elkan C. The foundations of cost-sensitive learning. In: *Proceedings of the Seventeenth International Conference on Artificial Intelligence*. Seattle, WA; August 4-10, 2001:1.
18. Lönnqvist A, Pirttimäki V. The measurement of business intelligence. *Inf Syst Manag*. 2006;23(1):32-40. doi:10.1201/1078.10580530/45769.23.1.20061201/91770.4.
19. Walls MR, Thomas ME, Brady TF. Improving system maintenance decisions: a value of information framework. *Eng Econ*. 1999;44(2):151-167. doi:10.1080/00137919908967514.
20. Davison L. Measuring competitive intelligence effectiveness: insights from the advertising industry. *Comp Int Rev*. 2001;12:25-38. doi:10.1002/cir.1029.
21. Dinger C. Are you aiming for the bullseye or just the target? Developing an outcome-based RBQM strategy. *Clinical Leader*. Published June 15, 2023. Accessed February 13, 2025. <https://www.clinicalleader.com/doc/are-you-aiming-for-the-bullseye-or-just-the-target-developing-an-outcome-based-rbqm-strategy-0001>.

CONTACT INFORMATION

Your comments and questions are valued and encouraged. Contact the authors at:

Author Name: Joe Lebers

Company: ThoughtSphere

Address: 99 S Almaden Blvd, San Jose, CA 95113

Email: joe.lebers@thoughtsphere.com

Website: thoughtsphere.com

Author Name: Christina Dinger

Company: ThoughtSphere

Address: 99 S Almaden Blvd, San Jose, CA 95113

Email: christina.dinger@thoughtsphere.com

Website: thoughtsphere.com

Brand and product names are trademarks of their respective companies.