

Automated 1-D Barcode Localization and Decode System for Warehouses

(Technical Paper)

Why We Need to Build Trust in AI

(STS Paper)

A Thesis Prospectus submitted to the

Faculty of the School of Engineering and Applied Sciences
University of Virginia • Charlottesville, VA

In Partial Fulfillment of the Requirements of the Degree
Bachelor of Science, School of Engineering

Kevin Chung
Spring 2024

Technical Project in Cooperation with
Coros, corp. • Mountain View, CA

On my honor as a University Student, I have neither given nor received
unauthorized aid on this assignment as defined by the Honor Guidelines
for Thesis-Related Assignments

Introduction

We are experiencing a technological revolution with the rapid development of artificial intelligence (AI) and machine learning (ML) (Tang et al., 2020). One of the biggest recent developments has been the rise of generative AI. With generative AI, there are now machine learning models that can generate images from a prompt, seemingly comprehend and respond to text, and even generate audio using someone else's voice. If you've been on the internet in the past year, you've likely heard of ChatGPT, a model capable of responding to or completing text prompts and seemingly able to understand the content of a sentence (OpenAI, 2023). Generative models like ChatGPT hold massive implications for anyone who needs to read or write, which is nearly everyone.

Another major application of AI and ML is in computer vision, which, according to IBM, one of the pioneers of artificial intelligence, "is a field of artificial intelligence that enables computers and systems to derive meaningful information from digital images, videos, and other visual inputs." (n.d.) It's one of the most heavily researched fields in computer science; it had a market share of 11.9 billion in 2022 (Lee, 2023). Additionally, similar to generative AI, it has applications in pretty much any industry. Computer vision has also already made its way into our daily lives, with things such as facial recognition to unlock your phone, video filters that change what you look like, and autonomous cars.

As these new technologies become more commonplace in our lives, automating things such as writing or driving, it's imperative that we discuss trust in these technologies. They need access to some level of our personal or proprietary information to function, information that is necessary to keep private. Additionally, trust is one of the most important factors when determining if a technology is going to be adopted (Hoff and Bashir, 2015). The problem is that

currently there is not much trust in AI and ML. According to a survey of over 17,000 people across 17 countries, 61% of participants were hesitant to trust AI systems. Some reasons that were cited were doubts about security, fairness, and safety of AI systems (Gillespie et al., 2023). I argue, however, that there are some AI models that we can and should trust. Without public trust in AI and ML, we cannot ethically implement these innovations into other systems, preventing us from taking advantage of these extremely powerful tools.

Technical Topic

Laser scanning technology for 1-dimensional barcode decoding has existed since the 1970s (“The History of the Bar Core”, n.d.). However, recent advances in computer vision techniques and computing power have resulted in the emergence of camera-based barcode decoding. Generative adversarial networks (GANs) (Goodfellow et al., 2020), a subset of generative AI, have been crucial in improving decode performance for vision-based approaches. GAN-based models are often used to “deblur” an image, enhancing the contrast on an image and sharpening the edges around objects in an image.

My technical project aims to combine DeblurGANv2 (Kupyn et al., 2019), a popular GAN-based model, with other computer vision techniques, such as Hough transforms (Duda and Hart, 1972) and localization models, to create a system for warehouses that can locate packages, shipping labels, and barcodes. It is also able to decode the located barcodes and read text either printed on the packages or shipping labels. In a fast-moving warehouse environment, images will often have motion blur and lighting conditions will vary. By using DeblurGANv2, our system should be able to control for various lighting conditions and improve the sharpness of the barcode image. A paper by Wang et al. (2022) shows that applying DeblurGANv2 to barcode images improves barcode decoding independent of the library used to decode. Working with my

team at Coros, a startup company based in Mountain View, I hope to create a system with acceptable performance using a vision-based approach rather than laser scanning.

My proposed method uses a pipeline that can be split into three stages: detection, enhancement, and decode. For the detection stage, the system will use a YOLO-family (You Only Look Once) model (Redmon, 2016), YOLOv8, to detect bounding boxes for parcels, shipping labels, and 1D barcodes. Parcel and shipping label detections can be fed through other pipelines, allowing a lot of flexibility with this approach. This paper will, however, only focus on the 1D barcodes. Once the barcodes have been detected, the system crops the barcodes and moves on to the enhancement stage. In this stage my method uses a Hough transform to align the barcode vertically, and DeblurGANv2 to enhance the sharpness of the barcode. I also experimented with various thresholding techniques to improve the contrast of the barcode. After the barcodes have been enhanced, the system uses Pyzbar, a Python library built to decode 1D barcodes.

