

Hierarchical Training Pipeline for Event-Based Robotic Perception Models for Autonomous Roof Bolting

Rik Banerjee

M3 Robotics Lab, Colorado School of Mines
Golden CO, USA

Akram Marseet

M3 Robotics Lab, Colorado School of Mines
Golden CO, USA

Andrew J. Petruska

M3 Robotics Lab, Colorado School of Mines
Golden CO, USA

ABSTRACT

Event cameras are used for their performance in high dynamic-range lighting conditions which are canonical to active mining environments. Direct labeling of event-based image data to train a model to perform semantic segmentation using traditional methods is slow and error-prone. This study proposes a framework to use roughly hand-labeled color images from a mine as an input to an intermediary probabilistic algorithm called alphasammatting to generate a ground-truth data set. These high-fidelity labels can be used to train a semantic segmentation model to differentiate the support strap from the roof. This model can then be leveraged to segment an event-based scene to enable autonomous roof bolting. This pipeline has been shown to achieve an accuracy of 88% with a false positive rate of 3%.

INTRODUCTION

Roof bolting is commonly regarded as one of the riskiest occupations in the United States [1]–[4]. This arises from several factors: the operator faces the danger of sustaining injuries from the bolting machine [5], and there exists a risk of becoming a victim of a roof collapse [6]. Furthermore, a significant portion of accidents occurs when operators with less than one year of drilling experience are engaged in bolting. Enhancing miner safety and productivity could be achieved by relieving operators of the need to be

underground to identify roof drilling areas, position, and operate the drill. The low-light conditions in underground mines have been studied to have a measurable adverse effect on the health of the miners [7]. A robotic solution can be used to automate this process.

As can be seen in Figure 1, automating this process faces a two-fold challenge. Firstly, the autonomous system needs to classify the visual scene into *Rock* and *Not Rock*. This can be achieved by a process called semantic segmentation [8]. Secondly, a three-dimensional (3D) representation of the surface needs to be computed by using either stereo-vision or a depth sensor. This paper will present an architecture to solve the issue of semantically recognizing rock in an underground mine during active drilling operations. This means that the system needs to withstand vibrational loads, dust clouds, and challenging lighting conditions [7]. These elements pose challenges for implementing ready-made computer vision products to address these issues. Additionally, this work will introduce the use of neuromorphic sensors called event cameras, to perform the semantic segmentation task [9].

These cameras provide numerous benefits compared to traditional ones, including low latency, high temporal resolution, and an exceptionally high dynamic range [10]. These features enable effective information capture in a dynamic mining setting, resisting issues such as motion blur from vibrating sensor platforms, shadows from single-point source lighting, and interference from dust clouds. However, due to the novelty of this sensor, there is a lack of

This work is supported by National Institute for Occupational Safety & Health | NIOSH/Project 75D30121C12206.



Figure 1. Shows a typical color image of the roof with a bolted strap, taken at the Edgar Mine in Idaho Springs, CO. It is imperative for any computer vision system to be able to differentiate between the strap and the underlying rock

a comprehensive and diverse associated dataset [11]. This poses an additional challenge when considering commercial vision products, as the limited availability of labeled data hinders the design and training of a generalized machine learning model. Typically, the solution to this issue involves

utilizing simulated data [12]. However, simulators often lack representation of visual features commonly found in mines, and the transition from simulation to reality is known to be problematic. The pipeline proposed in this paper can be used as an alternative by quickly and cheaply generating a labeled data set.

It is difficult to accurately label event-based images with semantic information directly. Therefore, a traditional color camera is used in conjunction with two event cameras. This is done so that a semantic segmentation network can be trained in optical color space. The trained network is then used to predict a semantically segmented image using an event-based image, called a mask. This mask is shown to be accurate enough to close the loop on a self-supervised approach [13] to automatically label event-based images.

All data pictured is collected at the Edgar Experimental Mine in Idaho Springs, Colorado. In order to direct the drill autonomously to perform roof bolting, areas of rock need to be segmented differentially from areas of support strap [14].

The primary contribution of this work is the training architecture to enable a self-supervised approach to the semantic segmentation task, while simultaneously using domain transfer techniques to reduce the cost and improve the fidelity of labeling data.

