



FEDERAL HOUSING FINANCE AGENCY

ADVISORY BULLETIN

AB 2022-02: ARTIFICIAL INTELLIGENCE/MACHINE LEARNING RISK MANAGEMENT

Purpose

This advisory bulletin (AB) provides Federal Housing Finance Agency (FHFA) guidance to Fannie Mae and Freddie Mac (collectively, the Enterprises)¹ on managing risks associated with the use of artificial intelligence and machine learning (AI/ML). This AB is intended to highlight key risks inherent in the use of AI/ML that are applied across a variety of business and operational functions, and considerations for effectively managing these risks. FHFA recognizes that AI/ML is an evolving field and encourages the responsible innovation and use of AI/ML that is consistent with the safe and sound operations of the Enterprises.

Background

For purposes of this AB, artificial intelligence broadly refers to the development and application of computational tools and computer systems able to perform tasks normally requiring human intelligence, and machine learning is a sub-category of AI described as algorithms that optimize automatically through experience and with limited or no human intervention.² The combined term, AI/ML, encompasses the sub-categories of AI, such as computer vision and natural language processing, as well as the various methods used in ML, such as supervised learning, unsupervised learning, reinforcement learning, deep learning, and neural networks. AI/ML can be leveraged in models, applications, tools, and systems throughout its lifecycle. Generally, the AI/ML lifecycle includes stages addressing proof-of-concept, development, implementation and deployment, production use, and retirement.

¹ Common Securitization Solutions, LLC (CSS) is an “affiliate” of both Fannie Mae and Freddie Mac, as defined in the Federal Housing Enterprises Financial Safety and Soundness Act of 1992, as amended. 12 USC 4502(1).

² There are no industry-wide definitions for AI/ML, but for purposes of this AB, definitions from the Financial Stability Board are used. See Financial Stability Board, *Artificial Intelligence and Machine Learning in Financial Services* (November 2017).

The use of AI/ML presents benefits and risks as it increases the opportunity for decisions to be made and relied upon with significantly less human involvement. With increases in computing power, AI/ML can be used by the Enterprises to process vast datasets, identify complex relationships, and improve efficiencies and operations with reduced error and cost. However, AI/ML applications can also expose the Enterprises to financial, compliance, reputational, model, and other risks. For example, AI/ML algorithms developed using incomplete or unrepresentative data with unclear relationships between model inputs and outputs could exacerbate existing risks and result in poor or costly business decisions. As AI/ML continues to advance, the associated risks will also evolve—posing challenges to existing risk management practices. For instance, as AI/ML becomes more automated and integrated into business processes within and across business lines, the interconnected nature of the risks can introduce more complexity in risk management. Reliance on AI/ML without sufficient risk oversight and transparency can create heightened risks for the Enterprises.

FHFA’s Prudential Management and Operations Standards (PMOS), Appendix to 12 CFR Part 1236, sets forth general responsibilities of the board and senior management, as well as specific responsibilities for management and operations relating to ten enumerated standards, adopted as guidelines. Standard 1 (Internal Controls and Information Systems) and Standard 8 (Overall Risk Management Processes) highlight the need for the Enterprises to establish risk management practices that identify, assess, control, monitor, and report risk exposures, and the need to have appropriate risk management policies, standards, procedures, controls, and reporting systems in place. These guidelines are especially relevant to the Enterprises’ use and risk management of AI/ML.

Guidance

The Enterprise should incorporate the following guidance to manage the risks posed by the use of AI/ML, taking into consideration existing laws, regulations, and other FHFA supervision guidance. The sophistication of the AI/ML risk management activity should be proportionate to each Enterprise’s size, complexity, and risk profile. The Enterprise should leverage enterprise-wide risk management and control frameworks, including those used for model, data, technology, information security, third-party, and compliance risk management, to the extent practicable. These frameworks, however, may need to be enhanced and adapted with the considerations highlighted in this guidance to address the heightened risks that AI/ML can pose to business operations. Given the evolving nature of AI/ML, risk management should be flexible to accommodate changes in the adoption, development, implementation, and use of AI/ML at the Enterprise. The degree and scope of risk management and controls addressing AI/ML should be risk-based and commensurate with the extent and complexity of AI/ML development and use at the Enterprise, as well as the level of risk exposure. For example, high-risk AI/ML use cases—such as those that affect the Enterprise’s critical business functions, invoke compliance with laws

and regulations, or involve highly complex and opaque methods—warrant more robust risk management considerations than AI/ML uses that are low risk or transparent.

