

ImageGrains 2.0: Improved precision and generalization for grain segmentation

David Mair¹, Guillaume Witz², Ariel Do Prado^{1,3}, Philippos Garefalakis¹, Amanda Wild^{4,5}, Fanny Ville⁶,
Bennet Schuster¹, Michael Horn², Jürgen Österle^{7,8}, Stefano C. Fabbri⁹, Camille Litty⁹, Stefan
5 Achleitner¹⁰, Sebastian Leistner¹⁰, Clemens Hiller^{10,11}, and Fritz Schlunegger¹

¹Institute of Geological Sciences, University of Bern, Bern, 3012, Switzerland

²Data Science Lab, University of Bern, Bern, 3012, Switzerland

³Institute of Geosciences, University of São Paulo, São Paulo, 05508-080, Brazil

⁴Institute of Physical Geography and Geoecology, RWTH-Aachen University, Aachen, 52062, Germany

10 ⁵GFZ Helmholtz Centre for Geosciences, 14473 Potsdam, Germany

⁶Fluvial Dynamics Research Group (RIUS), University of Lleida, Lleida, 25003, Spain

⁷School of Geography, Environment and Earth Sciences, Victoria University of Wellington, Wellington, 6012, New Zealand

⁸Amt der Vorarlberger Landesregierung, Bregenz, 6901, Austria

⁹Federal Office of Topography swisstopo, Wabern, 3084, Switzerland

15 ¹⁰Unit of Hydraulic Engineering, University of Innsbruck, Innsbruck, 6020, Austria

¹¹Natural Hazards and Risk Management, Geoconsult ZT GmbH, Puch bei Hallein, 5412, Austria

Correspondence to: David Mair (david.mair@unibe.ch)

Abstract. Recent advances in deep-learning-based image segmentation have enabled the development of automated approaches to detect individual grains and measure them for geoscientific applications. These methods facilitate the creation
20 of much larger and more precise datasets than traditional manual grain measurements. However, they typically perform best as specialized models trained on homogeneous, task-specific datasets, and often show reduced accuracy when used ~~to~~
~~generalize to~~ different data types.

Here, we present an updated framework, ImageGrains 2.0 that leverages Cellpose-SAM, a recently published next-generation deep-learning model originally developed for cell segmentation in biomedical research. It currently represents the state-of-
25 the-art for dense segmentation in 2D and 3D biomedical datasets, ~~and~~. It yields robust results and is capable to generalize across distinctly different image datasets. These properties allow us to re-train the model with geoscientific dataset comprising annotated images of fluvial gravel, coarse pro-glacial deposits, and X-ray computer tomography scans of glacial till and marine sand. We benchmark the segmentation performance of ~~the~~our method against ground-truth annotations, compare it to the performance of other segmentation methods, and we evaluate its measurement accuracy. Our results indicate that this approach
30 outperforms existing methods and confirm that the outstanding performance of Cellpose-SAM is transferable to segment sediment grains. We analyze the size and shape of these segmented grains and find that an increase in grain segmentation accuracy leads to more precise and realistic morphometric results, e.g., more accurate grain size distributions. Additionally, we introduce an interactive graphical user interface for image annotation and correction of model predictions, facilitating the use of the framework ~~in a broader~~for broad range of image settings. Furthermore, this study underscores the importance of

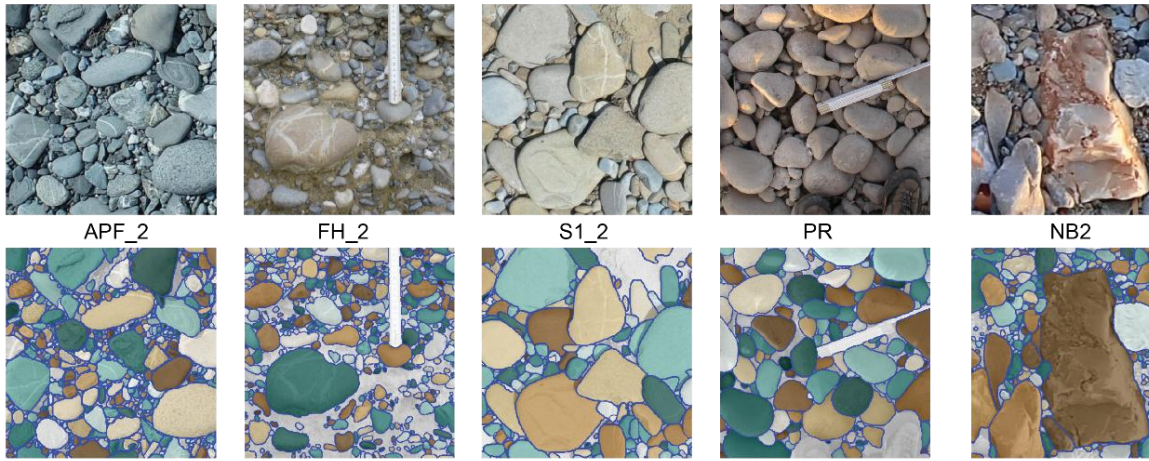
35 curating ~~of~~ more publicly available datasets, which could pave the way towards the generation of a foundation model for
segmenting granular particles in geoscientific imagery.

1 Introduction

Data on the size and shape of granular particles have been used across a broad range of geoscientific research areas, and such information has provided the basis for the quantification of the physical and chemical properties of clastic materials (e.g.,
40 Sklar, 2024; Israeli and Emmanuel, 2018). Grain morphometry, for instance, is essential for studying sediment production and transport dynamics in environments such as fluvial, marine, glacial, and hillslope systems (e.g., von Eynatten et al., 2012; DiBiase et al., 2017; Allen et al., 2017; Garefalakis et al. 2024). Such data on grain morphometry, particularly those collected in fluvial systems, have proven essential for understanding river hydraulics (e.g., Rickenmann and Recking, 2011; Dunne and Jerolmack, 2018), the mechanisms of sediment transport (e.g., Piégay et al., 2020; Tofelde et al., 2021), and the interplay
45 between fluvial hydraulics and bedload transport (e.g., Dietrich et al., 1989; Pfeiffer et al., 2017; Schlunegger et al., 2020; Deal et al., 2023). Moreover, the sediment particles' size and shape bear information on their transport through various geomorphic processes (e.g., Attal and Lavé, 2009; Miller et al., 2014; Novák-Szabó et al., 2018; Marc et al., 2021), and even on transport conditions on other planets (e.g., Szabo et al., 2015). Furthermore, distributions of the grains' size and shape as well as the geometric arrangement of the clasts in the stratigraphic record (e.g., imbrication) allow to reconstruct deposition
50 events and conditions (e.g., Spagnolo et al., 2016; Schlunegger and Garefalakis, 2018; Biguenet et al., 2021; Preusser et al., 2021).
Despite the importance of quantitative and robust data on grain size, shape and orientation, obtaining such data remains a challenge. Traditionally, such data on the grains' morphometries have been collected through laborious manual measurements of grains in the field or through sieving (e.g., Bunte and Abt, 2001) ~~or on~~; Watkins et al., 2019; Baynes et al., 2020; Garefalakis
55 et al., 2023), from topographic point clouds (e.g., Vázquez-Tarrio et al., 2017; Woodget and Austrums, 2017; Steer et al., 2022) or from imagery (e.g., Butler et al., 2001; Carbonneau et al., 2004; Detert & Weitbrecht, 2012; Buscombe, 2013; Sulaiman et al., 2014; Purinton & Bookhagen, 2019). Manual measurements are laborious, typically yield low numbers of observations (e.g., Eaton et al, 2019) and are prone to bias (e.g., Daniels and McCusker, 2010), especially when traditional geometric sampling techniques, e.g., Wolman counts or line sampling (Wolman, 1954; Fehr, 1987) are used. Obtaining grain
60 geometries from topographic point clouds is often limited by the technical limits of the acquisition method, grain occlusion or necessary geometric assumptions for grain segmentation (e.g., Steer et al., 2022; Woodget and Austrums, 2017). Image-based methods require field calibration for texture-based approaches and manual correction of individual grains for segmentation-based methods; both digital methods also tend to systematically over- or underestimate grain sizes (e.g., Chardon et al., 2022; Mair et al., 2022).
65 In ~~this context~~ the past years, machine-learning tools have been developed ~~more recently~~ in an effort to automate the measurements of grain size and shape ~~measurements in images~~, to improve ~~the~~ data quality, and to allow for an

~~increased~~increase in number of ~~observations~~observation. Among these, texture-based methods predict percentile values of grain size distributions if an unambiguous correlation between an image texture and a characteristic grain sized distribution exists, and if these were included in the training data (e.g., Buscombe, 2020; Lang et al., 2021). In contrast, segmentation-based methods delineate individual grains ~~through~~through object detection (Chen et al., 2022; Mair et al., 2024; Miazza et al., 2024; Sylvester et al., 2025) and facilitate the creation of large datasets, which allow for size and shape analysis down to ~~an~~the individual grain level. However, these segmentation models work best when trained as narrow specialist models on homogenous datasets, which often require task-specific, and sometimes site-specific, training and careful curation of the corresponding datasets (e.g., Chen et al., 2023; Prieur et al., 2023; Azzam et al., 2024; Miazza et al., 2024; Zegers et al., 2025; Schuster et al., 2025).

During recent years, a new generation of deep learning models that use a transformer architecture has become widely used in the field of computer vision for tasks related to the segmentation of objects in images (e.g., Dosovitskiy et al., 2020; Li et al., 2022). This resulted in the development of foundation models, such as the Segment Anything Model (SAM; Kirillov et al., 2023; Ravi et al., 2024), which are trained on very large and general datasets. These foundation models have proven effective at generalization, i.e., being able to predict outcomes for previously unseen data, which were not used for training. Due to a strong inductive bias, they are also considered as effective at out-of-distribution detection of objects, especially when fine-tuned on smaller datasets (Hendrycks et al., 2020; Fort et al., 2021). While these models perform well at segmenting numerous different types of objects in images, they are less efficient at accurately segmenting large quantities of a narrow range of objects, especially when they are not fine-tuned to specific datasets or when no specific prompts are used (Sylvester et al., 2025; Chan et al., 2025).



ImageGrains 2.0 (IG2) dataset: 243 image tiles with 29622 manually annotated grain masks

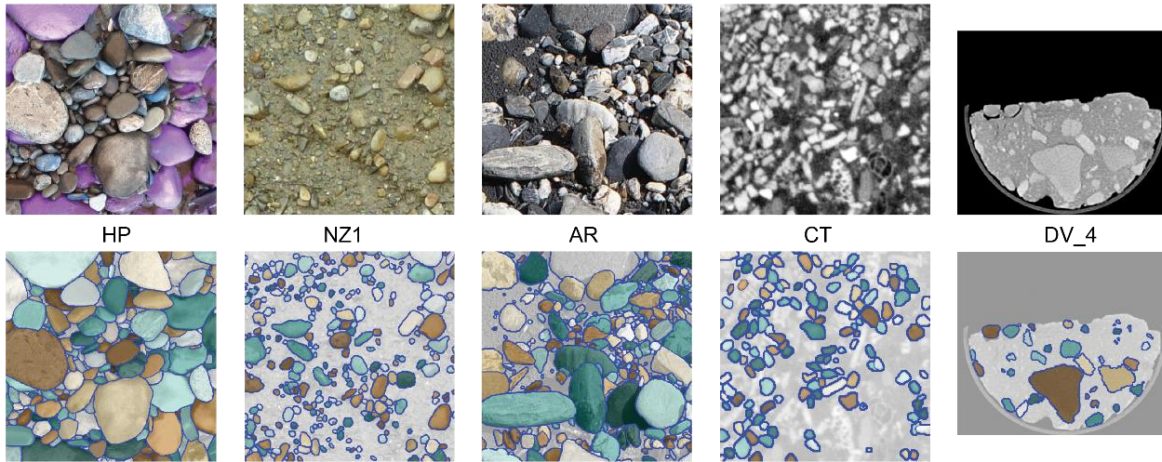
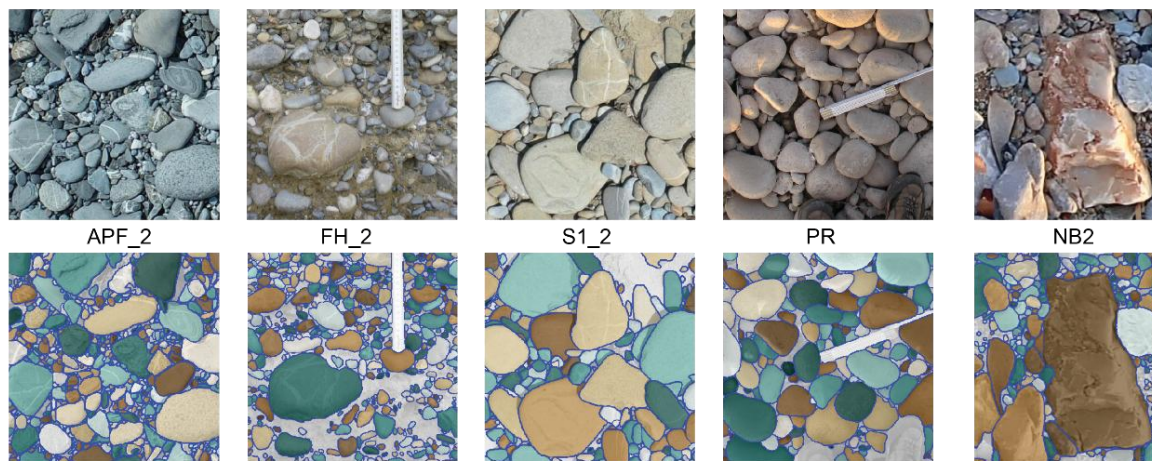


Figure 1: Example images and their manually annotated grain labels for various types of imagery and grains used in the IG2 dataset, and their respective indicated subset (Mair et al., 2025a; see Table S1 for more details). Individual clasts are shown in random colors with blue outlines.

90 Here, we present an updated update to the ImageGrains framework building on by Mair et al. (2024) that introduces two major
improvements. First, it leverages the strengths of Cellpose-SAM (Pachitariu et al., 2025), a recently published next-generation
 95 deep-learning model originally developed for cell segmentation in biomedical imagery. This model eliminates weaknesses of
 SAM and improves the performance for dense segmentation of many instances of the same object type with high accuracy,
 while maintaining the outstanding generalization ability. We utilize the new Cellpose-SAM model that was trained on large
 datasets of predominantly biomedical imagery of cells (for details, see Pachitariu et al., 2025) and retrained it to find grains in
 images of clastic sediment. This transfer learning approach allows us to utilize both the representations learned by SAM that
 enables the generalization across widely differing datasets and data types, and the Cellpose segmentation framework (Stringer
 et al., 2021) that facilitates the efficient and dense segmentation of grains with high accuracy without the necessity of prompt

engineering. To achieve this, we newly curated a dataset of 243,162 annotated image tiles, which we combine with the existing
 100 ImageGrains 1.0 dataset to obtain as set of 243 image tiles from various types of sediment grains (Fig. 4-1). This increase in
number of annotated images constitutes the second major improvement compared to the previous version. We then compare
 the segmentation results of our re-trained Cellpose-SAM model with the results of other state-of-the-art approaches, and we
 test the models' ability to generalize by using subsets of our dataset as unseen test splits. Finally, we highlight the potential of
 applying our workflow to 3D datasets of stacked images retrieved by X-ray computer tomography (CT) scans. Our results
 105 indicate that the new framework outperforms existing methods both in accuracy of the resulting segmentation, and in the
 capability to segment grains in new types of imagery.



ImageGrains 2.0 (IG2) dataset: 243 image tiles with 29622 manually annotated grain masks

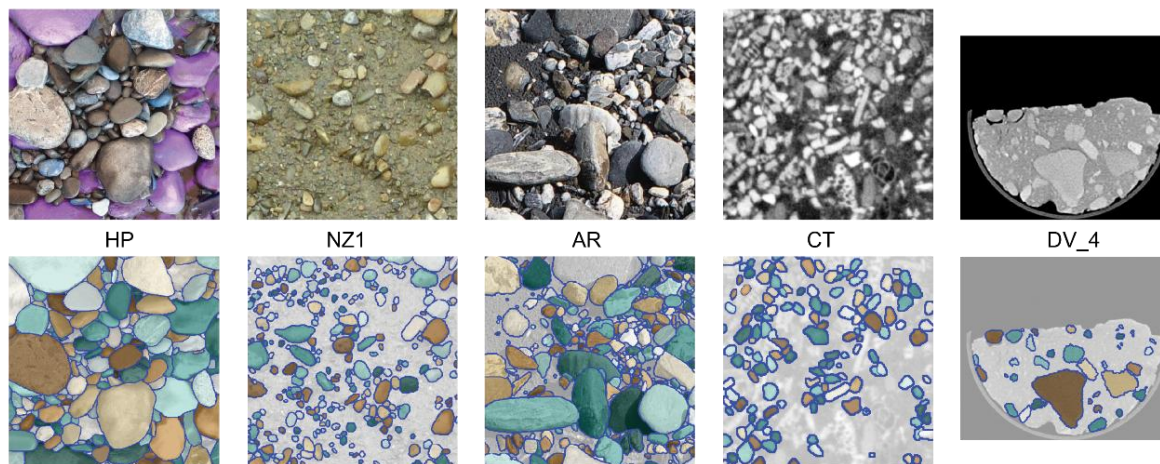


Figure 1: Example images and their manually annotated grain masks (labels) for various types of imagery and grains used in the IG2 dataset, and their respective indicated subset (Mair et al., 2025a; see Table 1 for more details). For visual examples of images with non-grain objects (e.g., vegetation) refer to Figure S1. Individual grains are shown in random colors with blue outlines. Please note that the examples for FH_2 and NZ1 are images taken from near vertical outcrops, where finer-grained background sediment cannot be resolved and therefore not annotated.

2 Methods

For ImageGrains 2.0, we employ the recently released Cellpose-SAM (Pachitariu et al., 2025) model architecture for segmenting biomedical images, which itself utilizes the Vision Transformer (ViT-transformer) of the Segment Anything Model (SAM; Kirillov et al., 2023) as backbone together with the gradient tracking of the original Cellpose framework (Stringer et al., 2021). Similar to the approach of Mair et al. (2024), we use a dataset consisting of images with annotated sediment grains (Fig. 1). In contrast to the dataset used by Mair et al. (2024), ImagegrainsImageGrains 2.0 (IG2; Mair et al., 2025a) is based on a much larger dataset including more image types (see Section 2.1). We use the IG2 dataset to train our model and to. We also evaluate its capability to quantify the size and shape of sediment grains (Section 2.2). In our approach, we apply transfer learning and retrain the pre-trained Cellpose-SAM foundation model to segment grains in images taken from clastic sediments (Fig. 2). In Section 2.3, we summarize key aspects of this foundation model and the adaptionsadaptations we made. Next, we describe how we set up other methods and models that we use to benchmark our approach (Section 2.4). We then proceed by quantifying and comparing the segmentation performances across all methods (Section 2.5). Finally, we obtain a set of aggregated metrics, which are based on the measured sizes and shapes of individual grains, to evaluate the effect of using segmented grain masks with varying precision. We note here that we use the term *method* for entire segmentation workflows, while the term *model* refers to a specific segmentation model. This distinction becomes important as some methods combine several models and we sometimes train several models ofusing the same method with different datasets.

