

Generative AI and Foundation Models in Medical Image

Masahiro Oda^{1,2*}

¹Scholarly Information Division, Information Technology Center,
Nagoya University, Furo-cho, Chikusa-ku, Nagoya 464-8601, Japan.

²Graduate School of Informatics, Nagoya University, Furo-cho,
Chikusa-ku, Nagoya 464-8601, Japan.

Corresponding author(s). E-mail(s): moda@mori.m.is.nagoya-u.ac.jp,
oda.masahiro.h6@f.mail.nagoya-u.ac.jp;

Abstract

In recent years, generative AI has attracted significant public attention, and its use has been rapidly expanding across a wide range of domains. From creative tasks such as text summarization, idea generation, and source code generation, to the streamlining of medical support tasks like diagnostic report generation and summarization, AI is now deeply involved in many areas. Today's breadth of AI applications is clearly distinct from what was seen before generative AI gained widespread recognition. Representative generative AI services include DALL·E 3 (OpenAI, California, USA) and Stable Diffusion (Stability AI, London, England, UK) for image generation, ChatGPT (OpenAI, California, USA), and Gemini (Google, California, USA) for text generation. The rise of generative AI has been influenced by advances in deep learning models and the scaling up of data, models, and computational resources based on the scaling laws. Moreover, the emergence of foundation models, which are trained on large-scale datasets and possess general-purpose knowledge applicable to various downstream tasks, is creating a new paradigm in AI development. These shifts brought about by generative AI and foundation models also profoundly impact medical image processing, fundamentally changing the framework for AI development in healthcare. This paper provides an overview of diffusion models used in image generation AI and large language models (LLMs) used in text generation AI, and introduces their applications in medical support. This paper also discusses foundation models, which are gaining attention alongside generative AI, including their construction

32 methods and applications in the medical field. Finally, the paper explores how
33 to develop foundation models and high-performance AI for medical support by
34 fully utilizing national data and computational resources.

35 **Keywords:** Generative AI, Diffusion model, Large language model, Foundation
36 model, National AI model

37 1 Introduction

38 In November 2022, ChatGPT (OpenAI, California, USA), powered by GPT-3.5 model,
39 was released. Due to its significantly improved performance compared to previous
40 chatbot AI models, it quickly gained global attention. Around the same time, Stable
41 Diffusion emerged, enabling high-quality image generation and signaling a technolog-
42 ical transformation in image generation. Since then, generative AI has become widely
43 recognized, leading to numerous research projects and commercial services leveraging
44 this technology. Today, generative AI is used in a wide range of applications, from
45 creative tasks such as text generation, summarization, idea generation, illustration,
46 digital art, and source code generation to the streamlining of medical support tasks
47 like diagnostic report generation and summarization. The breadth of generative AI
48 applications significantly departs from the pre-generative AI era. Notable generative
49 AI services include DALL-E 3 (OpenAI, California, USA), Stable Diffusion (Stability
50 AI, London, England, UK) for image generation, and ChatGPT and Gemini (Google,
51 California, USA) for text generation. The rapid rise of generative AI has been driven by
52 advancements in deep learning models and changes in training methodologies. These
53 innovations can potentially reshape the framework for AI development in medical
54 image processing.

55 This paper provides an overview of diffusion models and large language models
56 (LLMs) used in generative AI development, along with examples of their applications
57 in medical image processing. Additionally, this paper discusses foundational models,

58 which have garnered attention with generative AIs, and explores new AI development
59 approaches utilizing these models. Finally, this paper examines strategies to leverage
60 national computational and informational resources to develop internationally compet-
61 itive, nationally produced foundational models and high-performance AI for medical
62 image processing.

63 **2 Discriminative AI and generative AI**

64 The introduction of AlexNet [1] in 2012 marked the beginning of active research
65 in image processing using deep learning. At that time, most deep learning models
66 took data such as images as input and produced classification or detection results
67 as output. These models can be categorized as “Discriminative AI.” The main tasks
68 performed by Discriminative AI include image classification, object detection, and
69 segmentation. Among these, segmentation is a process that performs pixel-wise classi-
70 fication. A key characteristic of Discriminative AI is that the output data contains less
71 information than the input data. Additionally, models such as Convolutional Neural
72 Networks (CNNs) and Fully Convolutional Networks (FCNs) are commonly used in
73 its architecture.

74 Generative AI, on the other hand, takes inputs such as text and produces outputs
75 such as text or images using a generation model. A defining feature of Generative AI
76 is its ability to generate high-information outputs, such as long text or images, from
77 relatively low-information inputs, such as short sentences. Recent Generative AI mod-
78 els often adopt Diffusion Models or large-scale Transformers [2] as their architecture,
79 although some Generative AI models still use CNNs or FCNs. Details of diffusion
80 models are explained in 3.2.

81 A common way to interact with Generative AI is by providing instructions in the
82 form of text, leading to the widespread use of the term “prompt” to describe the
83 input given to Generative AI. In the past, due to AI’s limited capabilities, it was

84 seen as merely a machine that followed human instructions. However, with recent
85 advancements in AI performance, Generative AI is increasingly perceived as an entity
86 capable of intelligent reasoning, which has led to the adoption of the term “prompt,”
87 like how humans give instructions to one another.

88 The following chapters will discuss image generation AI and text generation AI,
89 both of which fall under the category of Generative AI.

90 **3 Image generation AI**

91 **3.1 Overview**

92 Previously, Variational Autoencoders (VAEs) [3] and Generative Adversarial Networks
93 (GANs) [4] were commonly used for image generation. VAEs offer stable training
94 and make it easier to evaluate training progress, but they struggle to create highly
95 expressive models. On the other hand, GANs can generate highly detailed images, as
96 seen in StyleGAN2 [5], which has demonstrated impressive facial image generation.
97 However, GANs suffer from unstable training and difficulty in evaluating training
98 progress.

99 Recent image generation services like DALL·E 3, Stable Diffusion, and Midjourney
100 can produce more visually appealing illustrations and photorealistic images than VAEs
101 or GANs. These services utilize diffusion models for image generation.