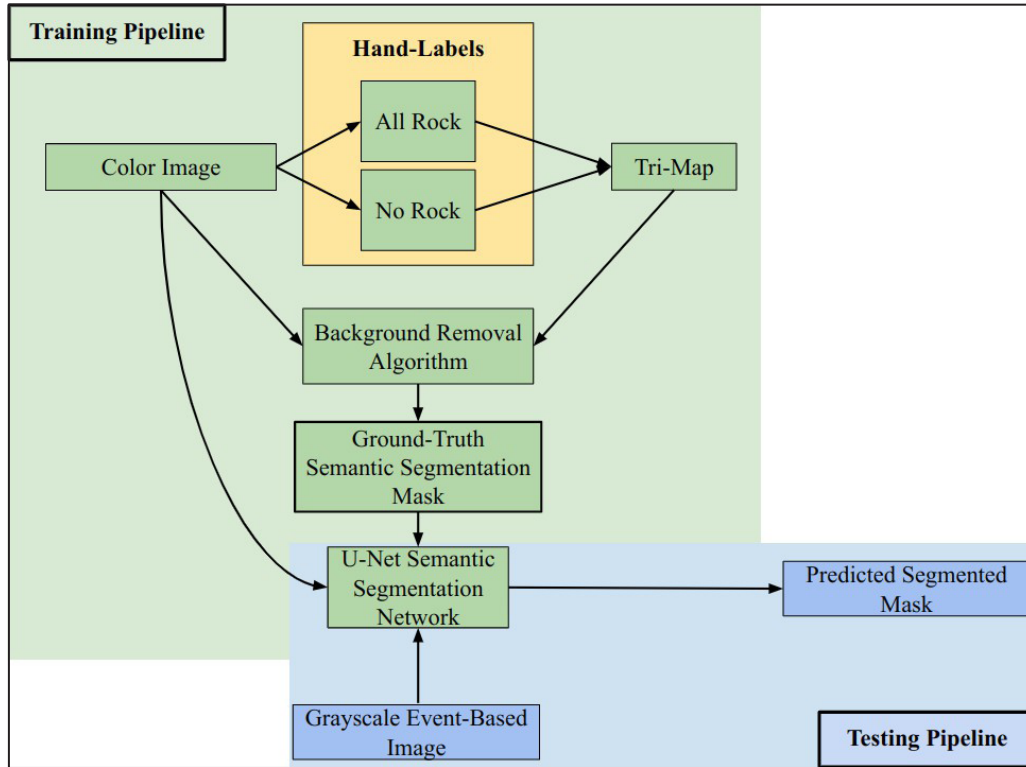


Figure 2. Shows both the training and testing pipelines intended for the development of event-based semantic segmentation

This paper is laid out with a discussion of the planned approach, a description of the experimental setup, a presentation of the results, a discussion grounding the results in the context of the project, and a conclusion.

METHODOLOGY

This effort is to enable the prediction of semantic segmentation masks given an event image taken during roof bolting as shown in Figure 2.

Since the goal is to minimize the per-pixel error between an event-based image and its corresponding ground-truth segmentation mask, any supervised detection or segmentation model can be used, given a diverse data set. To do so, color images are labeled twice. Once with all the *Rock* classification pixels removed and once with all the *Not Rock* pixels removed. This can be combined into a trimap. Using a background removal algorithm, a ground-truth mask can be computed. This combined with the original color image is used to train a supervised segmentation network. To test, time-synced event and color images are input to the trained network. The color prediction can serve as a labeled image for the event image. This process can be iterated until the accuracy of the event images asymptotes to that of the color images.

A. Developing a Data Set for Color Classification

The color pictures, like the one in Figure 3 are hand-labelled into a pair of images. One with all *Rock* pixels and one with all *Not Rock* pixels. This can be done within a much larger margin of error than would be possible if this were to result in a ground-truth image directly.

Figure 4 shows that the human-generated labels are often a lower fidelity because it is difficult to accurately label each pixel. Due to the allowance of an area of uncertainty, even though this is a binary classification problem, a data set can be created more quickly.

These images are then combined using bitwise logical operations as in 1 to create an area of uncertainty. This is then combined with the all rock image to generate a tri-map as seen in Figure 5. This is an image with an area of certain rock, certain not rock, and a band of uncertainty between them.

$$T = R \cup (R \cap N) \quad (1)$$

where,

- T is the set of all pixels that create the trimap.
- R is the set of all white pixels belonging to Rock,
- N is the set of all black pixels belonging to *Not Rock*.



Figure 3. Shows an image of the roof from a typical point-of-view of a roof bolting drill. It is critical to be able to differentiate the strap from the rock

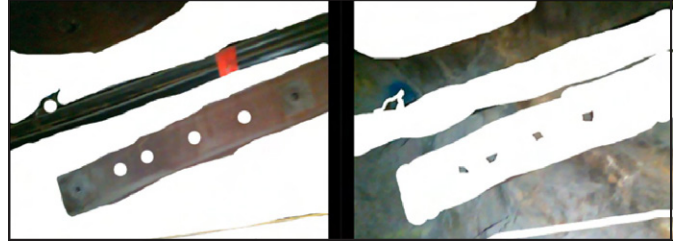


Figure 4. Left Shows a hand-labeled image of the roof with all the pixels belonging to the *Rock* classification removed. Right Shows a hand-labeled image of the roof with all the pixels belonging to the *Not Rock* classification removed