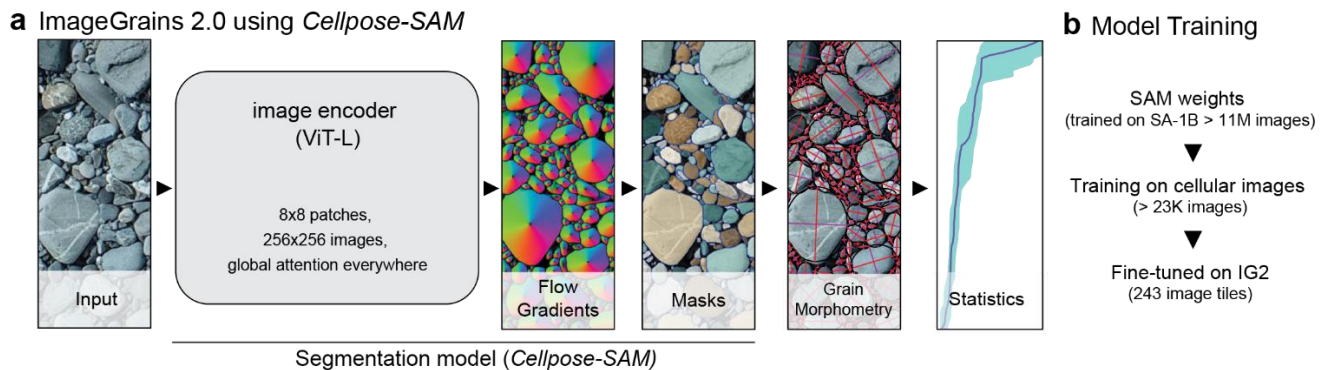


Figure 2: Overview of our workflow (a) that uses a re-trained Cellpose-SAM (Pachitariu et al., 2025) architecture for grain segmentation. The re-training (b) was done by fine-tuning Cellpose-SAM to the IG2 dataset (Mair et al., 2025a; see Table S1 for more details). ViT-L refers to the model weight used to initialize SAM, and SA-1B indicates the original training dataset of SAM (Kirillov et al., 2023). For details on the architecture of the segmentation model, refer to Pachitaru et al. (2025) and references therein.

2.1 The ImageGrains 2.0 (IG2) dataset

The IG2 dataset (Mair et al., 2025a) comprises over 29,000 manually annotated 2D masks of individual sediment grains captured in diverse image types, including RGB imagery taken from uncrewed aerial vehicles (UAVs), single-lens reflex cameras, compact digital cameras, and X-ray computed tomography (CT) slices (Table 1). Annotated grains cover a broad

range of sedimentary contexts, including fluvial gravels in outcrops (Garefalakis et al., 2023), fluvially transported pebbles, cobbles, and boulders on gravel bars and in river channels (Litty and Schlunegger, 2017; Mair et al., 2022; 2024), bioclastic sands from marine lagoons (Fabbri et al., 2024), and glaciofluvial deposits (Hiller et al., 2023; Schuster et al., 2025).

To create the IG2 dataset, we both expanded the existing subsets of the ImageGrains 1.0 dataset (IG1; Mair, 2023) and added new subsets. First, we added 5 additionally labelled image tiles to IG1 to improve the balance between the different image settings and data sources (Table 1) in the respective test and training splits. In particular, we added 1 image of a vertical gravel outcrop in the FH_2 subset, 1 image tile from the Swiss Sense River in the S1_2 subset, and 3 image tiles of fluvial pebbles from variable sources in the APF_2 subset. Second, we created 6 new subsets containing a total of 103 image tiles taken from fluvial sediments of rivers in Spain (with grains fully and partially painted in the field to track grain mobility and mobilization: HP, PP; Ville et al., 2023), Peru (PR), New Zealand (NZ2), Switzerland (AR), and Namibia (NB2). Furthermore, we included two more subsets of image tiles taken from coarse-grained and angular proglacial sediment (JF), and images retrieved from near-vertical outcrops of lithified conglomerates (NZ1). The images were acquired with different handheld and uncrewed aerial vehicle (UAV)-borne camera systems (Table 1) at varying resolutions. Finally, we completed the dataset by using X-ray CT (XRCT) scans taken from glacial tills (DV4), and micro-CT images of bioclastic marine sand (CT), which we annotated in two respective subsets. All these images were selected for variations regarding the objects displayed on the images. This includes - on purpose - the occurrence of vegetation and other objects that are not sediment grains to test the model against the possibility of false detections. We did so in order to challenge the models beyond variabilities in lithology, color, grain size and shape of the clasts. Specifically, we selected tiles that depict various objects such as scales, hands and equipment; tiles that featured different types of vegetation and water bodies (see Fig. S1 for examples of non-grain features), and that were acquired under variable light conditions. The combination of all subsets resulted in a total of 203 and 40 annotated image tiles that we used for training and testing, respectively.

Subset	Image Type	Setting	Image References	Annotator	Number of image tiles	Train	Test	Number of ROIs (grains)
APF_2	UAV, Handcamera	Fluvial Gravel	Mair et al. (2022), Mair et al. (2024), Chen et al. (2022), Brayshaw (2012), Litty & this study	DM	59	49	10	8090
AR	UAV, Handcamera	Fluvial gravel, debris flow deposits		DM	16	13	3	1011
CT	Micro-X-ray-CT	Bioclastic Sand	Fabbri et al. (2024)	DM	6	5	1	883
DV_4	X-ray-CT	Glaciogenic Diamict, drillcores	Schuster et al. (2025)	BS	19	16	3	849
FH_2	Handcamera	Outcrop gravel pit	Mair et al. (2024), Garefalakis et al. (2023)	DM	8	7	1	1426
HP	Handcamera	Fluvial gravel, fully painted particles	Ville (2024)	FV	7	6	1	2744
JF	UAV	Proglacial Sediments	Hiller et al. (2023)	DM	9	8	1	549
NB2	Handcamera	Fluvial Gravel	this study	AdP, AW, DM	42	35	7	2938
NZ1	Handcamera	Fluvial Gravel, near vertical outcrop	this study	DM	20	18	2	2971
NZ2	Handcamera	Fluvial Gravel	this study	DM	23	19	4	1809
PP	Handcamera	Fluvial gravel, partially painted particles	Ville (2024)	FV	8	7	1	3483
PR	Handcamera	Fluvial Gravel	Litty & Schlunegger (2017)	DM	7	5	2	733
S1_2	UAV	Fluvial Gravel Bar	Mair et al. (2024)	DM	19	15	4	2136
IG2 (all)					243	203	40	29622

Table 1: Dataset composition, image type, setting and summary statistics for the ImageGrains 2 (IG2; Mair et al., 2025a) dataset. ROI = region of interest.

To train and test segmentation models that are able to map a large variety of sediment grains on different image types, we ~~complemented and expanded~~ created labels for the ~~dataset (IG1; Mair, 2023)~~ dataset (IG1; Mair et al. (2024) by adding new ~~newly used~~ image data ~~and labels~~. For each image that was used in the expanded IG2 dataset, we chose subset tiles of varying sizes (ranging from 50×62 to 2750×2000 ~~pixel~~ pixels) that capture the full grain size variability and the complexity of the image content. ~~Specifically, we selected tiles that contained various objects such as scales, hands and equipment; tiles that featured different types of vegetation and water bodies, and that were acquired under variable light conditions.~~ This resulted in a large variety of tiles for each image (Mair et al., 2025; Table S1). These tiles were annotated manually using the LABKIT plugin (Arzt et al., 2022) for FIJI (Schindelin et al., 2012) and napari (v0.6; napari contributors, 2019), where each grain was labelled individually (i.e., dense labelling) as precisely as possible at the scale of individual pixels. In total, we added 162 such annotated image tiles from various sources and settings to the 81 tiles of fluvial sediment images compiled and annotated for the previous version (IG1). For all datasets, we manually generated representative subgroups of images, called stratified train and test splits (Table S1), to create internally balanced subsets for all imagery from different settings.

~~First, we proceeded by adding 5 additionally labelled image tiles to the original ImageGrains (IG1; Mair, 2023) dataset as a first task. The goal was to improve the balance between the different images and data sources (i.e., Brayshaw et al. 2012; Litty and Schlunegger, 2017; Mair et al., 2022; Chen et al. 2022; Garefalakis et al., 2023) in the respective test and training splits. In particular, we added 1 image of vertical gravel outcrops in the FH_2 subset, 1 image tile from the Swiss Sense River in the~~

S1_2 subset, and 3 image tiles of fluvial pebbles from variable sources in the APF_2 subset. In addition, we added 6 subsets containing 103 image tiles taken from fluvial sediment from rivers in Spain (with in the field painted elasts; HP, PP, n = 15), Peru (PR, n = 7; Litty and Schlunegger, 2017), New Zealand (NZ2, n = 23), Switzerland (AR, n = 16), and Namibia (NB2, n = 42). Furthermore, we included two more subsets of image tiles taken from coarse-grained and angular proglacial sediment (JF, n = 9; Hiller et al., 2023), and images retrieved from near-vertical outcrops of lithified conglomerates (NZ1, n = 20). The images were acquired with different handheld and unscrewed aerial vehicle (UAV) borne camera systems at varying resolutions (Mair et al., 2025a; see also Table S1). Finally, we completed the dataset by using X-ray CT (XRCT) scans taken from glacial tills (DV4, n = 19; Schuster et al., 2024; 2025), and micro-CT images of bio-elastic marine sand (CT, n = 6; Fabbri et al., 2024), which we annotated in two respective subsets. All these images were selected for variations regarding the objects displayed on the images. This includes on purpose the occurrence of vegetation and other objects that are not sediment particles to test the model against the possibility of false detections, in order to challenge the models beyond variabilities in the lithology, color, grain size and shape of the elasts. The combination of all subsets resulted in a total of 203 and 40 annotated image tiles that we used for training and testing, respectively.

2.2 2D Grain morphometry

Aside from quantifying the segmentation performance, we assessed the importance of precisely segmenting grain masks for yielding accurate results using grain size and shape metrics as benchmark information. For each grain mask, or region of interest (ROI), we used standard image analysis tools implemented in scikit-image (v0.25.2; van der Walt et al., 2014) ~~that have been successfully used~~ to represent the morphometry of ~~grains in geoscientific research (e.g., Szabó et al., 2015; Miller et al., 2024; Lepp et al., both grain size and shape as well as orientation.~~

~~To quantify grain sizes, 2024; Benet et al., 2024; Back et al., 2025).~~ Here we fitted ellipses to approximate the shape of the target grains for which we calculated the lengths of the minor and major axes (b- and a-axis, respectively, of an ellipse). This approach has been demonstrated to ~~well~~ capture well grain sizes of clastic material in 2D images (e.g., Purinton and Bookhagen, 2019; Chardon et al., 2022; Garefalakis et al., 2023; Mair et al., 2024; Sklar, 2024). The uncertainties of the grain size percentile values are quantified through bootstrapping, thereby resampling any grain size distribution (GSD) ~~a~~ 1000 times (for details, see Section 2.4 in Mair et al., 2022). We then calculated the differences between grains in the ground truth and predicted grains as difference for the percentile values. Furthermore, we tested if the GSDs of the ground truth and the predictions were statistically different ~~between predictions and ground truth~~ with a ~~we~~ two-sample Kolmogorov–Smirnov test.

Here, the two distributions being identical was the null hypothesis, which we ~~consider~~ rejected for $p > 0.05$. We note that due to the 2D nature of images, an inherent limit exists for detecting grains of small sizes, which itself is controlled by the image resolution (see Section 2.5 for details for our approach in this study). An important implication thereof is that the resulting GSDs are truncated at a specific threshold characterizing the smallest measurable size.

To quantify grain shape and orientation we used several approaches (Fig. 3) that have been successfully applied in recent studies across various geoscientific research fields (e.g., Szabó et al., 2015; Miller et al., 2024; Lepp et al., 2024; Benet et al.,

2024; Back et al., 2025a). We calculated the eccentricity ~~of the same ellipse fit to approximate the grain elongation in 2D, which is, i.e., the ratio of the focal distance over the length of the major axis, of the fitted ellipses to approximate the grain elongation in 2D, which is a classic metric for quantifying the rock fabric (Tucker, 1988).~~ Similarly, we used the convexity, sometimes also called solidity (e.g., in scikit-image; van der Walt et al., 2014), which is the ratio of pixels in the ROI to pixels within the convex hull, as proxy value for the 2D roughness of each grain. This parametrization of grain roughness has been frequently used (e.g., Cox and Budhu, 2008; Lepp et al., 2024) with relatively small differences compared to other parametrizations (Back et al., 2025b). Next, we obtained the isoperimetric ratio (IR) and normalized isoperimetric ratio (IR_n , Pokhrel et al., 2024; Quick et al., 2020) for each grain mask as indicator for the roundness of a grain. We note the selected approaches to compute IR and IR_n values can return values >1 in some cases, which could be the consequence of geometrically imperfect reconstructions (see supporting information of Quick et al., 2020). Finally, we ~~track~~measure the 2D grain orientation as azimuth angle of the b-axis of the above-described ellipse fit and the y-axis, i.e., the image height, of each image tile in degrees from 0° to 180° . Note that these metrics are an exemplary selection, and that any 2D shape and size metric can be calculated from the segmented grain ROIs.

We calculated all the aforementioned metrics for all ROIs in both ground truth and predicted masks that fall in the central ~~90%81% area~~ of each image tile by avoiding the outermost 5% from the edge of each image ~~edge~~. We did so to avoid a bias that could be introduced by considering grains – possibly cut ones – at the ~~border~~edge of image tiles. We then calculated differences between ground truth grains and predicted grains for all corresponding metrics. ~~By comparing these morphometric values across datasets, we can quantify how the segmentation quality affects the morphometric results.~~

2.3 Cellpose-SAM: re-training and inference

230 The Cellpose framework (Stringer et al., 2021) used a deep-learning model, which is based on a U-Net (Ronneberger et al., 2015) type of neural network with image style transfer (Gatys et al., 2016), a widely employed architecture that has proven very successful in image segmentation tasks (Siddique et al., 2021; Azad et al., 2024). This framework was combined with an equation modelled on heat diffusion to predict vector flows. From these flows, individual objects are segmented through gradient tracking. In Cellpose-SAM the previous backbone model was replaced with a modified version of the SAM transformer (Kirillov et al., 2023; see also Section 2.4.1 below for more details on SAM). Specifically, it used the image encoder module of SAM and replaced the decoder parts with a Cellpose’s vector representing the flow ~~representation~~ for prediction (Pachitariu et al., 2025). Moreover, the encoder itself was modified in several ways for the Cellpose-SAM architecture. First, the dimension of the input image was reduced to 256×256 pixels (from 1024×1024), and the patch size was reduced to 8×8 (from 16×16). This was done in an effort to decrease the processing time and to decrease computational costs. Accordingly, the position and patch embeddings were also ~~down-sampled~~downsampled, while global attention was used for all layers. This approach differed from using a global attention in only some layers in the original SAM architecture. It aims to improve the performance of tasks where dense segmentation is required, i.e., where multiple objects of the same class are segmented (for more details, refer to Pachitariu et al., 2025). ~~The Pachitariu et al. (2025) initialized the~~ Cellpose-SAM

model was initialized with the SAM ViT-L (large) model weights, which ~~itself allow for faster inference while maintaining a~~
245 high segmentation quality compared to ViT-H (huge; Kirillov et al., 2023). Both ViT-L and ViT-H had been trained on the
SA-1B dataset (Kirillov et al., 2023). Pachitariu et al. (2025) then trained Cellpose-SAM on an updated dataset of 22826
cell and cell nuclei images with > 3.3 million labelled objects. Notably, the updated architecture is much larger (> 304 million
trainable parameters compared to > 6.6 million trainable parameters in the old backboneU-net model). Furthermore, the
improved model can use multi-channel images, i.e., color, images, because of its training with random ~~channel~~-permutations
250 of channels. This was not possible with the previous models that converted the images to single-channel greyscale images
before the segmentation. As a result, multi-channel images are now the default image input.

We retrained the Cellpose-SAM model on our ImageGrains 2.0IG2 dataset using default settings and follow recommendations
of Pachitariu et al. (2025) for the custom re-training to obtain our new default model for ImageGrains. This included Our setup
from that of Pachitariu et al. (2025) in terms of the training length (we trained our model for 500 epochs with a instead of 100).
255 We did so, in an effort to maximize the learning rate of 1e-5. Here the model. For training, we used all 203 image tiles of the
train split in every epoch. ~~Training, and it~~ was accomplished during < within < 1.5 hours on a NVIDIA A100 GPU with 80 GB
memory at the UBELIX HPC cluster maintained by the University of Bern. The image tiles from the test split were used for
the validation of every 10 epochs. By default, the learning rate increased linearly from zero to 1e-5 over the first 10 epochs,
and then decreased by a factor of 10 every ten epochs over the last 50 epochs. The loss function was the default Cellpose
260 segmentation loss, which is the mean squared error between the 2D flows (in the XY plane; Pachitariu et al., 2025). This error
is calculated for the ground truth and the predicted flows, to which the cross-entropy between the probabilities of the ground-
truth and predicted objects is added. During the training, the images were randomly flipped, rotated and resized using a
uniformly distributed scaling factor between 0.5 and 1.5 before they were randomly cropped to 224×224 pixels. By default,
all image tiles were normalized to image intensity percentiles between 1 and 99 for each channel, and the AdamW optimizer
265 (Loshchilov and Hutter, 2019) together with a weight decay factor of 0.1.

Upon evaluating the model on 2D images, we ~~employed~~ again employed the default settings of Cellpose-SAM, which include
a block tile size of 256 pixels, a fractional tile overlap of 0.1, a threshold value of 0.0 for the object probability, and 0.4 for the
flow error, respectively. Again, the image intensity was normalized to the percentile range between 1 to 99 for each input
channel. All of these values were kept at default values of the algorithm, due to their demonstrated effectiveness (Pachitariu
270 et al., 2025). Contrary to previous Cellpose versions, no rescaling of image tiles was applied. For 3D segmentation, we used
the dedicated 3D approach of Cellpose (Stringer et al., 2021), which computes 2D flows and probabilities for slices in the XY,
YZ, and XZ planes: (Stringer and Pachitariu, 2025). The resulting values are then averaged to create 3D flow vectors. For the
construction of 3D segmentation masks, which themselves are generated from the 3D flow vectors, we used the default
computation, which considers a 3D smoothing factor of 1.0. This default setup allows for a high-quality 3D segmentation
275 despite upscaling from 2D segmentations (e.g., Zhou et al., 2025). For the demonstration of our 3D segmentation
demonstration performance, we used a stack of 400 TIFF images generated with XR-CT ~~from of~~ a drill core (from site
5068_1_C from 4-5m depth) ~~of taken from a~~ glacio-fluvial sediment infill in a glacially over-deepened valley in southern

Germany (Schuster et al., 2024). For details on the XR-CT scanning and image reconstruction, we refer to Schuster et al. (2025).

