



ISSN: 2454-9940



**INTERNATIONAL JOURNAL OF APPLIED
SCIENCE ENGINEERING AND MANAGEMENT**

E-Mail :
editor.ijasem@gmail.com
editor@ijasem.org

www.ijasem.org

The Use of Text-to-Image Generators in the Design of a User-Friendly Interface for Website Development

¹Mrs. G Swarnalata,²Pakalapati Satish Varma, ³Tadisetti Ravi Teja Sai Kumar, ⁴Nelli Prem Praneeth,

¹Associate Professor, Department of CSE, Rajamahendri Institute of Engineering & Technology, Bhoopalapatnam, Near Pidimgoyyi, Rajahmundry, E. G. Dist. A.P 533107.

^{2,3,4}Student, Department of CSE, Rajamahendri Institute of Engineering & Technology, Bhoopalapatnam, Near Pidimgoyyi, Rajahmundry, E. G. Dist. A.P 533107.

Abstract-

Diffusion Training and Dreambooth technology enabled the propagation of varied diffusion models, which in turn made it easier to use Text-to-image Generators to create high-quality images with defined styles. The main objective is to enhance the efficiency of creating a Usability Interface (UI) for a website. They enhanced Stable Diffusion models so that users may make images in the "FacebookAlegria" (Alegria style) and "blacksilhouette icons" (icon style) styles. The team was able to reduce development timelines because of the ease and speed of image creation. Out of the overall quantity of pictures generated, four were chosen for the Alegria style and seventeen for the icon style for the UI. The results of the research showed a significant change following the adoption compared to previous work done by the same company. Cuts of 81.65% in development time and 22.80% in development budget were made to the front end. The results of the test demonstrate, therefore, that the development was much more effective. Using text-to-image generations is an effective approach to reduce the period of UI development.

Keywords-Text-to-image, Stable diffusion, Efficiency, User Interface, DeepLearning.

INTRODUCTION

Experts believe that 80 million jobs will be eliminated by 2025 due to the rise of AI, which can now do tasks that used to need physical labor [15]. This would happen since there are many benefits to automating and optimizing processes using AI. In this setting, businesses would have to adopt this innovation if they want to reap the benefits of its many competitive advantages. The term "prompt"

refers to a program that enables users to make graphics just by adding text [12]. No matter the creative approach, the resulting photographs must be of high enough quality to show objects and components clearly [6]. How much of an imitation there is [1] is of sufficient quality to mix in with both hand created and computer-generated photos, as shown when the average observer cannot tell that the image is from an AI source. The number ten. Regarding the legal side, it has to be stated that in the Peruvian legal setting, all AI-generated photographs are gratuite [3]. Consequently, the project won't have to deal with the problem of licensing purchases. In order to make the most of these present prospects, a website was developed using text-to-image generators. So, the development team could get the particular user interface graphics they needed without relying on a tertiary University of San Andres de Lima, Peru, Wilfreda Ticona-B, Faculty of Engineering the ability to draw them, or who is an illustrator (wmamani@esan.edu.pe). Improving development efficiency via cutting costs and processtimes is the primary goal of this study. Using diffusion models [2] that were enhanced using DreamBooth technology [13] to accomplish the learning of the styles needed by the site was the approach that was used to produce the pictures. Presented below are several styles: The first is Facebook's aesthetic, sometimes known as Alegria, which consists of geometric pictures with flat colors. Black silhouettes without details will make up the second style. It's for icons. The two diffusion models that came out of this process were then uploaded to a Google Collab server along with the code needed to produce the photos.

RELATEDWORKS

Due to its user-friendliness and versatility, Stable Diffusion technology has recently exploded in popularity for producing high-quality creative

pictures of many kinds [12]. The pictures are created via word input, which are called prompts, as Oppenleander explains. [11]. The text-to-picture image production process relies on a trial-and-error approach called prompt engineering, which makes it easy for people without particular knowledge or abilities to handle [12]. In order to customize and create new pictures, the general public may learn how to utilize prompts appropriately using this learning approach. To create pictures using the book titles inputted into the system as input, a library incorporated a text-to-image generator using GAN technology in 2022 [9]. Nevertheless, there are defects, aberrations, and deformities in the output. That's why it's possible to use the FOmetric to objectively assess picture quality by looking at how realistic the image looks [8]. However, by comparing the outcomes of both algorithms using the measure, stable diffusion technology has shown that it can create pictures of equal or better quality than GAN [4]. In 2023, one author used Dreamwooth technology to teach a diffusion model using pseudo-words, which allowed for efficient and rapid training [16]. By modifying it in this manner, I am able to create new styles and objects with only one word [15]. I can then retrain it with very little resources and a particular phrase. [13]. Using Dreambooth technology, over 1 GB of memory, and 1000 steps, the estimated training time is 5 minutes [14], while this can vary based on the base model used, the amount of pictures utilized for retraining, and the number of steps. According to these numbers, Dreamboothretraining

Part A: Generative Pictures
Two models were obtained after the fine-tuning procedure. Iterates many photos with the necessary style using these models. We had to produce a lot of photos and choose the ones that didn't have major distortions in the elements since the technology is unpredictable.