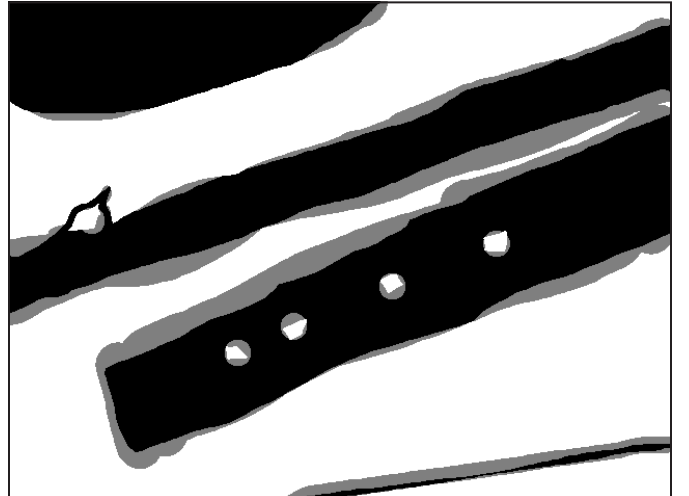


Figure 5. Shows the generated tri-map. Here the white pixels represent *Rock* labels, the black pixels represent *Not Rock* label and the grey pixels represent the area of uncertainty

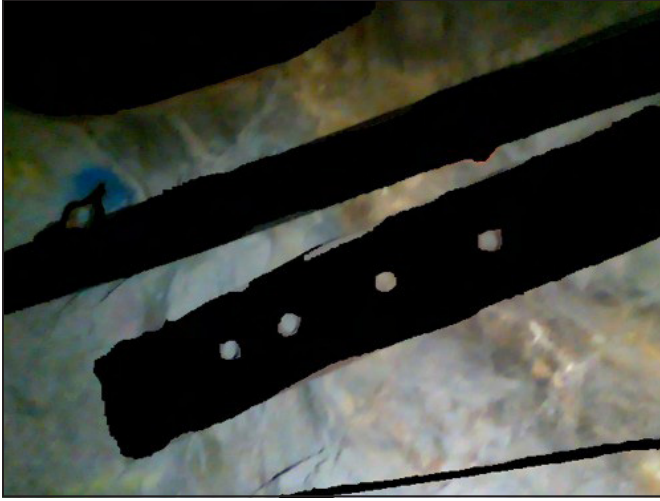


Figure 6. Shows the result of the alpha matting algorithm where the boundary between the rock and strap has been drawn. This mask can be processed to eventually be ground truth

Now what remains is for the decision boundary to be drawn using a general background removal algorithm called alphasatting [15]. But any other algorithm could be used. This is the auto-labeling step in the process. The mask is thresholded to suit a binary segmentation architecture and cleaned up using morphological operations. The resulting binary image is treated as ground truth as shown in Figure 6.

As can be seen, this process results in a high-fidelity mask. However, the auto-labeling engine used currently is a numerical one that requires a color image as well as the generation of a trimap. All of these factors result in a system that is infeasible for online, real-time rock segmentation in the mine. The pipeline can be used to generate a large, diverse dataset with training labels for a binary, semantic segmentation network.

B. Developing the Color Segmentation Network

The semantic segmentation network is built using a U-Net architecture [16]. The network uses an encoder-decoder structure with four successive convolutional layers. This extracts the salient features from the image and the ground truth mask. These are combined to minimize a binary cross-entropy loss [17] using the Adam optimizer [18].

The data set is divided into batches of four and trained for 15 epochs using an NVIDIA GeForce RTX 3070 graphical processing unit. A validation set is set aside to test the model at the end of each epoch. The final validation accuracy of the trained network is 98%. The accuracy and loss

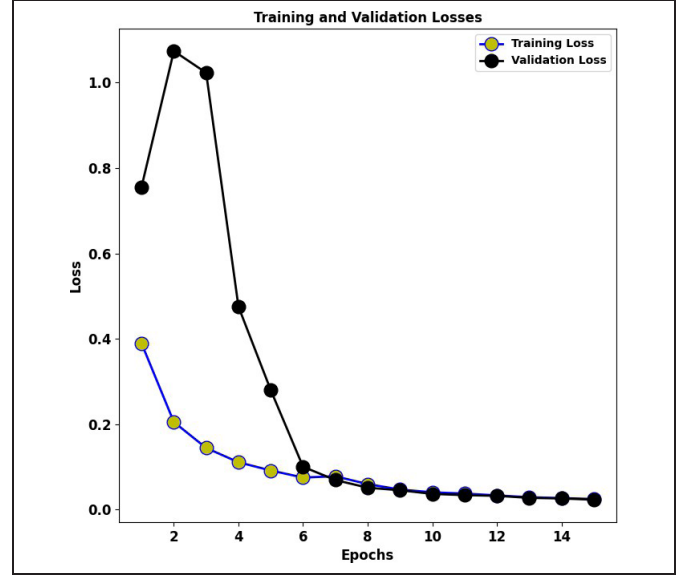


Figure 7. Shows the loss curves generated during the training of the network. An early stopping strategy was utilized so as not to overtrain the network

curves are plotted in Figure 7. Several strategies were used so as to prevent over-fitting [19]. Firstly an early-stopping strategy was used to stop the training process before validation variance started rising as can be seen by the validation curve in Figure 7. Secondly, the training and validation set was shuffled at each epoch.

C. Domain Transfer

The next step of the process is to be able to leverage this trained network to predict masks for event-based images. This will allow for real-time, low latency semantic segmentation as shown in Figure 8.

It can be seen that the mask predicted from the color image is of higher fidelity, with the holes in the strap being recognized by the network. Therefore, this can be used to validate the predictions from unseen event images.

