

PART ONE — Purpose, Scope and Fundamental Principles

Purpose

ARTICLE 1 — The purpose of this Framework is to establish the fundamental principles governing the design, implementation and monitoring of AI-powered products and services developed, provided or operated by Next4biz in a trustworthy, controlled, transparent and human-centered manner.

This Framework sets out the corporate practices and approach adopted by Next4biz to enhance the benefits provided by AI systems, to mitigate foreseeable risks, and to ensure that such systems are operated in accordance with their defined intended purposes.

Scope

ARTICLE 2 — This Framework covers the entirety of machine learning, deep learning, natural language processing, generative AI, large/small language models, forecasting, classification, recommendation, anomaly detection and AI-supported decision-making mechanisms developed by Next4biz or used within the scope of Next4biz products.

Within this scope, Issue Intelligence, PRIME, Chatbot, Auto Responder, Marketplace Responder, CSM Responder, sentiment analysis, InsightX, ComplaintXtract,

AnomalyNet, toXt2RPM, workload forecasting services, sentiment analysis agent

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

SETTINGS

REJECT ALL

ACCEPT ALL

determined by taking into account the system's intended purpose, its technical architecture, customer requirements, and the risks it may pose.

Nature of the Framework

ARTICLE 3 — This Framework describes Next4biz's general approach to the development and use of AI systems. It does not constitute a performance commitment, a service level guarantee, or a contractual guarantee of outcomes with respect to any specific model, service or customer project.

Product- and project-specific technical requirements, acceptance criteria, performance targets, customer responsibilities and service terms are separately regulated in the relevant contract, project document, technical specification or solution design report.

Risk-proportionate application

ARTICLE 4 — The principles set forth in this Framework are not applied with the same method and intensity for every AI system. The level of control to be applied is determined by considering whether the system communicates directly with the customer, whether it produces automated decisions or transactions, the nature of the data it processes, and the impact that an erroneous output may cause.

Systems that send responses to customers, automatically route a record, initiate transactions in external third-party systems, or generate recommendations that may produce significant consequences for individuals are subject to additional control, verification and oversight processes commensurate with their risk and impact level.

Not all of the technical methods and controls specified in this document are applied together in every AI system or customer project. The methods to be used are determined by taking into account the system's intended purpose, technical

architecture, data structure, risk level and customer requirements, and are

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

Human-centricity and human oversight (human-in-the-loop)

ARTICLE 5 — Next4biz AI systems are developed for the purpose of supporting human decisions and business processes. The level of automation granted to AI systems is determined in proportion to the nature of the use case and its potential impacts.

Outcomes that are ambiguous, high-impact, or that the customer requires to be subject to human review may be routed to the relevant user, agent, or pre-review pool. It is essential that authorized users are able to review, modify, or reject the AI output, or, where necessary, deactivate the system.

Purpose and scope limitation

ARTICLE 6 — For each AI system, the intended purpose, function, target users, and circumstances in which it must not be used are defined. Systems are not used to make decisions or perform transactions outside the defined scope.

Customer-specific restrictions on topics, products, brands, communications, recommendations, and transactions may be defined within the scope of the project. These restrictions are enforced through system instructions, knowledge sources, business rules, access controls, and output controls.

Reliability and measurability

ARTICLE 7 — AI systems are validated using datasets, test scenarios, and evaluation methods appropriate to their intended purposes. The data used in the training or model configuration process and the validation and test data on which performance is measured are separated from one another to the extent possible. Care is taken to ensure that test data represent the classes, user expressions, transaction types, and exceptional situations that the system will encounter in the

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

used as the test set. Final performance is reported by jointly evaluating the average of all repetitions and the variability among the results. In classification problems with imbalanced class distributions, the stratified k-fold method, which ensures that class proportions are preserved in each subset, is preferred.

In datasets containing related records belonging to the same customer, user, campaign, conversation, process, or document group, the distribution of samples belonging to the same group across both the training and the test data is prevented. In such cases, group k-fold or similar group-based splitting methods are used to reduce the risk of data leakage. For time-dependent data, chronological splitting, rolling window, or time-series cross-validation methods are applied to prevent future records from being used in training on past data.

