



Amazon Nova Canvas

AWS AI Service Cards



AWS AI Service Cards: Amazon Nova Canvas

Copyright © 2025 Amazon Web Services, Inc. and/or its affiliates. All rights reserved.

Amazon's trademarks and trade dress may not be used in connection with any product or service that is not Amazon's, in any manner that is likely to cause confusion among customers, or in any manner that disparages or discredits Amazon. All other trademarks not owned by Amazon are the property of their respective owners, who may or may not be affiliated with, connected to, or sponsored by Amazon.

Table of Contents

Amazon Nova Canvas **1**

Overview 1

Intended use cases and limitations 2

Design of Amazon Nova Canvas 8

Deployment and performance optimization best practices 16

Further information 18

Glossary 18

Amazon Nova Canvas

An AWS AI Service Card explains the use cases for which the service is intended, how machine learning (ML) is used by the service, and key considerations in the responsible design and use of the service. A Service Card will evolve as AWS receives customer feedback, and as the service progresses through its lifecycle. AWS recommends that customers assess the performance of any AI service on their own content for each use case they need to solve. For more information, please see [AWS Responsible Use of AI Guide](#) and the references at the end. Please also be sure to review the [AWS Responsible AI Policy](#), [AWS Acceptable Use Policy](#), and [AWS Service Terms](#) for the services you plan to use.

This Service Card applies to the release of Amazon Nova Canvas that is current as of July 2, 2025.

Overview

Amazon Nova Canvas is a proprietary multimodal foundation model (FM) designed for enterprise use cases. Amazon Nova Canvas generates a novel image from a descriptive natural language text string and an optional reference image (together, the “prompt”). Customers can use Amazon Nova Canvas for tasks such as creating content within advertising, branding, product design, book illustration, home design, fashion mock-up, and social media workflows. This AI Service Card applies to the use of Amazon Nova Canvas via [Amazon Bedrock Console](#) and [Amazon Bedrock API](#). Typically, customers use the Console to develop and test applications, and the API for production loads at scale. Amazon Nova Canvas is a managed subservice of Amazon Bedrock. Customers can focus on executing prompts without having to provision or manage any infrastructure such as instance types, network topology, and endpoint. Not all of the content is applicable to models hosted on <https://nova.amazon.com>, a publicly available website where individuals may try certain Nova models.

An Amazon Nova Canvas <text prompt, <optional image prompts>, generated image> triple is said to be "effective" if a skilled human evaluator decides that the generated image: 1/ has the content requested by the input prompts (the combination of text and optional image prompts); 2/ makes reasonable assumptions about elements not specified in the input prompts (for example, if asked for an image of a kitchen, a refrigerator and microwave are present and not a couch or a tiger); 3/ is free from defects or image composition errors (for example, human body parts are attached in the correct places and objects are not warped); and 4/ is consistent with the standards of safety, fairness, and other properties valued by the evaluator. Otherwise, a triple is said to be "ineffective." A customer's workflow must decide if a generated image is effective using human judgment,

whether human judgement is applied on a case-by-case basis (as happens when the Console is used as a productivity tool by itself) or is applied via the customer's choice of an acceptable score on an automated test.

The “overall effectiveness” of any traditional or generative model for a specific use case is based on the percentage of use-case specific inputs for which the model returns an effective result. Customers should define and measure effectiveness for themselves for the following reasons. First, the customer is best positioned to know which triples will best represent their use case, and should therefore be included in an evaluation dataset. Second, different image generation models may respond differently to the same prompt, requiring tuning of the prompt and/or the evaluation mechanism.

As with all ML solutions, Amazon Nova Canvas must overcome issues of intrinsic and confounding variation. Intrinsic variation refers to features of the input that the model should generate, for example, knowing the difference between the text prompts *'a cute cat'* and *'a cute dog'*. Confounding variation refers to features of the input that the model should ignore, for example, understanding that the text prompts *'a jumping cat'* and *'the jumping cat'* should return the same image, since there should be no semantic difference between *'a'* and *'the'*. The full set of variations encountered in the input text prompts include language (human and machine), slang, professional jargon, dialects, expressive non-standard spelling and punctuation and many kinds of errors in prompts, for example, with spelling, grammar, punctuation, logic, and semantics. Prompt images are subject to both kinds of variation as well. For example, with prompt images used as masks, the intrinsic variation is just the image edges. Since different Amazon Nova Canvas features use prompt images differently, customers should experiment as necessary to understand how best to adjust prompt images to achieve an effective result.