2.4 Other methods and models

We explored how well our approach compares to other methods that were publicly available at the time of writing this article and that were either used as a foundation model for general object detection or that were specifically tailored to segment grains. Note that we did not include the methods of Mörtl et al. (2022), Chen et al. (2024), or Soloy et al. (2020), because their models or code were not publicly available. Furthermore, we did not include the method of Chen et al. (2022) because of its relatively
weak performance in previous studies (Mair et al., 2024). ~~In particular, we~~
We first compared the outcome of our segmentation with that of the SAM (Vit-H) in its basic ~~mask generator~~-configuration to
generate masks. We viewed the performance of SAM as a baseline benchmark that every dedicated method should exceed, due to its ~~segmentation~~-capability to segment grains on a broad range of image datasets and object types (Kirillov et al. 2023) without fine-tuning for specific data, such as sediment grains. Next, we compared our results to the results of
Segmenteverygrain that uses prompt engineering for improving segmentations by SAM (Sylvester et al., 2025). Here, we used both their default prompt engineering model, SEG-SAM (default) and one that we trained on our IG2 dataset, SEG-SAM (IG2). Finally, to evaluate the relative improvement in segmentation performance with our new default model, we compared its segmentation results with those of the best performing model of Mair et al. (2024). In a second step, we evaluated~~Note that we did not include the methods of Mörtl et al. (2022), Chen et al. (2024), or Soloy et al. (2020), because their models or code~~
~~were not publicly available. Furthermore, we did not include the method of Chen et al. (2022) because of its relatively weak~~
~~performance in previous studies (Mair et al., 2024).~~ In a second step, we tested our default model’s ability to generalize to data not used during training with a setup where the S1_2 and PR subsets were ~~not used during~~excluded from training. We selected these two subsets for this test because for these subsets, the performance of both our fine-tuned Cellpose-SAM model and most benchmark models was highest amongst all ~~subsets~~subsets with heterogeneous image tiles of fluvial pebbles ~~under natural~~
~~conditions~~. Hence, we anticipated the largest impact on the segmentation performance if we left ~~these~~ out these data from the training split. Particularly, we compared the performance of our default model to that of all other models including specialized Cellpose v2 models, which were trained only on subset datasets used in this generalization.
In the following section, we briefly describe how we set up all benchmark models.

2.4.1 The Segment Anything Model (SAM)

SAM is a foundation segmentation model with a vision model transformer (Dosovitskiy et al., 2020; Li et al., 2022) ~~that has~~ been successful (Na et al., 2024; Archit et al., 2025). SAM itself was pre-trained with images from a large dataset of annotated images (11 million images with over 1 billion annotation masks; SA-1B) that was created with a custom data engine (Kirillov et al., 2023). This model can thus be used for segmenting a broad range of objects, and it is adaptable to more specific requirements related to various downstream tasks via prompt engineering, inspired by similar advances in Natural Language

310 Processing (Brown et al., 2020). The model itself consists of an image encoder, a mask decoder for inference, and a prompt
encoder that is employed for flexible and prompt handling (Kirillov et al., 2023). We used the ~~default~~largest available model
checkpoint ~~{ViT-H}~~ (huge), which demonstrated the highest segmentation precision (Kirillov et al., 2023), together with its
default mask decoder to predict grain masks. For this zero-shot instance segmentation (i.e., segmenting objects in images that
315 quality and duplicate masks. We used the predictions resulting from this model setup as baseline benchmark for grain
segmentation without custom method developments because they are based on data that is openly ~~available and~~accessible. In
addition, these predictions can be achieved without any fine-tuning or supervised training on images that display sediment
grains.

2.4.2 Segmenteverygrain

320 Segmenteverygrain combines SAM (i.e., the ViT-H checkpoint; see Section 2.4.1 for details) with a U-Net style convolutional
neural network for the prompt engineering upon segmenting grains in images (Sylvester et al., 2025). The default U-Net model
was trained on 66 different images displaying grains. The images themselves were split into 44,533 patches of 256×256 pixels.
We used both the default model and a model, which we fine-tuned with the entire IG2 dataset. Here, we employed the default
train/test splits of our IG2 dataset and the test split for validation- to ensure comparability with the other fine-tuned models.
325 Aside from this, we used the default configuration of Segmenteverygrain and followed the recommendation for fine-tuning
Segmenteverygrain-it (Sylvester et al., 2025). This configuration included the Adam optimizer and image augmentation. We
trained the refined model for 500 epochs and set the minimum object size to 15 pixels in order to match similar values of other
models. We used the predictions of Segmenteverygrain as benchmark for the approach referred to as prompt-based
segmentation.

330 2.4.3 Cellpose 2 models

To evaluate the impact of considering both the expanded datasets~~dataset~~ (IG2) and the new backbone architecture, we
compared the segmentation results of Cellpose SAM with those of older Cellpose (v2.3) models. We started by using the
IG1_full_set model of Mair et al. (2024), Cellpose 2 (IG1), which was trained on their original dataset that roughly comprised
a third of the images of the IG2 dataset (i.e., subsets S1_2, APF and FH; see also Section 2.1 for details on the datasets). We
335 then trained a Cellpose 2 model (Stringer and Pachitariu, 2022) with the same architecture on the full IG2 dataset, Cellpose 2
(IG2), using the same hyper-parameters and configuration as in the original publication (Mair et al., 2024). This included
training any Cellpose 2 model for 1000 epochs, a learning rate of 0.2 with a step-wise reduction of the learning rate by a factor
of two for every 10 epochs during the last 100 epochs and a batch size of 8 single-channel images, thereby employing the
default Cellpose implementation for image augmentation. Furthermore, Cellpose models can be trained as ~~a~~ specialist models
340 if fine-tuned to a specific dataset (Stringer et al., 2021; Mair et al., 2024). Therefore, we trained two more Cellpose 2 models

on the IG2 ~~data~~dataset but without S1_2 and PR image tiles. For further fine-tuning, we re-trained them only on the respective subsets S1_2 and PR.

2.5 Evaluating segmentation performance

We quantified the segmentation performance by first evaluating how closely individual predicted masks are representing ground truth annotations on an object (ROI) basis followed by calculating a precision score based on the number of grain masks that exceed an accuracy threshold (Everingham et al., 2015). This was done by comparing the predicted grain masks to the best-matching masks in the ground truth labels using the approach of Stringer et al. (2021). ~~This was done by~~ calculating average precision (AP) scores, evaluated at different intersection-over-union (IoU) thresholds. Specifically, we calculated the IoU metric for each grain mask with its closest ground truth match. We ~~use~~used an IoU threshold of >0.5 (for AP@0.5), and the increasingly stricter range of 0.5 to 0.9 (to calculate the average of AP values, i.e., mAP) to determine which grains were considered as true positives (TP). Grain masks in the ground truth that were not matched by a predicted mask with an IoU value above the aforementioned thresholds were counted as false negative (FN). Likewise, predicted grains with no corresponding grain mask in the ground truth that met the IoU quality criteria were considered as false positive (FP). The average precision ~~is~~was then calculated as the ratio between TP, and the sum of TP, FN and FP, i.e., $TP/(TP + FN + FP)$.

We used these standard metrics in object detection (e.g., Padilla et al., 2020; Yang et al., 2023; Zhou et al., 2025) to quantify the quality of the segmented grains, based on individual ROI quality, and to compare the results with those of the other methods. We chose to use this object-based ~~metric~~metrics because ~~it combines~~they combine both false negative and false positive detections in a single step, making it a more stringent metric than traditional metrics, such as precision, recall, or simple intersection over union (IoU) scores- (e.g., Minace et al., 2022). Furthermore, the average precision and mean average precision scores are ubiquitous metrics used to evaluate object-detection models in the field of computer vision, which includes models such as SAM (Kirilov et al., 2023), YOLO (Redmon et al., 2016) or Mask R-CNN (He et al., 2018)), among others. Similar to the approach of Mair et al. (2024), we excluded grains for which the minor axis of a simple ellipsoidal fit was <8 pixels both in the ground truths and in the segmented grain masks. The reason for this is that for most image types displaying sediments, we find it difficult to consistently distinguish between grains that are smaller than ~~these~~8 pixels during image annotation, which might render any predictions of smaller grains unstable. We acknowledge that ~~this~~threshold value can vary across different image settings, e.g., it is usually easier to identify very small grains in single-channel CT images with a high contrast than in multi-channel color images taken from fluvial sediments with a coarser resolution. This is in line with similar but larger thresholds (i.e., 20 or more pixels) determined by other approaches on similar fluvial sediment imagery (e.g., Chen et al., 2022; Purinton & Bookhagen, 2019; Chan et al., 2025).

3.1 Grain size and shape in ground truth ROIs

We first calculated standard 2D grain morphometry metrics (Fig. 3) for more than 18,500 manually labeled masks (ROIs) that remained after filtering for minimum grain size and distance to image tile edge, representing edges. These manually labeled masks represent about 63% of all labeled ROIs (for train and test split combined; Table S41). The resulting grain sizes vary substantially across the dataset, with b-axis lengths ranging from 8.0 to 481.0 pixels and a-axis lengths from 9.1 to 713.2 pixels (Table S2S1). The mean grain sizes reflect this variation of more than one order of magnitude, with average b-axis lengths ranging from 10.8 ± 2.7 pixels (CT) to 100.3 ± 111.3 pixels (NZ2). Within each data subset, grain sizes are highly variable, and systematic differences in mean grain size are observed between subsets for both a- and b-axes (Table S2S1).

The measured grain shapes vary among the data subsets. For grain roughness, expressed by the convexity values, the average of the unitless ration for the overall dataset is 0.93 ± 0.05 , which is consistent aeross-all with the averages of the respective data subsets, despite the subsets exhibiting a their high within-subset variability (Table S2S1). For example, the convexity shows a strong variability in some individual image tiles, with values ranging between 0.62 and 1.00. In general, a greater variation is observed for grain roundness values, expressed by the normalized isoperimetric ratio (IR_n , or circularity). Whereas the respective values generally range from 0.38 to >1.0 , the average IR_n values across data subsets range from 0.83 ± 0.07 (DV_4) to 0.97 ± 0.04 (CT), indicating systematic differences in roundness values between subsets (Table S2S1). The average value of IR_n of 0.89 ± 0.09 calculated in on the basis of the image tiles basis also reflects this broad variability. The data representing the grain elongation shows the largest variability with where the eccentricity values for grain ellipse-approximations ranging range from 0.25 to 0.97. Despite this broad range, the average eccentricity values aeresso of the data subsets are consistent showing with each other and yielding an average of 0.73 ± 0.14 , again demonstrating a strong within-subset variation (Table S2S1).

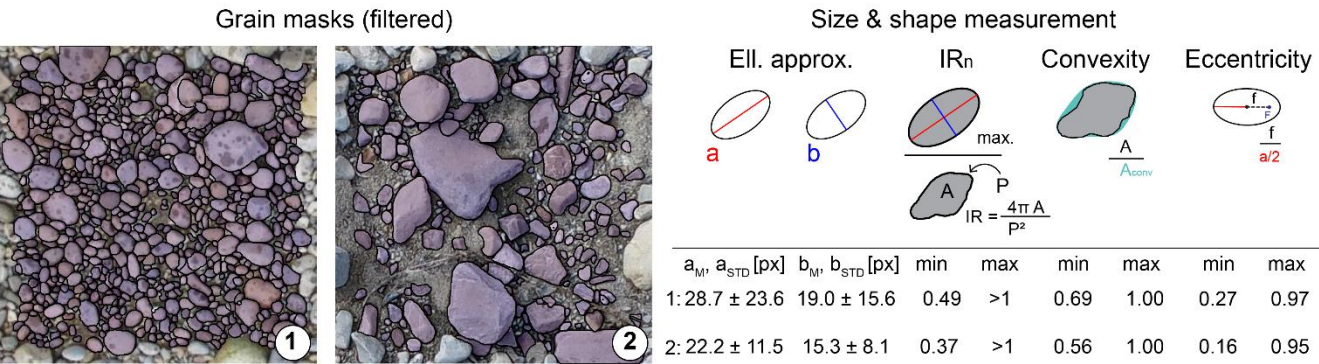


Figure 3: SelectedTwo examples of how grain size (with mean and 1 sigma standard deviation values), and shape measurementswere measured based on grain mask area and outlines, displayed here for two examples of annotated images tiles (labels) from subset APF_2. Ell. approx. = Ellipse approximation, a = major axis, b = minor axis, f = focal length, F = focal point, A = area, A_{conv} = area of the convex hull, P = perimeter, IR_n = normalized isoperimetric ratio (IR_n, or circularity; Pokhrel et al., 2024), px = pixel. For details on the metrics, refer to Section 2.2.

3.2 Segmentation performance

Overall, our default segmentation model, Cellpose-SAM (IG2), achieves a high accuracy in grain segmentation across all image types and most subsets (AP@0.5 >0.6; mAP >0.5 for train and test splits combined; Table 42; Fig. 4). Notably, the performance was largely independent from the dataset balance, i.e., the number of image tiles in a subset and its relative share of the entire dataset. For example, the segmentation performance for image tiles from subset DV_4 scores highest among all subsets (with an AP@0.5 value of 0.84; Table 2), despite their image type of XR-CT images accounting for only around 10% of all the image tiles in the dataset (Table 1).

Compared with alternative approaches, Cellpose-SAM (IG2) consistently outperforms all other models if applied to both the full IG2 dataset and across subsets (Table 42). This advantage is maintained in both training and test splits (Table S3S2; Fig. S2). Specifically, the median AP@0.5 across all test image tiles is 18% higher than that of the second-best method (0.71 vs. 0.3253 for Cellpose 2 trained on IG2; Fig. 4a; Table S3S2), and 20% higher across all image tiles (0.72 vs. 0.52 for Cellpose 2 trained on IG2; Table 42). The performance of the fine-tuned CellposedCellpose-SAM (IG2) model remains robust even for challenging image tiles, with almost no prediction-scoring AP@0.5 values below 0.4 (Figs. 4a, S1).

Upon comparing the performance of the methods other than our fine-tuned Cellpose-SAM (IG2), three observations can be made. First, the second-best model (using both the test and the full datasets) were trained with the IG2 dataset. Second, without fine-tuning to IG2, both Cellpose-SAM and Segmenteverygrain perform poorly (Figs. 4a, S1), which is expected since they were both fine-tuned to different image data (see Section 2.4). Finally, SAM (Vit-H) achieves moderate performance without fine-tuning to IG2, comparable to some other methods in methods in specific subsets (Table 42), highlighting its out-of-the-box segmentation capability. However, it does not match the performance of the best performing model.

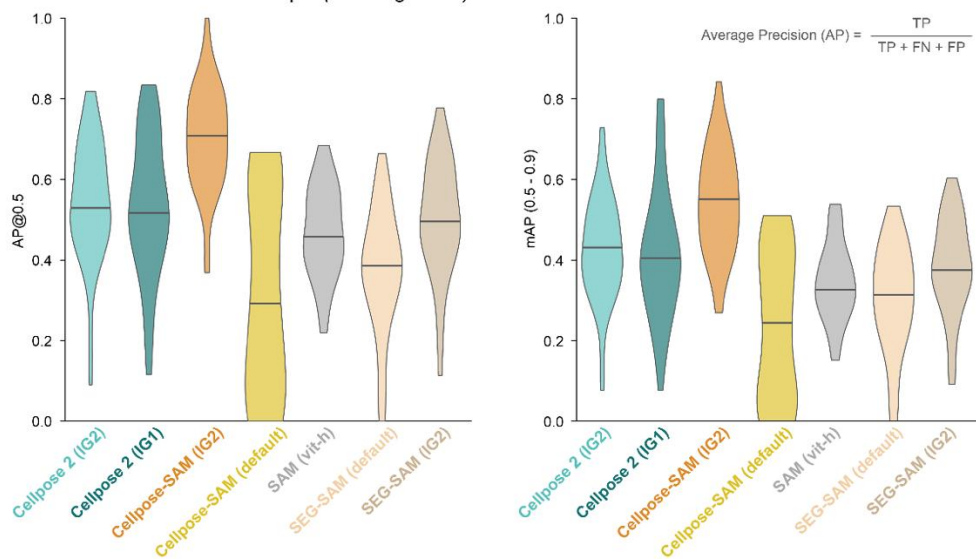
Metric	Method/Model	IG2 (all)	Data Subset													
			S1_2	PR	NZ2	FH_2	NZ1	NB2	APF_2	AR	JF	CT	DV_4	PP	HP	
AP@0.5	Cellpose 2 (IG2)	0.52	0.65	0.66	0.57	0.49	0.42	0.40	0.52	0.44	0.37	0.09	0.52	0.76	0.65	
	Cellpose 2 (IG1)	0.50	0.70	0.70	0.46	0.55	0.39	0.40	0.57	0.40	0.40	0.19	0.38	0.66	0.59	
	Cellpose-SAM (IG2)	0.72	0.80	0.74	0.73	0.69	0.58	0.68	0.69	0.63	0.72	0.68	0.84	0.83	0.76	
	Cellpose-SAM (default)	0.17	0.38	0.41	0.15	0.15	0.02	0.14	0.22	0.14	0.27	0.60	0.45	0.60	0.43	
	SAM (Vit-H)	0.44	0.56	0.49	0.44	0.48	0.37	0.37	0.43	0.38	0.42	0.54	0.49	0.60	0.54	
	SEG-SAM (default)	0.33	0.47	0.52	0.40	0.34	0.18	0.32	0.32	0.30	0.26	0.24	0.21	0.23	0.53	
	SEG-SAM (IG2)	0.51	0.58	0.68	0.63	0.58	0.44	0.41	0.48	0.50	0.34	0.07	0.56	0.79	0.71	
mAP	Cellpose 2 (IG2)	0.38	0.49	0.57	0.47	0.36	0.30	0.31	0.38	0.33	0.27	0.06	0.40	0.54	0.46	
	Cellpose 2 (IG1)	0.37	0.53	0.62	0.39	0.41	0.29	0.30	0.42	0.28	0.30	0.15	0.25	0.46	0.41	
	Cellpose-SAM (IG2)	0.55	0.62	0.63	0.58	0.50	0.42	0.52	0.51	0.49	0.51	0.54	0.74	0.60	0.55	
	Cellpose-SAM (default)	0.12	0.27	0.35	0.13	0.12	0.01	0.10	0.17	0.09	0.16	0.41	0.34	0.42	0.31	
	SAM (Vit-H)	0.31	0.41	0.41	0.35	0.34	0.25	0.25	0.30	0.26	0.29	0.42	0.36	0.45	0.38	
	SEG-SAM (default)	0.26	0.38	0.47	0.35	0.27	0.13	0.25	0.26	0.24	0.20	0.18	0.19	0.18	0.39	
	SEG-SAM (IG2)	0.38	0.43	0.55	0.50	0.43	0.28	0.30	0.35	0.37	0.23	0.05	0.43	0.57	0.48	

Table 2: Segmentation performance of all methods and models for the annotated grains in the IG2 dataset (test and train splits combined), and its subsets with the best performing model indicated in bold. All values are mean AP@0.5 or mAP values for all image tiles in the

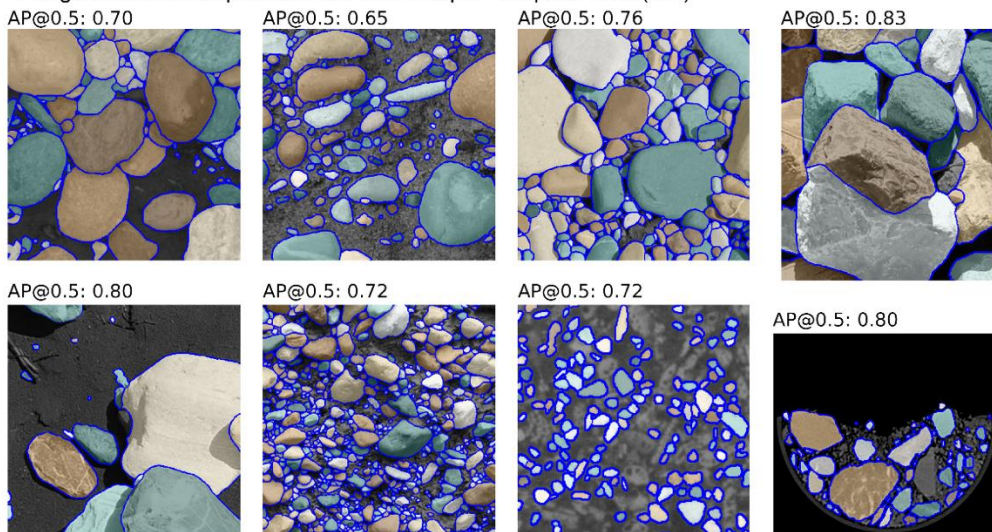
420 respective subsets, while for the entire dataset (IG2 –all), we report the median performance across all image tiles. Please note that the values for IG2 (all) are calculated on an image basis and therefore they are not the average of the respective values reported for the data subsets in the right part of the table. AP@0.5 = average precision evaluated at the intersection-over-union (IoU) threshold of 0.5; mAP = mean average precision for IoU thresholds ranging from 0.5 to 0.9.