I. Governance

AI/ML tools and systems can support a range of functions across the Enterprise, such as customer engagement, risk analysis, credit decision-making, fraud detection, and information security. The use of AI/ML can also expose the Enterprise to heightened risks, including compliance, financial, operational, and model risks. Effective governance of AI/ML should address these varied, cross-sectional risks in the context of the complexity and sophistication of the AI/ML methods used and the extent and materiality of each AI/ML use case.

The Enterprise should develop an enterprise-wide strategy for responsible AI/ML adoption that identifies the goals, benefits, and risks of AI/ML and clearly documents the corresponding risk management approach and framework for ensuring the application of appropriate risk governance. This strategy should be consistent with a risk culture and applicable risk appetite that integrates AI/ML core ethical principles into business processes and operations.³ The existing enterprise-wide risk management framework and governance processes should be leveraged to the extent practicable and updated to incorporate AI/ML concepts and risk management considerations.

The Enterprise should consider the following as foundational components when establishing a safe and sound AI/ML governance structure:

- **AI/ML Core Ethical Principles** – A set of core ethical principles should guide the Enterprise’s use of AI/ML and facilitate consistent governance across various business activities and functions, taking into consideration legal and compliance risks as well as how humans should interact with AI/ML systems. Personnel should be trained and aware of when and how these principles apply. These principles can include, but are not limited to, the following:
 - Transparency – Provide adequate clarity regarding how and why AI/ML is used, in addition to sufficient understanding, interpretability,⁴ and explainability,⁵ allowing for objective assessment and conceptual soundness validation.

³ See FHFA Advisory Bulletin 2020-06, *Enterprise Risk Management Program* (Dec. 11, 2020).

⁴ Interpretability refers to the extent to which a human can understand the choices taken by a model in the algorithmic decision-making process.

⁵ Explainability refers to how an AI/ML approach uses inputs to produce outputs (i.e., can the outcome be explained).

- Accountability – Assign appropriate human responsibility for AI/ML outcomes with adequate explanation and justification throughout each lifecycle stage in order to avoid and mitigate adverse outcomes.
- Fairness and Equity – Implement processes that drive fair and equitable AI/ML outcomes across different groups. Fairness is evaluated in consideration of the conditions and objectives of the AI/ML activity, and when applicable, in light of social, economic, political, or cultural biases.
- Reliability – Design AI/ML capabilities to operate as intended throughout each lifecycle stage, taking into account purpose, values, accuracy, and safety.
- Privacy and Security – Respect and protect privacy rights and data used for development and use of AI/ML throughout each lifecycle stage using industry best practices, as applicable.
- **AI/ML Definitions and Taxonomy** – An enterprise-wide definition and taxonomy for AI/ML terms and capabilities fosters a common vocabulary and understanding across the enterprise in a field that is rapidly evolving. Examples of capabilities include, but are not limited to, techniques such as prediction, classification, natural language, vision, web scraping. Examples of AI/ML terms include, but are not limited to, techniques such as supervised learning, unsupervised learning, reinforcement learning, neural networks, and deep learning. A taxonomy with clear definitions of AI/ML terms and capabilities should facilitate the effective identification and management of AI/ML risks. This taxonomy should include what the Enterprise is and is not classifying as an AI/ML model.
- **AI/ML Inventory** – A comprehensive inventory that captures the Enterprise’s AI/ML use cases across business lines, can provide the Enterprise with a holistic view of how to best manage its AI/ML associated risks. The Enterprise should determine the degree to which it needs to identify and document AI/ML techniques in addition to use cases, understanding that AI/ML can be embedded in models, applications, systems, platforms, tools, and services—either developed in-house or procured from third-party vendors. The AI/ML inventory should be appropriate for the Enterprise’s size, complexity, and risk profile, and include AI/ML use cases that range from proof-of-concept through production. To the extent practicable, the AI/ML inventory should be aligned with

existing enterprise-wide inventory systems, such as those used for models, IT assets, and third parties.

