

Model Documentation Form

This Form includes all the information to be documented as part of Measure 1.1 of the Transparency Chapter of the Code of Practice. Crosses on the right indicate whether the information documented is intended for the AI Office (AIO), national competent authorities (NCAs) or downstream providers (DPs), namely providers of AI systems who intend to integrate the general-purpose AI model into their AI systems. Whilst information intended for DPs should be made available to them proactively, information intended for the AIO or NCAs is only to be made available following a request from the AIO, either ex officio or based on a request to the AIO from NCAs. Such requests will state the legal basis and purpose of the request and will concern only items from the Form strictly necessary for the AIO to fulfil its tasks under the AI Act at the time of the request, or for NCAs to exercise their supervisory tasks under the AI Act at the time of the request, in particular to assess compliance of high-risk AI systems built on general-purpose AI models where the provider of the system is different from the provider of the model.

Any elements of information from the Model Documentation Form shared with the AIO and NCAs shall be treated in accordance with the confidentiality obligations and trade secret protections set out in Article 78 AI Act.

Date this document was last updated: 21/05/2026

Document version number: v1

General information						AIO	NCAs	DPs
Legal name for the model provider:	Nasjonalbiblioteket					☒	☒	☒
Model name:	borealis-open-27b					☒	☒	☒
Model authenticity:	We provide a package that proves that Nasjonalbiblioteket has signed the files and that Hellenic Academic and Research Institutions CA have approved that we are indeed Nasjonalbiblioteket. Verification can be done following https://ai.nb.no/verify . f4b9f176eb6839f67065f09ce880a28854ff3c71a9f0686a8fe3b06fc5b01021 model-00001-of-00002.safetensors 9a4ce8b7225d05a07b7d4fc3175d43e7179b3eb0dc61c7c2069ce8bb02c27195 model-00002-of-00002.safetensors					☒	☒	☐
Release date:	01/06/2026 Date when the model was first released through any distribution channel.					☒	☒	☒
Union market release date:	01/06/2026 Date when the model was placed on the Union market.					☒	☒	☒
Model dependencies:	google/gemma-3-27b-it					☒	☒	☒
Model properties						AIO	NCAs	DPs
Model architecture:	The borealis-open-27b is a 27b-parameter instruction-tuned model. This model is based on google/gemma-3-27b-it, and fine-tuned on textual instructions only.					☒	☒	☒
Design specifications of the model:	This model is a supervised fine-tuned adaptation of the Gemma 3 27B decoder-only Transformer architecture, trained on the aurora-sft-open dataset.					☒	☒	☐
Input modalities:	☒ Text	☒ Images	☐ Audio	☐ Video	If any other please specify:	☒	☒	☒
For each selected modality please include maximum input size or write 'N/A' if not defined	Maximum size: 4096 tokens	Maximum size: 896x896 pixels	Maximum size: N/A	Maximum size: N/A	Maximum size: N/A	☒	☐	☒
Output modalities:	☒ Text	☐ Images	☐ Audio	☐ Video	If any other please specify:	☒	☒	☒
For each selected modality please include maximum output size or write 'N/A' if not defined	Maximum size: 8192 tokens	Maximum size: N/A	Maximum size: N/A	Maximum size: N/A	Maximum size: N/A	☐	☐	☒
Total model size:	27 Billion					☒	☐	☐
The range within which the total number of parameters falls.	☐ 1—500M	☐ 500M—5B	☐ 5B—15B	☒ 15B—50B		☐	☒	☒
	☐ 50B—100B	☐ 100B—500B	☐ 500B—1T	☐ >1T				
Methods of distribution and licenses						AIO	NCAs	DPs
Distribution channels:	The model is available on HuggingFace at https://huggingface.co/NbAiLab/borealis-open-27b					☒	☒	☐
	The model is available on HuggingFace at https://huggingface.co/NbAiLab/borealis-open-27b					☐	☐	☒