We next evaluated the generalization capability of segmenting grains on image subsets that were not included in the fine-tuning for some model versions, specifically S1_2 and PR (Fig. 5). ~~Here,~~In this test of how well the model can deal with unseen
425 data, we find that a Cellpose-SAM model trained on all ~~IG+IG2~~ image tiles except S1_2 and PR outperforms all other methods that were not fine-tuned on these subsets (Figs. 5 a, b). Notably, for both subsets, this model achieves a performance that is comparable to some models that were trained with images from the respective subsets, including both versions of Segmenteverygrain and both versions of Cellpose 2 in PR (Figs. 5 a, b). Across both subsets, the overall best-performing model remains Cellpose-SAM that was trained on the full IG2 dataset (Tables ~~1, S22~~; Fig. 5). However, for S1_2, the older
430 Cellpose 2 architecture trained as a specialist model achieves a performance close to that of the best-performing model (Fig. 5).

a Performance on IG2 test split (40 image tiles)



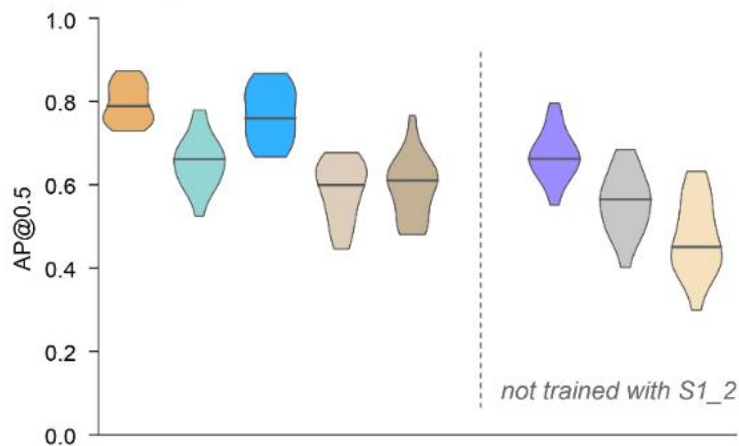
b Segmentation examples from the IG2 test split - Cellpose-SAM (IG2)



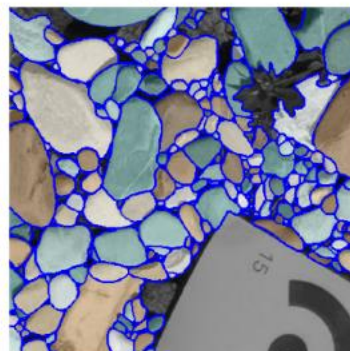
the IG2 dataset (test and train splits combined), and its subsets with the best performing model indicated in bold. All values are mean AP@0.5 or mAP values for all image tiles in the respective subsets, while for the entire dataset (IG2-all), we report the median performance across all image tiles. Please note that values for IG2 (all) are calculated on an image basis and therefore they are not the average of the respective values reported for the data subsets on the right.

Figure 4: Segmentation results for the IG2 test split calculated on an image-tile-basis with performance of the tested methods (a), and examples of predicted grain masks (b). AP@0.5 = average precision evaluated at the intersection-over-union (IoU) threshold of 0.5; mAP = mean average precision for IoU thresholds ranging from 0.5 to 0.9; TP = true positive, FP = false positive, FN = false negative. SEG-SAM = Segmenteverygrain.

a S1_2: Segmentation Performance

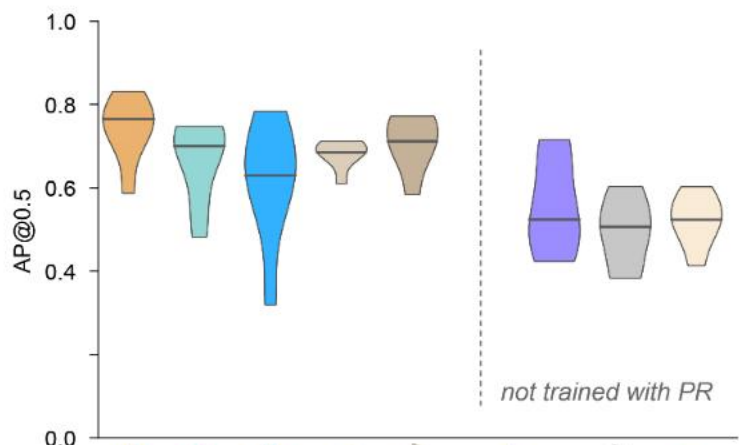


AP@0.5: 0.80

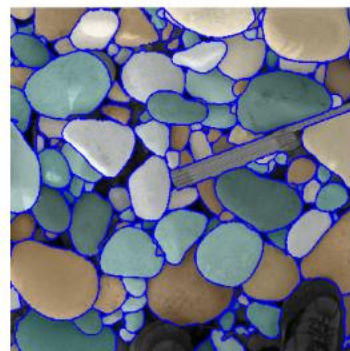


Cellpose-SAM (IG2) example

b PR: Segmentation Performance



AP@0.5: 0.79



Cellpose-SAM (IG2) example

Cellpose-SAM (IG2)
Cellpose 2 (IG2)
Cellpose 2 (specialist)
SEG-SAM (IG2)
SEG-SAM (specialist)
Cellpose-SAM (IG2 without S1_2, PR)
SAM (vit-h)
SEG-SAM (default)

Figure 5: Segmentation results for the S1_2 (a) and the PR (b) subsets. Specialist models were fine-tuned to the respective subset (see Section 2 for details). AP@0.5 = average precision evaluated at intersection over union (IoU) threshold of 0.5; SEG-SAM = Segmenteverygrain.

3.3 Size and shape accuracy of predicted grains

In a second step, we calculated the same 2D metrics (Fig. 3) for predictions from all ~~tested~~evaluated methods and compare them to the respective ground truth ~~for each~~ image tile. We first evaluate the differences (Δ) between predicted and ground

truth ROI masks for each model, aggregated across the full IG2 dataset (Fig. 6). Overall, the predictions derived from our default model (Cellpose-SAM trained on the full IG2 dataset) are the most accurate, showing the closest agreement with the ground truth for 9 out of 11 size and shape metrics (Fig. 6). For the two remaining metrics, mean ΔIR_n and mean $\Delta Eccentricity$, our model's predictions are still highly very similar to the best-performing alternative models (Fig. 6).

In more detail, the total number of detected grains of our default model, Cellpose-SAM (IG2) is also very close to the ground truth, achieving a 93% recovery rate (Fig. 6). Most notably, mean differences ~~(including the $\pm$ one sigma standard deviation range)~~ in grain size are around 3% or below ~~(with $\pm \leq 12\%$ standard deviation for the one sigma range)~~ for both the a- and b-axes, averaged across all detected grains. The resulting grain size distributions (GSDs) ~~, characterized by,~~ based on the lengths of both axes, are statistically identical to the ground truth (within 95% confidence, $p \geq 0.05$ for a two-sample Kolmogorov–Smirnov test) in 88% and 82% of cases for the a- and b-axes, respectively - compared to 54% and 57%, respectively, for the second-best model in this comparison. Our default model also yields the lowest ~~the average differences to the percentile differences values of the GSDs in grain size~~ the ground truth with mean ~~A values reaching values~~ differences of 0 ~~(pixels for the b-axis),~~ and 2.5 ~~(pixels for the a-axis) pixels, respectively~~ (Fig. 6). For all shape metrics, the differences between the predicted and ground truth grains are generally relatively small across all models. For our default model, mean Δ values are consistently below 2% (Fig. 6).

We next ~~examine~~ examined how the predicted grain masks from our default model compare to the ground truth ROIs across different data subsets. For grain size metrics (mean diameter and percentile differences), most subsets are close to the overall dataset average, with mean Δ a- and Δ b-axis differences within $\pm 10\%$ (Fig. S2S3). However, two ~~subsetssubsets~~ show a larger variability: AR and NZ2 have average relative differences of -12.4% and 16.4% for the a-axis, and -14.6% and 11.5% for the b-axis, respectively (Table S4S3). A similar variability between subsets is also observed for the average percentile differences in some subsets (Fig. S2S3, Table S4S3). Consequently, the GSDs for both a- and b-axes are statistically identical to the ground truth in five subsets (100% of image tiles). For another three subsets, the GSDs ~~remains of both axes are~~ identical ~~forto the ground truth in~~ more than 75% of the image tiles (Table S4S3). For four of the remaining five subsets, the accuracy of the GSD ~~differdiffers~~ particularly when the lengths of the a- and b-axes are considered separately, with 63–100% of GSDs matching the ground truth. Only subset AR shows a lower agreement where only 50% of GSDs match the ground truth. Yet this is still comparable to the best results of the second-best performing method and consistent with the average value of the full IG2 dataset (Fig. 6).

Considering grain shape, deviations in the shape metrics, i.e., mean IR_n , mean convexity, mean eccentricity and azimuth, of our default model from the ground truth are generally small. Here, the deviations ~~remain below are less than~~ 5% (Fig. S2S3) even in subsets with the largest differences: between predictions and ground truth. The only notable exception is the mean IR_n value for DV_4, which deviates by more than 10% (Fig. S2S3; Table S4S3). Finally, the inter-image variability contributes to higher relative standard deviations for several mean difference values in both grain size and shape metrics, resulting in a broader spread in Fig. S2S3.

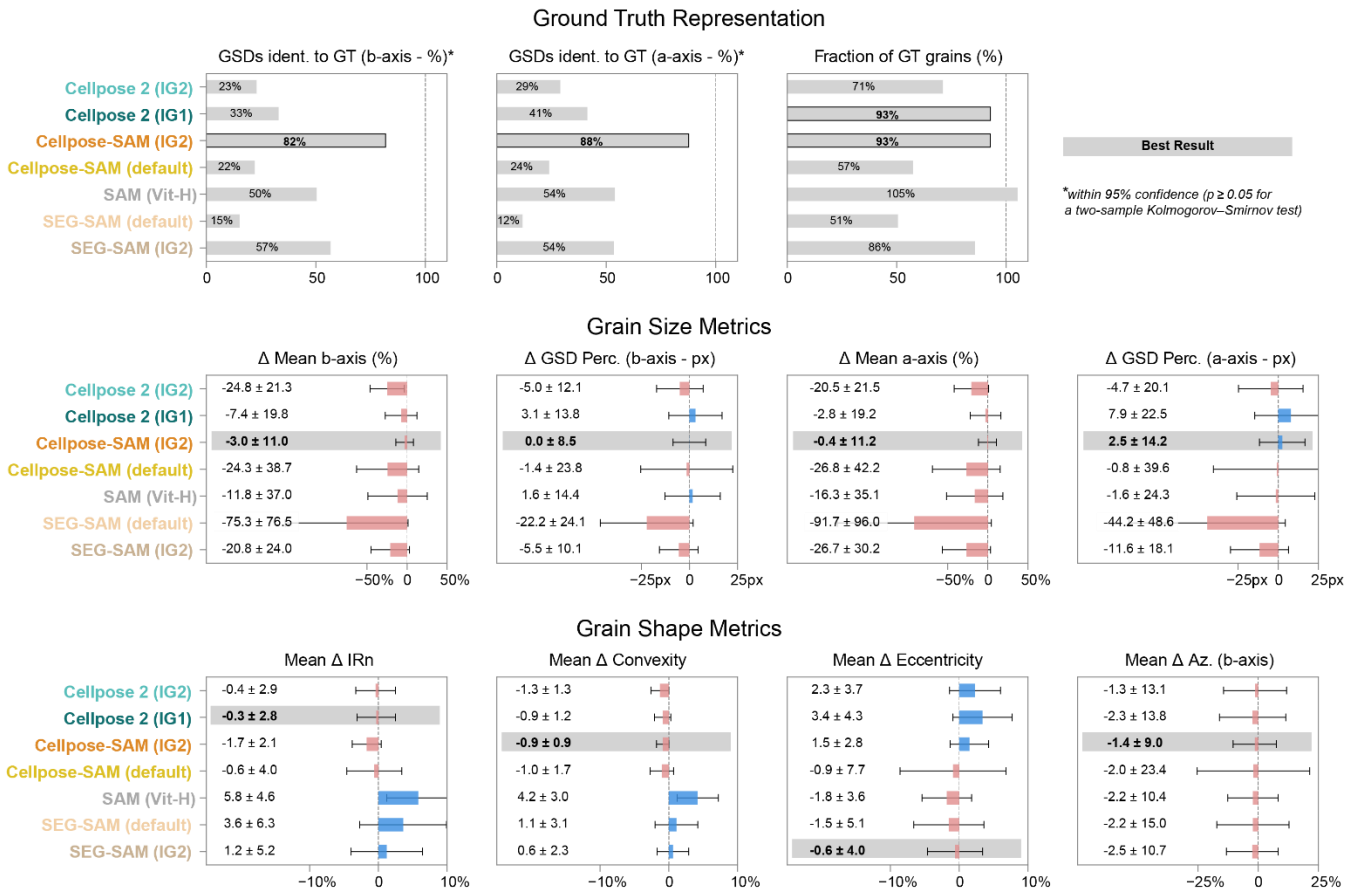


Figure 6: Summary of differences in 2D grain morphometry metrics calculated in relation to manually labelled grain masks (ground truth) across all data subsets. Mean and average standard deviation (1 sigma) values are calculated for image-averaged values. Values for the best performance in each metric are indicated in bold, while the dashed line indicates perfect representations of the ground truth. GT = ground truth, GSD = grain size distribution, Perc. = percentile, IRn = normalized isoperimetric ratio (IRn, Pokhrel et al., 2024; Quick et al., 2020), Az. = azimuth. SEG-SAM = Segmenteverygrain. For details on individual metrics refer to Section 2.2.

4 Discussion

The results show that our IG2 dataset can be used to successfully train and evaluate deep learning models for segmenting individual sediment grains in a broad variety of images ~~and taken from a~~ large range of depositional settings (Section 4.1). By re-training a state-of-the-art segmentation model (Cellpose-SAM) originally developed for bio-medical research with images of sediment grains, we obtain a model that significantly outperforms other current methods for the same task (Section 4.2), despite using the same backbone architecture and starting weights from the Segment Anything Model (SAM). The high-quality segmentation masks generated by our approach allow us to quantify how well grain size and shape are reconstructed relative to the ground truth, which we directly relate to the excellent segmentation performance (Section 4.3). Finally, we discuss the limits of our approach and avenues for future development (Section 4.4).

4.1 IG2 ~~Dataset~~dataset composition and characteristics

The IG2 dataset comprises over 29,000 manually annotated 2D masks of individual sediment grains captured in diverse image types, including RGB imagery taken from uncrewed aerial vehicles (UAVs), single lens reflex cameras, compact digital cameras, and X-ray computed tomography (CT) slices (Table S1). ~~Annotated grains cover a broad range of sedimentary contexts, including fluvial gravels in outcrops (Garefalakis et al., 2023), fluvially transported pebbles, cobbles, and boulders on gravel bars and in river channels (Mair et al., 2022; 2024; Litty and Schlunegger, 2017), bioclastic sands from marine lagoons (Fabbri et al., 2024), and glaciofluvial deposits (Schuster et al., 2025; Hiller et al., 2023).~~

The IG2 dataset was designed to encompass the broadest possible range of grain types (lithology, shape), bedding characteristics (imbrication, fine-material patches), image acquisition conditions (lighting, shadows, brightness), and background elements or non-grain objects (vegetation, scale objects, water bodies). Non-grain objects were excluded from annotation. Where shadows were present, only visible grain boundaries were traced across the shadows. Consequently, the trained Cellpose-SAM model effectively ignores a variety of non-grain features, such as sieves, scales, shoes, ground control point (GCP) markers (e.g., Fig. 5), partial shadows (e.g., Figs. 3, 4), and vegetation that have previously hampered automated grain detection (e.g., Chan et al., 2025; Miazza et al., 2024; Mair et al., 2022). However, the ~~robustness of the~~ grain segmentation ~~may/might~~ only generalize to be robust for those objects similar in shape with shapes, size-ranges, and color colors that are similar to those in the training data. Users should also note that grains partially obscured by other objects can introduce biases in quantifying the size and shape of grains, as their reconstructed outlines may deviate from the true boundaries. Yet, this is a problem that is associated with all image-based data collections.

The broad range of imagery and settings ~~resulted in~~yielded a substantial variability ~~in grain size of sediment grains with different grain sizes~~ and ~~shapes~~shapes, with a-axis values ranging from 9 pixels to over 700 pixels (Table ~~S2~~S1). This order-of-magnitude range exceeds that of many object-detection tasks and has previously hindered a robust segmentation across the ~~full~~entire size-spectrum (e.g., Chan et al., 2024; Mair et al., 2024). Variations in grain roundness and elongation (Fig. 3, Table ~~S2~~S1) surpass those observed along major terrestrial rivers (e.g., Quick et al., 2020; Pokhrel et al., 2024) and even Martian systems (e.g., Szabo et al., 2015). This large range in grain size and shape makes our dataset ideal for evaluating the capability of segmentation models for ~~grain shapes~~-reconstructions of grain shapes. We emphasize that these ground-truth sizes and shapes are not intended to represent specific geomorphic conditions, but they are rather considered to serve as a benchmark for evaluating the fidelity of model-predicted grain masks on images taken under highly variable conditions (see Section 3.2; Fig. 3).

Due to the high variability in grain size, shape as well as general image content and the relatively small number of 243 image tiles, some dataset imbalance persists between training and test splits. Image tiles were carefully selected to minimize this, but divergent segmentation performance in certain cases (e.g., HP; Table ~~S3~~S2) suggests the occurrence of some remaining imbalance in specific subsets.

4.2 Capabilities of Cellpose-SAM

4.2.1 Segmentation performance and generalization ability

Our results demonstrate that the high segmentation accuracy achieved by the Cellpose-SAM architecture on biomedical images (Pachitariu et al., 2025) can be effectively transferred to the segmentation of sediment grains. On average, our default model ~~that, which~~ was trained on ~~the~~ IG2 ~~dataset~~, correctly segments a larger number of grains and achieves a higher precision than all benchmark models considered in this study, with an average improvement of $\Delta AP@0.5 = 0.18$ compared to the second-best model (Fig. 4; Table ~~1~~2). It outperforms both earlier Cellpose 2 models, which employed a U-Net backbone, and workflows that also incorporated the SAM, such as SAM itself (ViT-H backbone; Kirilov et al., 2023) and Segmenteverygrain (Sylvester et al., 2025). In contrast, the ~~not fine-tuned~~ Cellpose-SAM model not fine-tuned with the IG2 dataset exhibits the lowest performance (Fig. 4), underscoring the effectiveness of re-training and fine-tuning even with comparatively small datasets, such as ours. Notably, training the other benchmark models on the IG2 dataset leads to moderate performance gains for Segmenteverygrain, but little to no improvement for Cellpose 2 (Fig. 4). This suggests that the default Segmenteverygrain model was originally trained on imagery substantially~~markedly~~ different from our IG2 dataset, whereas the Cellpose 2 model previously used by Mair et al. (2024) seems to have had already reached its performance limit for images of coarse-grained fluvial sediments. Across the benchmark models used in this study, the median segmentation performances of SAM, the trained Segmenteverygrain, and Cellpose 2 were broadly comparable, with differences mainly at the upper and lower end of the performance distributions. However, ~~they~~these benchmark models did not achieve the same level of segmentation performance as the fine-tuned Cellpose-SAM model (Fig. 4). This outcome aligns with the results of Pachitariu et al. (2025), who demonstrated similar advantages of Cellpose-SAM over other SAM-based architectures (Na et al., 2024; Israel et al., 2024) in biomedical segmentation tasks.