102 This chapter discusses the mechanism of diffusion models and their applications
103 in medical image processing.

104 **3.2 Diffusion models**

105 This section provides a brief explanation of diffusion models. Diffusion models are
106 inspired by thermodynamics, where the concentration of substances, temperature,
107 and energy differences tend to equalize over time. This phenomenon occurs because
108 molecules within a substance undergo random motion due to thermal energy, gradually

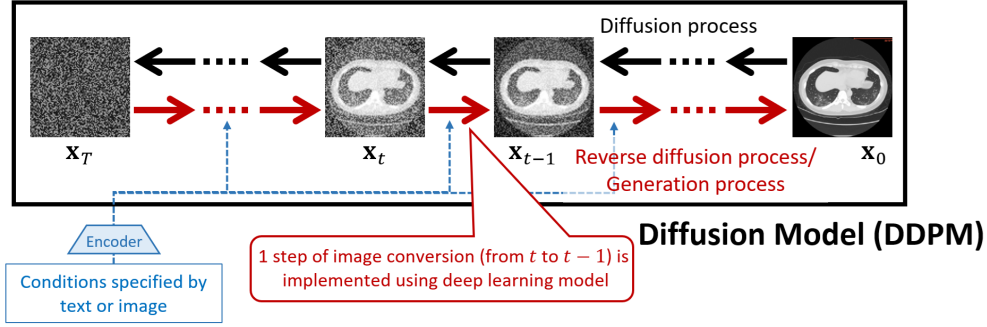


Fig. 1 Conceptual diagram of task-specific AI development using foundation models

109 mixing until they achieve a uniform distribution. The Diffusion Probabilistic Model
 110 (DPM) [6], which forms the foundation of diffusion models, was developed based on
 111 this diffusion process. In DPM, noise is gradually added to the data distribution,
 112 transforming it into a completely noisy distribution (diffusion process). The model then
 113 reconstructs the data distribution by reversing this process (reverse diffusion process).
 114 Please note that the assumption of the existence of reverse transformations of the
 115 diffusion process in DPM is different from the diffusion process in thermodynamics.

116 For image generation, a commonly used variant of DPM is the Denoising Diffusion
 117 Probabilistic Model (DDPM) [7]. In DDPM, an image undergoes a diffusion process,
 118 where gradual noise is added until the image is completely transformed into pure noise.
 119 Then, through the reverse diffusion process, the model progressively removes the noise
 120 to reconstruct an image. Since this reverse diffusion process creates an image, it is also
 121 called the generation process. This process can be understood as a transformation from
 122 a noise distribution to the target image distribution. Figure 1 illustrates the diffusion
 123 and reverse diffusion processes in DDPM.

124 In the implementation of DDPM, the reverse diffusion process, removing noise and
 125 reconstructing an image, is typically performed using a deep learning model such as
 126 U-Net [8]. The U-Net is an encoder-decoder style FCN model, which is commonly used
 127 to perform medical image segmentation. This model takes a noisy image as input and
 128 outputs an image with slightly reduced noise. The final generated image is obtained

129 by repeatedly applying this noise reduction process. Since the deep learning model
130 needs to reduce noise at each iteration appropriately, it receives both the step count
131 and generation conditions as input. The reverse diffusion process typically involves
132 about 1,000 steps to generate high-quality images.

133 In image generation, text or other images are often used to control the output. This
134 control is achieved by providing embedding representations of generation conditions
135 to the deep learning model at each step of the reverse diffusion process. For exam-
136 ple, in DALL-E 3, which generates images based on text descriptions, a Transformer
137 converts text into an embedding representation, which is then used as a generation
138 condition in the reverse diffusion process. Similarly, in diffusion model-based segmen-
139 tation methods, images are converted into embedding representations and used as
140 generation conditions.

141 3.3 Medical image generation using diffusion models

142 Many methods utilizing diffusion models have emerged in medical image processing,
143 and this section, along with the following sections, introduces some of the most rep-
144 resentative examples. Soon after diffusion models gained attention, medical image
145 generation techniques using diffusion models were proposed.

146 Methods for generating 2D brain MR images [9], 3D brain MR images [10], and 4D
147 cardiac MR images [11] using diffusion models have been introduced, and these papers
148 indicate that they achieve better results than GAN-based image generation. Addi-
149 tionally, diffusion models have been applied to generate laparoscopic images during
150 surgery [12, 13]. These laparoscopic image generation methods introduced techniques
151 to control the presence of surgical instruments and organs in the generated images.
152 Zhou et al. [12] used text descriptions as generation conditions, while Venkatesh et al.
153 [13] used label images. In these approaches, text or label images are first converted

154 into embedding representations, which are then fed into the diffusion model for image
155 generation.

156 Here, this paper examines how generated medical images can benefit computer-
157 aided medical applications. While diffusion models can generate realistic medical
158 images, these generated images are not actual patient images and, therefore, cannot
159 be used directly for diagnostic or therapeutic support. However, this does not mean
160 that generated medical images are useless. Image generation using diffusion models
161 allows the incorporation of text descriptions, label images, and other generation con-
162 ditions, making it possible to produce images that accurately reflect real anatomical
163 structures, diseases, and intraoperative conditions when appropriately controlled. As
164 a result, generated medical images play a valuable role in data augmentation for train-
165 ing deep learning models used in diagnostic and therapeutic support. Compared to
166 traditional data augmentation techniques such as geometric transformations, color
167 transformations, and Mixup [14], which have been widely used, image generation using
168 diffusion models can create more realistic and diverse variations in datasets.

169 A study [15] explored the use of diffusion-generated images for training deep-
170 learning models in mammogram-based diagnostic support. This study used Stable
171 Diffusion to generate mammogram images with embedded tumors. During image gen-
172 eration, text-based prompts were used to control the position and size of the embedded
173 tumors. By incorporating these generated images into the training of diagnostic sup-
174 port deep-learning models, the study successfully developed a higher-performance
175 model compared to training with only real images. Thus, generated images using dif-
176 fusion models can serve as an advanced data augmentation technique, significantly
177 enhancing the effectiveness of medical AI applications.

178 3.4 Medical image segmentation using diffusion models

179 One of the most effective applications of diffusion models in medical image processing
180 is segmentation. Among the early methods utilizing diffusion models for segmenta-
181 tion, MedSegDiff [16] was one of the first proposed approaches. This method generates
182 segmentation result images through the reverse diffusion process, where embedding
183 representations extracted from medical images are provided as generation conditions.
184 This process enables the generation of segmentation results as images corresponding
185 to the input medical images. Compared to previously proposed FCN-based methods,
186 such as SegNet and nnU-Net, and FCN+Vision Transformer (ViT) approaches, such
187 as TransUNet, MedSegDiff demonstrated higher segmentation accuracy. In addition to
188 MedSegDiff and its improved version, MedSegDiff-V2 [17], many diffusion model-based
189 segmentation methods have emerged and gained attention at MICCAI, a leading inter-
190 national conference on computer-aided medical applications. This section introduces
191 several diffusion model-based medical image segmentation methods.

192 The reverse diffusion process used for segmentation in diffusion models can be
193 implemented using a U-Net-based model. However, U-Net may sometimes fail to gen-
194 erate images that faithfully correspond to the medical images given as generation
195 conditions. To address this issue, Chowdary et al. [18] proposed a Multi-sized Trans-
196 former U-Net (MT U-Net), a deep learning model designed explicitly for diffusion
197 model-based segmentation. In diffusion models, the deep learning model takes two
198 inputs: the output image from the previous reverse diffusion step and the original med-
199 ical image. MT U-Net enhances image generation by extracting effective features from
200 these two inputs using Cross Attention between image feature representations. Addi-
201 tionally, it incorporates a Multi-size Transformer module, which utilizes information
202 from various spatial positions and scales for image generation. These enhancements
203 enabled MT U-Net to achieve better segmentation results than traditional FCN-based
204 methods, FCN+ViT-based methods, and MedSegDiff and MedSegDiff-V2.