Simultaneously, this architecture can be used to validate these predictions. The extrinsic calibration parameters, discussed in the following section, are used to map the color images into the event camera frames. Both the event and warped color images are input to the network. The prediction from the input color image is used as ground-truth labels and compared to the semantic prediction from the event image. Now, the domain transfer loop can be closed, as event-based images can be chosen and added to the dataset along with their predicted masks. Additional training regiments can be run with this multi-domain dataset so that the network becomes more generalized over time.

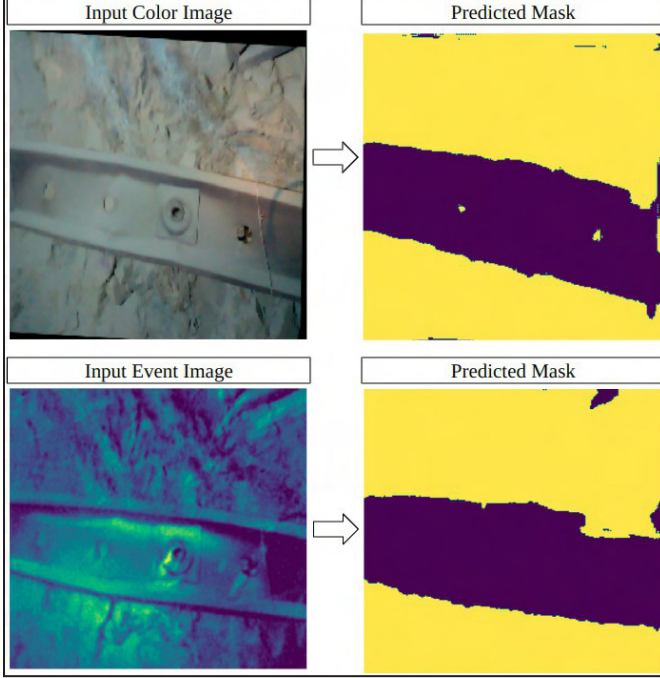


Figure 8. Shows the result of predicting both semantic and color masks using the network trained and validated on the color data set

D. Evaluation Metrics

To analyze the performance of the predictions of the network, a Bayesian approach [20] is used.

$$P(R_p | N) = \frac{P(R_p) \times P(N)}{P(N | R_p)} \quad (2)$$

$$P(N_p | R) = \frac{P(R_p) \times P(R)}{P(R | R_p)} \quad (3)$$

where $P(R_p | N)$ is the probability of a pixel being predicted as Rock given that it is *Not Rock*. This is also called a False Positive. $P(R_p)$ is the probability of a pixel being predicted as Rock. $P(N)$ is the probability of a pixel in the test image belonging to the *Not Rock* classification. $P(N | R_p)$ is the probability of a pixel being *Not Rock* given that it was predicted to be Rock. $P(N_p | R)$ is the probability of a pixel being predicted as *Not Rock* given that it is Rock. This is also called a False Negative.

A confusion matrix is computed for each image. This allows for the categorization of each pixel into True Positive (TP), True Negative (TN), False Positive (FP), and False Negative (FN). This is appropriate for a binary classification task as the eventual goal is to autonomously identify the appropriate areas in the rock to drill.

The accuracy of the eventual prediction can be computed as

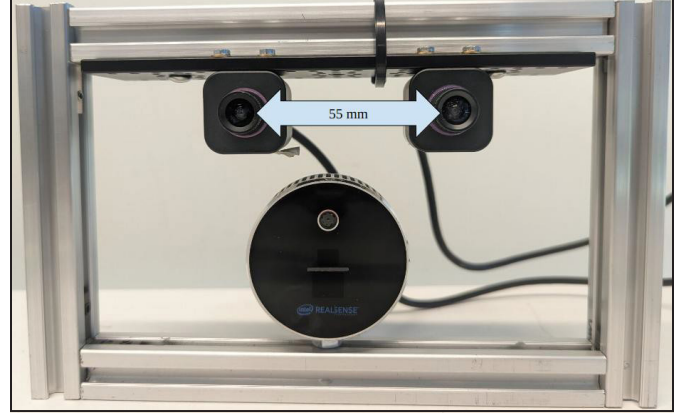


Figure 9. Shows the sensor rig used during mobile data collection operations in the mine

$$\text{Accuracy} = \frac{\text{True Positive} + \text{True Negative}}{\text{Number of Pixels}} \quad (4)$$

DATA SET COLLECTION

The sensors used are two iniVation DAVIS (Dynamic and Active-pixel Vision Sensor) event cameras and one Intel RealSense L515 (Lidar-Camera). These are mounted rigidly as shown in Figure 9.

The pipeline was evaluated by collecting both training and testing data in the Edgar Mine in Idaho Springs, Colorado. The hypothesis is that semantic masks can be predicted for event-based images in severe environmental conditions during drilling activities.

A. Extrinsic and Intrinsic Calibration

The sensors are calibrated using computer vision techniques to compute the relative positions of the sensor axes. From this, we can get the distance between the stereo event cameras as well as the homogeneous transform between the L515 color optical frame and each of the event cameras' frames. This is the result of the extrinsic calibration [21]. The relative position of each sensor is computed with reference to the geometric center of the mobile rig using pose computation using a fiducial marker [22] as in Figure 10.