We have evaluated the generalization capability of Cellpose-SAM by training a model that excluded the PR and S1_2 subsets. When compared with Segmenteverygrain and SAM (ViT-H), Cellpose-SAM achieves the highest segmentation performance across both datasets, though with notable differences between them (Fig. 5). For the S1_2 image tiles, the model performs on par with several benchmark models ~~that~~ were explicitly trained on the S1_2 data. In contrast, the model performance for the PR images is more variable, with median $AP@0.5$ values similar to those of SAM and Segmenteverygrain when these models were not trained with the PR data (Fig. 4). Overall, for the PR tiles, the best-performing model ~~remained~~was our default ImageGrains model, which included PR images in its training data.

These results demonstrate the strong generalization capability of Cellpose-SAM, while also indicating that its performance depends, to some extent, on the type of images and ~~how close these are, in regard to both~~on the similarity of the objects of interest and non-grain objects; to the data that were used during training. Additionally, our findings show that Cellpose 2 models trained as dataset-specific specialists, i.e., on smaller and more homogeneous image sets, can achieve a segmentation accuracy, which is comparable to, or even exceeding that of the more generalist SAM-based architectures. This was particularly evident for the S1_2 data subset (Fig. 5). Furthermore, our default Cellpose-SAM was the best performing model in all subsets

(Table 2) if the data was included in training irrespective of the relative share of the respective subset of the entire IG dataset (Table 1). Particularly the performance in subsets with few images of different types (e.g., DV_4, CT) suggests that the inter-subset balance has little effect on the model's generalization ability. We hypothesize that the ability of the SAM architecture to learn a multitude of representation categories, even with class imbalance (Kirilov et al., 2023), is the underlying cause for this performance. Therefore, the deep architecture, combined with the very general initial training of SAM and the two iterations of fine-tuning (first by Pachitariu et al. 2025 and then in our workflow; cf. Fig. 2), should effectively eliminate the need for an inter-subset balance during model training.

4.2.2 3D Segmentation

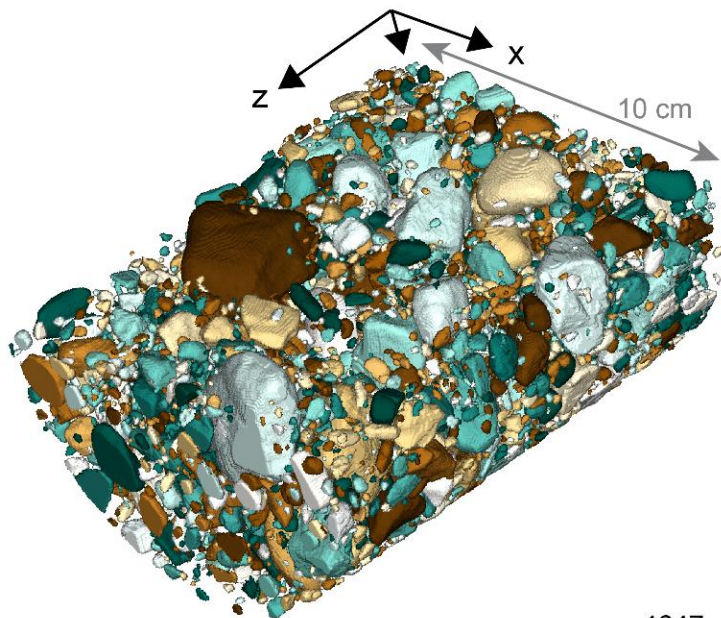
Data about grain size and shape achieved through image-based measurements in 2D have been successfully used to investigate sedimentary systems across a broad range of research (e.g., Garefalakis et al., 2024; Allen et al., 2017; Williams et al., 2013; Marchetti et al., 2022). However, 2D data do not always provide an accurate representation of the size and morphology of grains in 3D (e.g., Garefalakis et al., 2023; Steer et al., 2022; Bunte and Abt, 2001). Consequently, the segmentation of sedimentary grains in 3D remains a critical objective, even though several methodological advancements have been made in recent years to address this challenge (e.g., Rheinwalt et al., 2025; Steer et al., 2022; Walicka and Pfeifer, 2022; Domokos et al., 2024; Kettler et al., 2023). Most of these approaches are designed for segmenting grains from datasets approximating the grains' surfaces in 3D, such as topographic point clouds or meshed grids derived from LiDAR (e.g., Brodu and Lague, 2012) or structure-from-motion (SfM) photogrammetry (e.g., Eltner et al., 2016; Woodget et al., 2018). Grains from such datasets are typically only partially visible due to occlusion (Rheinwalt et al., 2025), which often necessitates the fitting of predefined geometric models during segmentation (e.g., Steer et al., 2022).

In contrast, XR-CT scans can, despite some inherent technical limitations due to scanning contrast and resolution, provide complete volumetric representations of entire grains (e.g., Cnudde and Boone, 2013; Houston et al., 2013; Mitra et al., 2024). These XR-CT data are analogous to the 3D image stacks of microscopic samples used to train and evaluate the 3D-segmentation capabilities of Cellpose (Stringer et al., 2021) and Cellpose-SAM (Pachitariu et al., 2025). Schuster et al. (2025) demonstrated that models with a Cellpose 2 architecture can be successfully trained to segment coarse grains in 3D XR-CT stacks of images taken from glacio-~~fluvi~~ally transported sediment, fluvial gravels. In our study, we incorporated the annotated 2D images from Schuster et al. (2025) as subset DV_4, along with annotated micro-XR-CT imagery from Fabbri et al. (2024), to leverage the dedicated 3D capabilities of Cellpose-SAM. We note here that we used these datasets to test the 3D segmentation quality of our default segmentation model and that we refer to cited literature for details on and a discussion of 3D XR-CT scanning itself.

The 3D segmentation of the used example of stacked images yielded 4,647 visually well-defined coarse grains from ~~glaciofluvial~~glacio-fluvial diamictic sediment (Fig. 7). Among all IG2 data subsets, the segmentation performance on the 2D image tiles was highest for DV_4, with a mean AP@0.5 of 0.84, whereas the CT dataset achieved an intermediate performance (Table 42). Notably, for the DV_4 images, the new Cellpose-SAM-based approach produced a net increase in mean AP@0.5

595 of more than 0.2 compared to the Cellpose 2 model of Schuster et al. (2025). These results indicate that the new default model implemented in ImageGrains 2.0 is well suited for such datasets, with its 3D segmentation capability being particularly promising. It should be noted that, ideally, 3D ground-truth labels would be required to rigorously benchmark the segmentation performance in 3D. However, to our knowledge, the manual effort required to annotate large numbers of image slices in 3D XR-CT stacks has so far impeded the creation of such a reference dataset.

3D Segmentation example



n = 4647
glacio-fluvial diamictic sediment
400 XY - images of XR-CT scan
for details, see Schuster et al. (2025)

600 **Figure 7:** Example where grains were segmented with the default Cellpose-SAM-based segmentation model of ImageGrains 2.0 in 3D from a stack of 400 XR-CT scans taken from of a drill core (drill site 5068_1_C from 4-5m depth; Schuster et al., 2024) made up of coarse-grained glacio-fluvial sediment.

4.3 Relating size and shape accuracy to segmentation performance

605 Segmentation-based approaches for measuring grain size and shape have long been limited by inaccuracies arising from over-segmentation, under-segmentation, and imprecise grain boundaries (Chardon et al., 2022; Mair et al., 2022; Steer et al., 2022). The introduction of deep-learning models has substantially improved the accuracy and precision of automated grain segmentation in 2D imagery (Mair et al., 2024; Miazza et al., 2024), with recent developments additionally leveraging the capability of SAM (Chan et al., 2025; Sylvester et al., 2025).

Our new default segmentation model within ImageGrains — a fine-tuned Cellpose-SAM model — further improves both the accuracy and the precision across the entire IG2 dataset, achieving ~~up to 20%an~~ improvement of up to a 20% in AP@0.5 and mean average precision (mAP) metrics compared to previous models (Section 3.2; Fig. 4). Moreover, the reconstructed grain masks produced by this new default model most closely match the ground-truth ~~regions of interest (ROIs)~~ in both size and shape measurements among all benchmark models (Section 3.3; Fig. 6). These results confirm the inferences (e.g., Mair et al., 2024) where an improvement in the segmentation performance results in a more accurate quantification of the size and shape of individual grains.

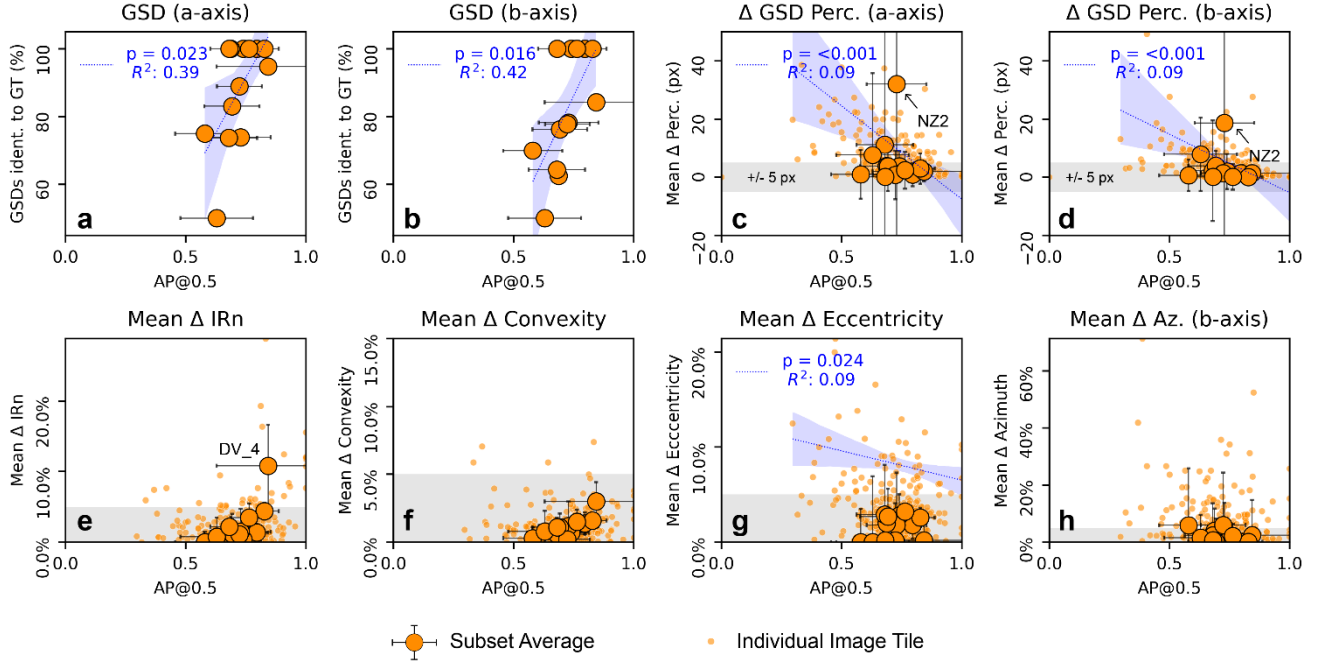


Figure 8: Comparison of segmentation performance metric (AP@0.5) with relative differences in grain size (a-d) and shape (e-h) between predicted grains and the ground truth ROIs for our new default model (Cellpose-SAM IG2). Grey areas indicate very low differences between predicted grain masks and ground truth ROIs, with differences within ± 5 pixels (c, d) and $<5\%$ (e-h), respectively. Only statistically significant correlations ($p \leq 0.05$, $R^2 \geq 0.05$) for individual images are indicated. For shape metrics (e-h) only values with Δ values $>5\%$ were considered for correlation. R^2 = coefficient of determination. Please note that the y-axes in panels c, d are cropped for a better visualization of the bulk of results.

Despite this strong overall performance, the model yields returns results withof variable quality across data subsets and individual images (Tables 1, S4; Figs. 8, S2). This variability provides an opportunity to assess how informative the model predictions are for individual size and shape metrics (Fig. 8). In this context, we observe the proportion of grain-size distributions (GSDs) of predicted grains that are statistically indistinguishable from the ground truth ($p \geq 0.05$; two-sample Kolmogorov–Smirnov test) to increase significantly with segmentation performance for all data subsets (Fig. 8). We observe similar trends in the results of other tested methods, with an even higher statistical significance due to their generally lower and more variable segmentation performance (Fig. S3S4). Although the relatively large number of statistically

indistinguishable GSDs prevents the definition of an AP threshold for perfect correspondence, a 100% match between corresponding ~~grain size distributions~~GSDs in the ground truth and the ~~results~~predictions is only achieved for cases with average AP@0.5 values exceeding 0.68 (Figs. 8, S4). For some ~~dataset subsets~~, higher average AP values ~~yield~~do not necessarily ~~yield~~ perfect GSD matches, underscoring the importance of effects related to dataset-specific variabilities; however, no perfect match is achieved below this average AP@0.5 value of 0.68.

This seemingly minimum threshold of 0.68 AP@0.5 IoU (or c. 0.5 mAP@50-90) for statistically representative GSDs depends on three factors: i) the specific segmentation results (i.e., whether FP, FN are caused by over/under-segmentation or missed grains), ii) the grain size sorting, and iii) the shape of the GSD itself. These factors are interdependent. For example, mis-segmented grains have a much smaller effect on the obtained ~~size distribution~~GSD for a very well sorted sediment with a narrow and Gaussian-shaped GSD than for a poorly sorted sediment with log-shaped or bi-modal GSDs. Therefore, it is unlikely that such a minimum AP threshold exists as ~~a single, generally an universally~~ applicable value. This implies that efforts to improve ~~GSD the~~ accuracy ~~best of the GSD should be directed with an~~ aim ~~at improving to improve~~ the segmentation accuracy, ~~i.e., This can be achieved~~ through ~~model~~ fine-tuning ~~the model~~ to custom annotations and ~~through~~ adhering to classic ~~guidelines for~~ image acquisition ~~guidelines, such as using. This includes the use of~~ suitable sensors with sufficiently high resolutions ~~and, thereby~~ maintaining stable image acquisition conditions (e.g., ~~Mair et al., 2022~~Bunte and Abt, 2001; Baptista et al., 2012; Bertin and Friedrich, 2016; Purinton and Bookhagen, 2019; ~~Baptista~~Mair et al., ~~2012~~2022).

The average differences between ~~the~~ percentile values of GSDs ~~is~~are generally small, typically < 5 pixels (both a-, and b-axes; Table ~~S4~~S3) for a majority of the data subsets and for most individual images (Fig. 8). In general, correlations ~~—~~, where significant ~~—~~, indicate that higher AP@0.5 values are associated with smaller differences in percentile-based GSD metrics relative to the ground truth (Figs. 8, S3).

For most grain shape metrics, which include roughness (mean IR_n), roundness (mean convexity), elongation (mean eccentricity), and orientation (mean azimuth), the corresponding differences between predicted masks and the ground truths are minor (< 5% on average) for our default model (Fig. 8; see also Fig. 6 and Table ~~S4~~S3), both at the data-subset and individual-image levels. We therefore conclude that segmentation performances with $AP@0.5 \geq 0.6$ enable sufficiently accurate and precise reconstructions of the grains' shapes for nearly all images, with only a few outliers (Fig. 8). The other tested methods exhibit larger differences between predicted grain masks and the ground truths, and correlations between segmentation performance and reduced deviations are more frequently detectable (Fig. ~~S3~~S4). Overall, the models that were fine-tuned with our IG2 dataset tend to ~~yield~~return results that are more reliable in representing the shape than ~~in representing~~ the size of grains.

In summary, the segmentation performance, expressed by AP scores combining false-negative and false-positive detections, correlates with the degree of agreement between predicted and ground-truth grain properties. Therefore, improving the segmentation performance indeed reduces the differences between predicted grain masks and ground truths until they are statistically identical. On our data, our new default model reaches that performance threshold for some data subsets in all metrics (Fig. 8).

665 4.4 Limitations, ~~Applicability~~applicability, and outlook

4.4.1 Limitations of image-based segmentation

Segmentation-derived grain size and shape data from 2D imagery are widely used across geoscientific disciplines, but the 2D nature of images itself imposes some limits on the applicability of segmentation-based approaches. First, ~~image data have size limits for grain detection~~the sizes of grains that can be detected and measured are controlled by ~~image~~the resolution ~~and of the~~
670 image and the content: displayed on it. In the IG2 dataset, our default segmentation model applies a filter with a minimum size ~~filter~~ of 8 pixels for the minor axis of the fitted ellipses. This threshold is close to the technical lower limit of the Cellpose-SAM backbone, which can theoretically detect circular objects ~~as small as~~with a diameter $\geq$ 5 pixels ~~in diameter~~ (Stringer et al., 2021). The reported lower detection limits ~~for~~for grain segmentation vary considerably among studies, depending on image type and resolution: 6 pixels (Schuster et al., 2025), 12 pixels (Mair et al., 2024), 20 pixels (Chen et al., 2022; Purinton and
675 Bookhagen, 2019), and up to 30 pixels (Chan et al., 2025). At the upper end of the size spectrum, the size of detectable objects is limited by the image extent. By default, Cellpose-SAM detects only objects that are ~~not larger~~smaller than 40% of the entire image area (Pachitariu et al., 2025). In addition, although the SAM-based architecture can handle a greater range of size variability than previous models (Pachitariu et al., 2025), an extremely large variability in grain size, i.e., spanning more than an order of magnitude, may require ~~combining that several~~ masks predicted are combined, which themselves are based on
680 predictions from multiple segmentation runs using differently rescaled images (Chan et al., 2025). We note that the size-limits of objects due to the 2D nature of images always result in truncated GSDs. This effect can complicate comparisons with data collected through other methods, e.g., when comparing grid- or areal-derived GSDs (e.g., Bunte and Abt, 2001; Steer et al., 2022).

Second, our approach is well suited for segmenting objects in 3D in stacked imagery data (see section 4.2.2 above). However,
685 it is not well suited for segmenting ~~3D objects based on surface~~ data ~~of surfaces, which are routinely alone or incomplete 3D representations of grains, such as those~~ obtained from topographic point clouds. ~~Usually, such data do not contain the complete 3D representation of grains, and therefore.. Therefore, they~~ require a geometric extrapolation to ~~successfully~~ segment those grains successfully (Steer et al., 2022; Rheinwalt et al., 2025).