205 Other segmentation approaches have improved accuracy by refining the noise-
206 generation process in diffusion models. Typically, diffusion models add Gaussian-
207 distributed noise to images. However, Chen et al. [19] suggested that, for segmentation
208 tasks that involve binary images, it is more effective to use Bernoulli-distributed noise,
209 which takes only two discrete values. Their approach, which incorporated Bernoulli
210 noise into a diffusion model for segmentation, demonstrated superior results when
211 applied to MRI and CT images. Their method outperformed traditional FCN-based
212 and ViT-based approaches and diffusion model-based methods that use Gaussian
213 noise.

214 Some medical images accompany text descriptions of their contents. When train-
215 ing segmentation models, utilizing ground-truth label images and textual information
216 enables efficient learning even with limited training data. A segmentation method pro-
217 posed based on this idea is TextDiff [20]. The TextDiff has a diffusion model-based
218 image encoder and a Clinical BioBERT as a text encoder. It combines embed-
219 ding representations from both modalities using cross-modal attention. The resulting
220 multi-scale attention map is upsampled and integrated to obtain the final segmenta-
221 tion output. TextDiff uses pre-trained image and text encoders to improve learning
222 efficiency while training only the cross-modal attention module and subsequent hid-
223 den layers. This enables multi-modal segmentation models to be trained effectively
224 even with limited data. This approach successfully developed a high-performance
225 segmentation model using a small dataset of just 150 X-ray images.

226 Diffusion models are also being applied to video-based segmentation. Lu et al.
227 proposed Diff-VPS [21], a method for polyp segmentation in colonoscopy videos.
228 This approach improves diffusion model-based segmentation accuracy by incorporat-
229 ing multi-task learning and temporal information. The Diff-VPS introduces two key
230 innovations. One is the multi-task learning. The model solves image classification
231 and object detection as additional subtasks within the segmentation diffusion model,

232 enabling it to utilize high-level contextual information. The other one is the temporal
233 reasoning module. To incorporate temporal information from video frames, this mod-
234 ule is trained using a task that estimates the target frame based on previous frames.
235 This training method allows the model to extract and leverage temporal information
236 in the diffusion-based segmentation process. Experimental results demonstrated that
237 Diff-VPS achieved higher segmentation accuracy than conventional methods for polyp
238 segmentation in colonoscopy videos.

239 **4 Text generation AI**

240 **4.1 Overview**

241 Text generation AI is designed to generate text based on given conditions. Notable
242 commercial services utilizing text generation AI include ChatGPT and Gemini. Chat-
243 GPT, introduced in November 2022, attracted worldwide attention because it could
244 generate text more naturally and human-likely compared to previous text generation
245 AIs.

246 This section briefly introduces natural language processing (NLP), which underlies
247 text generation AI, and discusses LLMs, which have significantly focused on NLP
248 research in recent years.

249 **4.2 Advances in NLP and emergence of LLMs**

250 NLP aims to enable computers to process natural language, which humans use daily,
251 and solve various language-related tasks. The core tasks in NLP include text classifi-
252 cation and generation, and research has been conducted to enhance these capabilities.
253 Some practical applications of NLP include machine translation, kana-kanji conver-
254 sion, search engines, and dialogue systems, many of which had already been studied
255 and commercially applied even before ChatGPT emerged.

256 Similar to image processing, NLP models have increasingly adopted deep learning
257 techniques. To handle sequential data such as text, early models relied on Recurrent
258 Neural Networks (RNNs) and Long Short-Term Memory (LSTM) networks [22], which
259 improved RNNs’ ability to process long sequences. However, recent high-performance
260 text generation AI models predominantly utilize Transformer-based NLP architec-
261 tures [23]. The Transformer model first converts extracted morphemes from text into
262 embedding representations, then applies Multi-Head Self-Attention to learn the rela-
263 tionships between embeddings, ultimately generating the output. This mechanism has
264 also been adapted for image processing in models such as ViT [24].

265 The performance of Transformer-based natural language processing models has
266 improved significantly, leading to their widespread adoption in commercial services.
267 OpenAI’s ChatGPT uses the Generative Pre-trained Transformer (GPT), while
268 Google’s Gemini is based on PaLM 2 and a Transformer model also named Gemini.
269 Additionally, Meta’s Llama is another example of a Transformer-based model. The
270 rapid performance gains of these NLP models are driven by the scaling up of data,
271 model size, and computational resources.

272 In 2020, OpenAI published research on Scaling Laws for Neural Language Models
273 [25], reporting that the performance of language models improves as the amount of
274 training data, the number of model parameters, and the computational resources used
275 during training increase. Based on these findings, OpenAI developed the GPT mod-
276 els using massive datasets collected from the internet, Transformer-based models with
277 many parameters, and extensive GPU resources. The first GPT model (GPT-1), intro-
278 duced in 2018, was trained on 4.5GB of text data, had 117 million parameters, and
279 was trained using eight GPUs [26]. OpenAI rapidly expanded the scale of its models,
280 leading to GPT-2 in 2019, which utilized 40GB of text data and had 1.5 billion param-
281 eters. GPT-3, released in 2020, was trained on 570GB of text data, had 175 billion

Table 1 LLMs and numbers of model parameters, scales of training data, and number of GPUs used for their development. ‘-’ indicates that information has not been disclosed

LLM	Num. of parameters	Training data	Num. of GPUs	Released
GPT-1[26] (OpenAI)	117M	4.5GB (text)	8 GPUs	June, 2018
GPT-2 (OpenAI)	1.5B	40GB (text)	-	Feb., 2019
GPT-3 (OpenAI)	175B	570GB (text)	10,000 GPUs[27]	June, 2020
GPT-3.5 (OpenAI)	355B	-	-	March, 2022
GPT-4 (OpenAI)	-	-	-	March, 2023
Llama 2 (Meta)	70B	2T (token)	-	July, 2023
Llama 3 (Meta)	405B	15.6T (token)	-	July, 2024
PaLM 2 (Google)	540B	-	-	May, 2023
Gemini (Google)	1.6T	-	-	Dec., 2023

parameters, and was trained using 10,000 GPUs (NVIDIA V100) [27]. By 2022, GPT-3.5 had increased in scale to 355 billion parameters, and ChatGPT was built using a customized version of GPT-3.5. These details are summarized in Table 1. Models with such many parameters are categorized as LLMs, which achieve high performance by leveraging vast amounts of data and large-scale computational environments.

Since LLMs require an enormous amount of text data for training, it is impractical to annotate correct answers for such large datasets manually. Instead, Self-Supervised Learning (SSL) is used for training. One example of SSL in Transformer-based large-scale models is BERT [28], which is trained using a masked language modeling (MLM) approach. In this method, a sentence is converted into a fill-in-the-blank problem. For example, given the sentence “The capital of Aichi Prefecture is Nagoya City”, the training model sees “The capital of Aichi Prefecture is ?” and learns to predict the missing word “Nagoya City”. GPT models also utilize a similar SSL approach, allowing them to learn from massive datasets without relying on manually annotated data. After the pre-training phase, the model undergoes fine-tuning and human-in-the-loop adjustments to optimize its performance for specific tasks.