A. Roles and Responsibilities

Consistent with the Enterprise's overall enterprise risk management (ERM) program,⁶ the board of directors (board) is responsible for overseeing enterprise-wide risk management and fostering an effective risk culture. An enterprise-wide approach to managing AI/ML risks should be incorporated into the Enterprise's ERM program and managed within the Enterprise's risk appetite and applicable risk limits framework. Senior management is responsible for executing the AI/ML strategy and the specific risk management practices for AI/ML. Senior management should consider an interdisciplinary approach to AI/ML business decision-making, risk management, and risk oversight that includes sufficient representation from first-line business functions and second-line oversight functions when developing, implementing, and using AI/ML.

Effective AI/ML risk management includes the following considerations commensurate with the risk and complexity involved in the Enterprise's use of AI/ML:

- Assigned AI/ML risk management roles that are clearly defined and include accountability;
- Clear reporting lines and communication protocols for reporting relevant AI/ML metrics and escalating conflicts;
- Appropriately allocated resources for AI/ML that are in line with business needs and consider the benefits and risks;
- The sufficiency of technical expertise and appropriateness of resources for the complexity and scope of AI/ML techniques;
- The ability of designated personnel to provide current and appropriate guidance on AI/ML adoption and use strategy;
- The training of personnel across the three lines of defense on AI/ML applications, risks, and controls;
- The regular updating of AI/ML related policies, standards, and procedures and the appropriate integration of these into business lines; and
- The timely remediation of issues or concerns identified by FHFA or internal audit, or self-identified by the business.

B. Policies, Standards, and Procedures

⁶ See FHFA Advisory Bulletin 2020-06, *Enterprise Risk Management Program* (Dec. 11, 2020).

The Enterprise's risk policies, standards, and procedures should incorporate measures for identifying, assessing, controlling, monitoring, and reporting AI/ML risks. The Enterprise should develop and maintain processes that promote safe and sound practices throughout the AI/ML lifecycle, incorporating independent review and effective challenge of AI/ML by the second line. Policies, standards, and procedures should also clearly define roles and responsibilities, strategies, risk appetite, and documentation requirements. AI/ML core ethical principles, definitions, taxonomy, and inventory should also be incorporated into policies, standards, and procedures to ensure consistent application across the enterprise. To accommodate the rapidly changing nature of AI/ML, related policies, standards, and procedures may need to be updated on a more frequent basis than non-AI/ML related governing documents.

II. Risk Identification and Assessment

The Enterprise's decision of whether to develop, acquire, and use AI/ML should begin with effective and timely risk identification and risk assessment processes that capture the risks and benefits associated with AI/ML.⁷ This should include analyzing and addressing past incidents and lessons learned from the Enterprise's use of AI/ML. Given the rapid technological advancement of AI/ML and the ability of AI/ML models to dynamically update over time, the identification and assessment of AI/ML risks may need to be done frequently, as needed. For example, a risk assessment conducted when an AI/ML tool was in a proof-of-concept stage can quickly become outdated if the scope of use expands in production. As risks can manifest across the Enterprise beyond a single use case, it is critical to know whether an AI/ML approach that was independently reviewed initially has significantly evolved over time.

Whether AI/ML is developed in house or procured from a third party, risk identification and assessment of AI/ML risks should be incorporated in a timely manner into existing risk management processes. This includes identifying when AI/ML meets the definition of a model⁸ and determining the appropriate risk management processes that apply. This process should follow clear criteria and document the Enterprise's rationale to pursue a particular use case.

Risk identification and assessment should incorporate cross-collaboration and review among stakeholders across divisions, business lines, and risk teams to comprehensively capture AI/ML risks. The Enterprise should have personnel with adequate AI/ML and data analytics subject matter expertise in key positions across all three lines of defense to accurately identify and assess AI/ML risks at appropriate junctures in the AI/ML lifecycle. For instance, AI/ML may be embedded into third-party software and hardware used in customer decisioning or interface that is not readily apparent but influences performance. In this example, stakeholders in technology,

⁷ Consistent with FHFA Advisory Bulletin 2020-06, *Enterprise Risk Management Program* (Dec. 11, 2020), and FHFA Advisory Bulletin 2014-02, *Operational Risk Management* (Feb. 18, 2014).