Third, some limitations arise from image type and content. While modern deep-learning models, particularly Cellpose-SAM,
690 can be applied to a broad variety of imagery because of the strong capability for generalization (Pachitariu et al., 2025; see also Section 4.2.1), the segmentation can be hampered, and, thus, the performance be limited ~~by~~due to the occurrence of complex contents in the imagery itself. Examples include unrelated objects or vegetation, challenging lighting conditions (e.g., shadows, reflections, or glare from water bodies), motion blur, and color imbalance. To counter these effects, we composed the IG2 dataset with images encompassing a broad range of ~~image~~ conditions at which the images were taken, thereby
695 enhancing the model's robustness. Consequently, our default model is able to ~~effectively~~ handle a variety of adverse imagery conditions and un-related objects effectively. However, due to the finite size of the dataset ~~size~~, the performance may decline when applied to ~~image types~~imagery that ~~are~~ substantially different from those differs from the contents displayed on the

imagery or type of imagery represented in the IG2 dataset. Examples of such differences include: objects that are very different from clasts (such as man-made objects), images taken in very different vegetation zones and on other planets, or entirely new types of imagery such as images of thin sections. In such cases, fine-tuning with a small number of additional training tiles (as few as seven or fewer) can yield substantial improvements, as demonstrated for the CT, HP and FH_2 data splits (Table S11). Here, the combination of a deep transformer architecture with SAM's generalist encoder weights effectively reduces the need for extensive dataset balancing, which was essential for earlier, shallower architectures (e.g., Mair et al., 2024).

Finally, an insufficient segmentation performance likely poses the major limiting factor for some applications. The definition of an exact threshold that constitutes a sufficient level of segmentation performance remains an application- and dataset-specific task (see section 4.3; Zaidi et al., 2022). In case of insufficient segmentation (over- or under-segmentation, as well as missed grains), a manual post-processing, i.e., correcting the mis-segmented grains, is recommended. This could be done through the graphical user interface or any software tool that allows the creation and modification of mask ROIs.

4.4.2 Hardware requirements and environmental costs

Cellpose-SAM and similar transformer-based architectures require considerably more computational resources and dedicated GPU support than earlier, shallower segmentation models, both during training and inference. Nonetheless, training a Cellpose-SAM model is feasible on a standard desktop equipped with a mid-range GPU, such as an NVIDIA GeForce RTX 3070 with 8 GB of RAM. Under this configuration, training with our dataset required more than 40 hours, compared to ~~under~~ ≤ 1.5 hours on an NVIDIA A100 GPU with 80 GB of memory (see Section 2.3 for details). For inference, dedicated GPUs (e.g., NVIDIA or Apple M2 and newer chips) with at least 3 GB of RAM are required to segment large images ($> 1000 \times 1000$ pixels) within a few seconds. Detailed performance benchmarks are provided in Pachitariu et al. (2025; Tables S2, S3). Measurements of grain size and shape in ImageGrains operate at comparable speeds, enabling the automated analysis of thousands of grains within minutes.

The recent increases in the depths of models and the sizes of training data for deep learning led to a large increase of computational resources that consume substantial amounts of energy, which exerts a high toll on the environment (e.g., Strubell et al., 2019; van Wynsberghe, 2021). It is difficult to precisely quantify the environmental costs of such models due to the question of how to factor in the initial training of SAM (which is reported to have cost c. 6963 kWh or equivalent to 2.8 metric tons of CO₂ released; Kirilov et al. 2023), and the open question of how to account for secondary effects beyond the computational cost and energy consumption (e.g., Yu et al., 2024; Bouza et al., 2024). Furthermore, quantifying the environmental impact of energy consumption for training or inference requires dedicated tests and often-unavailable information, such as where a specific piece of hardware was used and what energy source powered it (e.g., Lacoste et al., 2019; Patterson et al., 2021; Jay et al., 2023). However, we can estimate the computational cost (and therefore the energy cost) of our approach, based on the results of Pachitariu et al. (2025; cf. Tables S2 and S3 therein). We note that the Cellpose-SAM architecture runs on local desktop computers with mid-range consumer hardware configurations (see above), both for inference and for retraining smaller models (not more than a few hundred images). Therefore, despite the rather costly initial training

and release of SAM (Kirilov et al., 2023), using or fine-tuning our model in ImageGrains has a comparatively small computational cost and thus a low environmental impact, comparable to running any other software for similar durations on a desktop PC.

Finally, the ImageGrains library is distributed as an installable Python package, allowing both local and cloud-based deployment across platforms (see code availability section below). This enables an efficient execution even on free online platforms, such as Google Colab, for users without access to suitable local hardware.

4.4.3 Applicability

The introduction of a fine-tuned Cellpose-SAM model as the default segmentation engine in the ImageGrains 2.0 library significantly enhances the performance of segmenting grains relative to previous versions (Mair et al., 2024). This improvement directly benefits not only the direct applications of this workflow (e.g., Patel et al., 2025; Rezwan et al., 2025; Zegers et al., 2025) but also a broad range of applications that rely on the segmentation of grains in 2D, such as conducted in studies of coarse-grained fluvial sediments (e.g., especially when high accuracy is required across large areas (e.g., Guerit et al., 2018; Purinton and Bookhagen, 2021; Marchetti et al., 2022). This is particularly the case when grain size accuracy had previously been a challenge (e.g., Chardon et al., 2022; Miazza et al., 2024), or when the spatial variability in grain size is high (e.g., Rice and Church, 1998; Guerit et al., 2014) (Patel et al., 2025; Rezwan et al., 2025; Zegers et al., 2025). The expanded IG2 dataset enables the application of our default model to additional sedimentary contexts, such as XR-CT imagery of glaciofluvial-glacio-fluvial clasts (Schuster et al., 2025), bioclastic marine sands (e.g., Fabbri et al., 2024) and proglacial angular sediments (e.g., Hiller et al., 2024). For CT-based image stacks, the segmentation can be extended to fully capture the 3D nature of sediment grains. The high precision of the predicted grain masks allows for a robust analysis of the shape of grains, with the metrics used in this study provided as default outputs in ImageGrains. Furthermore, the availability of individual 2D grain masks as output enables the analysis of shapes tailored and customized to specific research requirements. Moreover, the high segmentation performance and generalization capability enable robust and precise segmentation of sediment grains in a broad range of image datasets and applications. These images of sediment other than coarse, clastic material. Our workflow should be adaptable, with minimal annotation effort, for similar segmentation tasks of other types of sediment (e.g., Dawson and John, 2024; Mitra et al., 2024; Liu et al., 2026), or even for segmenting other image-based objects that are of geoscientific relevance (e.g., Hsiang et al., 2019; Kloster et al., 2023). In general, the approach can be generalized to any category of segmentable objects in geoscientific imagery, where dense segmentation is needed to create ROI-based outputs for object morphometry. Such outputs could then be used as inputs for other machine learning tasks, e.g., for classification tasks. Furthermore, our publicly available dataset can be used in combination with our own labels for fine-tuning to other image types and settings. Finally, the manually annotated masks for individual grains can be used to train and test other segmentation approaches.

4.4.4 Future directions

This study demonstrates that state-of-the-art deep learning models originally developed for biomedical image segmentation can be successfully adapted for segmenting sediment grains. Similar to the segmentation of biomedical images, domain-specific architectures optimized for grain imagery outperform generic computer vision approaches, despite sharing foundational components such as SAM encoders and pre-trained weights (see Section 4.2.1). This suggests that the same underlying principles apply for these tasks.

However, a notable difference between the two application domains lies in absolute performance of the segmentation. The median AP@0.5 for our full IG2 dataset (0.72) is lower than that reported for Cellpose-SAM on the Cellpose biomedical dataset (>0.85 ; Pachitariu et al., 2025). We attribute this difference primarily to the smaller size of the IG2 dataset (243 image tiles compared to over 1000 microscopy images in the Cellpose nuclei dataset), the greater variability of the imagery regarding the size distribution and texture of the ~~grain~~grains displayed in this imagery, and the conditions at which they were taken in the field. This interpretation is supported by the variable model performance across IG2 subsets, which correlates with grain size accuracy—~~and~~ and to a lesser extent—~~with~~ grain shape accuracy (Fig. 8; Table ~~S4S3~~). Thus, future improvements in measuring the size and shape of grains will depend strongly on enhancing the performance of segmenting grains.

Progress in this direction is likely to come from the creation of larger and more variable annotated datasets and from the adoption of standardized image acquisition protocols that reduce the variability in image content. Given the effort of labelling/annotating large amounts of images and the fact that only a handful of annotated images are required for fine tuning the model, we expect that these larger and more variable datasets will be compiled from a collection of smaller datasets from different studies, if published by the community (e.g., Saddi et al., 2026). Moreover, obtaining annotations from multiple experts for the same images would enable calculation of values reflecting an inter-annotator consensus-level, a common benchmark for assessing absolute segmentation quality in other fields (e.g., Braylan et al., 2022; Yang et al., 2023; Zhou et al., 2025). Current SAM-based segmentation models appear capable of achieving such inter-annotator consensus-level accuracy for 2D images when trained on suitable datasets (Pachitariu et al., 2025; Na et al., 2025; Israel et al., 2025). This suggests that future advances in grain size and shape measurement will be driven less by architectural refinements but more by the development of larger, ~~high-quality~~ and representative training datasets of high quality.

5 Conclusions

We present a collection of 243 manually annotated images of sediment, the IG2 dataset, designed to enable systematic assessment and training of deep-learning architectures for segmenting sediment grains. Furthermore, we introduce ImageGrains 2.0, an updated open-source framework for automated measurement of grain size and shape that integrates ~~a fine-tuned~~ Cellpose-SAM as a fine-tuned default segmentation model for this task. Across the IG2 dataset, the model improves segmentation performance by up to 20% in AP@0.5 and mAP compared to ~~previous~~selected benchmark workflows, yielding grain masks that most closely match the ground-truth regions of interest for both size and shape metrics. The model's strong

generalization capability enables accurate segmentation across multiple image types and settings, including XR-CT imagery, and across variable grain textures and imaging conditions. The segmentation performance correlates directly with the accuracy of the derived grain size and shape metrics, confirming that improved segmentation performance translates to more robust geomorphic measurements. Additionally, the model can be fine-tuned with only a few additional image tiles to adapt to new sediment types or imaging conditions, making it applicable for a broad range of applications.

Code availability

All code is available as open-source code in the ImageGrains library (<https://github.com/dmair1989/imagegrains>), which is also installable as Python package (<https://pypi.org/project/imagegrains>). A graphical user interface and Jupyter notebooks are provided, enabling the use of ImageGrains 2.0 without the need to write custom code.

Data availability

The image and annotations of the IG2 dataset are available in the dedicated Zenodo repository (Mair et al., 2025a; <https://doi.org/10.5281/zenodo.17866827>). The model weights for the fine-tuned Cellpose-SAM default segmentation model, and the other Cellpose-2-based models, are available in a separate Zenodo repository (Mair et al., 2025b; <https://doi.org/10.5281/zenodo.15309323>).

Author contribution

DM conceptualized the research and developed the code together with GW. The data were curated by DM, with image annotations performed by DM, AdP, AW, BS, and FV, and images contributed by AdP, AW, PG, FV, BS, JÖ, SF, CL, SA, SL, CH, and FS. DM interpreted the results with scientific inputs from GW, MH, and FS. DM prepared the manuscript and figures with contributions from all authors.

Competing interests

The authors declare that they have no conflict of interest.

Disclaimer

Copernicus Publications remains neutral with regard to jurisdictional claims made in the text, published maps, institutional affiliations, or any other geographical representation in this paper. While Copernicus Publications makes every effort to include

appropriate place names, the final responsibility lies with the authors. Views expressed in the text are those of the authors and do not necessarily reflect the views of the publisher.

Acknowledgements

The images for NB2 are part of a larger dataset, not yet published in full, whose collection was funded by EU Horizon 2020 grant No. 860383. JÖ acknowledges funding from the MBIE Science Whitinga Fellowship (21-VUW-029), which supported fieldwork in Canterbury, New Zealand. SA, SL and CH acknowledge the funding by the Austrian Academy of Sciences, Earth System Sciences research initiative (ESS), Hidden.ice project in which data for subset JF (JamtalFerner) was obtained. Subsets PP and HP data subset were acquired by FV with a grant funded by the Ministry of Economy, Industry and Competitiveness, Spain (BES-2017-081850) and under financial support from the MorphHab (PID2019104979RBI00/AEI/10.13039/501100011033) and MorphPeak (CGL201678874-R/AEI/10.13039/501100011033) research projects funded by the Spanish State Research Agency (Ministry of Science and Innovation) and the European Regional Development Fund Scheme (FEDER). We acknowledge access to high-performance GPUs for model training through the UBELIX HPC cluster maintained by the University of Bern.

References

- Allen, P. A., Michael, N. A., D'Arcy, M., Roda-Boluda, D. C., Whittaker, A. C., Duller, R. A., and Armitage, J. J.: Fractionation of grain size in terrestrial sediment routing systems, *Basin Research*, 29, 180–202, <https://doi.org/10.1111/bre.12172>, 2017.
- Archit, A., Freckmann, L., Nair, S., Khalid, N., Hilt, P., Rajashekar, V., Freitag, M., Teuber, C., Buckley, G., von Haaren, S., Gupta, S., Dengel, A., Ahmed, S., and Pape, C.: Segment Anything for Microscopy, *Nat. Methods*, 22, 579–591, <https://doi.org/10.1038/s41592-024-02580-4>, 2025.
- Arzt, M., Deschamps, J., Schmied, C., Pietzsch, T., Schmidt, D., Tomancak, P., Haase, R., and Jug, F.: LABKIT: Labeling and Segmentation Toolkit for Big Image Data, *Front Comput Sci*, 4, <https://doi.org/10.3389/fcomp.2022.777728>, 2022.
- Attal, M. and Lavé, J.: Pebble abrasion during fluvial transport: Experimental results and implications for the evolution of the sediment load along rivers, *J. Geophys. Res. Earth Surf.*, 114, <https://doi.org/10.1029/2009JF001328>, 2009.
- Azad, R., Aghdam, E. K., Rauland, A., Jia, Y., Avval, A. H., Bozorgpour, A., Karimijafarbigloo, S., Cohen, J. P., Adeli, E., and Merhof, D.: Medical Image Segmentation Review: The Success of U-Net, *IEEE Trans. Pattern Anal. Mach. Intell.*, 46, 10076–10095, <https://doi.org/10.1109/TPAMI.2024.3435571>, 2024.
- Azzam, F., Blaise, T., and Brigaud, B.: Automated petrographic image analysis by supervised and unsupervised machine learning methods, *Sedimentologia*, 2, <https://doi.org/10.57035/journals/sdk.2024.e22.1594>, 2024.
- Back, A. L., Kana Tepakbong, C., Bédard, L. P., and Barry, A.: From rocks to pixels: a comprehensive framework for grain shape characterization through image analysis of roundness and roughness descriptors, *Front. Earth Sci. (Lausanne)*, 13, <https://doi.org/10.3389/feart.2025.1634237>, 2025a.
- Back, A. L., Kana Tepakbong, C., Bédard, L. P., and Barry, A.: From rocks to pixels: a comprehensive framework for grain shape characterization through the image analysis of size, orientation, and form descriptors, *Front. Earth Sci. (Lausanne)*, 13, <https://doi.org/10.3389/feart.2025.1508690>, 20252025b.
- Back, A. L., Kana Tepakbong, C., Bédard, L. P., and Barry, A.: From rocks to pixels: a comprehensive framework for grain shape characterization through the image analysis of size, orientation, and form descriptors, *Front Earth Sci (Lausanne)*, 13, <https://doi.org/10.3389/feart.2025.1508690>, 20252025a.

- Baptista, P., Cunha, T. R., Gama, C., and Bernardes, C.: A new and practical method to obtain grain size measurements in sandy shores based on digital image acquisition and processing, *Sediment. Geol.*, 282, 294–306, <https://doi.org/10.1016/j.sedgeo.2012.10.005>, 2012.
- 860 Baynes, E. R. C., Lague, D., Steer, P., Bonnet, S., and Illien, L.: Sediment flux-driven channel geometry adjustment of bedrock and mixed gravel–bedrock rivers, *Earth Surf. Process. Landf.*, 45, 3714–3731, <https://doi.org/10.1002/esp.4996>, 2020.
- Benet, D., Costa, F., Widiwijayanti, C., Pallister, J., Pedreros, G., Allard, P., Humaida, H., Aoki, Y., and Maeno, F.: VolcAshDB: a Volcanic Ash DataBase of classified particle images and features, *Bull Volcanol.*, 86, <https://doi.org/10.1007/s00445-023-01695-4>, 2024.
- 865 Bertin, S. and Friedrich, H.: Field application of close-range digital photogrammetry (CRDP) for grain-scale fluvial morphology studies, *Earth Surf. Process. Landf.*, 41, 1358–1369, <https://doi.org/10.1002/esp.3906>, 2016.
- Biguenet, M., Sabatier, P., Chaumillon, E., Chagué, C., Arnaud, F., Jorissen, F., Coulombier, T., Geba, E., Cordrie, L., Vacher, P., Develle, A. L., Chalmin, E., Soufi, F., and Feuillet, N.: A 1600 year-long sedimentary record of tsunamis and hurricanes in the Lesser Antilles (Scrub Island, Anguilla), *Sediment. Geol.*, 412, <https://doi.org/10.1016/j.sedgeo.2020.105806>, 2021.
- 870 Braylan, A., Alonso, O., and Lease, M.: Measuring Annotator Agreement Generally across Complex Structured, Multi-object, and Free-text Annotation Tasks, in: WWW 2022 - Proceedings of the ACM Web Conference 2022, 1720–1730, <https://doi.org/10.1145/3485447.3512242>, 2022.
- 875 Brayshaw, D. D.: Bankfull and effective discharge in small mountain streams of British Columbia, <https://doi.org/10.14288/1.0072555>, 2012.
- Brodu, N. and Lague, D.: 3D terrestrial lidar data classification of complex natural scenes using a multi-scale dimensionality criterion: Applications in geomorphology, *ISPRS Journal of Photogrammetry and Remote Sensing*, 68, 121–134, <https://doi.org/10.1016/j.isprsjprs.2012.01.006>, 2012.
- 880 Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., Hesse, C., Chen, M., Sigler, E., Litwin, M., Gray, S., Chess, B., Clark, J., Berner, C., McCandlish, S., Radford, A., Sutskever, I., and Amodei, D.: Language Models are Few-Shot Learners, 2020.
- Bunte, K. and Abt, S. R.: Sampling Surface and Subsurface Particle-Size Distributions in Wadable Gravel-and Cobble-Bed Streams for Analyses in Sediment Transport, Hydraulics, and Streambed Monitoring, 2001.
- 885 Buscombe, D.: *SediNet: a configurable deep learning model for mixed qualitative and quantitative optical granulometry*, *Earth Surf Process Landf.*, 45, 638–651, <https://doi.org/10.1002/esp.4760>, 2020.
- Buscombe, D.*: Transferable wavelet method for grain-size distribution from images of sediment surfaces and thin sections, and other natural granular patterns, *Sedimentology*, 60, 1709–1732, <https://doi.org/10.1111/sed.12049>, 2013.
- 890 *Buscombe, D.*: *SediNet: a configurable deep learning model for mixed qualitative and quantitative optical granulometry*, *Earth Surf Process Landf.*, 45, 638–651, <https://doi.org/10.1002/esp.4760>, 2020.
- Butler, J. B., Lane, S. N., and Chandler, J. H.: Automated extraction of grain-size data from gravel surfaces using digital image processing, *Journal of Hydraulic Research*, 39, 519–529, <https://doi.org/10.1080/00221686.2001.9628276>, 2001.
- 895 Carbonneau, P. E., Lane, S. N., and Bergeron, N. E.: Catchment-scale mapping of surface grain size in gravel bed rivers using airborne digital imagery, *Water Resour Res.*, 40, <https://doi.org/10.1029/2003WR002759>, 2004.
- Chan, V., Rheinwalt, A., and Bookhagen, B.: OrthoSAM: Multi-Scale Extension of the Segment Anything Model for River Pebble Delineation from Large Orthophotos, <https://doi.org/10.5194/egusphere-2025-4003>, 29 August 2025.
- Chardon, V., Piasny, G., and Schmitt, L.: Comparison of software accuracy to estimate the bed grain size distribution from digital images: A test performed along the Rhine River, *River Res. Appl.*, 38, 358–367, <https://doi.org/10.1002/rra.3910>, 2022.
- 900 *Chardon, V., Piasny, G., and Schmitt, L.*: *Comparison of software accuracy to estimate the bed grain size distribution from digital images: A test performed along the Rhine River, River Res Appl*, 38, 358–367, <https://doi.org/10.1002/rra.3910>, 2022.
- 905 Chen, X., Hassan, M. A., and Fu, X.: Convolutional neural networks for image-based sediment detection applied to a large terrestrial and airborne dataset, *Earth Surface Dynamics*, 10, 349–366, <https://doi.org/10.5194/esurf-10-349-2022>, 2022.