Where the dataset is limited, the leave-one-out cross-validation method may be used. In this method, a single sample is set aside for testing in each iteration, all remaining samples are used for training, and the process is repeated for each sample in the dataset. While the leave-one-out method allows the majority of the available data to be used in training, it is applied taking into account data size, model complexity, and class distribution, due to its high computational cost and the variability that may arise in the results.

Cross-validation results do not automatically substitute for an independent final test set. In projects where the amount of data is sufficient, model and hyperparameter selection is performed using training and validation data, and final performance is measured on an independent test set not previously used in the model development process. The data partitions used for model selection and for performance reporting are kept separate from one another, thereby reducing the risk of overfitting to the test data.

The testing process evaluates not only standard samples for which the system is expected to produce correct results, but also incomplete or ambiguous inputs,

prohibited answers or transactions are examined through automated measurements and, where necessary, expert evaluation.

System performance is evaluated not solely on the basis of an overall success rate, but by taking into account the error types required by the use case, class-level results, false positive and false negative rates, confidence scores, the need for human intervention, and end-to-end business outcomes. The performance indicators and acceptance thresholds to be used are determined before testing is carried out, having regard to the intended purpose of the system and the potential impact of erroneous results.

Test results are recorded together with the dataset used, the number of samples, the class distribution, the data splitting method, the number of folds, the evaluation period, and the model and configuration version. Where cross-validation is used, the standard deviation among fold results or similar measures of variability are reported alongside the average performance value. Results obtained from samples that are limited, imbalanced, outdated, or insufficiently representative of live usage conditions are not presented as findings that definitively represent the system's overall and ongoing performance.

Transparency and explainability

ARTICLE 8 — Next4biz endeavors to ensure that the intended purpose, basic mode of operation, and known limitations of its AI systems are understandable to the relevant parties.

Depending on the use case and the structure of the model used, explainable AI methods such as SHAP, LIME, ELI5, Layerwise-Relevance Propagation (LRP), permutation feature importance, coefficient analysis, or tree-based feature importance levels may be utilized to identify the variables affecting the system output. Through these methods, the indicator variables influencing a specific

sources, similarity scores, confidence values, and applied decision rules that influence the classification or response may be recorded.

To the extent required by the use case, the information underlying the class, score, recommendation, or response produced by the system, together with the confidence level, influential variables, decision rules, and related system records, may be provided to the user or to the authorized review unit.

Explainability outputs assist in interpreting the model's decision, but do not mean that the model has identified a causal relationship or that the variables highlighted in the explanation alone determine the outcome. The level of technical explainability may be limited taking into account information security, intellectual property, trade secrets, and the need to prevent misuse of the system.

Fairness and prevention of discrimination

ARTICLE 9 — Care is taken to ensure that AI systems do not produce systematically and unjustifiably adverse outcomes with respect to specific persons or groups.

Depending on the nature of the use case and on the lawful availability of the required data, model performance is evaluated separately for the relevant user groups. In this evaluation, measurements such as group-level accuracy, precision, recall, false positive rate, false negative rate, selection rate, and confidence score calibration may be used. Where necessary, differences in outcomes between groups are examined using fairness metrics such as demographic parity, equal opportunity, and equalized odds.

In the analyses, not only individual attributes but also, where sufficient samples are available, intersections of variables such as age group, language, region, channel, or customer segment may be taken into account. Minimum sample size, confidence intervals, bootstrap analysis, or similar statistical methods are used to assess whether observed differences result from sample insufficiency

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

Significant performance differences that are identified are evaluated with respect to data distribution, labeling quality, sample balance, model structure, confidence threshold, business rules, and manner of use. Where necessary, bias mitigation methods such as resampling, class or sample weighting, data enrichment, feature selection, group-based threshold optimization, or retraining of the model may be applied.

Not every performance difference measured between groups in itself indicates the existence of discrimination. Results are evaluated together with the intended purpose of the system, the representation of the groups within the data, the impact of error types, and the applied business rules. If the data distribution or group-level performance changes during live use, the relevant measurements are periodically re-examined.

Privacy and data protection

ARTICLE 10 — Data used in AI systems are processed in accordance with the applicable data protection obligations and Next4biz information security rules. Data protection controls are applied so as to cover the entire data lifecycle, including the ingestion of data into the system, its preparation, its transfer to the model, its storage, its logging, and its deletion.