4.3 LLMs for medical applications and Japanese language

LLMs built using general text data may not achieve sufficient performance for specialized tasks such as medical support. To address this, many LLMs trained on medical

Table 2 LLMs developed using medical text corpora. Numbers of model parameters, scales of training data, and number of GPUs used for their development are also shown. ‘-’ indicates that information has not been disclosed

LLM	Num. of parameters	Training data	Num. of GPUs	Released
BioBERT	110M	18G (token)	8 GPUs	Sep., 2019
BioMedLM	2.7B	110GB	-	March, 2024
GatorTronGPT	20B	277B (token)	-	Nov., 2023
PMC-LLaMA	13B	79B (token)	32 GPUs	Aug., 2023
Med-PaLM (Google)	540B	-	-	Dec., 2022
Med-PaLM 2 (Google)	-	-	-	May, 2023

text corpora have been proposed [29]. One example is BioBERT, a model with 110 million parameters trained on texts from PubMed and PMC. In addition, larger-scale models such as BioMedLM, GatorTronGPT, and PMC-LLaMA have also been introduced. These models are summarized in Table 2.

LLMs built using Japanese-language corpora have also been introduced. These are summarized in Table 3. ELYZA Inc. has released ELYZA-japanese-Llama-2-7b and Llama-3-ELYZA-JP-70B, versions of Meta’s Llama 2/3 models fine-tuned on Japanese text. In addition, Preferred Networks and Preferred Elements have released PLaMo-13B and PLaMo-100B, while Japan’s National Institute of Informatics (NII) has published LLM-jp-13B and LLM-jp-172B. Among them, LLM-jp-172B is particularly notable. It has 172 billion parameters, making it comparable in scale to OpenAI’s GPT-3. Other notable examples include LLMs fine-tuned using data from Japan’s National Examination for Medical Practitioners, such as Llama3-Preferred-MedSwallow-70B and Preferred-MedLLM-Qwen-72B, both released by Preferred Networks.

4.4 Medical support using LLMs

Research on the medical applications of LLMs is rapidly advancing, enabling capabilities such as automatic generation and summarization of radiology reports, structuring medical findings, and automatic anonymization of clinical text. Numerous research

Table 3 LLMs developed using Japanese text corpora. Numbers of model parameters, scales of training data, and number of GPUs used for their development are also shown. ‘-’ indicates that information has not been disclosed

LLM	Num. of parameters	Training data	Num. of GPUs	Released
PLaMo-13B (Preferred Networks)	13B	1.4T (token)	480 GPUs	Sep., 2023
PLaMo-100B (Preferred Elements)	100B	2T (token)	-	Sep., 2024
ELYZA-japanese-Llama-2-7b (ELYZA)	7B	18B (token)	32 GPUs	Aug., 2023
Llama-3-ELYZA-JP-70B (ELYZA)	70B	-	-	June, 2024
LLM-jp-13B (NII)	13B	300B (token)	96 GPUs	Oct., 2023
LLM-jp-3 172B (NII)	172B	2.1T (token)	-	Sep., 2024
Llama3-Preferred-MedSwallow-70B (Preferred Networks)	70B	-	-	July, 2024
Preferred-MedLLM-Qwen-72B (Preferred Networks)	72B	-	-	March, 2025

320 results have been reported, and several companies have released LLM-based commer-
321 cial medical support services.

322 As an example of commercial service using LLMs, Therapixel (France) offers
323 MammoScreen [30], which automatically detects tumors in mammogram images and
324 generates drafts of radiology reports, including findings and impressions. In the United
325 States, RADPAIR [31] provides tools for radiology report input via dictation and
326 automated report generation support. In the future, more companies providing image-
327 based diagnostic support are expected to begin incorporating LLMs to assist in
328 radiology report creation.

329 Research studies using LLMs in medical support are also increasing rapidly. One of
330 the early studies applying LLMs in this context [32] explored the automatic generation
331 of descriptive text from medical images. This approach, called ChatCAD, performs
332 tasks such as tumor detection on medical images and then uses an LLM to generate
333 explanatory text based on the results. Another example, introduced in Section 3.4, is
334 TextDiff [20], which applies LLMs to medical image segmentation. This method uses
335 Clinical BioBERT, a specialized LLM for clinical text, as a text encoder and combines
336 textual and image information to generate segmentation results.

5 Foundation models

5.1 Generalized performance of large-scale models

In Section 4.2 in the chapter on text generation AI, this paper introduced the Scaling Laws [25], which state that the performance of neural language models improves with increases in the amount of training data, the number of model parameters, and the scale of computation during training. This principle applies not only to language models but also to image processing models.

A report by Zhai et al. [33] demonstrated that in ViT-based image processing models, increasing the amount of training data, the number of model parameters, and the scale of computation similarly leads to performance improvements. One well-known example of an image processing model built using massive datasets and large-scale computation is Meta’s Segment Anything Model (SAM) [34]. SAM was trained on 11 million images and over one billion region annotations, resulting in a ViT-based model with over 600 million parameters. Another example is Flamingo by DeepMind [35], which also leverages large-scale data and models to achieve high performance. Meta has also released an improved version, Segment Anything Model 2 (SAM 2) [36], which enhances SAM’s accuracy and supports video processing. Compared to SAM, SAM 2 offers better accuracy and faster processing speed, incorporating temporal attention into the model. SAM 2 was trained on approximately 51,000 videos, enabling it to segment objects in video sequences accurately.

Large-scale image processing models such as SAM, sometimes called Large Vision Models (LVMs), possess general-purpose capabilities that are not limited to specific tasks. While SAM is primarily trained on natural images and performs well in segmenting natural scenes, it has also been shown to work for medical image segmentation. Several studies have reported the applications of SAM to medical images, including tumor segmentation in pathology images [37], cardiac segmentation in ultrasound

images [19], and segmentation across various imaging modalities [38]. Although SAM’s segmentation accuracy is generally lower than that of task-specific models, it can still perform segmentation across various scenarios, regardless of differences in object scale and shape or whether the image is color or grayscale. SAM can recognize and segment images without additional fine-tuning using the target task data (i.e., zero-shot learning), even for medical images such as pathology images. These results indicate that large-scale image processing models exhibit strong generalization capabilities.

5.2 Foundation models and changes in AI development

In traditional machine learning-based AI development, data collection and model development were carried out separately for each task, resulting in task-specific AI systems. However, this approach requires repeated efforts for every new task and is not efficient.

In contrast, training relatively large-scale models using cross-task datasets makes it possible to create models with general-purpose performances that can be applied to multiple tasks. These models are known as foundation models. Foundation models have the following characteristics:

- They are trained on large-scale, cross-task datasets.
- They can be transferred to a wide range of downstream tasks.

SAM is a foundation model in image processing, while GPT is an example in natural language processing.