- Chen, Y., Bao, J., Chen, R., Li, B., Yang, Y., Renteria, L., Delgado, D., Forbes, B., Goldman, A. E., Simhan, M., Barnes, M. E., Laan, M., McKeever, S., Hou, Z. J., Chen, X., Scheibe, T., and Stegen, J.: Quantifying Streambed Grain Size, Uncertainty, and Hydrobiogeochemical Parameters Using Machine Learning Model YOLO, *Water Resour Res*, 60, <https://doi.org/10.1029/2023WR036456>, 2024.
- 910 Chen, Z., Scott, C., Keating, D., Clarke, A., Das, J., and Arrowsmith, R.: Quantifying and analysing rock trait distributions of rocky fault scarps using deep learning, *Earth Surf Process Landf*, 48, 1234–1250, <https://doi.org/10.1002/esp.5545>, 2023.
- Daniels, M. D. and McCusker, M. H.: Operator bias characterizing stream substrates using Wolman pebble counts with a standard measurement template, *Geomorphology*, 115, 194–198, <https://doi.org/10.1016/j.geomorph.2009.09.038>, 2010.
- 915 Deal, E., Venditti, J. G., Benavides, S. J., Bradley, R., Zhang, Q., Kamrin, K., & Perron, J. T.: Grain shape effects in bed load sediment transport. *Nature*, 613(7943), 298–302, <https://doi.org/10.1038/s41586-022-05564-6>, 2023.
- Detert, M. and Weitbrecht, V.: Automatic object detection to analyze the geometry of gravel grains – a free stand-alone tool, in: *River Flow 2012*, 595–600, 2012.
- 920 DiBiase, R. A., Lamb, M. P., Ganti, V., and Booth, A. M.: Slope, grain size, and roughness controls on dry sediment transport and storage on steep hillslopes, *J Geophys Res Earth Surf*, 122, 941–960, <https://doi.org/10.1002/2016JF003970>, 2017.
- Dietrich, William E., et al. "Sediment supply and the development of the coarse surface layer in gravel-bedded rivers." *Nature* 340.6230: 215-217, 1989.
- 925 Domokos, G., Jerolmack, D. J., Sipos, A. Á., and Török, Á.: How river rocks round: Resolving the shape-size paradox, *PLoS One*, 9, <https://doi.org/10.1371/journal.pone.0088657>, 2014.
- Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., Dehghani, M., Minderer, M., Heigold, G., Gelly, S., Uszkoreit, J., and Houshy, N.: An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale, 2021.
- 930 Dunne, K. B. J. and Jerolmack, D. J.: Evidence of, and a proposed explanation for, bimodal transport states in alluvial rivers, *Earth Surface Dynamics*, 6, 583–594, <https://doi.org/10.5194/esurf-6-583-2018>, 2018.
- Eaton, B. C., Dan Moore, R., and Mackenzie, L. G.: Percentile-based grain size distribution analysis tools (GSDtools)- Estimating confidence limits and hypothesis tests for comparing two samples, *Earth Surface Dynamics*, 7, 789–806, <https://doi.org/10.5194/esurf-7-789-2019>, 2019.
- 935 Eltner, A., Kaiser, A., Castillo, C., Rock, G., Neugirg, F., and Abellán, A.: Image-based surface reconstruction in geomorphometry-merits, limits and developments, *Earth Surface Dynamics*, 4, 359–389, <https://doi.org/10.5194/esurf-4-359-2016>, 2016.
- Everingham, M., Eslami, S. M. A., Van Gool, L., Williams, C. K. I., Winn, J., and Zisserman, A.: The Pascal Visual Object Classes Challenge: A Retrospective, *Int. J. Comput. Vis.*, 111, 98–136, <https://doi.org/10.1007/s11263-014-0733-5>, 2015.
- 940 ~~von Eynatten, H., Tolosana Delgado, R., and Karius, V.: Sediment generation in modern glacial settings: Grain size and source rock control on sediment composition, *Sediment Geol*, 280, 80–92, <https://doi.org/10.1016/j.sedgeo.2012.03.008>, 2012.~~
- Fabbri, S. C., Sabatier, P., Paris, R., Falvard, S., Feuillet, N., Lothoz, A., St-Onge, G., Gailler, A., Cordrie, L., Arnaud, F., Biguenet, M., Coulombier, T., Mitra, S., and Chaumillon, E.: Deciphering the sedimentary imprint of tsunamis and storms in the Lesser Antilles (Saint Martin): A 3500-year record in a coastal lagoon, *Mar Geol*, 471, <https://doi.org/10.1016/j.margeo.2024.107284>, 2024.
- 945 Fort, S., Ren, J., and Lakshminarayanan, B.: Exploring the Limits of Out-of-Distribution Detection, <https://doi.org/10.48550/arXiv.2106.03004>, 2021.
- Garefalakis, P., do Prado, A. H., Mair, D., Douillet, G. A., Nyffenegger, F., and Schlunegger, F.: Comparison of three grain size measuring methods applied to coarse-grained gravel deposits, *Sediment Geol*, 446, <https://doi.org/10.1016/j.sedgeo.2023.106340>, 2023.
- 950 Garefalakis, P., do Prado, A. H., Whittaker, A. C., Mair, D., and Schlunegger, F.: Quantification of sediment fluxes and intermittencies from Oligo–Miocene megafan deposits in the Swiss Molasse basin, *Basin Research*, 36, <https://doi.org/10.1111/bre.12865>, 2024.

- 955 Gatys, L. A., Ecker, A. S., and Bethge, M.: Image Style Transfer Using Convolutional Neural Networks, in: Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition, 2414–2423, <https://doi.org/10.1109/CVPR.2016.265>, 2016.
- Gordon Wolman, M.: A METHOD OF SAMPLING COARSE RIVER-BED MATERIAL, American Geophysical Union, 1954.
- 960 Guerit, L., Barrier, L., Liu, Y., Narteau, C., Lajeunesse, E., Gayer, E., and Métivier, F.: Uniform grain-size distribution in the active layer of a shallow, gravel-bedded, braided river (the Urumqi River, China) and implications for paleo-hydrology, Earth Surface Dynamics, 6, 1011–1021, <https://doi.org/10.5194/esurf-6-1011-2018>, 2018.
- Guerit, L., Barrier, L., Narteau, C., Métivier, F., Liu, Y., Lajeunesse, E., Gayer, E., Meunier, P., Malverti, L., and Ye, B.: The grain-size patchiness of braided gravel-bed streams - Example of the Urumqi River (northeast Tian Shan, China), Advances in Geosciences, 37, 27–39, <https://doi.org/10.5194/adgeo-37-27-2014>, 2014.
- 965 He, K., Gkioxari, G., Dollár, P., and Girshick, R.: Mask R-CNN, 2017.
- Hendrycks, D. and Dietterich, T.: Benchmarking Neural Network Robustness to Common Corruptions and Perturbations, 2019.
- Hiller, C., Leistner, S., Helfricht, K., and Achleitner, S.: Robust estimations of areal grain size distribution from geometric surface roughness in a proglacial outwash area, Geomorphology, 439, <https://doi.org/10.1016/j.geomorph.2023.108857>, 2023.
- 970 Houston, A. N., Schmidt, S., Tarquis, A. M., Otten, W., Baveye, P. C., and Hapca, S. M.: Effect of scanning and image reconstruction settings in X-ray computed microtomography on quality and segmentation of 3D soil images, Geoderma, 207–208, 154–165, <https://doi.org/10.1016/j.geoderma.2013.05.017>, 2013.
- Ibbeken, H. and Schleyer, R.: Photo-sieving: A method for grain-size analysis of coarse-grained, unconsolidated bedding surfaces, Earth Surf Process Landf, 11, 59–77, <https://doi.org/10.1002/esp.3290110108>, 1986.
- 975 Israel, U., Marks, M., Dilip, R., Li, Q., Yu, C., Laubscher, E., Iqbal, A., Pradhan, E., Ates, A., Abt, M., Brown, C., Pao, E., Li, S., Pearson-Goulart, A., Perona, P., Gkioxari, G., Barnowski, R., Yue, Y., and Van Valen, D.: CellSAM: A Foundation Model for Cell Segmentation, <https://doi.org/10.1101/2023.11.17.567630>, 20 November 2023.
- Jay, M., Ostapenko, V., Lefevre, L., Trystram, D., Orgerie, A.-C., and Fichel, B.: An experimental comparison of software-based power meters: focus on CPU and GPU, in: 2023 IEEE/ACM 23rd International Symposium on Cluster, Cloud and Internet Computing (CCGrid), 106–118, <https://doi.org/10.1109/CCGrid57682.2023.00020>, 2023.
- 980 Kettler, C., Phillips, E., Pichler, K., Smrzka, D., Vandyk, T. M., and Le Heron, D. P.: 3D macro- and microfabric analyses of Neoproterozoic diamictites from the Valjean Hills, California (United States), Front Earth Sci (Lausanne), 11, <https://doi.org/10.3389/feart.2023.929011>, 2023.
- 985 Kirillov, A., Mintun, E., Ravi, N., Mao, H., Rolland, C., Gustafson, L., Xiao, T., Whitehead, S., Berg, A. C., Lo, W.-Y., Dollár, P., and Girshick, R.: Segment Anything, 2023.
- Kumar, N., Verma, R., Anand, D., Zhou, Y., Onder, O. F., Tsougenis, E., Chen, H., Heng, P. A., Li, J., Hu, Z., Wang, Y., Koohbanani, N. A., Jahanifar, M., Tajeddin, N. Z., Gooya, A., Rajpoot, N., Ren, X., Zhou, S., Wang, Q., Shen, D., Yang, C. K., Weng, C. H., Yu, W. H., Yeh, C. Y., Yang, S., Xu, S., Yeung, P. H., Sun, P., Mahbod, A., Schaefer, G., Ellinger, I., Ecker, R., Smedby, O., Wang, C., Chidester, B., Ton, T. V., Tran, M. T., Ma, J., Do, M. N., Graham, S., Vu, Q. D., Kwak, J. T., Gunda, A., Chunduri, R., Hu, C., Zhou, X., Lotfi, D., Safdari, R., Kascenas, A., O’Neil, A., Eschweiler, D., Stegmaier, J., Cui, Y., Yin, B., Chen, K., Tian, X., Gruening, P., Barth, E., Arbel, E., Remer, I., Ben-Dor, A., Sirazitdinova, E., Kohl, M., Braunewell, S., Li, Y., Xie, X., Shen, L., Ma, J., Baksi, K., Das, Khan, M. A., Choo, J., Colomer, A., Naranjo, V., Pei, L., Iftekharuddin, K. M., Roy, K., Bhattacharjee, D., Pedraza, A., Bueno, M. G., Devanathan, S., Radhakrishnan, S., Koduganty, P., Wu, Z., Cai, G., Liu, X., Wang, Y., and Sethi, A.: A Multi-Organ Nucleus Segmentation Challenge, IEEE Trans. Med. Imaging, 39, 1380–1391, <https://doi.org/10.1109/TMI.2019.2947628>, 2020.
- 995 Lang, N., Irniger, A., Rozniak, A., Hunziker, R., Dirck Wegner, J., and Schindler, K.: GRAINet: Mapping grain size distributions in river beds from UAV images with convolutional neural networks, Hydrol Earth Syst Sci, 25, 2567–2597, <https://doi.org/10.5194/hess-25-2567-2021>, 2021.
- 1000 Lepp, A. P., Miller, L. E., Anderson, J. B., O’regan, M., Winsborrow, M. C. M., Smith, J. A., Hillenbrand, C. D., Wellner, J. S., Prothro, L. O., and Podolskiy, E. A.: Insights into glacial processes from micromorphology of silt-sized sediment, Cryosphere, 18, 2297–2319, <https://doi.org/10.5194/tc-18-2297-2024>, 2024.

- Li, Y., Mao, H., Girshick, R., and He, K.: Exploring Plain Vision Transformer Backbones for Object Detection, 2022.
- 1005 Litty, C. and Schlunegger, F.: Controls on pebbles' size and shape in streams of the Swiss alps, *Journal of Geology*, 125, 101–112, <https://doi.org/10.1086/689183>, 2017.
- Li, Y., Mao, H., Girshick, R., and He, K.: Exploring Plain Vision Transformer Backbones for Object Detection, 2022.
- Loshchilov, I. and Hutter, F.: Decoupled Weight Decay Regularization, 2019.
- ~~Mair, D.: Dataset and model weights for ImageGrains (Version v1), Zenodo, <https://doi.org/10.5281/zenodo.8005771>, 2023.~~
- 1010 Mair, D., Do Prado, A. H., Garefalakis, P., Lechmann, A., Whittaker, A., and Schlunegger, F.: Grain size of fluvial gravel bars from close-range UAV imagery – uncertainty in segmentation-based data, <https://doi.org/10.5194/esurf-2022-19>, 23 May 2022.
- Mair, D., do Prado, A., Garefalakis, P., Wild, A. L., Ville, F., Schuster, B., Österle, J., Fabbri, S. C., Litty, C., Achleitner, S., Leistner, S., Hiller, C., & Schlunegger, F.: ImageGrains 2.0 dataset (Version v2), Zenodo, <https://doi.org/10.5281/zenodo.17866827>, 2025a.
- 1015 ~~Mair, D., Witz, G., do Prado, A., Garefalakis, P., Wild, A. L., Ville, F., Schuster, B., Horn, M., Österle, J., Fabbri, S. C., Litty, C., Achleitner, S., Leistner, S., Hiller, C., & Schlunegger, F.: ImageGrains 2.0 models, Zenodo, <https://doi.org/10.5281/zenodo.15309323>, 2025b.~~
- Mair, D., Witz, G., Do Prado, A. H., Garefalakis, P., and Schlunegger, F.: Automated detecting, segmenting and measuring of grains in images of fluvial sediments: The potential for large and precise data from specialist deep learning models and transfer learning, *Earth Surf Process Landf*, 49, 1099–1116, <https://doi.org/10.1002/esp.5755>, 2024.
- 1020 Mair, D.: Dataset and model weights for ImageGrains (Version v1), Zenodo, <https://doi.org/10.5281/zenodo.8005771>, 2023.
- Marc, O., Turowski, J. M., and Meunier, P.: Controls on the grain size distribution of landslides in Taiwan: The influence of drop height, scar depth and bedrock strength, *Earth Surface Dynamics*, 9, 995–1011, <https://doi.org/10.5194/esurf-9-995-2021>, 2021.
- 1025 Marchetti, G., Bizzi, S., Belletti, B., Lastoria, B., Comiti, F., and Carbonneau, P. E.: Mapping riverbed sediment size from Sentinel-2 satellite data, *Earth Surf. Process. Landf.*, 47, 2544–2559, <https://doi.org/10.1002/esp.5394>, 2022.
- ~~do Prado, A., Garefalakis, P., Wild, A. L., Ville, F., Schuster, B., Österle, J., Fabbri, S. C., Litty, C., Achleitner, S., Leistner, S., Hiller, C., & Schlunegger, F.: ImageGrains 2.0 dataset (Version v2), Zenodo, <https://doi.org/10.5281/zenodo.17866827>, 2025a.~~
- 1030 ~~Mair, D., do Prado, A., Garefalakis, P., Wild, A. L., Ville, F., Schuster, B., Österle, J., Fabbri, S. C., Litty, C., Achleitner, S., Leistner, S., Hiller, C., & Schlunegger, F.: ImageGrains 2.0 dataset (Version v2), Zenodo, <https://doi.org/10.5281/zenodo.17866827>, 2025b.~~
- 1035 Marchetti, G., Bizzi, S., Belletti, B., Lastoria, B., Comiti, F., and Carbonneau, P. E.: Mapping riverbed sediment size from Sentinel-2 satellite data, *Earth Surf Process Landf*, 47, 2544–2559, <https://doi.org/10.1002/esp.5394>, 2022.
- Miazza, R., Pascal, I., and Ancey, C.: Automated grain sizing from uncrewed aerial vehicles imagery of a gravel-bed river: Benchmarking of three object-based methods, *Earth Surf Process Landf*, 49, 1503–1514, <https://doi.org/10.1002/esp.5782>, 2024.
- 1040 Miller, K. L., Szabó, T., Jerolmack, D. J., and Domokos, G.: Quantifying the significance of abrasion and selective transport for downstream fluvial grain size evolution, *J. Geophys. Res. Earth Surf.*, 119, 2412–2429, <https://doi.org/10.1002/2014JF003156>, 2014.
- ~~Miller, K. L., Szabó, T., Jerolmack, D. J., and Domokos, G.: Quantifying the significance of abrasion and selective transport for downstream fluvial grain size evolution, *J Geophys Res Earth Surf*, 119, 2412–2429, <https://doi.org/10.1002/2014JF003156>, 2014.~~
- 1045 Minae, S., Boykov, Y., Porikli, F., Plaza, A., Kehtarnavaz, N., and Terzopoulos, D.: Image Segmentation Using Deep Learning: A Survey, *IEEE Trans. Pattern Anal. Mach. Intell.*, 44, 3523–3542, <https://doi.org/10.1109/TPAMI.2021.3059968>, 2022.
- Mitra, S., Paris, R., Bernard, L., Abbal, R., Charrier, P., Falvard, S., Costa, P., and Andrade, C.: X-ray tomography applied to tsunami deposits: Optimized image processing and quantitative analysis of particle size, particle shape, and sedimentary fabric in 3D, *Mar. Geol.*, 470, <https://doi.org/10.1016/j.margeo.2024.107247>, 2024.
- 1050