STS Topic

Discussing AI in today's age almost always comes with a discussion of trust. Some people have already adopted generative AI into their normal workflow, using it to answer questions, generate writing for them, or write code, with the most commonly used model being ChatGPT. However, those who are users of or informed about ChatGPT know that everything the model outputs will not necessarily be true. The model is inherently generative, as it's in the name, meaning that the output it generates will look like things it has seen in the past, but is not fundamentally grounded in truth. As a result generative AI generates realistic but fabricated content, which could potentially lead to the spread of misinformation or fake news (Tredinnick and Laybats, 2023, p. 47). Generative AI models can also propagate biases and prejudices in

their training data, which can have serious consequences when used in applications such as image generation. In an article by Moss (2023), it's stated that a survey showed that 73% of employees do not trust generative AI and believe it poses security risks. If we want to take full advantage of generative AI, we need to understand how we can build trust in it.

Another large reason for the distrust of AI is that many AI models are “black boxes.” This is a different type of trust, but very closely linked with the distrust in reliability. There are methods such as looking at the interior of the model to see what part of the input is being used or changing the input methodically to determine how different perturbations affect the model output. However, these both fail on generative AI models such as ChatGPT or Dall-E, due to the sheer size of these models (Bischoff, 2023). In order to maximize our benefits of generative AI, we change the stigma around generative AI, or we develop a method to peer into the model and understand what is going on behind the scenes. Innovating methods to analyze the inner workings of a ML model requires an extensive understanding of ML architecture (which I do not possess), so I will be discussing the former.

Changing a stigma of distrust around AI requires understanding the antecedents of trust in technology. According to Jacovi et al., one of the most important antecedents of trust for AI models is the ability to evaluate its reliability. While non-experts cannot easily evaluate the reliability of most models, they can gain trust by proxy of expert evaluations (2021, p. 629). Experts evaluating models should be able to communicate what different metrics mean in non-technical terms, and provide recommendations on circumstances that maximize the models effectiveness and accuracy. If we place a larger emphasis on peer evaluations, or form some sort of other certification for more reliable models, the ability to determine if a model is reliable should become more accessible.

Prioritizing reducing bias and discrimination can also build trust in AI. This is especially important when bias in AI limits accessibility to a technology, or it acts in a manner that is discriminatory towards a subset of people. Bias most often arises from the dataset used to train a model, since ML is designed to propagate patterns from the training data. However, in a study conducted by Ferreira et al., they found that some of the most popular facial recognition models produce more error when tested on Black faces, despite the training data having a majority of Black faces (2021). The onus is on the engineers that select and train these models to ensure that they do not result in unfair biases. The public can also call attention to software that uses AI that discriminates against groups of people, which should incentivize the engineers of the software to more carefully inspect their models. There have already been great strides towards prioritizing fairness in models, such as IBM releasing an open source tool called “AI Fairness 360” that allows developers to easily check for bias in their models, and the EU forming a commission to set forward strict ethical guidelines for AI development (Rossi, 2018, pp. 127–34).

Conclusion

AI is one of the fastest-growing industries right now, meaning we are going to see rapid improvements in the capabilities of generative AI in the next few years. If we do not trust the results from generative AI, or it continues to propagate bias or prejudices, it will be nearly impossible to apply it in the real world, as both often lead to lack of acceptance and adoption. It is unfair to place the responsibility of evaluating trustworthiness of AI on the public, as most people are not educated enough on the subject to evaluate a model. The responsibility falls on experts to clearly communicate to the public the reliability and bias of various models. If experts are able to accomplish this, it should allow everyone to have some idea of how reliable each model is.

Trust is extremely complicated and nuanced, and it may very well be the case that people never truly trust AI. There are many, many facets of trust, each which need an extensive effort to address and build. However, we need to take small steps in building public trust in AI and ML, as our technology is headed in the direction of AI being intertwined with nearly everything. Without trust, we cannot morally use these technologies to their full potential. AI and ML have proven to have the capability to revolutionize, we just need to trust them to do so.

References

Bischoff, Manon. “New Tool Reveals How AI Makes Decisions.” *Scientific American*.

Accessed October 17, 2023.