The principle of data minimization is observed in the design of systems. It is ensured that personal or sensitive data not necessary for the intended purpose are not transferred to the model. Regular expressions, dictionary-based checks, Named Entity Recognition (NER) models, and data loss prevention mechanisms may be used in combination to detect personal data such as first name, surname, telephone number, e-mail address, identification number, address, financial information, and similar data in model inputs.

Detected personal data are transformed, depending on the use case, through

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

It is essential that data are protected with up-to-date encryption methods during transfer and in the environments in which they are stored, that access is restricted on a role- and authorization-based basis, and that data belonging to different customers are logically or physically separated. Data access may be recorded together with user, service, time, and transaction information, depending on the nature of the operation performed.

Training, validation, test, and live-use data are separated from one another to the extent possible. The use of live environment data for model training, fine-tuning, or retraining purposes is carried out taking into account the project scope, data usage authorization, and the relevant approval processes. In test and development environments, the use of masked, de-identified, or synthetic data instead of real personal data is preferred.

Where third-party models or services are used, the scope of the data to be transferred to the model is technically restricted. The service provider's data retention, model training, logging, and data processing conditions are assessed. Where possible, configurations that prevent customer data from being used by the service provider for model development purposes are preferred.

Input and output logs retain only the data necessary for error analysis, traceability, and operational monitoring. Personal or sensitive data in log records are masked, de-identified, or removed prior to recording where required. Access rights to logs are restricted, and records are deleted at the end of the designated retention period or de-identified in an irreversible manner.

Data sources, applied transformations, access authorizations, data transfer points, and retention periods are recorded. When the purpose of use of a data source ceases to exist or its retention period expires, processes are operated for the deletion of the relevant data from active systems, temporary storage areas, and, to the extent applicable, backups.

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

ineffective, the creation of unauthorized transactions, the use of services outside their intended purpose, and similar risks are assessed during the design, development, and operation phases of the system.

Within this scope, technical controls such as authentication, role- and authorization-based access control, segregation of customer data, input and output validation, secure configuration management, protection of secret keys, transaction limitation, request rate control, retention of security logs, and monitoring of anomalous use may be applied.

The output produced by an AI model is not directly accepted as a trusted or authorized command. Before any transaction is performed in external systems, model outputs are validated at the application layer with respect to authorization, parameters, data type, business rules, and transaction scope. Additional verification or human approval may be applied for transactions that are irreversible, that have financial consequences, or that create significant business impact.

It is essential that systems transition to a safe state in the event of error, outage, security suspicion, or uncertainty. In such cases, the system may be caused to refrain from generating transactions, to suppress the output, to provide a limited response, to inform the user, to record the transaction, or to transfer the process to an authorized person.

The effectiveness of the applied security controls is assessed, depending on the nature of the system, through security tests, access control tests, malicious input scenarios, API tests, vulnerability scans, and controlled attack simulations. Identified security vulnerabilities are prioritized according to their risk levels and remediated.

Accountability

ARTICLE 12 — For each AI system, a business owner, a technical owner, and, where necessary, approval authorities are designated. Responsibilities relating to the

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

PART THREE — Governance and Responsibilities

AI governance

ARTICLE 13 — AI systems developed within Next4biz or integrated into its products shall be recorded together with their intended purposes of use and risk levels.

During the development of a new AI system, a significant modification of an existing system, or its opening to a new area of use, the technical, operational, information security and, where necessary, legal impacts shall be assessed.

Use scenarios assessed as carrying significant risk may be reviewed with the participation of the relevant product, R&D, information security, legal and management units.

Responsibilities of the units

ARTICLE 14 — Product and project teams are responsible for clearly defining the system's intended purpose of use and customer expectations.

R&D and software teams ensure the implementation of data, model, technical controls, testing, versioning and monitoring mechanisms.

The information security unit assesses the security risks relating to the system's architecture and data flow. Legal and compliance units provide opinions on the relevant obligations and customer documentation where the use scenario so requires.

The customer is responsible for the accuracy, currency and authorization of use of the content it provides to the system, and for the correct and complete communication of customer-specific business rules.

ARTICLE 15 — Before an AI system is developed or deployed in a new use case, an assessment shall be conducted of the system's purpose, its intended users, the data to be used, the outcomes it is expected to produce, and the impacts of potential errors.