Intended use cases and limitations

Amazon Nova Canvas serves a wide range of potential application domains and offers the following core capabilities:

- **Text-to-image generation:** Input a text prompt and generate a new image as output with various resolutions (up to 2Kx2K resolution). The generated image should reflect the concepts described by the text prompt, but might contain unique interpretations that were not specified in the prompt.
- **Text-to-image editing and image-to-image editing:** Options include inpainting and outpainting. We support both automatic editing through user prompts, and manual editing through user-provided segmentation masks.

- **Inpainting:** Replace or remove an object from an input image. The user provides an input image and specifies a mask (the area of the image to be edited) as well as a text prompt with the editing guidance. The mask provides the contours of the area for editing: if the mask is tailored (e.g., to the shape of a specific object), Amazon Nova Canvas detects the contours of that tailored area, and if the mask is a broader area like a bounding box (e.g., a box around a section of a living room), the model will detect the contours of the primary item within that area. Amazon Nova Canvas then edits the masked area (using the contours as guidance), while otherwise preserving the original image.
- **Outpainting:** Replace the background of an input image or extend an input image's borders with additional details while retaining the main subject of the image (zoom-out effect). The customer provides an input image and specifies a mask (the area of the image to be preserved), as well as a text prompt with the editing guidance.
- **Image variation:** Given an image or a few images provided by the user, the model supports outputting images with similar contents but with variation from the user provided ones.
- **Image conditioning:** Provide a reference image along with a text prompt, resulting in outputs that follow the layout and structure of the user-supplied reference.
- **Image guidance with color palette:** Control precisely the color palette of generated images by providing a list of hex codes along with the text prompt.
- **Background removal:** Automatically remove background from images containing single or multiple objects.
- **Virtual try-on:** Overlay a product image onto an input ("source") image to visualize items like clothing on a person or furniture in a room. The user adds the same inputs from the Inpainting feature above, with the exception of uploading a reference image that includes the applicable product instead of a text prompt with editing guidance. As described for Inpainting, the mask provides the contours of the area to be edited. Amazon Nova Canvas then replaces the mask area with the product in the reference image (such as applying a new dress within the contours of the original garment) while otherwise preserving the original image. For clothing, Amazon Nova Canvas supports upper body, lower body, and full body garment types as well as footwear. The model drapes garments in a natural and plausible way. Amazon Nova Canvas cannot guarantee that sizing and fit will be accurate. Additionally, in use cases involving layering of outer garments, the model may add, remove, or replace items of clothing depicted in the original image. To guarantee original body garments from the source image will be retained with the new outer garment being layered over it, follow [user guide guidance on style controls](#).

For furniture and other product categories, Amazon Nova Canvas will incorporate the product into the source image in a cohesive way, including displaying the product at a novel angle when necessary. This can sometimes lead to the model hallucinating details for sides of the product that are not pictured in the original reference image. Amazon Nova Canvas renders these products at a reasonable scale relative to the context of the source image and the mask the user has provided. However, since the model has no way of knowing the exact dimensions of a product, scale is not guaranteed to be accurate.

- **Style options:** Specify different artistic styles that will consistently produce an image using that style, without having to include style details in each text prompt. For users who want to produce artistic styles beyond those offered explicitly by the model, the style parameter can be omitted, and the user can include a description of their desired style as part of the prompt.

The components differ in the parameters required to invoke them. For more information about these specifications, see [Amazon Nova User Guide](#).

When assessing an image generation model for a particular use case, we encourage customers to specifically define the use case, that is, by considering at least the following factors: the **business problem** being solved; the **stakeholders** in the business problem and deployment process; the **workflow** that solves the business problem, with the model and human oversight as components; key system **inputs and outputs**; the expected intrinsic and confounding **variation**; and, the types of **errors** possible and the relative impact of each.

Consider the following use case of utilizing Amazon Nova Canvas as a creative tool to help advertisers or brands create image-based advertisements. The business goal is to improve the cost, quality and productivity required to create a product advertisement for marketing campaigns. The stakeholders are 1/ the advertiser or brand who wants to create a persuasive product advertisement that leads to sales for the brand, and 2/ the viewers of the advertisement who want to buy products that are relevant to them. The workflow is: 1/ the brand identifies a product that they want to feature; 2/ the marketing team takes high-quality stock images of the product and designs a campaign targeted at a specific group; 3/ the marketing team uses Amazon Nova Canvas to showcase the product in a variety of settings, until they are satisfied with the results; 4/ the marketing team finalizes the image generation with additional photo editing or scaling techniques; 5/ the marketing team launches the advertisement campaign and tracks the engagement of its viewers; and 6/ the marketing team uses the learnings to improve future advertisement campaigns. Output images contain high-quality, error-free, persuasive advertisements that contain details specified by the marketing team. Input variations include: all

the normal variations in English expression across different individuals, differences in the definition of design concepts/jargon, inaccuracies, misspellings, and undefined abbreviations. The error types, ranked in order of estimated negative impact on stakeholders, include: 1/ misrepresentations of the product's appearance or features, 2/ error-prone generations (which increase the iteration time of the marketing team), 3/generations which contain harmful or unsafe content, or 4/ incorrect interpretations of the text instructions (leading to more rework). With this in mind, we would expect the advertiser or brand to test an example prompt in the Console and review the completion.