Homogenous transforms can be chained together to map a depth image from the Lidar to be centered around each of the event camera axes. This is used as the ground-truth depth image, from the perspective of the event cameras, ${}^E I$.

$${}^E I = {}^E T_L \times {}^L I \quad (5)$$

where L_I is the depth image taken by the LiDAR, ${}^E I$ is the ground-truth depth image from the perspective of each of the event cameras, and ${}^E T_L$ is the transform between the

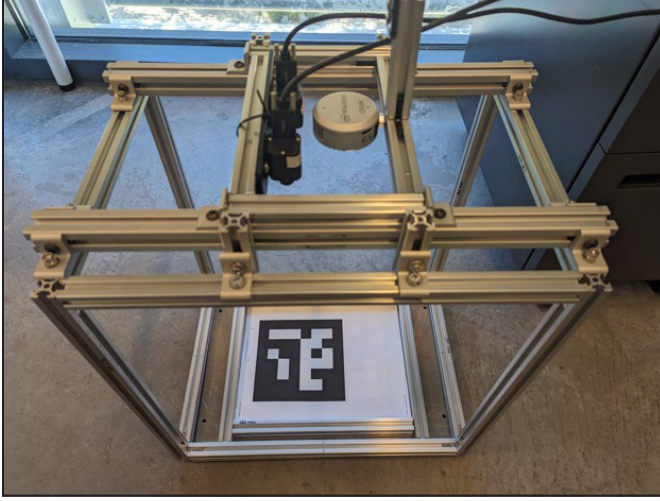


Figure 10. Shows the calibration rig used to extract the relative sensor poses

depth camera and each of the event cameras found from extrinsic calibration.

The event image E_I , is time synced with the color image, I which is then input to the trained network. The prediction is a real-time label for the event image and can be added to the data set to further diversify it.

RESULTS

The testing set is validated by the time-synced and mapped corresponding color image prediction as shown in Figures 11, 12, and 13.

Color Segmentation

The performance of the trained model on color images is evaluated using the metrics outlined above. The validation set has 612 test image and ground-truth label pairs available but was not used to train the network.

		Ground-truth Label	
		Rock	Not Rock
Predicted Label	Rock	0.997	0.003
	Not Rock	0.069	0.931
Accuracy		0.964	

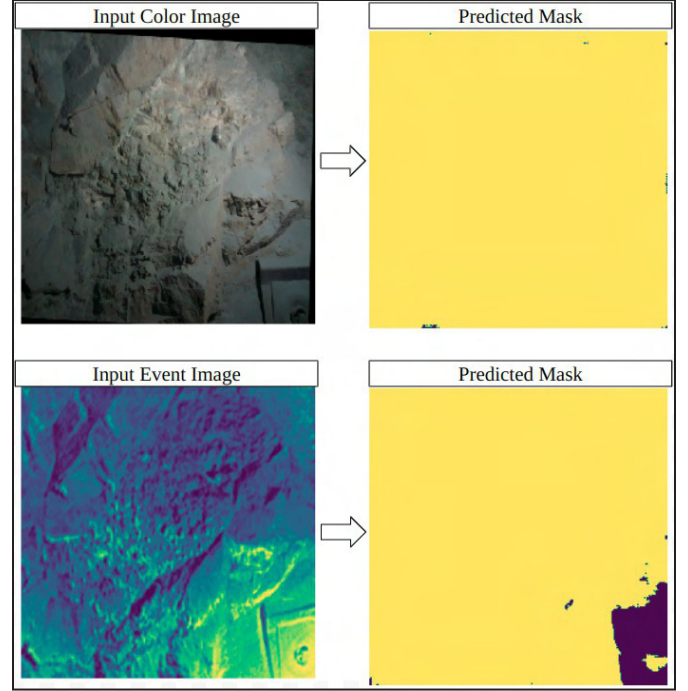


Figure 11. Shows a severe lighting scene with shadows. The model predicts a support strap from an event image, while the prediction from the color image misses it

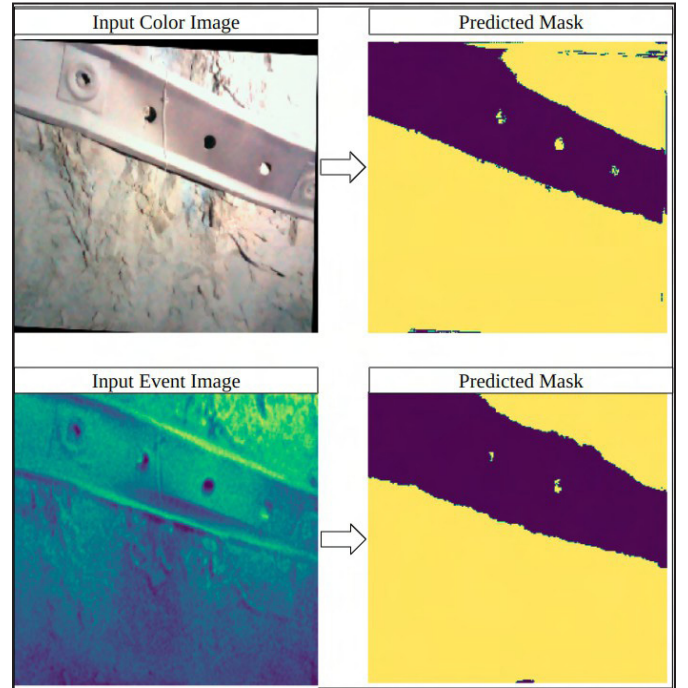


Figure 12. Shows the result of predicting a semantic mask using the network trained and validated on the color data set