The assessment shall take into account considerations such as whether the system merely provides recommendations, makes automated decisions, communicates directly with the customer, or initiates transactions in an external system.

As a result of the risk assessment, controls such as human approval, confidence thresholds, non-response, a pre-review queue, content filtering, transaction limits, or additional monitoring may be applied.

Data governance

ARTICLE 16 — The suitability of training, validation, testing, and live-use data for the system shall be assessed. The source, currency, and scope of the data, and the degree to which it represents the intended purpose of use, shall be taken into account.

It is acknowledged that significant changes in data distribution may affect system performance. Data that has become outdated or that does not represent current customer behavior shall not be used for model development or performance validation purposes without the necessary assessments being carried out.

Information and documents provided by the customer shall be processed solely within the scope of the relevant project and contract.

Testing and acceptance

ARTICLE 17 — AI systems shall undergo functional, technical, and scenario-based testing before being put into production. The scope of testing shall be determined in

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

Customer acceptance tests shall be conducted where necessary for customer-specific configurations. The decision to go live shall be made taking into account the assessment results and the accepted residual risk.

Change management

ARTICLE 18 — Significant changes to the model, knowledge sources, system instructions, threshold values, business rules, or integrations shall be implemented in a controlled manner.

Where a change may significantly affect system behavior, the relevant tests shall be repeated. Where necessary, the ability to roll back to the previous version or to disable the relevant feature shall be provided.

PART FIVE — Specific Principles Regarding AI Systems

Classification and routing systems

ARTICLE 19 — In classification and routing systems such as Issue Intelligence and PRIME, class definitions, category hierarchies, confidence scores, decision thresholds and workflow rules are determined according to the use case. Care is taken to ensure that classes are separable from one another, operationally meaningful and capable of being represented with a sufficient number of examples. Classes that overlap with one another, are ambiguous or have insufficient examples are examined separately during the data preparation and model development stages.

Class distributions are analyzed in model training and evaluation. In imbalanced datasets, methods such as class weighting, resampling, data augmentation or class-

The extent to which the confidence scores produced by the model reflect prediction accuracy is examined through calibration analyses. Decision thresholds are determined on validation data, taking into account the cost of misrouting, human review capacity and class-based error impacts. Where necessary, different thresholds on a category or class basis may be used instead of a single general threshold. Records whose confidence score falls below the determined thresholds, that produce close results across multiple classes, or that do not show sufficient similarity to the defined classes may be transferred to human review, a pre-assessment pool or an alternative workflow instead of being routed directly.

In systems such as PRIME, where ambiguous records are routed to a preliminary review pool, the pool is treated as a separate and valid system output. Records sent to the preliminary review pool are not automatically deemed incorrect predictions. System performance is measured across all defined classes — active, passive and the preliminary review pool — and with an evaluation method appropriate to the use case.

The initial result suggested by the model, the confidence score, alternative classes, the applied threshold value, the business rule and the final routing are recorded in a manner that keeps them distinct from one another. In asynchronously operating systems, a transaction or process identifier is created for each prediction, ensuring that the content ultimately saved by the user is matched with the relevant model prediction. This prevents intermediate predictions generated before the text is completed from being treated as the final system result, and prevents model performance from being confused with the outcome arising at the application layer.

Classification performance is not evaluated solely on the basis of overall accuracy. Depending on the use case, measurements such as precision, recall, F1 score, macro F1, weighted F1, balanced accuracy, false positive rate, false negative rate and the confusion matrix are used. Results are reported on a category and class basis. Particularly for classes with high operational impact, error types and the

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

performance decline in certain classes, the emergence of new forms of expression or a change in the category structure, the dataset, threshold values, class definitions and the need to retrain the model are reviewed.

Generative AI and dialogue systems

ARTICLE 20 — Generative AI systems such as Chatbot, Auto Responder, Marketplace Responder, CSM Responder and similar systems are configured in accordance with the purpose of use, target users, communication channels, information sources, customer rules and the operations the system is permitted to perform, as defined within the scope of the project.

At the start of the project, a Chatbot Design Document or an equivalent technical analysis document is prepared together with the customer. This document specifies the topics the chatbot will answer, the areas to be excluded from scope, the information sources to be used, target user profiles, the language and tone of communication, the conditions for transfer to an agent, prohibited responses, product recommendation rules, sensitive topics, external service calls and the operations requiring human approval. Customer-specific rules are enforced through system instructions, information retrieval filters, business rules, output controls and integration permissions.