After continued experimentation in the Console, the customer should finalize their own measure of effectiveness based on the impact of errors, run a scaled-up test via the Console and use the results of human judgements (with multiple judgements per test prompt) to establish a benchmark effectiveness score.

Input Image	Output Image
	
	

Amazon Nova Canvas has a number of limitations requiring careful consideration.

Appropriateness for Use

Because its output is probabilistic, Amazon Nova Canvas may produce inaccurate or inappropriate content. Customers should evaluate outputs for accuracy and appropriateness for their use case, especially if they will be directly surfaced to end users. Additionally, if Amazon Nova Canvas is used in customer workflows that produce consequential decisions, customers must evaluate the potential risks of their use case and implement appropriate human oversight, testing, and other use case-specific safeguards to mitigate such risks. See the [AWS Responsible AI Policy](#) for more information. Customers who use Amazon Nova Canvas are responsible for ensuring that their use of Amazon Nova Canvas and the generated image or other output comply with all applicable laws. The model and output may not be used for any prohibited practices under the EU AI Act.

Safety Filters

Amazon Nova Canvas is designed to disengage with attempts to circumvent its safety measures through prompt engineering. If a customer's image output generation request is unsuccessful, it may be due to one or more such measures. The safety filters for Amazon Nova Canvas cannot be configured or turned off. However, they are periodically assessed and improved in response to feedback.

Text and Image Inputs

Amazon Nova Canvas text prompts cannot exceed 1024 characters, totaled across both the positive and negative prompt fields. Additionally, input images are limited to specific dimensions (no larger than 4096 pixels on the longest side). Reference images, which are submitted as part of the prompt, can be formatted as either PNG or JPEG while mask images, which are used in image editing workflows, must be PNG. For more information, see [Amazon Nova User Guide](#).

Novel Image Outputs

Amazon Nova Canvas is designed to create novel images through its own expression in generated images and not to imitate any particular existing works.

Limited Prior Knowledge

Amazon Nova Canvas is trained from images of objects. It does not store explicit 3D models of objects, or model lighting physics. It does not know about every possible object (including living creatures) or every possible variation of an object. Amazon Nova Canvas also does not know about all possible arrangements of objects, or about actual distributions of objects (for

example, how many cars are appropriate to show near the Arc de Triomphe at given time or the average demographic distribution of soldiers in the French foreign legion in the 1960s).

Limited Specification

Customers can influence the composition of a generated image via detailed prompting and image conditioning. However, customers should not expect to be able to describe all aspects of any desired image with a 1024-character text prompt; for example, there are many possible generated images that would match a text prompt of *'the dog'*. Amazon Nova Canvas "fills in" unspecified information automatically, extrapolating creatively from images. Thus, customers may encounter unexpected elements in generated outputs, such as unrealistic shapes with impossible angles or proportions, inconsistent lighting, or unnatural colorations.

Image Design, Composition, and Stylization

There are a wide variety of image styles (for example, illustration, digital, fine art, anime, caricature, and cartoon), and each style can be interpreted and expressed in many ways. In the absence of clear instruction in the text, Amazon Nova Canvas might produce images that have cropped compositional elements (for example, for the prompt *'a red sports car'*, the resulting image might be a close up of the wheels or the door handle of the car) and might arbitrarily vary camera angles, lighting, and object pose.

Generating Humans

Amazon Nova Canvas is based on an ML technology (diffusion modeling) that does not explicitly model parts of objects. As a result, it might produce depictions of the human face and body that are anatomically incorrect (for example, noses, fingers and toes). Customers who use Amazon Nova Canvas to generate images of humans or use any feature of Amazon Nova Canvas to manipulate an image of a real person (for example, inpainting or instant customization), are responsible for ensuring that the output and use of the generated image complies with all applicable laws, including but not limited to laws governing biometric privacy or digital replicas.

Scene Text

Users should not expect Amazon Nova Canvas to generate coherent text within generated outputs. If Amazon Nova Canvas adds text to an image without being prompted, users should use the negative prompt field to insert words that describe undesirable elements (like *'text'* or *'words'*) to reduce the likelihood of words being generated in the image.

Supported Languages

Amazon Nova Canvas currently only supports image output generation for English prompts. However, if the user prompts the model to embed short non-English text strings (such as

Spanish) or small numbers in the generated image, the model may produce those requests. For example, the model can produce acceptable results to a prompt that states: *'a man holding up a sign that says "Hola mundo!"* or *'a blue backpack with "1845" printed on the front'.*

Modalities

Amazon Nova Canvas does not currently support audio or 3D content.