<https://www.scientificamerican.com/article/new-tool-reveals-how-ai-makes-decisions/>.

Duda, Richard O., and Peter E. Hart. “Use of the Hough Transformation to Detect Lines and Curves in Pictures.” *Communications of the ACM* 15, no. 1 (January 1, 1972): 11–15.

<https://doi.org/10.1145/361237.361242>.

Ferreira, Marcos Vinícius, Ariel Almeida, João Paulo Canario, Matheus Souza, Tatiane Nogueira, and Ricardo Rios. “Ethics of AI: Do the Face Detection Models Act with Prejudice?” In *Intelligent Systems*, edited by André Britto and Karina Valdivia Delgado, 89–103. Lecture Notes in Computer Science. Cham: Springer International Publishing, 2021. https://doi.org/10.1007/978-3-030-91699-2_7.

Gillespie, Nicole, Steven Lockey, Caitlin Curtis, Javad Pool, and Ali Akbari. “Trust in Artificial Intelligence: A Global Study.” Brisbane, Australia: The University of Queensland; KPMG Australia, February 16, 2023. <https://doi.org/10.14264/00d3c94>.

Goodfellow, Ian, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. “Generative Adversarial Networks.”

Communications of the ACM 63, no. 11 (October 22, 2020): 139–44.

<https://doi.org/10.1145/3422622>.

Hoff, Kevin Anthony, and Masooda Bashir. “Trust in Automation: Integrating Empirical Evidence on Factors That Influence Trust.” *Human Factors* 57, no. 3 (May 1, 2015): 407–34. <https://doi.org/10.1177/0018720814547570>.

IBM. “What Is Computer Vision?” Accessed October 12, 2023.

<https://www.ibm.com/topics/computer-vision>.

Jacovi, Alon, Ana Marasović, Tim Miller, and Yoav Goldberg. “Formalizing Trust in Artificial Intelligence: Prerequisites, Causes and Goals of Human Trust in AI.” In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 624–35. Virtual Event Canada: ACM, 2021.

<https://doi.org/10.1145/3442188.3445923>.

Kupyn, Orest, Tetiana Martyniuk, Junru Wu, and Zhangyang Wang. “DeblurGAN-v2: Deblurring (Orders-of-Magnitude) Faster and Better.” arXiv, August 10, 2019.

<https://doi.org/10.48550/arXiv.1908.03826>.

Lee, Eric. “Computer Vision Market Size Worth USD 22.76 Billion in 2032 | Emergen Research,” March 9, 2023.

<https://www.prnewswire.com/news-releases/computer-vision-market-size-worth-usd-22-76-billion-in-2032--emergen-research-301767807.html>.

Magazine, Smithsonian. “The History of the Bar Code.” Smithsonian Magazine. Accessed October 5, 2023.

<https://www.smithsonianmag.com/innovation/history-bar-code-180956704/>.

Moss, Tim. “Dangers of Generative AI - Why Some Companies Are Banning It.”

Simplishow (blog), July 24, 2023.

<https://simplishow.com/blog/dangers-of-generative-ai/>.

OpenAI. “GPT-4 Technical Report.” arXiv, March 27, 2023.

<https://doi.org/10.48550/arXiv.2303.08774>.

Redmon, Joseph, Santosh Divvala, Ross Girshick, and Ali Farhadi. “You Only Look Once:

Unified, Real-Time Object Detection.” arXiv, May 9, 2016.

<https://doi.org/10.48550/arXiv.1506.02640>.

Rossi, Francesca. “Building Trust in Artificial Intelligence.” *Journal of International Affairs*

72, no. 1 (2018): 127–34.

Tang, Xuli, Xin Li, Ying Ding, Min Song, and Yi Bu. “The Pace of Artificial Intelligence

Innovations: Speed, Talent, and Trial-and-Error.” arXiv, September 3, 2020.

<http://arxiv.org/abs/2009.01812>.

Tredinnick, Luke, and Claire Laybats. “The Dangers of Generative Artificial Intelligence.”

Business Information Review 40, no. 2 (June 1, 2023): 46–48.

<https://doi.org/10.1177/02663821231183756>.

Wang, Chaoxin, Nicolais Guevara, and Doina Caragea. “Using Deep Learning to Improve

Detection and Decoding Of Barcodes.” In *2022 IEEE International Conference on*

Image Processing (ICIP), 1576–80, 2022.

<https://doi.org/10.1109/ICIP46576.2022.9897371>.