⁸ See FHFA Advisory Bulletin 2013-07, *Model Risk Management Guidance* (Nov. 20, 2013).

modeling, and third-party risk management should be involved in order to adequately identify and assess risks.

The financial, compliance, legal, reputational, and operational risks that are typically assessed for any business activity should be evaluated with respect to the use of AI/ML. Risks may become heightened given the complexity and speed of AI/ML innovation and use, which can manifest in unfamiliar ways, thus making AI/ML risks harder to identify in an effective and timely manner. Key risk considerations are discussed in more detail below.

A. Model Risks

For AI/ML, the following are heightened model risks:

- **Black Box Risk** – There can be an inherent tradeoff between model complexity, accuracy, and transparency when using AI/ML models. Complex AI/ML models may not offer clear relationships between model inputs and outputs that are readily understandable by humans. A lack of interpretability, explainability, and transparency – or “black box risk” – can translate into higher levels of uncertainty about the conceptual soundness and suitability of the AI/ML approach. Related to this is the risk of a lack of expertise among model developers in building and users in applying AI/ML models.
- **Overfitting** – Model out-of-sample performance may be significantly worse than in-sample performance when a model learns from idiosyncratic patterns in the training data that is not representative of the population being modeled.⁹ While overfitting is a common risk with traditional models, the risk is heightened with the use of AI/ML models. Undetected overfitting could result in incorrect predictions or categorizations.
- **Model Drift** – The risk of model performance degradation over time is also heightened with the use of AI/ML models. This can be driven by data drift—which occurs when there are changes in the population being modeled thereby affecting the representativeness of input data—or concept drift, which occurs when the relationships between model inputs and outputs change.
- **Model Calibration and Feedback** – Dynamic model calibration, self-updating, and continuous feedback with the use of certain AI/ML models can present heightened model risks, as these models may create a feedback loop that is not well understood. The accuracy of the AI/ML model’s results may degrade rapidly if compromised feedback is not detected in a timely manner. More opaque and complex AI/ML models can also present challenges in understanding why a particular approach experiences performance degradation due to a lack of transparency.

⁹ In-sample performance is model performance based on the training sample, while out-of-sample performance is model performance generated using data excluded from the training sample.

- Bias¹⁰ – Bias in AI/ML models contributes to poor predictability and can lead to discriminatory or unfair outcomes that benefit or harm some individuals, groups, or communities disproportionately. Bias can arise from the data used and can be amplified by the algorithm itself.
- Model Misuse – Business users may lack an adequate level of understanding of the AI/ML model’s output and limitations. Model misuse may also be driven by misalignment between the model methodology or algorithm and the business problem to be addressed and quantified by the model.
- Vendor Models – The use of vendor AI/ML models may heighten existing vendor model risks because of increased model and data complexity and lack of transparency due to the proprietary nature of such models.

B. Data Risks

The quality and appropriateness of data used in AI/ML is crucial in producing reliable decisions or predictions. Large and diverse datasets drive many AI/ML algorithms. Unrepresentative and unsuitable data reduces the accuracy and utility of AI/ML. The following data risks are heightened with the use of AI/ML:

- Appropriateness and suitability of data for purpose (e.g., data source and selection of data).
- Appropriateness and suitability of the dataset for a particular stage of use (e.g., data for training versus production, testing, and validation).
- Accuracy and quality of data used in training and production.
- Appropriateness of data sampling techniques used that could result in imbalanced datasets.
- Bias in selection of data such as omission bias or stereotype bias, and bias in data processing.
- Complex, high-dimensional data, and new, unfamiliar data sources, such as third-party data or unstructured data.
- Time and cost associated with acquiring, curating, and preparing data.
- Lack of data lineage preservation and the failure to identify root causes of errors or risks associated with the storage and movement of data that could affect data integrity.
- Security of data from unintentional and intentional manipulation of data, such as data poisoning.