The value of foundation models lies in their ability to reduce the burden of AI development for individual tasks. Using a foundation model minimizes the effort required for task-specific data collection and rapidly builds models with performance comparable to task-specific AI systems. This capability has the potential to transform traditional AI development methodologies. In recent years, various approaches to developing AI

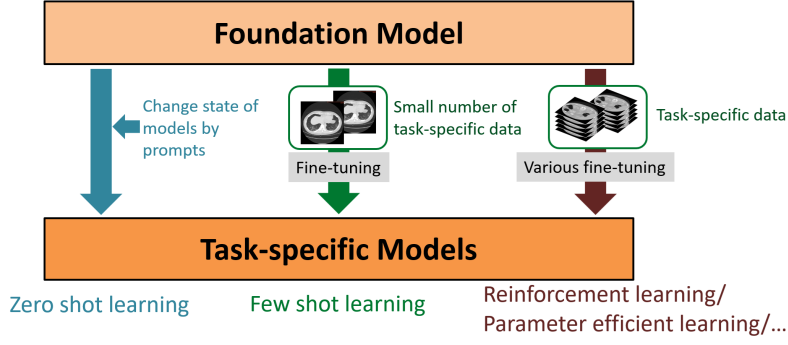


Fig. 2 Conceptual diagram of task-specific AI development using foundation models

for downstream tasks using foundation models have been proposed. Figure 2 illustrates some of these. Zero-shot learning is the method of adapting a foundation model to a task without using any task-specific data by adjusting its state with prompts. Few-shot learning is a method of fine-tuning a foundation model with a small amount of task-specific data. Additionally, methods such as parameter-efficient learning have been proposed, which enable efficient training of large foundation models on limited computational resources.

Various methods of AI development using foundation models are also being explored in medical image processing. Two core challenges in medical image processing are the significant effort required to collect large-scale datasets and the limited data availability for rare diseases. In this context, foundation models' ability to support downstream task development using only small amounts of data is extremely valuable. Further research advancements in this area are highly anticipated. In his keynote talk at the SPIE Medical Imaging 2024 international conference, Shuo Li discussed one vision for the ideal form of a medical foundation model [39]. His idea involved building a general-purpose medical foundation model that could be transferred to various downstream tasks in healthcare support, thereby enabling the development of AI systems for various medical tasks. He described a hierarchical structure of downstream tasks in medical support, where tasks are organized by organ, and beneath each organ task, there are disease-specific tasks. This concept is illustrated in Fig. 3.

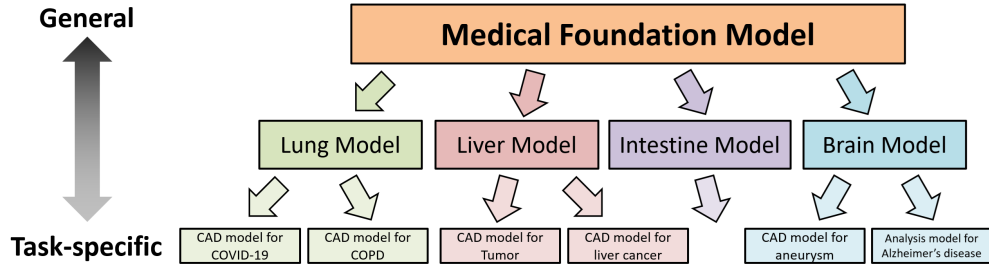


Fig. 3 Relationship between medical foundation model and downstream tasks as proposed by Shuo Li at SPIE Medical Imaging 2024 [39]

5.3 Methods for developing vision foundation models

This section introduces methods used to develop foundation models for image processing. In foundation model development, self-supervised learning (SSL) is typically used for pretraining, using many unlabeled images [40]. After pretraining, transfer learning or fine-tuning is performed using labeled images for the target downstream task, resulting in a task-specific model. The most commonly used SSL techniques for pretraining vision foundation models include Contrastive Learning (CL), Masked Image Modeling (MIM), and a combination of CL and MIM.

Contrastive Learning (CL) generates two augmented versions of the same image through data augmentation. The model extracts feature representations using an encoder, and training is performed to increase similarity between positive pairs (images augmented from the same original image) and decrease similarity between negative pairs (images augmented from different images). Several CL-based methods have been proposed. Among them, SimCLR [41] is a fundamental contrastive learning method. MoCo [42] reuses negative samples efficiently during training. SimSiam [43] and DINO [44] perform CL without using negative samples. SimCLR2 [45] and MoCov2 [46] are enhanced versions of SimCLR and MoCo, respectively. A multimodal version of CL that utilizes image-text pairs is CLIP [47], which is widely used for developing multimodal foundation models.

Masked Image Modeling (MIM) divides an image into patches, masks a portion of them, and trains the model to reconstruct the original image using the unmasked patches. MIM is commonly used as an SSL technique when employing ViT. A fundamental approach in MIM is the Masked Autoencoder [48]. It has been shown that masking up to 75% of patches during training still leads to high accuracy in downstream tasks. Variants such as Attention-guided MIM [49] have also been proposed, which adjust the masking process using attention weights to focus on semantically meaningful regions.

Several methods combine CL and MIM to enhance model performance. For example, Contrastive Masked Autoencoders (CMAE) [50] create positive and negative image pairs using pixel shifts and then apply masking and reconstruction in the encoder similar to a masked autoencoder. The model is trained using loss functions from both CL and MIM. Other notable approaches include “What Do Self-Supervised Vision Transformers Learn?” [51], which further explores SSL for ViT-based models.

5.4 Foundation models for medical image processing

One significant advantage of foundation models is their applicability across a wide range of downstream tasks. However, achieving high performance, even with transfer learning, becomes difficult if there is a large domain gap between the image data used to train the foundation model and the downstream task’s image domain. To achieve better transfer learning results in specific image domains, using a foundation model built using data closely aligned with the target domain is more effective. For example, SAM was trained on general-purpose natural images and may not provide sufficient accuracy when directly applied to medical image processing. Therefore, a variety of foundation models have been proposed that are specifically trained using single- or multi-modal medical images. This section introduces such foundation models for medical image processing.

453 MedSAM [38] is a segmentation-oriented foundation model applicable to 2D medi-
454 cal images from various modalities. It was created by training the SAM architecture on
455 medical images, including CT, MR, X-ray, ultrasound, mammography, optical coher-
456 ence tomography (OCT), endoscopy, dermoscopy, fundus, and pathology images, with
457 over 1.57 million images used for training. Unlike SAM, which allows various prompts,
458 MedSAM uses only bounding boxes as prompts. It has been shown to outperform
459 SAM in tasks like organ and tumor segmentation from 2D medical images.

460 MedSAM-2 [52] extends SAM 2 to support 3D medical images and video segmen-
461 tation. By treating one spatial axis of a 3D image as a temporal axis, MedSAM-2
462 leverages SAM 2’s temporal modeling capabilities to enable 3D segmentation.

463 BiomedCLIP [53] is a multimodal foundation model that combines images and
464 text. Built on the CLIP framework, it incorporates PubMedBERT [54] as its text
465 encoder to improve performance in the medical domain. It was trained on PMC-15M,
466 a 15 million figure-caption pairs dataset extracted from 4.4 million scientific articles
467 in PubMed Central (PMC). BiomedCLIP supports tasks such as image classification
468 and visual question answering (VQA).