In order for the system to produce accurate and up-to-date responses about the customer's processes and products, the necessary information, documents and content must be provided by the customer. This content may include product and service descriptions, frequently asked questions, user manuals, procedures, campaign information, product catalogs, business rules, prohibited topics, contractual limitations and situations requiring transfer to an agent. The accuracy, currency, consistency and usage authorization of the content provided by the customer are the customer's responsibility.

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

combining keyword search with semantic search, metadata filtering and re-ranking methods may be utilized.

When user input is received, content related to the query is retrieved from the knowledge base and transmitted to the large language model only together with the necessary context. The number of content items to be retrieved, the similarity threshold, the text segment size, the context length and the re-ranking settings are determined by testing on project data. During retrieval, filters relating to the customer, product, language, channel or authorization level are applied, preventing irrelevant content or content for which no access authorization exists from being passed to the model.

Documents added to the knowledge base do not directly constitute a retraining of the model. In systems using RAG, information is used by presenting the relevant document segments to the model as context during response generation. This approach enables information to be updated on a per-source basis, outdated content to be removed from the system and the generated response to be associated with the sources used.

In order to reduce the risk of hallucination and responses outside the sources, the model is required to answer on the basis of information sources approved by the customer. Rules are defined in the system instructions requiring the model not to present information outside the provided context as definitive information, not to speculate when information is unavailable and not to produce unverifiable content. Where the relevance of the retrieved documents to the query, the source coverage and the level of support for the response are found insufficient, the system may refrain from generating a response, may give a limited answer, may request additional information or may transfer the process to an agent.

The consistency of responses with sources may be examined through automated similarity, semantic consistency, source coverage and supportedness checks. For

Rules such as not recommending products designated by the customer, not giving advice about specific product components, not producing directive responses in sensitive areas such as health, law, finance or similar fields, not comparing specific brands or products, and not answering questions of a personal-advice nature such as "which product is best for me" are defined in the Customer AI Control Plan. These rules may be enforced through topic classification, prohibited expression and concept lists, product metadata, output filters and business rule controls.

When the system detects that the user input relates to a topic whose answering is prohibited or which requires human assessment, it may produce a standard informational message, suppress the response or transfer the conversation to an agent. In such cases, the reason the system did not produce a response may be recorded in association with the relevant customer policy or business rule identifier.

Against prompt injection, jailbreak attempts and requests aimed at neutralizing system instructions, user inputs and system instructions are separated from one another. Content provided by the user is not accepted as a trusted system instruction or an authorized operation command. To ensure that instructions contained within texts retrieved from information sources do not alter model behavior, the content and instruction layers are separated, known attack patterns are examined through input controls and the precedence of system instructions is preserved.

A tool call or operation request generated by the generative AI model is not executed directly. The tool name, operation authorization, parameter types, mandatory fields, user permissions and customer business rules are validated at the application layer. Operations that modify data, produce financial consequences, are difficult to reverse or create a significant impact for the customer may be subject to additional verification or human approval.

Model inputs and outputs may be checked with respect to personal or sensitive

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

Before the system goes live, it is tested not only with normal usage scenarios but also with off-topic questions, missing information, contradictory documents, prohibited product recommendations, requests for health advice, prompt injection, inputs containing personal data, questions with no answer in the sources and unauthorized operation requests. Based on the test results, the system instructions, information retrieval thresholds, document chunking method, re-ranking settings, output controls and conditions for transfer to an agent are updated.

Additions, removals and updates made to the knowledge base are recorded with version information. Content that has lost its currency, is contradictory or has been reported by the customer as withdrawn from use is removed from the retrieval index. After significant changes made to the model, system instructions, information retrieval method or customer rules, the relevant test scenarios are re-executed.

These controls are intended to reduce the risk of producing content that is outside the sources, inappropriate, unauthorized or contrary to fact. Due to the probabilistic nature of the model used, the scope and currency of the information provided by the customer, the nature of user inputs and the behavior of third-party models, they do not constitute a guarantee that all errors will be eliminated under all circumstances.