License:		Gemma (https://ai.google.dev/gemma/terms)				☒	☒	☐	
						☐	☐	☒	
						☒	☐	☒	
Use						AIO	NCA's	DPs	
Acceptable Use Policy:		- The model may hallucinate or produce incorrect information. - Safety alignment reduces but does not eliminate harmful or inappropriate outputs. - Performance outside Norwegian and English use cases has not been fully characterized. As the model is a derivative of Gemma 3, the same use restrictions apply: - Illegal Acts & Violence: Generating content that violates laws, promotes violence, or details how to build weapons. - Severe Personal Harm: Facilitating suicide, self-harm, child exploitation, human trafficking, or cyberattacks. - Malicious Cyber Activities: Writing malware, automating spam, or engineering phishing campaigns. - Hate Speech & Harassment: Generating slurs, organizing targeted bullying, or inciting hatred against protected groups. - Deception & Impersonation: Deploying scams, committing fraud, or misleadingly pretending to be a real person. - Explicit Content: Generating sexually explicit text, pornography, or unauthorized adult chatbots. - Regulated Advice: Providing automated medical, financial, or legal advice without certified oversight. - Critical Decisions: Making automated judgments on high-stakes matters like credit scores, housing, or hiring.				☒	☒	☒	
Intended uses:		- Norwegian-centric assistant-style tasks (e.g., drafting, summarization, Q&A, light reasoning). - Assessment of Norwegian writing style and quality. - Early evaluation of behavior, language coverage (Norwegian / Bokmål / Nynorsk / Sami), and quality.				☒	☒	☒	
Type and nature of AI systems in which the general-purpose AI model can be integrated:		This model is intended for integration into Norwegian-centric conversational assistants, drafting and summarization tools, question-answering and light reasoning systems, and language quality assessment applications, with a focus on evaluating behavior, writing quality, and language coverage across Norwegian Bokmål, Nynorsk, and Sami in human-supervised workflows.				☒	☒	☒	
Technical means for model integration:		The model is available on HuggingFace and it is also availabe in quantized formats for llama.cpp and ollama and MLX formats for Apple devices.				☐	☐	☒	
Required hardware:		GPU VRAM 24 GB, System RAM 64 GB				☐	☐	☒	
Required software:		Transformers: https://github.com/huggingface/transformers@v4.49.0-Gemma-3				☐	☐	☒	
Training process						AIO	NCA's	DPs	
Design specifications of the training process:		The model is a supervised fine-tuned version of the Gemma 3 27B initialized from google/gemma-3-27b-it and trained using full-parameter fine-tuning on Norwegian instruction datasets. Training uses cosine learning rate scheduling (5e-6 initial LR), and a 0.1 warmup ratio over 3 epochs, with sequences up to 4096 tokens.				☒	☒	☐	
Decision rationale:		The model is an adaption of Gemma 3 to Norwegian contexts for instruction-following behaviour. SFT seems sufficient for a first release.				☒	☒	☐	
Information on the data used for training, testing, and validation						AIO	NCA's	DPs	
Data type/modality: <i>Select all that apply.</i>		☒ Text	☐ Images	☐ Audio	☐ Video	If any other please specify:	☒	☒	☒
Data provenance: <i>Select all that apply</i>		☒ Web crawling		☐ Private non-publicly available datasets obtained from third parties		☐ User data	☒	☒	☒
For definitions of each listed category, see the Template for the Public Summary of the Training Content of General-Purpose AI models provided by the AI Office		☒ Publicly available datasets				☐ Data collected through other means	☒	☒	☒
		☒ Synthetic data that is not publicly accessible (when created directly by or on behalf of the provider)				If any other please specify:			

How data was obtained and selected:	The model was trained on NbAiLab/aurora-sft-open. NbAiLab/aurora-sft-open consist of a machine-translated version of English instruction data from allenai/tulu-3-sft-mixture, a curated instruction dataset developed as part of the Mimir project, xsample/tulu-3-mig-50k, a parallel corpus for Bokmål–Nynorsk translation, a dataset containing of single and multi-turn translation conversations constructed between several Sámi languages and Norwegian Bokmål, syntetic data using deepseek-ai/DeepSeek-V3.2.	☒	☒	☐
Number of data points:	The dataset consists of 882 263 documents and 674 704 598 tokens (gemma3 tokenizer).	☐	☒	☐
	The dataset consists of 882 263 documents and 674 704 598 tokens (gemma3 tokenizer).	☒	☐	☐
Scope and main characteristics:	The training, validation, and test data consist of high-quality supervised instruction datasets focused on Norwegian (Bokmål, Nynorsk), Sámi languages, and English. The data covers conversational AI, instruction-following, translation, summarization, and news-related generation tasks such as title and ingress creation. The dataset is split into ~95% training, 2.5% validation, and 2.5% test.	☒	☒	☐
Data curation methodologies:	Strict filtering based on content and structural quality scores using deepseek-ai/DeepSeek-V3.2 to ensure high-quality samples for fine-tuning.	☒	☒	☒
Measures to detect unsuitability of data sources:	The data used for the training consists in its majority of synthetic data automatically assessed by existing frontier LLMs, or humanly supervised machine translations of heavy-filtered existing datasets in English. CSAM filters were also applied to filter out toxic and unsuitable content.	☒	☒	☐
Measures to detect identifiable biases:	We manually verified a subset of translated and synthetically generated samples to inspect biases. We removed biases against model self-identification and political orientation.	☒	☒	☐
Computational resources (during training)		AIO	NCA s	DP s
Training time:	Less than 1 month	☐	☒	☐
	6.671 days wall clock time; 53.368 NVIDIA H200 GPU-days (8×6.671 days)	☒	☐	☐
Amount of computation used for training:	10 ²⁰ FLOPs	☐	☒	☐
	2.897 × 10 ²⁰ FLOPs	☒	☐	☐
Measurement methodology:	We estimated the number of FLOPS per token using the OpenAI scaling law paper. We then multiply that number by the number of tokens trained on.	☒	☒	☐
Energy consumption (during training and inference)		AIO	NCA s	DP s
Amount of energy used for training:	4.883 × 10 ⁻¹ MWh	☒	☒	☐
Measurement methodology:	We train our model with Weights & Biases as logging system, this logs the power usage during training. We take the average power usage and then multiply by the compute time.	☒	☒	☐
Benchmarked amount of computation used for inference¹:	Model runs comfortably on a 64-80 GB VRAM.	☒	☒	☐
	Measured using the vLLM inference runtime.	☒	☒	☐

¹ This item relates to energy consumption during inference, which makes up the “energy consumption of the model” (Annex XI, 2(e), AI Act) together with energy consumption during training. Since energy consumption during inference depends on more than just the model itself, the information required for this item is limited to relevant information depending only on the model, namely computational resources used for inference.