469 BioViL-T [55] is another multimodal model that combines X-ray images with radi-
470 ology report text. It uses a custom multimodal architecture trained on about 174,100
471 image-text pairs from the MIMIC-CXR v2 dataset. BioViL-T can be used for image
472 classification and text generation from images.

473 PathAsst [56] is a multimodal foundation model for pathology images and text.
474 Based on the CLIP framework, it was trained on over 207,000 image-text pairs col-
475 lected from sources such as PubMed. PathAsst is suitable for image classification and
476 text generation.

477 Prov-GigaPath [57] is a single-modality foundation model tailored for pathology
478 images. It was pretrained on 1.3 billion image patches extracted from 171,000 whole-
479 slide images, using methods such as DINOv2. It is applicable to various downstream
480 tasks in pathology.

481 RETFound [58] is a foundation model developed for ophthalmology, supporting
482 fundus and OCT images. It was trained on 1.6 million images using the Masked
483 Autoencoder framework. RETFound has shown excellent performance across various
484 retinal image-based downstream tasks.

485 **6 Future directions in medical image processing**

486 **research**

487 **6.1 Overview**

488 In general image processing and natural language processing, the development of large-
489 scale models and foundation models has progressed by referencing the Scaling Laws.
490 Large models such as GPTs and Gemini are developed using the vast computational
491 resources of big tech companies, making it difficult for academic researchers to compete
492 in model development.

493 In contrast, even big companies face challenges in collecting massive datasets in
494 medical image processing. As a result, current efforts focus on developing task-specific
495 foundation models within this domain. Moreover, in developing AI for medical support,
496 it is essential to be mindful of data bias related to country or race. However, no
497 foundation model that is specifically suited to the Japanese patient population has
498 yet been realized. To develop AI tailored to medical support needs in each country,
499 foundation models that account for regional and demographic biases must be created.
500 Collaboration among academic researchers in medical image processing is necessary
501 to achieve this, particularly for data collection and model development.

502 What is required to make this a reality? Recalling the Scaling Laws, three key
503 factors are necessary to improve model performance, including a large amount of
504 training data, a large number of model parameters, and a significant scale of compu-
505 tational resources during training. While preparing these elements is not easy, it is not
506 impossible if we fully utilize national data and computational resources. Assuming the
507 implementation of large-scale models is possible through software development, the
508 following sections will explore how to prepare the necessary large-scale medical image
509 datasets and computational resources.

510 **6.2 Building large-scale medical image datasets**

511 There are already examples where large-scale medical image datasets have been devel-
512 oped. One notable case is the "Medical Image Big Data Cloud Platform" developed
513 by the Medical Big Data Research Center at Japan's National Institute of Informatics
514 (NII). This platform stores over 400 million medical images across various modalities,
515 including radiological, endoscopic, ophthalmologic, and pathology images. Since 2017,
516 the center has collected these images through the SINET network, resulting in one of
517 Japan's most significant medical image datasets.

518 Once a vast number of images is collected, the next major challenge becomes how
519 to annotate them efficiently. A notable report by Qu et al. [59] describes an approach
520 using a human-in-the-loop system to enable large-scale annotation. In their study, they
521 annotated multi-organ abdominal segmentation regions on 8,448 CT cases (approxi-
522 mately 3.2 million 2D slice images) in an impressive three-week period. Their workflow
523 involved AI-driven automatic segmentation, followed by manual correction by human
524 experts, which is a human-in-the-loop setup. They utilized three segmentation mod-
525 els, including Swin UNETR, nnU-Net, and U-Net, and leveraged model disagreement
526 (inter-model uncertainty) to prioritize and streamline human correction. Such a system
527 demonstrates that efficient annotation of large datasets is feasible within a reasonable
528 time.

529 Additionally, synthetic and simulation data are proving highly effective in building
530 large annotated datasets. Constructing AI models using easily generated synthetic or
531 simulation data and then fine-tuning them on real data can significantly reduce the
532 overall data collection cost. Due to advances in simulation technologies, the sim-to-
533 real approach, which is commonly used in robotics and autonomous driving, is now
534 being actively researched in biomedical applications. One such example is the 4D
535 eXtended CArdiac-Torso (XCAT) Phantom [60], which can generate virtual medical
536 images by simulating various biological characteristics such as age, gender, body size,
537 organ structure, and motion on a virtual human model. It can even simulate realistic
538 noise and imaging artifacts, making it possible to generate large-scale medical image
539 datasets at low cost. Another example is a study by Menten et al. [61], which used
540 vascular development simulation for training an AI for blood vessel segmentation.
541 They generated numerous 3D retinal vascular structures via simulation and created
542 annotated OCT angiography images, which were used to train a segmentation AI.
543 Their results showed that using simulated data led to higher segmentation accuracy
544 than using only real data. As these research efforts demonstrate, integrating synthetic
545 and simulation data into data collection and AI development is extremely promising
546 and valuable in the context of medical support applications.

547 **6.3 Utilizing large-scale computational resources**

548 Currently, AI development requires distributed training using many GPUs, espe-
549 cially for large-scale models. However, acquiring many high-performance GPUs is a
550 significant challenge. To overcome this, researchers can leverage supercomputers devel-
551 oped by universities and research institutions. Considering the growing demands of
552 AI development, many of these supercomputers are equipped with many GPUs and
553 support general development environments, including Python. With access to such
554 infrastructure, it is entirely feasible to develop large-scale AI models. Although using

555 supercomputers typically incurs usage fees, several support programs are available to
556 subsidize these costs. When effectively utilized, these programs can allow access to
557 powerful computational resources with little to no financial burden.

558 As discussed throughout this chapter, large-scale medical image datasets and com-
559 putational resources essential for developing foundation models for medical support
560 are realistically available. By maximizing the use of national data and computing
561 resources, we can anticipate the emergence of nationally developed foundation mod-
562 els and high-performance AI systems that contribute to the advancement of medical
563 support technologies.

564 7 Conclusions

565 In this paper, the mechanisms of image generation AI and text generation AI, along
566 with examples of their applications in the medical field are introduced. Furthermore,
567 this paper introduced the Scaling Laws, a concept that has gained prominence as gen-
568 erative AI continues to evolve, and introduced foundation models, which are expected
569 to significantly impact the future of AI development. Finally, this paper examined the
570 importance of preparing large-scale medical image datasets and securing large-scale
571 computing resources to realize foundation models for medical image processing. This
572 paper emphasized the necessity of developing foundation models for medical support
573 using national data and computing resources. As medical image processing enters a
574 new era, I hope that academic researchers will continue to play a key role and make
575 meaningful contributions to its advancement.

576 Declarations

577 **Conflict of interest.** The author have no conflicts of interest regarding this
578 manuscript.

579 **Ethics approval.** Ethics approval is not necessary because no human or animal
580 participants, specimens, or data were involved in this work.

581 **Data availability.** No data was used in this paper.