¹⁰ See, e.g., National Institute of Standards and Technology (NIST) research on identifying and managing bias in artificial intelligence.

C. Other Operational Risks

The use of AI/ML involves other operational risks, such as information technology, information security, third-party, and business resiliency risks. Depending on the scope and complexity of AI/ML use cases, the following are areas of potential risk:

- IT infrastructure – Legacy IT systems may not be able to support the storage, transfer, and processing of big datasets for AI/ML. Implementing AI/ML can also place a high demand on IT infrastructure and cloud-based services. Insufficient computing power and hardware can degrade network latency and performance standards per established key indicators. For example, AI/ML models that require reliable computing speed to handle model complexity and frequent recalibration needed for production readiness may be negatively impacted by ill-equipped IT systems.
- Information security – Adopting AI/ML systems may pose risks to existing processes that can compromise the confidentiality, integrity, and availability of information. Open source software or application program interfaces (APIs) embedded into AI/ML technology may also present susceptibility to adversarial attacks.
- Business continuity – Business functions supported by AI/ML can feed into downstream business processes or other AI/ML systems that can cause significant disruptions across the enterprise if AI/ML performance is degraded or compromised.
- Use of AI/ML through third-party providers – Third-party provided products and services—ranging from those with embedded AI/ML to cloud providers hosting AI/ML platforms—present potential business resiliency and concentration risks if AI/ML services are limited to a few vendors.¹¹

D. Regulatory and Compliance Risks

The use of AI/ML presents regulatory and compliance risks, such as compliance with consumer protection, fair lending, privacy, and employment discrimination laws and regulations. For example, the use of AI/ML-based credit underwriting models in credit decision-making can present compliance risks due to a lack of explainability of the model, interpretability of the model output, and adequacy of controls in the decision-making process that may be mandated by consumer protection and fair lending laws and regulations. Additionally, personal data used in AI/ML may be subject to complex data governance and privacy laws with requirements such as anonymizing data, securing consent to use the data, and maintaining a record of how data is used, accessed, and stored.

¹¹ See FHFA Advisory Bulletin 2018-08, *Oversight of Third-Party Provider Relationships* (Sept. 28, 2018).

III. Control Framework

The degree and scope of risk management and controls addressing AI/ML should be commensurate with the extent and complexity of AI/ML development and use at the Enterprise and level of risk exposure. The Enterprise should consider the evolving nature of AI/ML when evaluating, adjusting, or adding mitigating controls. Appropriate stakeholders should determine whether controls are in line with applicable risk appetite metrics. Controls mitigating AI/ML risk should be embedded in policies, standards, and procedures, and in the roles and responsibilities of all stakeholders throughout the AI/ML lifecycle. Key control considerations are discussed in more detail below.

A. Model Controls

While FHFA guidance for model risk management and model controls framework¹² applies to AI/ML models, the Enterprise should also consider:

- Whether model risk policies, standards, procedures, and practices sufficiently address AI/ML concepts such as—but not limited to—model interpretability, explainability, transparency, bias, fairness, dimensionality reduction, hyperparameter selection, feature engineering, and dynamic retraining and updating. Existing model risk management practices may need to be adapted to address non-traditional use cases, such as chatbots, cybersecurity, and human resources analytics.
- Whether the Enterprise has staff across all lines of defense with appropriate knowledge, skills, and experience in AI/ML data science, analytics, and modeling. For example, model owners and users should have a sufficient understanding of the underlying AI/ML model assumptions and limitations.
- Whether the Enterprise has an AI/ML model development process that guides initial determinations on data quality and suitability, model conceptual soundness, explainability, and appropriateness of use.
- Whether the Enterprise has tools and techniques to determine drivers of AI/ML model decisions and to assist in model interpretability, bias detection, and performance testing.
- Whether the frequency of AI/ML model performance tracking and ongoing monitoring is adequate to observe changes in model drift and degradation, dynamic updating, and the adequacy of corresponding model change management processes. For example, AI/ML models may update more frequently than traditional models, requiring recalibration and tuning as the algorithm learns from new data. To accommodate this more frequent update cycle, the AI/ML model should be dynamically monitored to detect changes in performance and impact on business usage.