- Mörtl, C., Baratier, A., Berthet, J., Duvillard, P. A., and De Cesare, G.: GALET: A deep learning image segmentation model for drone based grain size analysis of gravel bars, in: Proceedings of the IAHR World Congress, 5326–5335, <https://doi.org/10.3850/IAHR-39WC2521716X2022895>, 2022.
- 1055 Na, S., Guo, Y., Jiang, F., Ma, H., and Huang, J.: Segment Any Cell: A SAM-based Auto-prompting Fine-tuning Framework for Nuclei Segmentation, <https://doi.org/10.1109/TNNLS.2025.3611322>, 2024.
- Na, S., Guo, Y., Jiang, F., Ma, H., Gao, J., and Huang, J.: Segment Any Cell: A SAM-Based Auto-Prompting Fine-Tuning Framework for Nuclei Segmentation, IEEE Trans Neural Netw Learn Syst, 1–10, <https://doi.org/10.1109/TNNLS.2025.3611322>, 2025.
- 1060 Novák-Szabó, T., Sipos, A. Á., Shaw, S., Bertoni, D., Pozzebon, A., Grottoli, E., Sarti, G., Ciavola, P., Domokos, G., and Jerolmack, D. J.: Universal characteristics of particle shape evolution by bed-load chipping, 2018.
- Pachitariu, M., Rariden, M., and Stringer, C.: Cellpose-SAM: superhuman generalization for cellular segmentation, <https://doi.org/10.1101/2025.04.28.651001>, 1 May 2025.
- Padilla, R., Netto, S. L., and da Silva, E. A. B.: A Survey on Performance Metrics for Object-Detection Algorithms, in: 2020 International Conference on Systems, Signals and Image Processing (IWSSIP), 237–242, <https://doi.org/10.1109/IWSSIP48289.2020.9145130>, 2020.
- 1065 Patel, N. K., Schlunegger, F., Mair, D., Pati, P., do Prado, A. H., Garefalakis, P., and Choudhury, R. K.: Source-to-sink patterns of grain size along the Yamuna River in the Indian Himalaya, Geomorphology, 486, 109909, <https://doi.org/10.1016/j.geomorph.2025.109909>, 2025.
- 1070 Pfeiffer, A. M., Finnegan, N. J., and Willenbring, J. K.: Sediment supply controls equilibrium channel geometry in gravel rivers, Proc. Natl. Acad. Sci. U. S. A., 114, 3346–3351, <https://doi.org/10.1073/pnas.1612907114>, 2017.
- Piégay, H., Arnaud, F., Belletti, B., Bertrand, M., Bizzi, S., Carbonneau, P., Dufour, S., Liébault, F., Ruiz-Villanueva, V., and Slater, L.: Remotely sensed rivers in the Anthropocene: state of the art and prospects, Earth Surf. Process. Landf., 45, 157–188, <https://doi.org/10.1002/esp.4787>, 2020.
- 1075 Pokhrel, P., Attal, M., Sinclair, H. D., Mudd, S. M., and Naylor, M.: Downstream rounding rate of pebbles in the Himalaya, Earth Surface Dynamics, 12, 515–536, <https://doi.org/10.5194/esurf-12-515-2024>, 2024.
- Preusser, F., Büschelberger, M., Kemna, H. A., Miocic, J., Mueller, D., and May, J. H.: Exploring possible links between Quaternary aggradation in the Upper Rhine Graben and the glaciation history of northern Switzerland, International Journal of Earth Sciences, 110, 1827–1846, <https://doi.org/10.1007/s00531-021-02043-7>, 2021.
- 1080 Prieur, N. C., Amaro, B., Gonzalez, E., Kerner, H., Medvedev, S., Rubanenko, L., Werner, S. C., Xiao, Z., Zastrozhnov, D., and Lapôtre, M. G. A.: Automatic Characterization of Boulders on Planetary Surfaces From High-Resolution Satellite Images, J Geophys Res Planets, 128, <https://doi.org/10.1029/2023JE008013>, 2023.
- Purinton, B. and Bookhagen, B.: Introducing PebbleCounts: A grain-sizing tool for photo surveys of dynamic gravel-bed rivers, <https://doi.org/10.5194/esurf-7-859-2019>, 17 September 2019.
- 1085 Quick, L., Sinclair, H. D., Attal, M., and Singh, V.: Conglomerate recycling in the Himalayan foreland basin: Implications for grain size and provenance, GSA Bulletin, 132, 1639–1656, <https://doi.org/10.1130/B35334.1>, 2020.
- Ravi, N., Gabeur, V., Hu, Y.-T., Hu, R., Ryali, C., Ma, T., Khedr, H., Rädle, R., Rolland, C., Gustafson, L., Mintun, E., Pan, J., Alwala, K. V., Carion, N., Wu, C.-Y., Girshick, R., Dollár, P., and Feichtenhofer, C.: SAM 2: Segment Anything in Images and Videos, 2024.
- 1090 Redmon, J., Divvala, S., Girshick, R., and Farhadi, A.: You Only Look Once: Unified, Real-Time Object Detection, 2016.
- Rezwan, N., Mair, D., Whittaker, A. C., Schlunegger, F., do Prado, A. H., and Setu, S. H.: Assessing Flood-Influenced Sediment Dynamics Using UAV Photogrammetry and Machine Learning: Insights from River Sense, Switzerland, <https://doi.org/10.5194/egusphere-2025-5145>, 28 October 2025.
- Rheinwalt, A., Purinton, B., and Bookhagen, B.: Curvature-based pebble segmentation for reconstructed surface meshes, Earth Surface Dynamics, 13, 923–940, <https://doi.org/10.5194/esurf-13-923-2025>, 2025.
- 1095 Rice, S. and Church, M.: Grain size along two gravel-bed rivers: Statistical variation, spatial pattern and sedimentary links, Earth Surf. Process. Landf., 23, 345–363, [https://doi.org/10.1002/\(SICI\)1096-9837\(199804\)23:4<345::AID-ESP850>3.0.CO;2-B](https://doi.org/10.1002/(SICI)1096-9837(199804)23:4<345::AID-ESP850>3.0.CO;2-B), 1998.
- Rickenmann, D. and Recking, A.: Evaluation of flow resistance in gravel-bed rivers through a large field data set, Water Resour. Res., 47, <https://doi.org/10.1029/2010WR009793>, 2011.
- 1100

- Ronneberger, O., Fischer, P., and Brox, T.: U-Net: Convolutional Networks for Biomedical Image Segmentation, <https://doi.org/10.48550/arXiv.1505.04597>, 2015.
- Saddi, K. C., van Emmerik, T. H. M., Miglino, D., Poggi, M., Isgro, F., Tasserone, P. F., Daniele, L., and Manfreda, S.: Exploring the Transferability of Image-Based Algorithms for River Plastic Detection: The Value of Small Mixed Data Sets, *Water Resour. Res.*, 62, <https://doi.org/10.1029/2025WR040605>, 2026.
- Schlunegger, F. and Garefalakis, P.: Clast imbrication in coarse-grained mountain streams and stratigraphic archives as indicator of deposition in upper flow regime conditions, *Earth Surface Dynamics*, 6, 743–761, <https://doi.org/10.5194/esurf-6-743-2018>, 2018.
- Schindelin, J., Arganda-Carreras, I., Frise, E., Kaynig, V., Longair, M., Pietzsch, T., Preibisch, S., Rueden, C., Saalfeld, S., Schmid, B., Tinevez, J. Y., White, D. J., Hartenstein, V., Eliceiri, K., Tomancak, P., and Cardona, A.: Fiji: An open-source platform for biological-image analysis, <https://doi.org/10.1038/nmeth.2019>, July 2012.
- Schlunegger, F. and Garefalakis, P.: Clast imbrication in coarse-grained mountain streams and stratigraphic archives as indicator of deposition in upper flow regime conditions, *Earth Surface Dynamics*, 6, 743–761, <https://doi.org/10.5194/esurf-6-743-2018>, 2018.
- Schlunegger, F., Delunel, R., and Garefalakis, P.: Short communication: Field data reveal that the transport probability of clasts in Peruvian and Swiss streams mainly depends on the sorting of the grains, *Earth Surface Dynamics*, 8, 717–728, <https://doi.org/10.5194/esurf-8-717-2020>, 2020.
- Schuster, B., Gegg, L., Schaller, S., Buechi, M. W., Tanner, D. C., Wielandt-Schuster, U., Anselmetti, F. S., and Preusser, F.: Shaped and filled by the Rhine Glacier: the overdeepened Tannwald Basin in southwestern Germany, *Scientific Drilling*, 33, 191–206, <https://doi.org/10.5194/sd-33-191-2024>, 2024.
- Schuster, B., Mair, D., Schmid, T. C., Gegg, L., Buechi, M. W., Schaller, S., Anselmetti, F. S., and Preusser, F.: Automated X-ray computed tomography-based analysis of clast fabric in drill-cores of glacial diamicts using deep learning models, *Boreas*, <https://doi.org/10.1111/bor.70023>, 2025.
- Siddique, N., Paheding, S., Elkin, C. P., and Devabhaktuni, V.: U-Net and Its Variants for Medical Image Segmentation: A Review of Theory and Applications, *IEEE Access*, 9, 82031–82057, <https://doi.org/10.1109/ACCESS.2021.3086020>, 2021.
- Sklar, L. S.: Grain Size in Landscapes, *Annu Rev Earth Planet Sci*, 52, 663–692, <https://doi.org/10.1146/annurev-earth-052623-075856>, 2024.
- Soloy, A., Turki, I., Fournier, M., Costa, S., Peuziat, B., and Lecoq, N.: A deep learning-based method for quantifying and mapping the grain size on pebble beaches, *Remote Sens (Basel)*, 12, 1–23, <https://doi.org/10.3390/rs12213659>, 2020.
- Spagnolo, M., Phillips, E., Piotrowski, J. A., Rea, B. R., Clark, C. D., Stokes, C. R., Carr, S. J., Ely, J. C., Ribolini, A., Wysota, W., and Szuman, I.: Ice stream motion facilitated by a shallow-deforming and accreting bed, *Nat. Commun.*, 7, <https://doi.org/10.1038/ncomms10723>, 2016.
- Steer, P., Guerit, L., Lague, D., Crave, A., and Gourdon, A.: Size, shape and orientation matter: Fast and semi-automatic measurement of grain geometries from 3D point clouds, *Earth Surface Dynamics*, 10, 1211–1232, <https://doi.org/10.5194/esurf-10-1211-2022>, 2022.
- Stringer, C. and Pachitariu, M.: Cellpose3: one-click image restoration for improved cellular segmentation, *Nat. Methods*, 22, 592–599, <https://doi.org/10.1038/s41592-025-02595-5>, 2025.
- Stringer, C., Wang, T., Michaelos, M., and Pachitariu, M.: Cellpose: a generalist algorithm for cellular segmentation, *Nat Methods*, 18, 100–106, <https://doi.org/10.1038/s41592-020-01018-x>, 2021.
- Strubell, E., Ganesh, A., and McCallum, A.: Energy and Policy Considerations for Deep Learning in NLP, 2019.
- Sulaiman, M. S., Sinnakaudan, S. K., Ng, S. F., and Strom, K.: Application of automated grain sizing technique (AGS) for bed load samples at Rasil River: A case study for supply limited channel, *Catena (Amst.)*, 121, 330–343, <https://doi.org/10.1016/j.catena.2014.05.013>, 2014.
- Sylvester, Z., Stockli, D. F., Howes, N., Roberts, K., Malkowski, M. A., Poros, Z., Martindale, R. C., and Bai, W.: Segmenteverygrain: A Python module for segmentation of grains in images, *J Open Source Softw*, 10, 7953, <https://doi.org/10.21105/joss.07953>, 2025.
- Szabó, T., Domokos, G., Grotzinger, J. P., and Jerolmack, D. J.: Reconstructing the transport history of pebbles on Mars, *Nat. Commun.*, 6, <https://doi.org/10.1038/ncomms9366>, 2015.

- 150 Szabó, T., Domokos, G., Grotzinger, J. P., and Jerolmack, D. J.: Reconstructing the transport history of pebbles on Mars, *Nat Commun*, 6, <https://doi.org/10.1038/ncomms9366>, 2015.
- Tofelde, S., Bernhardt, A., Guerit, L., and Romans, B. W.: Times Associated With Source-to-Sink Propagation of Environmental Signals During Landscape Transience, <https://doi.org/10.3389/feart.2021.628315>, 27 April 2021.
- 155 Tucker Bsc, M., London, O., and Boston Melbourne, E.: *mm Techniques in Sedimentology* Blackwell Scientific Publications, 1988.
- Van Der Walt, S., Schönberger, J. L., Nunez-Iglesias, J., Boulogne, F., Warner, J. D., Yager, N., Gouillart, E., and Yu, T.: *Scikit-image: Image processing in python*, *PeerJ*, 2014, <https://doi.org/10.7717/peerj.453>, 2014.
- van Wynsberghe, A.: Sustainable AI: AI for sustainability and the sustainability of AI, *AI and Ethics*, 1, 213–218, <https://doi.org/10.1007/s43681-021-00043-6>, 2021.
- 160 Vázquez-Tarrio, D., Borgniet, L., Liébault, F., and Recking, A.: Using UAS optical imagery and SfM photogrammetry to characterize the surface grain size of gravel bars in a braided river (Vénéon River, French Alps), *Geomorphology*, 285, 94–105, <https://doi.org/10.1016/j.geomorph.2017.01.039>, 2017.
- Vázquez-Tarrio, D., Borgniet, L., Liébault, F., and Recking, A.: Using UAS optical imagery and SfM photogrammetry to characterize the surface grain size of gravel bars in a braided river (Vénéon River, French Alps), *Geomorphology*, 285, 94–105, <https://doi.org/10.1016/j.geomorph.2017.01.039>, 2017.
- 165 Ville F, Vericat D, Batalla RJ, Rennie C.: PhotoMOB: Automated GIS method for estimation of fractional grain dynamics in gravel bed rivers. Part 2: Bed Stability and Fractional Mobility, *EarthArXiv* (preprint), <https://doi.org/10.31223/X5Q108.2023>.
- 170 Ville F.: *Morpho-sedimentary dynamics of mountain rivers affected by hydropeaks*, PhD thesis, University of Lleida, 2024.
- von Eynatten, H., Tolosana-Delgado, R., and Karius, V.: Sediment generation in modern glacial settings: Grain-size and source-rock control on sediment composition, *Sediment Geol*, 280, 80–92, <https://doi.org/10.1016/j.sedgeo.2012.03.008>, 2012.
- Walicka, A. and Pfeifer, N.: Automatic Segmentation of Individual Grains From a Terrestrial Laser Scanning Point Cloud of a Mountain River Bed, *IEEE J Sel Top Appl Earth Obs Remote Sens*, 15, 1389–1410, <https://doi.org/10.1109/JSTARS.2022.3141892>, 2022.
- 1175 Watkins, S. E., Whittaker, A. C., Bell, R. E., Brooke, S. A. S., Ganti, V., Gawthorpe, R. L., McNeill, L. C., and Nixon, C. W.: Straight from the source’s mouth: Controls on field-constrained sediment export across the entire active Corinth Rift, central Greece, *Basin Research*, 32, 1600–1625, <https://doi.org/10.1111/bre.12444>, 2020.
- 1180 ~~Van Der Walt, S., Schönberger, J. L., Nunez-Iglesias, J., Boulogne, F., Warner, J. D., Yager, N., Gouillart, E., and Yu, T.: Scikit-image: Image processing in python, PeerJ, 2014, https://doi.org/10.7717/peerj.453, 2014.~~
- Williams, R. M. E., Grotzinger, J. P., Dietrich, W. E., Gupta, S., Sumner, D. Y., Wiens, R. C., Mangold, N., Malin, M. C., Edgett, K. S., Maurice, S., Forni, O., Gasnault, O., Ollila, A., Newsom, H. E., Dromart, G., Palucis, M. C., Yingst, R. A., Anderson, R. B., Herkenhoff, K. E., Le Mouélic, S., Goetz, W., Madsen, M. B., Koefoed, A., Jensen, J. K., Bridges, J. C., Schwenzer, S. P., Lewis, K. W., Stack, K. M., Rubin, D., Kah, L. C., Bell, J. F., Farmer, J. D., Sullivan, R., Van
- 1185 Beek, T., Blaney, D. L., Pariser, O., Deen, R. G., Kemppinen, O., Bridges, N., Johnson, J. R., Minitti, M., Cremers, D., Edgar, L., Godber, A., Wadhwa, M., Wellington, D., McEwan, I., Newman, C., Richardson, M., Charpentier, A., Peret, L., King, P., Blank, J., Weigle, G., Schmidt, M., Li, S., Milliken, R., Robertson, K., Sun, V., Baker, M., Edwards, C., Ehlmann, B., Farley, K., Griffes, J., Miller, H., Newcombe, M., Pilorget, C., Rice, M., Siebach, K., Stolper, E., Brunet, C., Hipkin, V., Lèveillé, R., Marchand, G., Sánchez, P. S., Favot, L., Cody, G., Steele, A., Flückiger, L., Lees, D.,
- 1190 Nefian, A., Martin, M., Gailhanou, M., Westall, F., Israël, G., Agard, C., Baroukh, J., Donny, C., Gaboriaud, A., Guillemot, P., Lafaille, V., Lorigny, E., Paillet, A., Pérez, R., Saccoccio, M., Yana, C., Aparicio, C. A., Rodríguez, J. C., Blázquez, I. C., et al.: Martian fluvial conglomerates at gale crater, *Science* (1979), 340, 1068–1072, <https://doi.org/10.1126/science.1237317>, 2013.
- Woodget, A. S. and Austrums, R.: Subaerial gravel size measurement using topographic data derived from a UAV-SfM approach, *Earth Surf. Process. Landf.*, 42, 1434–1443, <https://doi.org/10.1002/esp.4139>, 2017.
- 1195 Woodget, A. S. and Austrums, R.: Subaerial gravel size measurement using topographic data derived from a UAV-SfM approach, *Earth Surf. Process. Landf.*, 42, 1434–1443, <https://doi.org/10.1002/esp.4139>, 2017.

Woodget, A. S., Fyffe, C., and Carbonneau, P. E.: From manned to unmanned aircraft: Adapting airborne particle size mapping methodologies to the characteristics of sUAS and SfM, *Earth Surf Process Landf*, 43, 857–870, <https://doi.org/10.1002/esp.4285>, 2018.

Yang, F., Zamzmi, G., Angara, S., Rajaraman, S., Aquilina, A., Xue, Z., Jaeger, S., Papagiannakis, E., and Antani, S. K.: Assessing Inter-Annotator Agreement for Medical Image Segmentation, *IEEE Access*, 11, 21300–21312, <https://doi.org/10.1109/ACCESS.2023.3249759>, 2023.

Yang, F., Zamzmi, G., Angara, S., Rajaraman, S., Aquilina, A., Xue, Z., Jaeger, S., Papagiannakis, E., and Antani, S. K.: Assessing Inter-Annotator Agreement for Medical Image Segmentation, *IEEE Access*, 11, 21300–21312, <https://doi.org/10.1109/ACCESS.2023.3249759>, 2023.

Zaidi, S. S. A., Ansari, M. S., Aslam, A., Kanwal, N., Asghar, M., and Lee, B.: A survey of modern deep learning based object detection models, <https://doi.org/10.1016/j.dsp.2022.103514>, 30 June 2022.

Zegers, G., Hayashi, M., and Garcés, A.: Distributed estimation of surface sediment size in paraglacial and periglacial environments using drone photogrammetry, *Earth Surf Process Landf*, 50, <https://doi.org/10.1002/esp.70093>, 2025.

Zhou, F. Y., Marin, Z., Yapp, C., Zou, Q., Nanes, B. A., Daetwyler, S., Jamieson, A. R., Islam, M. T., Jenkins, E., Gihana, G. M., Lin, J., Borges, H. M., Chang, B.-J., Weems, A., Morrison, S. J., Sorger, P. K., Fiolka, R., Dean, K. M., and Danuser, G.: Universal consensus 3D segmentation of cells from 2D segmented stacks, *Nat. Methods*, 22, 2386–2399, <https://doi.org/10.1038/s41592-025-02887-w>, 2025.

Zhou, F. Y., Marin, Z., Yapp, C., Zou, Q., Nanes, B. A., Daetwyler, S., Jamieson, A. R., Islam, M. T., Jenkins, E., Gihana, G. M., Lin, J., Borges, H. M., Chang, B.-J., Weems, A., Morrison, S. J., Sorger, P. K., Fiolka, R., Dean, K. M., and Danuser, G.: Universal consensus 3D segmentation of cells from 2D segmented stacks, *Nat Methods*, 22, 2386–2399, <https://doi.org/10.1038/s41592-025-02887-w>, 2025.