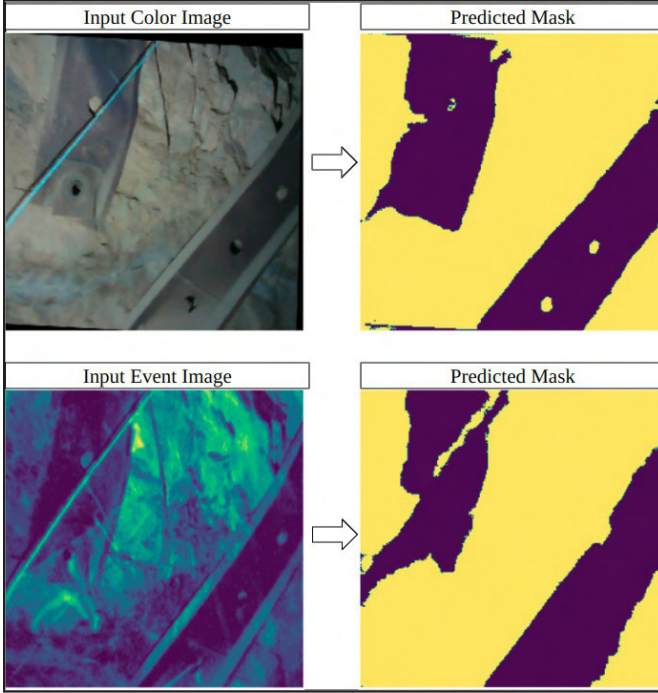


Figure 13. Shows the result of predicting a semantic mask using the network trained and validated on the color data set

Event Segmentation

The performance of the model on event images is evaluated using the labels from the associated color-based prediction. This test set is comprised of 8324 images from 2 different sections of the mine.

		Ground-truth Label	
		Rock	Not Rock
Predicted Label	Rock	0.962	0.038
	Not Rock	0.210	0.790
Accuracy		0.876	

DISCUSSION

This study demonstrates the efficacy of using transfer learning techniques to execute the semantic segmentation task in an active mining environment using event-based images. Notably, the false positive rate, denoted as $P(R_p|N)$, is approximately 3%. This finding is particularly encouraging

for the intended application, given that the primary failure mode involves drilling into areas other than rock. The network exhibits a tendency to predict false negatives, i.e., $P(R_p|N)$, about three times more frequently when event images are used as input compared to traditional color images. This characteristic contributes to a more conservative approach in the selection of suitable roof areas for drilling. The higher likelihood of false negatives results in a conservative strategy, emphasizing precision and minimizing the risk of drilling in inappropriate locations, thereby enhancing the overall reliability and safety of the drilling process in active mining scenarios.

The higher incidences of false negatives can be explained by lower-quality ground truth as seen in Figure 11. Statistically, the strap detected by the event image is classified as a false negative prediction. However, upon inspection, it can be seen that the lighting conditions adversely affected the color image and the prediction missed the strap altogether. The superior optical properties of the event camera allowed for the prediction to label the strap accurately and outperform a traditional camera. There were several instances in the dataset where the color of the support strap had asymptotized to that of the rock behind it. In each of these cases, the event-based prediction was able to capture more detail than the traditional color image prediction. This shows that this work can be valuable to the application of robotics in mining. Even in severe environmental and lighting conditions, computer vision can be applied to automate otherwise dangerous processes.

The training set includes labeled images of holes in the support strap, and Figure 12 shows the continued accurate prediction of these features even after transferring the learning to event images. This observation holds significant promise for several reasons. Firstly, it indicates the robustness and adaptability of the model to recognize and predict specific features, such as holes in the support strap, across different data modalities. This adaptability is crucial for practical applications where variations in imaging technologies might be encountered.

Moreover, the sustained accurate prediction of support strap holes in event images underscores the effectiveness of the transfer learning approach. It suggests that the knowledge gained from the initial training on traditional color images successfully transfers and generalizes to event-based images, showcasing the model's capacity to extract meaningful information and maintain predictive accuracy across diverse datasets. This adaptability and transferability are vital attributes for real-world scenarios where a model trained on one set of conditions needs to perform reliably in different, potentially challenging, environments.