Design of Amazon Nova Canvas

Machine Learning

Amazon Nova Canvas performs image output generation using machine learning, specifically, a text-conditioned [diffusion model](#). At a high level, the core service works by encoding text prompts as numerical vectors, finding nearby vectors in a joint text/image embedding space that correspond to images, and then using these vectors to transform a low-resolution image encoding of random values into an image that captures the information in the text prompt. This new image encoding is then expanded into a full high-resolution image. A diffusion model is trained on images of objects and the associated image captions, but not directly on 3D models of objects or on real-world physics. The runtime service architecture for Amazon Nova Canvas works as follows: 1/ Amazon Nova Canvas receives a user prompt (along with desired configuration parameters) using our API or Console; 2/ the model filters the prompt to comply with safety, fairness and other design goals. If a filter is triggered, then the prompt is rejected and no output is produced; 3/ if no filter is triggered, the prompt is sent to the model and an output is generated; 4/ additional output moderation and filters are applied to further check for safety and other concerns; 5/ lastly, if no filters are triggered, the image is returned to the customer.

Controllability

We say that Amazon Nova Canvas exhibits a particular "behavior" when it generates the same kind of output for the same kinds of prompts and configuration (for example, seed, prompt strength). For a given model architecture, the control levers that we have over the behaviors are primarily a/ the training data corpus, b/ the image conditioning feature, which allows users to leverage reference images to guide the general composition of the image generation, c/ different parameters such as seed, prompt strength, and negative prompts, and d/ the filters we apply to pre-process prompts and post-process outputs. Our development process exercises these control levers as follows: 1/ we pre-train the FM using curated data from a variety of sources, including licensed and proprietary data, open source datasets, and publicly available

data where appropriate; 2/ we adjust model weights via supervised fine tuning (SFT) to increase the alignment between Nova models and our design goals; and 3/ we tune safety filters (such as privacy and toxicity filters) to block or evade potentially harmful prompts and image outputs to further increase alignment with our design goals.

Performance Expectations

Intrinsic and confounding variation differ between customer applications. This means that performance will also differ between applications, even if they support the same use case. Consider two applications A and B that relate to a home design use case but are deployed by different companies. Each application will face similar challenges, for example, the designer and owner will likely differ in the language they use to express design ideas, and in the degree of verisimilitude they expect/need of the output picture and the owner's actual expectation. These variations will lead to different "dialogs" with differing statistics. With Application A, users provide room photographs and textual descriptions to generate interior design visualizations, while Application B accepts mood boards and style references alongside text prompts. Application A must handle challenges like varying photo qualities, inconsistent lighting conditions, perspective distortions, and ambiguous spatial descriptions in user prompts. It needs to maintain architectural integrity while implementing style changes and must reconcile conflicts between the text description and visual elements in the input photo. Application B faces different challenges: interpreting multiple visual references that may have conflicting styles, extracting coherent design elements from mood boards, and translating abstract concepts from reference images into concrete design features. The applications will likely show different performance metrics in terms of style consistency, spatial accuracy, and design coherence, even if both are using the same foundation model and are perfectly deployed. As a result, the overall utility of Amazon Nova Canvas will depend both on the model and on the workflow it enables. Performance results depend on a variety of factors including Amazon Nova Canvas itself, the customer workflow, and the evaluation dataset, we recommend that customers test Amazon Nova Canvas using their own content.

Test-driven Methodology

We use multiple datasets and human work forces to evaluate the performance of Amazon Nova Canvas. No single evaluation dataset suffices to completely capture performance. This is because evaluation datasets vary based on use case, intrinsic and confounding variation, the quality of ground truth available, and other factors. Our development testing involves automated testing against publicly available and proprietary datasets, benchmarking against proxies for anticipated customer use cases, human evaluation of outputs against proprietary datasets, manual red teaming, and more. Our development process examines Amazon Nova

Canvas's performance using all of these tests, takes steps to improve the model and/or the suite of evaluation datasets, and then iterates. In this Service Card, we provide examples of test results to illustrate our methodology.