Automatic response systems

ARTICLE 21 — In systems where an AI output is sent directly to the customer without agent review, additional verification and policy controls relating to the scope, content and communication rules determined by the customer are applied before the response is sent.

The types of topics, operations and customer requests that automatic response systems may answer are explicitly defined within the scope of the project. For requests that are not included in the scope of automatic response, that lack

Certain topics such as health, law, finance, product recommendations, product comparisons, campaign commitments, or price or stock verification may be excluded by the customer from the scope of automatic response. The non-recommendation of certain products or product components, the non-use of certain expressions, the provision of only a standard informational message on certain topics and the situations requiring transfer to an agent are defined within the customer rules.

These rules may be enforced through topic and intent classification, information source coverage control, product and category metadata, confidence thresholds, prohibited topic and expression lists, output validation mechanisms and rule-based suppression controls. For a response to be sent automatically, it is expected that the necessary information sources are present, that the response remains within the defined scope and that it does not violate the relevant customer rules.

Where a response is found unsuitable for automatic sending, the operation may be suppressed, a draft response may be created, additional information may be requested or the record may be routed to human review. The suppression or routing decision may be recorded in association with the relevant policy, topic, product, channel or business rule identifier.

The message length, format, forms of address, prohibited content, link usage and reply rules specific to the e-mail, web, WhatsApp, marketplace or other written communication channels through which the automatic response will be sent are included in the system configuration. Different reply and routing policies may be applied for the same information across different channels.

Automatically sent responses are recorded in a traceable manner together with the model version used, the information sources, the applied rules and the final sending result. User feedback, agent corrections, suppressed responses and erroneous automatic sendings are monitored, and the scope, threshold values and customer

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

FeedbackX360 are of a decision-support nature.

The models and methods to be used in these systems are determined taking into account the language of the text, its length, channel characteristics, class structure and purpose of use. In sentiment analysis, transformer-based language models, multi-class or multi-label classification methods and, where necessary, ensemble approaches combining the results of different models may be used. In addition to the overall sentiment class, sentiment analysis on a topic or product-feature basis may be applied where appropriate to the use case.

Model outputs are evaluated together with confidence scores, decision thresholds are determined on validation data and class-based performance results are monitored. Irony, implicit expression, short statements, spelling errors, texts containing more than one sentiment and industry-specific expressions are included in the test scenarios. Where changes in data and language use may affect model performance, the model, threshold values and training data are re-evaluated.

The results of these systems must not be used as a definitive determination on their own and independently of context, particularly in assessments that may produce significant consequences for employees or customers.

In systems performing summarization and root cause extraction, care is taken to ensure that critical information in the source text is preserved, that the produced result is consistent with the source content and that a cause not present in the source is not presented as a definitive finding. Where necessary, summary, topic or root cause results are evaluated through rule-based checks, a second model or user verification.

Anomaly (outlier) detection and forecasting systems

ARTICLE 23 — The results produced by AnomalyNet and similar anomaly detection systems are indicative in nature, pointing to unusual process instances that require

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

information such as time, user, role, transaction type, resource and duration as node or edge attributes. Multi-graph approaches that enable different characteristics of a process instance — such as control flow, activity, user and time — to be processed on separate graphs may be utilized.

In structures similar to GAMA+, graph representations may be processed with Graph Autoencoder, Graph Attention Network, GATv2 or similar graph neural network architectures. While graph encoders learn the node and edge relationships within the process structure, GRU, Transformer or hybrid decoders may be used to reconstruct the sequential and temporal characteristics of events. Events, attributes or process instances for which the model is unable to reconstruct normal process patterns may be identified as anomaly candidates via the reconstruction error.

In methods similar to SPECTRE, activity, user, time or other process attributes may be processed in separate channels. The temporal pattern of the relevant attribute may be learned through embedding layers created for each attribute, bidirectional GRU encoders and attention mechanisms. Through the joint optimization of the channels, the cleaned or expected version of the process record may be reconstructed, and the differences between the observed record and the reconstructed record may be examined at the event flow and attribute level. This approach may be used to reduce dependence on a single manually determined alarm threshold and to evaluate contaminated process records in conjunction with anomaly detection.