The blue wire as seen in Figure 13 is an unseen feature to the network. Despite the network’s generalization capability, where the color prediction correctly categorizes the feature as *Not Rock*, the event prediction exhibits some false positive pixels. This scenario is noteworthy for several reasons. Primarily, it emphasizes the model’s capacity to extrapolate knowledge from the training set and accurately predict features even when encountering previously unseen elements. The successful classification of the blue wire as *Not Rock* in the color prediction indicates the model’s ability to leverage feature-based information to make informed decisions about unfamiliar objects or structures. However, the occurrence of false positive pixels in the event prediction suggests that the network, when relying on event-based data, may encounter challenges in precisely identifying novel features that differ significantly from the characteristics present in the training set. This underscores the need for quick and inexpensive data augmentation techniques as presented by this work.

The incorporation of a self-supervised, hierarchical training pipeline makes a valuable contribution by quickly enhancing the event camera-based dataset. This approach involves a method where the system refines and improves its own learning through various stages, optimizing the dataset for better performance. The significance lies in the accelerated readiness it brings for automated operations in novel environments. By efficiently augmenting the dataset, not only saves time but also minimizes costs associated with initiating automated processes in new settings. Overall, the utilization of such a training pipeline facilitates a more seamless and cost-effective transition to automated operations across diverse environments.

This work was able to demonstrate a self-supervised approach to the semantic segmentation task which proved to be effective in the harsh lighting and environmental conditions of an active mine. In fact, the vibratory load expected on a drill rig will trigger the event cameras. While these vibrations would cause motion blur in traditional optical cameras, the event cameras generate images with more dense information. Additionally, the ability to use a model trained on color images, and then predict using event-based images is valuable as pre-existing data sets can be leveraged to accelerate the performance of novel sensors. While there are large data sets for urban driving like Cityscapes [23] and KITTI [24], semantic recognition networks struggle with generalization [25]. This leads to industry-wide specialization in urbanized environments in developed countries. This work can be used to fill in the gaps of representation at a much lower cost. In essence, this approach promotes a seamless transition from established imaging technologies

to emerging ones, ensuring a smoother integration of novel sensors into existing systems while optimizing the use of available data resources.

CONCLUSION

This work demonstrates a pipeline that uses inexpensive semantic hand labels and self-supervision methods to generate a generalized data set. Using this pipeline, this paper presents the capability to use event cameras to solve the semantic segmentation task in a mine during a drilling session. Additionally, the capability to benchmark against traditional color cameras is presented and shows that event cameras are able to perform the task.

REFERENCES

- [1] JJ Sammarco, A Podlesny, EN Rubinstein, and B Demich. An analysis of roof bolter fatalities and injuries in us mining. *Transactions of Society for Mining, Metallurgy, and Exploration, Inc*, 340(1):11, 2016.
- [2] Mark Aldrich. Engineers attack the “no. one killer” in coal mining: The bureau of mines and the promotion of roof bolting, 1947–1969. *Technology and culture*, pages 80–118, 2016.
- [3] Christopher Mark. The introduction of roof bolting to us underground coal mines (1948–1960): a cautionary tale. 2002.
- [4] Ronald C Althouse, Michael J Klishis, and R Larry Grayson. Microanalysis of roof bolter injuries. *Applied occupational and environmental hygiene*, 12(12):851–857, 1997.
- [5] Fred C Turin. Human factors analysis of roof bolting hazards in underground coal mines. 1995.
- [6] Syd Peng. Roof bolting and underground roof falls. *Geohazard Mechanics*, 1(1):32–37, 2023.
- [7] Jing Li, Yaru Qin, Cheng Guan, Yanli Xin, Zhen Wang, and Ruikang Qi. Lighting for work: a study on the effect of underground low-light environment on miners’ physiology. *Environmental Science and Pollution Research*, pages 1–10, 2022.
- [8] Yongqin Xian, Subhabrata Choudhury, Yang He, Bernt Schiele, and Zeynep Akata. Semantic projection network for zeroand few-label semantic segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2019.
- [9] Zhaoning Sun, Nico Messikommer, Daniel Gehrig, and Davide Scaramuzza. Ess: Learning event-based semantic segmentation from still images. In *European*