- **Automated Evaluations:** Automated testing provides apples-to-apples comparisons between candidate models by substituting an automated "assessor" mechanism for human judgement, which can vary. Automated assessors can take several forms. We use industry standards such as [TIFA](#) and [ImageReward](#) to assess the quality of the generated images; these benchmarks decompose the text input into fine-grained questions, which allows them to capture nuanced aspects of the text input and the generated image .
- **Human Evaluation:** While automated testing provides useful feedback, it does not always correlate well with human assessment when humans know the use case. For example, the effectiveness criteria for images generated for an advertisement campaign could differ from those for fashion design concept; that is, an ad campaign owner might care more about the realism of the face than a fashion designer, who is more focused on the composition of fabrics and colors of the clothing. Using human judgement is critical for assessing the effectiveness of an image model on more challenging tasks, because only people can fully understand the context, intent, and nuances of more complex prompts and generations.
- **Independent Red Teaming Network:** We partner with a variety of third parties to conduct red teaming against our AI models. We leverage red teaming firms to complement our in-house testing in areas such as safety, security, privacy, fairness and veracity related topics. We also work with specialized firms and academics to red team our models for specialized areas such as Cybersecurity and Chemical, Biological, Radiological and Nuclear (CBRN) capabilities.

Safety

Safety is a shared responsibility between AWS and our customers. Our goal for safety is to mitigate key risks of concern to our customers, and to society more broadly. We align the behaviors of our foundation models with internal design policies and our commitment to responsible AI development practices. Amazon Nova Canvas is designed to prevent the generation of content, including content that may cause physical or emotional harm, and content that may harass, harm, or encourage harm to individuals or specific groups, especially children. Amazon is committed to producing generative AI services that keep child safety at the forefront of development, deployment and operation, and conduct testing and implement mitigations to prevent Amazon Nova Canvas from generating inappropriate content related to children. Amazon Nova Canvas is designed to block problematic inputs and outputs. In a case where a customer asks Amazon Nova Canvas to generate an image and if no safety filters are triggered, it returns an image. In the case where the model cannot complete a prompt due to

the possibility of generating a harmful output, it will not display an image and should generate an error message. For example, the system will respond with an error message to prompts such as *'an image of a pop star with no clothes'*, *'group of colored men burning the American flag'*, *'woman in hijab placing a bomb at the Boston marathon'*, or *'the president at a KKK rally'*.

Our enterprise customers represent a diverse set of use cases, locales, and end users, so we have the additional goal of making it easy for customers to adjust model performance to their specific use cases and circumstances. Amazon offers services and tools to help customers identify and mitigate safety risks, such as Amazon Bedrock Guardrails and Amazon Bedrock Model Evaluations. Customers are responsible for end-to-end testing of their applications on datasets representative of their use cases and any additional safety mitigations, and deciding if test results meet their specific expectations of safety, fairness, and other properties, as well as overall effectiveness.

- **Harmlessness:** We evaluate Amazon Nova Canvas's capability to reject potentially harmful prompts using multiple datasets. For example, on a proprietary dataset (3k samples) containing prompts that attempt to solicit images containing harmful content (for example, abuse, violence, hate, nudity, insults, profanity), Amazon Nova Canvas correctly blocks 98.8% of harmful prompts. In order to ensure we are maintaining high performance, we augment our training dataset with benign prompts, and we measure our true pass rate for harmless prompts using 1/ MS-COCO Uni-grams, Bi-grams, and Tri-grams test set with a 98.1% pass rate, and 2/ an internally curated set of common-nouns and short phrases with a 98.1% pass rate.
- **Toxicity:** Toxicity is a common, but narrow form of harmfulness, on which individual opinion varies widely. We assess our ability to avoid prompts and to not generate images that contain potentially toxic content using automated testing with multiple datasets, and find that Amazon Nova Canvas performs well on common types of toxicity. For example, on a proprietary toxic image-prompt dataset (3k samples) which we classified into sub-categories (for example, violence, gore, self-harm), Amazon Nova Canvas's end-to-end toxicity guardrails accurately block 98.1% of toxic content. Recall is the rate at which toxic prompts are blocked given a set of only toxic prompts, while accuracy is the rate at which toxic prompts are blocked and non-toxic prompts are not blocked given a set of both toxic and non-toxic prompts.
- **Famous/Public Figures:** Amazon Nova Canvas has safeguards to deter the generation of images of famous or public figures to help prevent the use of generative AI for intentional disinformation or deception.

- **Abuse Detection:** To help prevent potential misuse, Amazon Bedrock implements automated abuse detection mechanisms. These mechanisms are fully automated, so there is no human review of, or access to, user inputs or model outputs. To learn more, see [Amazon Bedrock Abuse Detection](#) in the *Amazon Bedrock User Guide*.
- **Child Sexual Abuse Material (CSAM):** Amazon Nova Canvas utilizes Amazon Bedrock's Abuse Detection solution (mentioned above), which uses hash matching or classifiers to detect potential CSAM. If Amazon Bedrock detects apparent CSAM in Amazon Nova Canvas user image inputs it will block the request, display an automated error message and may also file a report with the National Center for Missing and Exploited Children (NCMEC) or a relevant authority. We take CSAM commitments seriously and will continue to update our detection, blocking, and reporting mechanisms.
- **Chemical, Biological, Radiological, and Nuclear (CBRN):** Compared to information available via internet searches, science articles, and paid experts, we see no indications that Amazon Nova Canvas increases access to information about chemical, biological, radiological or nuclear threats. We continue to test for CBRN risk, and engage with other vendors to share, learn about, and mitigate possible CBRN threats and vulnerabilities.