In cases where labeled anomaly data is limited, unsupervised, semi-supervised or synthetic anomaly generation-based training methods may be applied. Care is taken to ensure that synthetic anomalies represent situations that may be encountered in real processes, such as missing activities, unexpected activities, incorrect sequencing, repeated events, wrong users, unusual durations or attribute changes. Model results may be evaluated separately at the process, event and attribute level.

In systems performing workload and traffic forecasting, the trend, seasonality, periodicity, holiday, campaign and channel effects of the data are examined. Depending on the data structure, ARIMA, SARIMA, exponential smoothing, regression-based methods, LSTM, GRU, Temporal Convolutional Network or Transformer-based time series models may be used. External variables such as calendar information, campaign periods, channel changes and operational capacity may be included in the model as additional inputs.

Time series models are evaluated with a chronological training, validation and test split rather than random data splitting. The model's performance across different time periods is measured through rolling window or expanding window backtesting methods. Forecast results are evaluated with MAE, RMSE, WAPE or other error measurements appropriate to the use case, and where possible, prediction intervals are produced together with point forecasts.

During live use, changes in the process structure, anomaly scores and time series distributions are monitored. When new process types, changing user behaviors, sudden traffic increases or events not present in historical data are detected, the re-evaluation of the model, alarm thresholds, forecast variables and training data is ensured.

AI systems performing operations in external systems

ARTICLE 24 — Where AI systems call external services, create records, initiate workflows or modify existing data, the tools that may be used and the operation permissions are defined in advance.

These systems are granted only the permissions necessary to perform their tasks. Human approval or additional verification may be applied for operations that are irreversible, produce financial consequences or create a significant impact for the customer.

PART SIX — Customer-Specific Controls and Shared Responsibility

Customer control plan

ARTICLE 25 — In projects where the behavior of the AI system must be restricted in accordance with the customer's commercial, legal, or operational requirements, a Customer AI Control Plan may be prepared.

This plan sets out the topics on which the system may and may not respond, the information sources it will use, the products that may or may not be recommended, the conditions for transfer to a human agent, the transactions requiring human approval, and customer-specific communication rules.

The system cannot be expected to apply of its own accord a business rule that has not been communicated by the customer, has not been updated, or has not been technically defined in the system.

Content provided by the customer

ARTICLE 26 — The customer is responsible for the accuracy, currency, and authorization for use of the information, documents, rules, and product content it provides to the system.

Changes made to the content or new restrictions must be notified to Next4biz, and the necessary configuration work must be completed. The customer is informed that content that has become outdated may affect the system's results.

Limits of use

ARTICLE 27 — AI systems must not be used outside their defined intended purpose

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

PART SEVEN — Monitoring, Record-Keeping and Incident Management

Monitoring of live systems

ARTICLE 28 — AI systems operating in the live environment shall be monitored through technically and operationally appropriate indicators.

The scope of monitoring may include indicators such as system performance, user corrections, transfers to agents, non-response situations, policy violations, anomalous usage, service latencies and resource consumption.

In the event that a significant change in performance is observed, the data, model, knowledge source, threshold value or business rules may be re-evaluated.

Record-keeping and traceability

ARTICLE 29 — Significant operations of AI systems shall be recorded, taking into account the nature of the relevant system and data protection requirements.

In systems that require the matching of model output with the final business outcome, a transaction or process identifier may be used. In systems operating asynchronously, it must be traceable, to the extent possible, which model output and system version the final user transaction is associated with.

It is essential that only data necessary for operational needs be retained in log records and that sensitive data be protected.

Management of AI incidents

ARTICLE 30 — A systematic erroneous outcome, inappropriate content, an unauthorized transaction, suspected personal data leakage or a significant policy

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

Following the incident, a root cause analysis shall be conducted; causes originating from data, the model, rules, integrations or processes shall be evaluated, and the necessary corrective actions shall be implemented.

PART EIGHT — Third-Party Systems and Vendors

Use of third-party models and services

ARTICLE 31 — Third-party AI model services such as OpenAI, Google Vertex AI, Microsoft Azure, Amazon Bedrock, Anthropic, and similar services may be used in Next4biz solutions. In addition, open-source models, embedding models, re-ranking services, vector databases, content safety services, speech-to-text, text-to-speech, OCR, translation, observability, and model monitoring components may also be incorporated into the system architecture.