- Conference on Computer Vision*, pages 341–357. Springer, 2022.
- [10] Guillermo Gallego, Tobi Delbrück, Garrick Orchard, Chiara Bartolozzi, Brian Taba, Andrea Censi, Stefan Leutenegger, Andrew J Davison, Jörg Conradt, Kostas Daniilidis, et al. Event-based vision: A survey. *IEEE transactions on pattern analysis and machine intelligence*, 44(1):154–180, 2020.
 - [11] Alex Zihao Zhu, Dinesh Thakur, Tolga Özaslan, Bernd Pfrommer, Vijay Kumar, and Kostas Daniilidis. The multivehicle stereo event camera dataset: An event camera dataset for 3d perception. *IEEE Robotics and Automation Letters*, 3(3):2032–2039, 2018.
 - [12] Shital Shah, Debadepta Dey, Chris Lovett, and Ashish Kapoor. Airsim: High-fidelity visual and physical simulation for autonomous vehicles. In *Field and Service Robotics: Results of the 11th International Conference*, pages 621–635. Springer, 2018.
 - [13] Ashish Jaiswal, Ashwin Ramesh Babu, Mohammad Zaki Zadeh, Debapriya Banerjee, and Fillia Makedon. A survey on contrastive self-supervised learning. *Technologies*, 9(1):2, 2020.
 - [14] Xiaowei Feng, Fei Xue, Xiaotian Feng, and Tongyang Zhao. Failure characteristics of w strap in coal mine support. *Archives of Mining Sciences*, 67(2), 2022.
 - [15] Yağız Aksoy, Tunç Ozan Aydın, and Marc Pollefeys. Information-flow matting. 2019.
 - [16] Olaf Ronneberger, Philipp Fischer, and Thomas Brox. Unet: Convolutional networks for biomedical image segmentation. In *Medical Image Computing and Computer-Assisted Intervention–MICCAI 2015: 18th International Conference, Munich, Germany, October 5–9, 2015, Proceedings, Part III* 18, pages 234–241. Springer, 2015.
 - [17] I. J. Good. Rational decisions. *Journal of the Royal Statistical Society. Series B (Methodological)*, 14(1):107–114, 1952.
 - [18] Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization, 2017.
 - [19] Xue Ying. An overview of overfitting and its solutions. In *Journal of physics: Conference series*, volume 1168, page 022022. IOP Publishing, 2019.
 - [20] James Joyce. Bayes’ theorem. 2003.
 - [21] Fumio Itami and Takaharu Yamazaki. A simple calibration procedure for a 2d lidar with respect to a camera. *IEEE Sensors Journal*, 19(17):7553–7564, 2019.
 - [22] S. Garrido-Jurado, R. Muñoz-Salinas, F. J. Madrid-Cuevas, and M. J. Marín-Jiménez. Automatic generation and detection of highly reliable fiducial markers under occlusion. *Pattern Recognition*, 47(6):2280–2292, June 2014.
 - [23] Marius Cordts, Mohamed Omran, Sebastian Ramos, Timo Rehfeld, Markus Enzweiler, Rodrigo Benenson, Uwe Franke, Stefan Roth, and Bernt Schiele. The cityscapes dataset for semantic urban scene understanding. In *Proc. of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016.
 - [24] Moritz Menze and Andreas Geiger. Object scene flow for autonomous vehicles. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2015.
 - [25] Dewant Katare, Nicolas Kourtellis, Souneil Park, Diego Perino, Marijn Janssen, and Aaron Yi Ding. Bias detection and generalization in ai algorithms on edge for autonomous driving. In *2022 IEEE/ACM 7th Symposium on Edge Computing (SEC)*, pages 342–348. IEEE, 2022.

Karl Dufresne

Market Development Manager, DynaEnergetics
Europe GmbH Troisdorf, Germany

D. Scott Scovira

Global Manager Blasting Science, BME USA Inc.
Chatham, MA

ABSTRACT

Blasting in geothermally high temperature ground and/or extremely reactive ground is an increasing common challenge for surface and underground mines. Many current PETN based non-electric and electronic detonators have manufacturer published in-hole temperature limits ranging from 55°C to 65°C (131°F to 150°F). When confronting blasthole temperatures above the manufacturer published limits, operations may request site-specific guidance and qualification for use. The advised safe loading window for in-hole PETN based detonators is often too short to be practicable for production blasting.

To overcome limitations of PETN based initiation systems, surface mines with high temperature ground and/or reactive ground are known to load blastholes with a qualified high temperature bulk emulsion that are top primed with a cartridge of temperature qualified detonator sensitive emulsion plus an RDX detonating cord downline. As there is no in-hole delay detonator in this system, blast-hole timing is accomplished with surface delays. Although this system is safe, top priming and surface delaying does influence floor quality and top size fragmentation, and safe loading windows times reduce blast block size and increase blasting frequency.

A mining company with a site characterized by geothermally high temperature ground plus highly reactive ground employed a blasting system like the one described above. The site desires to improve blasting performance by moving to an in-hole delay initiation system that can withstand the

site ground conditions. At the request of the mining house, DynaEnergetics developed an electronic initiation system for application in geothermally high temperature ground and/or extremely reactive ground conditions.

The DynaEnergetics IGNEO high temperature electronic initiation system is rated for in-hole use in ground up to 150°C (302°F) for 48 hours. The detonator consists of a high strength detonator with an HMX base charge and thermally resistant downline. To make up a primer, the detonator is inserted into a booster charge assembly that is fitted with a small HMX shape charge. This high energy primer is designed to be compatible and reliable for firing bulk emulsions, including those with high percentages of urea inhibition to combat reactive ground. Up to 1000 detonators may be programmed in 1 ms increments up to 1500 ms. The system is controlled by a digital firing panel that can be used with either a wired or wireless trigger.

The system enables high temperature and/or reactive ground sites to utilize in-hole delay initiation and realize the associated performance benefits.

TRADITIONAL BLASTING METHODS IN HIGH TEMPERATURE AND/OR EXTREMELY REACTIVE GROUND

Reactive ground is a term used in the mining industry to describe the condition where an exothermic chemical reaction can occur between sulphide bearing rock and blast holes loaded with nitrate-based explosives. Reactive ground conditions increase the risk of explosives deflagration and/or unplanned detonations.