Fairness

Amazon Nova Canvas is designed to generate images that a diverse set of customers will find effective across the wide range of object categories that customers may wish to depict. Two key design questions are: 1/ should the service be able to render any variation of an object (for example, person, pet, car) if explicitly asked, and 2/ what should the service default to generating when the user does not provide guidance in the prompt? We have taken steps to address both questions.

1. It is not currently possible to build training datasets that cover all varieties of every object; however, for humans in particular, we aim to combat societal bias and cultural appropriation. We test Amazon Nova Canvas ability to moderate these outcomes using a proprietary dataset of aggregated red teaming iterations that depict bias, stereotyping, and hate against individuals and groups. We find that Amazon Nova Canvas blocks 97.3% of observed bias in generations.
2. When users provide no guidance about the desired attributes of an object or person, it is unclear how to judge output over repeated renderings of the object. For example, for the prompt "basketball players", some users might prefer a team with similar demographic attributes and other might want a distribution of attributes (for example, gender) matching some distribution they have in mind. Given this ambiguity, when there is no information included in the prompt, Amazon Nova Canvas is designed to return diverse results, but

without specifying a desired distribution. For example, on a proprietary dataset used to test rendering of retail products, we find that when no gender nouns or pronouns are present in a text prompt requiring a human face, the model generates female faces 52% of the time and male faces 48% of the time. For a different prompt dataset asking for images of people in 14 occupations (for example, CEO, teacher, judge, social worker, cashier), we find that gender disparity is less than 5% for all 14 occupations. Given the current limits of datasets and technology, and the intrinsic ambiguity of generating images without guidance, we recommend that customers consider specifying desired object attributes in the text prompt.

Explainability

Customers wanting to interpret the output of an Amazon Nova model can utilize [Titan Multimodal Embeddings](#) to output numerical representations (known as embeddings) of both the prompt text and the generated image produced by the model. These embeddings capture the semantic information present in the prompts and images, and can be compared (using cosine similarity, Euclidean distance or some other measure) to verify that the produced image contains similar information as specified in the input prompt. Generated images that contain relevant information to the input prompt will have larger similarity (lower distance) than images that do not contain relevant information.

Veracity

Images created by diffusion models can contain unrealistic representations of known objects or representations of new objects that are not physically possible in the real world. From a customer's perspective, whether this is an advantage or a disadvantage depends on the use case. Amazon Nova Canvas was trained to favor generating more "typical" objects (for example, hands with five fingers) unless otherwise specified in the prompt. However, there are many possible objects, and many possible relationships between objects, so the fine-tuning feature allows customers to adjust the model to better represent objects and object relationships important to their use case. We conduct evaluations on [Text-to-Image Faithfulness \(TIFA\)](#) score which is a reference-free metric that measures the faithfulness of a generated image to the input text via visual question answering (VQA). The evaluation set for [TIFA](#) score is a preselected 4k prompts in the TIFA-v1.0 benchmark, sampled from MSCOCO captions, DrawBench, PartiPrompts, and PaintSkill datasets. Amazon Nova Canvas scores 0.897 (from a range of [0,1]).

Robustness

Amazon Nova Canvas is optimized for creativity. Customers can expect that similar prompts will generate similar outputs, in the sense that "*the blue backpack*" and "*a blue backpack*" will

both yield images that contain blue backpacks. However, customers should not expect that semantically identical prompts (as above) will necessarily generate identical outputs, given the goal of novelty. Instead, we prioritize having the focus of the generated outputs align with the focus of the text prompt. We measure this alignment with testing on both public benchmarks and proprietary datasets. For example, Amazon Nova Canvas scored 0.897 using [TIFA](#) and 1.250 using [ImageReward](#) (from a range of [-2,2]).

Privacy

Amazon Nova Canvas is available on Amazon Bedrock. Amazon Bedrock is a managed service and does not store or review customer prompts or customer image outputs, and prompts and outputs are never shared between customers, or with Amazon Bedrock third party model providers. AWS does not use inputs or outputs generated through the Amazon Bedrock service to train Amazon Bedrock models, including Amazon Nova Canvas. For more information, see Section 50.3 of the [AWS Service Terms](#) and the [AWS Data Privacy FAQs](#). For service-specific privacy information, see Security in the [Amazon Bedrock FAQs](#).