582 References

- 583 [1] Krizhevsky A, Sutskever I, Hinton GE. ImageNet Classification with Deep Con-
584 volutional Neural Networks. Proceedings of the 26th International Conference on
585 Neural Information Processing Systems (NIPS 12). 2012;1:1097–1105.
- 586 [2] Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, et al.
587 Attention Is All You Need. arXiv:1706.03762. 2017;.
- 588 [3] Kingma DP, Welling M. Auto-Encoding Variational Bayes. International
589 Conference on Learning Representations (ICLR) 2014. 2014;.
- 590 [4] Goodfellow I, Pouget-Abadie J, Mirza M, Xu B, Warde-Farley D, Ozair S, et al.
591 Generative adversarial networks. Communications of the ACM. 2020;63(11):139–
592 144.
- 593 [5] Karras T, Laine S, Aittala M, Hellsten J, Lehtinen J, Aila T. Analyzing and
594 Improving the Image Quality of StyleGAN. 2020 IEEE/CVF Conference on
595 Computer Vision and Pattern Recognition (CVPR). 2020;p. 8107–8116.
- 596 [6] Sohl-Dickstein J, Weiss EA, Maheswaranathan N, Ganguli S. Deep unsupervised
597 learning using nonequilibrium thermodynamics. Proceedings of the 32nd Interna-
598 tional Conference on International Conference on Machine Learning (ICML 15).
599 2015;37:2256–2265.
- 600 [7] Ho J, Jain A, Abbeel P. Denoising Diffusion Probabilistic Models. Proceedings
601 of the 34th International Conference on Neural Information Processing Systems

- (NIPS 20). 2020;p. 6840–6851.
- [8] Ronneberger O, Fischer P, Brox T. U-Net: Convolutional Networks for Biomedical Image Segmentation. Medical Image Computing and Computer Assisted Intervention (MICCAI 2015). 2015;9351:234–241.
- [9] Pinaya WHL, Tudosiu PD, Dafflon J, Costa PFD, Fernandez V, Nachev P, et al. Brain Imaging Generation with Latent Diffusion Models. DGM4MICCAI 2022. 2022;13609:117–126.
- [10] Dorjsembe Z, Odonchimed S, Xiao F. Three-Dimensional Medical Image Synthesis with Denoising Diffusion Probabilistic Models. MIDL 2022. 2022;.
- [11] Kim B, Ye JC. Diffusion Deformable Model for 4D Temporal Medical Image Generation. Medical Image Computing and Computer Assisted Intervention (MICCAI 2022). 2022;13431:539–548.
- [12] Zhou Y, Towning R, Awad Z, Giannarou S. Image synthesis with class-aware semantic diffusion models for surgical scene segmentation. Healthcare Technology Letters. 2025;12(1).
- [13] Venkatesh DK, Rivoir D, Pfeiffer M, Kolbinger F, Speidel S. Synthesizing Multi-class Surgical Datasets with Anatomy-Aware Diffusion Models. arXiv:241007753. 2024;.
- [14] Zhang H, Cisse M, Dauphin YN, Lopez-Paz D. mixup: Beyond Empirical Risk Minimization. International Conference on Learning Representations (ICLR) 2018. 2018;.
- [15] Abe K, Takeo H, Nawano S. Artificial Case Image Generation of Breast Cancer Mass by Stable Diffusion and Its Application to Differentiation between Benign

- 625 and Malignant. Proceedings of JAMIT 2023. 2023;p. 142–148.
- 626 [16] Junde Wu RF, Fang H, Zhang Y, Yang Y, Xiong H, Liu H, et al. MedSegDiff:
627 Medical Image Segmentation with Diffusion Probabilistic Model. Medical Imaging
628 with Deep Learning. 2024;227:1623–1639.
- 629 [17] Wu J, Ji W, Fu H, Xu M, Jin Y, Xu Y. MedSegDiff-V2: diffusion-based medical
630 image segmentation with transformer. Proceedings of the Thirty-Eighth AAAI
631 Conference on Artificial Intelligence and Thirty-Sixth Conference on Innovative
632 Applications of Artificial Intelligence and Fourteenth Symposium on Educational
633 Advances in Artificial Intelligence (AAAI’24/IAAI’24/EAAI’24). 2024;p. 6030–
634 6038.
- 635 [18] Chowdary GJ, Yin Z. Diffusion Transformer U-Net for Medical Image Segmenta-
636 tion. Medical Image Computing and Computer Assisted Intervention (MICCAI
637 2023). 2023;14223:622–631.
- 638 [19] Chen T, Wang C, Shan H. BerDiff: Conditional Bernoulli Diffusion Model for
639 Medical Image Segmentation. Medical Image Computing and Computer Assisted
640 Intervention (MICCAI 2023). 2023;14223:491–501.
- 641 [20] Feng CM. Enhancing Label-Efficient Medical Image Segmentation with Text-
642 Guided Diffusion Models. Medical Image Computing and Computer Assisted
643 Intervention (MICCAI 2024). 2024;15008:253–262.
- 644 [21] Lu Y, Yang Y, Xing Z, Wang Q, Zhu L. Diff-VPS: Video Polyp Segmenta-
645 tion via a Multi-task Diffusion Network with Adversarial Temporal Reasoning.
646 Medical Image Computing and Computer Assisted Intervention (MICCAI 2024).
647 2024;15006:165–175.

- [22] Hochreiter S, Schmidhuber J. Long Short-Term Memory. *Neural Computation*. 1997;9(8):1735–1780.
- [23] Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, et al. Attention is all you need. *Proceedings of the 31st International Conference on Neural Information Processing Systems (NIPS 17)*. 2017;p. 6000–6010.
- [24] Dosovitskiy A, Beyer L, Kolesnikov A, Weissenborn D, Zhai X, Unterthiner T, et al. An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale. *ICLR 2021*. 2021;.
- [25] Kaplan J, McCandlish S, Henighan T, Brown TB, Chess B, Child R, et al. Scaling Laws for Neural Language Models. *arXiv:200108361*. 2020;.
- [26] Radford A, Narasimhan K, Salimans T, Sutskever I.: Improving language understanding with unsupervised learning. Available from: <https://openai.com/index/language-unsupervised/>.
- [27] Patterson D, Gonzalez J, Hölzle U, Le Q, Liang C, Munguia LM, et al. The Carbon Footprint of Machine Learning Training Will Plateau, Then Shrink. *Computer*. 2022;55(7):18–28.
- [28] Devlin J, Chang MW, Lee K, Toutanova K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. *NAACL-HLT 2019*. 2019;1:4171–4186.
- [29] Zhou H, Liu F, Gu B, Zou X, Huang J, Wu J, et al. A Survey of Large Language Models in Medicine: Progress, Application, and Challenge. *arXiv:231105112*. 2024;.
- [30] MammoScreen.: Available from: <https://www.mammoscreen.com/>.