The selection of a third-party model or service is not made solely on the basis of the model's general capabilities or comparative benchmark results. The selection jointly considers suitability for the use case, performance on Turkish and domain-specific data, context length, support for structured output and tool calling, response time, capacity limits, cost, service continuity, security features, data processing terms, regional deployment options, model version management, and the applicability of the required technical controls.

Before a third-party service is put into use, the following are examined: the scope of the inputs and documents to be transmitted to the service, whether the generated outputs are retained, whether the data is used by the provider for model training or service improvement purposes, the retention period of abuse-monitoring records, the regions in which the data is processed and stored, sub-service providers, and data deletion capabilities. This assessment is carried out on the basis of the API endpoint, feature, account type, and project configuration used.

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

or similar controls. The fact that a service supports these features in general does not mean that all models and API functions used are subject to the same terms.

Only the data necessary for performing the relevant operation is transmitted to third-party services. Personal, sensitive, or commercial data that is not necessary for the purpose of use is masked, de-identified, or removed from the scope of the input prior to transmission. API keys and access credentials are not kept in application code or model inputs; they are stored in secure secret management systems, and access to service accounts is granted in accordance with the principle of least privilege.

Third-party model outputs are not directly accepted as reliable or verified information. Depending on the use case, model outputs are checked at the Next4biz application layer with respect to consistency with the source, scope, format, content safety, customer rules, and transaction authorization. While the provider's own content filters may be used, reliance of customer rules and Next4biz controls solely on the safety mechanisms offered by the provider is prevented.

In model or service selection, alternative providers that meet the same use case may be compared. In such comparisons, indicators such as accuracy, consistency with the source, hallucination rate, policy violations, latency, availability, and cost are measured on a common test data set and common scenarios. Regression tests are applied on the same test set in order to determine the effect of model changes on system behavior.

If the provider changes the model version, safety policy, data processing terms, API behavior, or usage limits, the effect of the change on the Next4biz solution is assessed. Where possible, pinned model versions are used, automatic version transitions are monitored, and acceptance tests are re-run after significant model changes.

In projects where dependence on a single model or provider poses a risk to service

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

Where open-source models are used, the model card, disclosures regarding the training data, known limitations, commercial use conditions, license provisions, dependencies, and security risks are examined. Model files and third-party libraries are obtained from trusted sources and subjected to integrity checks and security scans. Running an open-source model in the Next4biz or customer environment does not eliminate the requirement to perform separate verification with respect to the accuracy and safety of the results produced by the model.

The use of third-party services is recorded in the technical documentation together with the provider name, the model and version used, the API feature, the scope of data transfer, retention settings, deployment region, security configurations, and the alternative service plan.

PART NINE — Training, Audit and Continuous Improvement

AI awareness

ARTICLE 32 — It shall be ensured that employees involved in the development, sale, implementation, or operation of AI systems possess knowledge of AI risks, data protection, information security, and system limitations at a level commensurate with their roles.

Customer users may be provided with the necessary information regarding the manner of use and the limitations of the relevant product.

Review and improvement

ARTICLE 33 — This Framework and the controls applied to AI systems shall be reviewed taking into account technological developments, customer feedback,

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

PART TEN — Final Provisions

Exceptions

ARTICLE 34 — Where a control envisaged under this Framework cannot be implemented for technical or operational reasons, the resulting risk shall be assessed and, where possible, an equivalent control shall be applied.

Exceptions involving significant risk shall be subject to the approval of the relevant product and management officers.

Related documents

ARTICLE 35 — The details concerning the implementation of this Framework shall be governed by the following supplementary documents:

The Next4biz Responsible AI Standard, the AI Risk Assessment Form, the AI System Card, the Customer AI Control Plan, the AI Incident Management Procedure, and product- or service-specific technical documents.

These documents shall be prepared in accordance with the principles established in the Framework and shall be updated independently as required.

Project-specific technical controls, performance targets, information sources, and usage limitations shall be defined in the relevant project document and in the Customer AI Control Plan. In the event of any discrepancy between this Framework and the contract or project documents, the provisions of the relevant contract shall prevail.

Entry into Force

ARTICLE 36 — This Framework shall enter into force on the date of its approval by

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.

We use cookies in accordance with legal regulations to improve our services and your experience on our site. Click 'Accept All' to allow all cookies, 'Reject All' to refuse optional cookies, or manage your preferences in settings.