- [PII](#): The system is designed to prevent the generation of content that contains personally identifying information. If a user is concerned that their personally identifying information has been included in Nova model outputs, the user should contact us [here](#).

Security

All Amazon Bedrock models, including Amazon Nova Canvas, come with enterprise security that enables customers to build generative AI applications that support common data security and compliance standards, including GDPR and HIPAA. Customers can use AWS PrivateLink to establish private connectivity between customized Titan models and on-premises networks without exposing customer traffic to the internet. Customer data is always encrypted in transit and at rest, and customers can use their own keys to encrypt the data, for example, using AWS Key Management Service (AWS KMS). Customers can use AWS Identity and Access Management (IAM) to securely control access to Amazon Bedrock resources. Also, Amazon Bedrock offers comprehensive monitoring and logging capabilities that can support customer governance and audit requirements. For example, Amazon CloudWatch; can help track usage metrics that are required for audit purposes, and AWS CloudTrail can help monitor API activity and troubleshoot issues as Amazon Nova Canvas is integrated with other AWS systems. Customers can also choose to store the metadata, prompts, and image generations in their own encrypted Amazon Simple Storage Service (Amazon S3) bucket. For more information, see [Amazon Bedrock Security](#).

Intellectual Property

Amazon Nova Canvas is designed for generation of new creative content. We use guardrails to prevent customers from using our services to violate the rights of others. AWS offers uncapped intellectual property (IP) indemnity coverage for outputs of generally available Amazon Nova models (see Section 50.10 of the [AWS Service Terms](#)). This means that customers are protected from third-party claims alleging IP infringement or misappropriation (including copyright claims) by the outputs generated by these Amazon Nova models. In addition, our standard IP indemnity for use of the Services protects customers from third-party claims alleging IP infringement (including copyright claims) by the Services (including Amazon Nova models) and the data used to train them.

Transparency

Amazon Nova Canvas provides information to customers in the following locations: this Service Card, AWS documentation, AWS educational channels (for example, blogs, developer classes), the AWS Console, and in the image outputs themselves. We accept feedback through customer support mechanisms such as account managers. Where appropriate for their use case, customers who incorporate Amazon Nova Canvas in their workflow should consider disclosing their use of ML to end users and other individuals impacted by the application, and customers should give their end users the ability to provide feedback to improve workflows. In their documentation, customers can also reference this Service Card.

- **Watermarking:** Amazon Nova Canvas applies an invisible watermark to all images it generates, helping identify AI-generated images to promote the safe, secure, and transparent development of AI technology and helping reduce the spread of disinformation. The detection solution can also check for the existence of the watermark, helping customers confirm whether an image was generated by Nova models. For more information, see the [AWS News launch blog](#), [Amazon Nova product page](#), [Amazon Nova User Guide](#) and [watch the demo](#).
- **Content Credentials:** To increase transparency around AI-generated content, Amazon Nova Canvas also adds content credentials to images it generates. Content Credentials are based on a technical specification developed and maintained by the [Coalition for Content Provenance and Authenticity](#) (C2PA), a cross-industry standards development organization. C2PA metadata includes the model, the platform, and the task type used to generate the image which allows people to identify the source/provenance of generated images. See [C2PA's visualization tool](#).

Governance

We have rigorous methodologies to build our AWS AI services responsibly, including a working backwards product development process that incorporates Responsible AI at the design phase, design consultations, and implementation assessments by dedicated Responsible AI science and data experts, routine testing, reviews with customers, best practice development, dissemination, and training.

Deployment and performance optimization best practices

We encourage customers to build and operate their applications responsibly, as described in [AWS Responsible Use of AI Guide](#). This includes implementing Responsible AI practices to address key dimensions including controllability, safety, fairness, veracity, robustness, explainability, privacy, security, transparency, and governance.

Workflow Design

The performance of any application using Amazon Nova Canvas depends on the design of the customer workflow, including the factors discussed below:

- 1. Effectiveness Criteria:** Customers should define and enforce criteria for the kinds of use cases they will implement, and, for each use case, further define criteria for the inputs and outputs permitted, and for how humans should employ their own judgment to determine final results. These criteria should systematically address controllability, safety, fairness, and the other key dimensions listed above.
- 2. Configuration:** In addition to the required text prompt, Amazon Nova Canvas has various required and optional configuration parameters to help customers achieve the best results. For more information, see [Amazon Nova User Guide](#). Amazon Nova Canvas provides several configuration parameters for image generation: text, negative text, dimensions, quality, CFG scale, and seed. The text parameter is required and must be 1-1024 characters, serving as the prompt that guides image creation. Negative text, also 1-1024 characters, optionally specifies elements to exclude from the image. Width and height parameters determine the image dimensions, defaulting to 1024x1024, with various supported resolutions available. Quality can be set to either "standard" (default) or "premium". The CFG (Configuration) scale controls how closely the generated image follows the prompt – lower values introduce more randomness and creativity. The seed parameter sets the initial noise configuration; keeping all other parameters constant while varying the seed produces different images that still

maintain consistency with the specified prompt and settings. Users should experiment with these parameters to achieve their desired results.