¹² See FHFA Advisory Bulletin 2013-07, *Model Risk Management Guidance* (Nov. 20, 2013).

- Whether the frequency and scope of model validation and effective challenge processes is adequate to sufficiently address AI/ML models and related concepts. For example, point-in-time independent model risk management and model validation approaches may need to be adapted as AI/ML models may not be static between reviews.
- All AI/ML models are expected to go through model validation. This includes AI/ML models used by internal audit and other functions that may not traditionally use model output such as the information technology functions. In all cases, the second line model risk management function should perform the validation, or contract with a third party for the validation should additional expertise be necessary.
- Model risk management processes for identification of material model changes may need to be enhanced, given the more frequent AI/ML model change management cycle.
- Whether model documentation requirements and frequency of update are adequate to reflect current AI/ML model input and output relationships and model operation.
- Whether consideration of ethical principles, such as fairness and bias, are adequately addressed throughout all lifecycle stages.
- Whether an adequate independent assessment of third-party AI/ML models is performed to evaluate the conceptual soundness, security, and integrity of the AI/ML model's development and performance.

Challenger Models

Challenger models are developed as an alternative to a champion or production model, allowing for testing of alternative theoretical or estimation methodologies. Challenger models may be developed internally or by external vendors, subject to the same principles as internally developed challenger models. The criteria for determining champion and challenger models should be clear and measurable, and provide adequate support for why one model is chosen to be the champion model along with analysis of model performance and related assumptions. The Enterprise should take a risk-based approach with regard to the intensity and frequency of a challenger model's validation and effective challenge and, to the extent AI/ML techniques are utilized, ensure heightened risk management considerations as described in this AB are considered.

B. Data Controls

Data risk management strategies, governance, policies, procedures, and standards may need to be enhanced to address increased data risks associated with the use of AI/ML.¹³ The Enterprise should consider the following when evaluating the data risks associated with AI/ML:

¹³ See FHFA Advisory Bulletin 2016-04, *Data Management and Usage* (Sept. 29, 2016).

- The adequacy of data risk management roles and responsibilities such as data ownership and management. For example, there may need to be more frequent and robust data accountability roles and approval processes to address data quality, relevance, and compliance concerns.
- The strength of practices and processes to mitigate the sources of data bias, such as data proxies and use of over- or under-represented data.
- The efficacy of each stage of data management, including the acquisition and sourcing of data, data preparation and processing, data quality review, and data sampling to address data bias, appropriateness, quality, and preservation.
- The adequacy of documentation requirements for each stage of data management, such as usage rights and data permissions.
- The strength of data lineage practices with all types of data formats, such as unstructured data, that adequately captures the transformations and modifications to data.
- The adequacy of enterprise-wide data architecture and systems to accommodate the storage, processing, and movement of vast, complex data sets and various data types used for AI/ML while ensuring business operations are not adversely affected.
- The degree and frequency of monitoring data at each stage of use to identify risks such as data drift and data anomalies.
- The adequacy of data testing measures and remediation to ensure data issues are resolved.
- The sufficiency of data security measures from internal and external threats and compromises to data.

C. Other Operational Controls

To address other operational risks raised with the use of AI/ML, the Enterprise should consider the following risk mitigation solutions:

- Scalable infrastructure to support data storage and computing power necessary to meet operational and business needs.
- Business continuity plans and incident response plans that are adapted to AI/ML tools, systems, and applications, including third-party AI/ML products and services.
- Contingency plans, including manual override functions, when automated AI/ML dependent processes become skewed.
- Workarounds that address interconnectivities and dependencies of data.
- Sufficient and consistent testing of in-house and third-party AI/ML tools, applications, and systems to assess integrity, security, and business resiliency.
- Appropriate change management practices and procedures to accommodate evolving AI/ML techniques.
- Security measures to monitor and protect cloud-based AI/ML models and data.
- Open-source software controls.

- Contractual requirements with third-party providers of AI/ML models and data that ensure transparency and accountability with use.