- [31] RADPAIR.: Available from: <https://www.radpair.com/>.
- [32] Wang S, Zhao Z, Ouyang X, Liu T, Wang Q, Shen D. Interactive computer-aided diagnosis on medical image using large language models. *Communications Engineering*. 2024;3.
- [33] Zhai X, Kolesnikov A, Houlsby N, Beyer L. Scaling Vision Transformers. 2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). 2022;p. 1204–1213.
- [34] Kirillov A, Mintun E, Ravi N, Mao H, Rolland C, Gustafson L, et al. Segment Anything. *arXiv:230402643*. 2023;.
- [35] Alayrac JB, Donahue J, Luc P, Miech A, Barr I, Hasson Y, et al. Flamingo: a visual language model for few-shot learning. *Proceedings of the 36th International Conference on Neural Information Processing Systems (NIPS 22)*. 2022;p. 23716–23736.
- [36] Ravi N, Gabeur V, Hu YT, Hu R, Ryali C, Ma T, et al. SAM 2: Segment Anything in Images and Videos. *arXiv:240800714*. 2024;.
- [37] Deng R, Cui C, Liu Q, Yao T, Remedios LW, Bao S, et al. Segment Anything Model (SAM) for Digital Pathology: Assess Zero-shot Segmentation on Whole Slide Imaging. *arXiv:230404155*. 2023;.
- [38] Ma J, He Y, Li F, Han L, You C, Wang B. Segment anything in medical images. *Nature Communications*. 2024;15.
- [39] Li S. Integrating vision and language: revolutionary foundations in medical imaging AI. *Proceedings of SPIE Medical Imaging 2024*. 2024;12926.

- 693 [40] Jing L, Tian Y. Self-Supervised Visual Feature Learning With Deep Neu-
694 ral Networks: A Survey. *IEEE Transactions on Pattern Analysis and Machine*
695 *Intelligence*. 2020;43(11):4037–4058.
- 696 [41] Chen T, Kornblith S, Norouzi M, Hinton G. A simple framework for con-
697 trastive learning of visual representations. *Proceedings of the 37th International*
698 *Conference on Machine Learning (ICML 20)*. 2020;p. 1597–1607.
- 699 [42] He K, Fan H, Wu Y, Xie S, Girshick R. Momentum Contrast for Unsupervised
700 Visual Representation Learning. *Proceedings of the IEEE/CVF Conference on*
701 *Computer Vision and Pattern Recognition (CVPR)*. 2020;p. 9729–9738.
- 702 [43] Chen X, He K. Exploring Simple Siamese Representation Learning. *Proceed-*
703 *ings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*
704 *(CVPR)*. 2021;p. 15750–15758.
- 705 [44] Caron M, Mairal J, Touvron H, Misra I, Bojanowski P, Jégou H, et al. Emerging
706 Properties in Self-Supervised Vision Transformers. *Proceedings of the IEEE/CVF*
707 *International Conference on Computer Vision (ICCV)*. 2021;p. 9650–9660.
- 708 [45] Chen T, Kornblith S, Swersky K, Norouzi M, Hinton G. Big Self-Supervised
709 Models are Strong Semi-Supervised Learners. *Proceedings of the 34th Interna-*
710 *tional Conference on Neural Information Processing Systems (NIPS 20)*. 2020;p.
711 22243–22255.
- 712 [46] Chen X, Fan H, Girshick R, He K. Improved Baselines with Momentum
713 Contrastive Learning. *arXiv:200304297*. 2020;.
- 714 [47] Radford A, Kim JW, Hallacy C, Ramesh A, Goh G, Agarwal S, et al. Learning
715 Transferable Visual Models From Natural Language Supervision. *Proceedings of*
716 *the 38th International Conference on Machine Learning*. 2021;p. 8748–8763.

- 717 [48] He K, Chen X, Xie S, Li Y, Dollár P, Girshick R. Masked Autoencoders Are
718 Scalable Vision Learners. 2022 IEEE/CVF Conference on Computer Vision and
719 Pattern Recognition (CVPR). 2022;.
- 720 [49] Kakogeorgiou I, Gidaris S, Psomas B, Avrithis Y, Bursuc A, Karantzas K,
721 et al. What to Hide from Your Students: Attention-Guided Masked Image Mod-
722 eling. Proceedings of the European Conference on Computer Vision (ECCV).
723 2022;13690:300–318.
- 724 [50] Huang Z, Jin X, Lu C, Hou Q, Cheng MM, Fu D, et al. Contrastive Masked
725 Autoencoders are Stronger Vision Learners. IEEE Transactions on Pattern
726 Analysis and Machine Intelligence. 2023;46(4):2506–2517.
- 727 [51] Park N, Kim W, Heo B, Kim T, Yun S. What Do Self-Supervised Vision
728 Transformers Learn? arXiv:230500729. 2023;.
- 729 [52] Zhu J, Hamdi A, Qi Y, Jin Y, Wu J. Medical SAM 2: Segment medical images
730 as video via Segment Anything Model 2. arXiv:240800874. 2024;.
- 731 [53] Zhang S, Xu Y, Usuyama N, Xu H, Bagga J, Tinn R, et al. BiomedCLIP: a
732 multimodal biomedical foundation model pretrained from fifteen million scientific
733 image-text pairs. arXiv:230300915. 2023;.
- 734 [54] Gu Y, Tinn R, Cheng H, Lucas M, Usuyama N, Liu X, et al. Domain-specific
735 language model pretraining for biomedical natural language processing. ACM
736 Transactions on Computing for Healthcare (HEALTH). 2021;3(1):1–23.
- 737 [55] Bannur S, Hyland S, Liu Q, Pérez-García F, Ilse M, Castro DC, et al. Learning to
738 Exploit Temporal Structure for Biomedical Vision-Language Processing. Proceed-
739 ings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition
740 (CVPR). 2023;p. 15016–15027.

- 741 [56] Sun Y, Zhu C, Zheng S, Zhang K, Sun L, Shui Z, et al. PathAsst: A Generative
742 Foundation AI Assistant Towards Artificial General Intelligence of Pathology.
743 Proceedings of the Thirty-Eighth AAAI Conference on Artificial Intelligence
744 (AAAI-24). 2024;p. 5034–5042.
- 745 [57] Xu H, Usuyama N, Bagga J, Zhang S, Rao R, Naumann T, et al. A whole-
746 slide foundation model for digital pathology from real-world data. Nature.
747 2024;630:181–188.
- 748 [58] Zhou Y, Chia MA, Wagner SK, Ayhan MS, Williamson DJ, Struyven RR, et al. A
749 foundation model for generalizable disease detection from retinal images. Nature.
750 2023;622:156–163.
- 751 [59] Qu C, Zhang T, Qiao H, Liu J, Tang Y, Yuille A, et al. AbdomenAtlas-
752 8K: Annotating 8,000 CT Volumes for Multi-Organ Segmentation in Three
753 Weeks. Proceedings of the 37th International Conference on Neural Information
754 Processing Systems (NIPS 23). 2023;p. 36620–36636.
- 755 [60] Segars WP, Tsui B, Cai J, Yin FF, Fung GS, Samei E. Application of the 4D
756 XCAT Phantoms in Biomedical Imaging and Beyond. IEEE Transactions on
757 Medical Imaging. 2017;37(3):680–692.
- 758 [61] Menten MJ, Paetzold JC, Dima A, Menze BH, Knier B, Rueckert D. Physiology-
759 based simulation of the retinal vasculature enables annotation-free segmentation
760 of OCT angiographs. 2022;13438:330–340.