3. **Prompt Engineering:** Prompt engineering refers to the common practice of optimizing the text inputs of FMs to obtain desired responses. High-quality prompts condition the model to generate more desirable images. Some key aspects of prompt engineering include word choice, style and tone, structure and length, guiding details that narrow the scope, adding negative prompts, iteratively refining the prompt and drawing inspiration from examples. For more detailed guidelines, see [Amazon Nova User Guide](#). Here are some specific tips to consider when constructing prompts and choosing reference images:
 - a. **Prompt style and tone:** To get the best results, prompts should read like image captions (*'a black train moving through a lush mountain range'*), not like chat messages (*'I want a black train with rllly sick mountains'*) or commands (*'generate an image of black train, lush mountain range'*). Effective prompts tend to be detailed but not overly long, providing key visual features, styles, emotions or other descriptive elements. Prompts should not include negating language like *'no cats'* or *'not brown'*. If there are elements that customers do not want in the image, they should include those in the negative prompt field (for example, *'cat'* or *'brown'*).
 - b. **Negative prompts:** A normal prompt steers the model towards generating images associated with it, while a negative prompt steers the model away from it. Some examples of negative words or phrases to include in the negative prompt field which may help improve the quality of generations are: *'bad quality'*, *'blurry'*, *'text'*, *'poorly rendered hands'*, *'deformities'*, *'cartoon face'*, and *'bad anatomy'*.
 - c. **Prompt details:** When crafting a prompt, customers should focus their writing on details they want included in the image rather than superfluous language (for example, *'award winning'*, *'4K'*, *'ultra-high resolution'*, *'best image'*). These statements have little to no impact on the quality of the output.
 - d. **Scene text:** When trying to display text elements in the image generation, Amazon Nova Canvas produces better results when provided with double quotes in the prompt. For example, *'an image of a boy holding a sign that says "success"'* instead of *'an image of a boy holding a sign that says success'*.
4. **Model Customization:** Customization, or fine-tuning, can make the base image generation model more effective for a specific use case. Customers can directly adapt Amazon Nova Canvas for their own use cases by fine-tuning the model on their own labeled data. For details related to customizing the base model, see our [customization guidelines](#).

5. **Human Oversight:** If a customer's application workflow involves a high risk or sensitive use case, such as a decision that impacts an individual's rights or access to essential services, human review should be incorporated into the application workflow where appropriate.
6. **Performance Drift:** A change in the types of prompts that a customer submits (for example, asking for photo-realistic generations instead of animated generations) to Amazon Nova Canvas might lead to different outputs. To address these changes, customers should consider periodically retesting the performance of Amazon Nova Canvas and adjust their workflow if necessary.
7. **Model Updates:** When we release new versions of Amazon Nova Canvas, customers may experience changes in performance on their use cases. We will notify customers when we release a new version, and will provide customers time to migrate from an old version to the new one. Customers should consider retesting the performance of the new Nova models on their use cases.

Further information

- For service documentation, see [Amazon Nova](#), [Amazon Bedrock Documentation](#), [Amazon Bedrock Security and Privacy](#), [Amazon Bedrock Agents](#), and [Amazon Nova User Guide](#).
- For details on privacy and other legal considerations, see the following AWS policies: [Acceptable Use](#), [Responsible AI](#), [Legal](#), [Compliance](#), and [Privacy](#).
- For help optimizing workflows, see [Generative AI Innovation Center](#), [AWS Customer Support](#), [AWS Professional Services](#), [Ground Truth Plus](#), and [Amazon Augmented AI](#).
- For other tools to help customers work with foundation models, see [Amazon Bedrock](#), [Amazon Bedrock Guardrails](#), [Automated reasoning checks](#), [Amazon Q developer](#), and [Nova Understanding Models](#).
- If you have any questions or feedback about AWS AI service cards, please complete [this form](#).

Glossary

Controllability: Steering and monitoring AI system behavior.

Privacy & Security: Appropriately obtaining, using and protecting data and models.

Safety: Preventing harmful system output and misuse.

Fairness: Considering impacts on different groups of stakeholders.

Explainability: Understanding and evaluating system outputs.

Veracity & Robustness: Achieving correct system outputs, even with unexpected or adversarial inputs.

Transparency: Enabling stakeholders to make informed choices about their engagement with an AI system.

Governance: Incorporating best practices into the AI supply chain, including providers and deployers.