D. Regulatory and Compliance Controls

The Enterprise may need to adapt its existing regulatory and compliance risk management practices and controls to accommodate AI/ML associated risks, including the following:

- Revising policies, procedures, and standards to address AI/ML explainability, interpretability, and transparency, and compliance with applicable laws and regulations.
- Designing a compliance risk management program,¹⁴ that includes analysis of relevant consumer protection, employment discrimination, privacy, and other laws and regulations as they apply to the use of personal and alternative data.
- Involving qualified compliance personnel during AI/ML development and implementation to ensure data and methodologies comply with applicable laws and regulations.
- Integrating fair lending reviews and testing, as appropriate, through all lifecycle stages.

IV. Risk Monitoring, Reporting, and Communication

The Enterprise should establish appropriate key risk indicators (KRIs) and key performance indicators (KPIs) for monitoring and analyzing AI/ML risks and risk management practices in line with risk appetite. These KRIs and KPIs can indicate whether existing risk management practices are effective or need to be modified. AI/ML related risk and performance metrics should be reported and communicated to the appropriate stakeholders across the enterprise. Reporting and communication protocols may need to be reviewed and adjusted more frequently to optimally capture and timely convey AI/ML associated risks as they evolve and change. The Enterprise should consider the following when monitoring, reporting, and communicating AI/ML risks within and across business lines:

- The degree and frequency of monitoring needed to adequately capture the scope of AI/ML risks, including model, data, compliance, information security, and other operational risks.
- The relevancy and effectiveness of KPIs and KRIs in measuring changes to the risk profile associated with AI/ML risks, and the frequency to which they need to be evaluated and reviewed for changes. Such metrics should also reveal the comparative business advantages or disadvantages of using AI/ML.

¹⁴ See FHFA Advisory Bulletin 2019-05, *Compliance Risk Management* (Oct. 3, 2019).

- The benefits and risks associated with AI/ML powered monitoring applications and the appropriate level of human involvement and discretion needed for monitoring AI/ML risks.
- The adequacy of reporting within and across business units, lines, and the enterprise, including board and senior management, to effectively communicate AI/ML risks.
- The type of information regarding AI/ML performance and risks that needs to be conveyed to different stakeholders across the enterprise and escalated to senior management and the board. For example, first line data scientists and modelers may rely on granular AI/ML metrics while second line risk management may utilize broader, aggregated AI/ML data.

Related Guidance and Regulations

12 CFR Part 1239, Responsibilities of Boards of Directors, Corporate Practices, and Corporate Governance Matters.

12 CFR Part 1236, Appendix, Prudential Management and Operations Standards.

12 CFR Part 1223, Minority and Women Inclusion.

Model Risk Management Guidance, Federal Housing Finance Agency Advisory Bulletin 2013-07, November 20, 2013.

Operational Risk Management, Federal Housing Finance Agency Advisory Bulletin 2014-02, February 18, 2014.

Data Management and Usage, Federal Housing Finance Agency Advisory Bulletin 2016-04, September 29, 2016.

Internal Audit Governance and Function, Federal Housing Finance Agency Advisory Bulletin 2016-05, October 7, 2016.

Information Security Management, Federal Housing Finance Agency Advisory Bulletin 2017-02, September 28, 2017.

Cloud Computing Risk Management, Federal Housing Finance Agency Advisory Bulletin 2018-04, August 14, 2018.

Oversight of Third-Party Provider Relationships, Federal Housing Finance Agency Advisory Bulletin 2018-08, September 28, 2018.

Business Resiliency Management, Federal Housing Finance Agency Advisory Bulletin 2019-01, May 7, 2019.

Compliance Risk Management, Federal Housing Finance Agency Advisory Bulletin 2019-05, October 3, 2019.

Enterprise Risk Management Program, Federal Housing Finance Agency Advisory Bulletin 2020-06, December 11, 2020.

Enterprise Fair Lending and Fair Housing Compliance, Federal Housing Finance Agency Advisory Bulletin 2021-04, December 20, 2021.

FHFA has statutory responsibility to ensure the safe and sound operations of the regulated entities and the Office of Finance. Advisory bulletins describe FHFA supervisory expectations for safe and sound operations in particular areas and are used in FHFA examinations of the regulated entities and the Office of Finance. Questions about this advisory bulletin should be directed to: SupervisionPolicy@fhfa.gov.