Here are the results for the "StyleAlegria" pseudo-word prompt:



Fig.3.ImagesinAlegriaStyle



Fig.4.ImagesinIconStyle

In all, 260 photos were produced, with 4 images returned by the searches. With intervals ranging from one minute to one and a half minutes, the time between queries was unpredictable.

RESULTS

At <http://teamlimonagrio.com/>, you can see the outcomes of the installation. After that, both the pre- and post-tests were used to determine efficiency. You can find the indicators that will be used to analyze the outcomes in the table that follows.

TABLE
II OBSERVATION GUIDE
INDICATORS

Nº	Indicators	Unit of measurement
1	Percentage of money spent on indirect costs	Percentage(%)
2	Percentage of Front End Budget	Percentage(%)
3	Rate of valid images made x day	Images/Days
4	Percentage of images used	Percentage(%)

Previous projects from the same firm are shown in the following table. The findings will be averaged so that they may be compared to the post-test.

TABLE III
RESULT PRE-TEST OBSERVATION GUIDE

Start Date/Indicators	1	2	3	4
01/02/2015	10.00%	16.66%	0.2	55.55%
21/06/2021	2.50%	50.00%	0.13	14.28%
01/06/2022	1.11%	16.66%	0.2	33.33%
12/12/2022	4.38%	25.00%	0.2	28.57%
01/01/2023	3.44%	18.75%	0.5	33.33%
Average	4.29%	25.41%	0.24	33.01%

Presently underway projects make up the post-test. The measurement would become quite old if lead times were longer to cover more projects, since the Stable Diffusion technology is updated so quickly. This table displays the outcomes of the post-test.

TABLE IV

RESULT POST-TEST OBSERVATION GUIDE

Start Date/Indicators	1	2	3	4
08/05/2023	2.37%	2.70%	1.5	7.47%

Indirect expenses as a percentage of total expenditure: 2.37 percent of the budget went into paying for the Google Drive and Google Collab services so that the team may use the tool throughout the technological deployment. On the other hand, prior initiatives necessitated the purchase of online tools, leading to expenditures that averaged 4.29 percent of the budget. This metric is reduced by 1.92 percent as a result of the implementation.

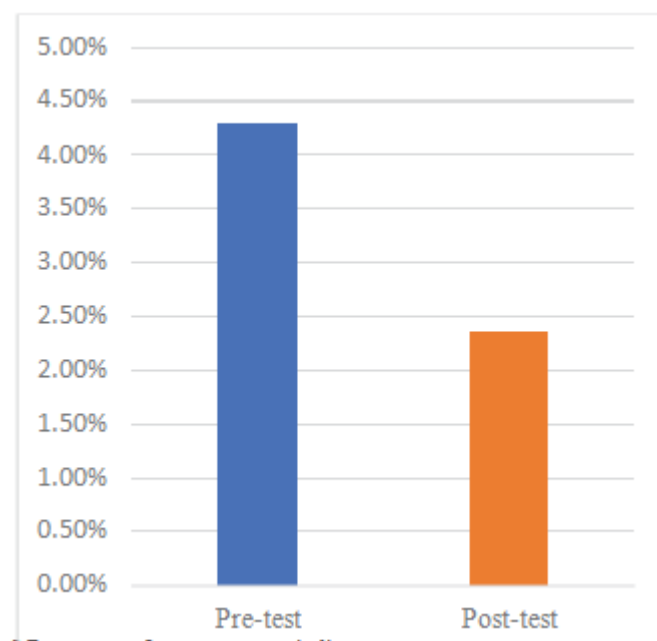


Fig.5. Percentage of money spent on indirect costs

Front-End Budget Percentage: On previous projects, we had to hire a draftsman to do the requisite drawings, which may add another minimum salary. The additional payment to the workers for the time in months brought the overall cost of the project to up to 25.41%. Time saved and money saved were the two most important factors in the post-test. Hence, it barely covered 2.70 percent of the whole expenditure. A 22.71% decrease in this metric has been achieved as a result of the implementation.

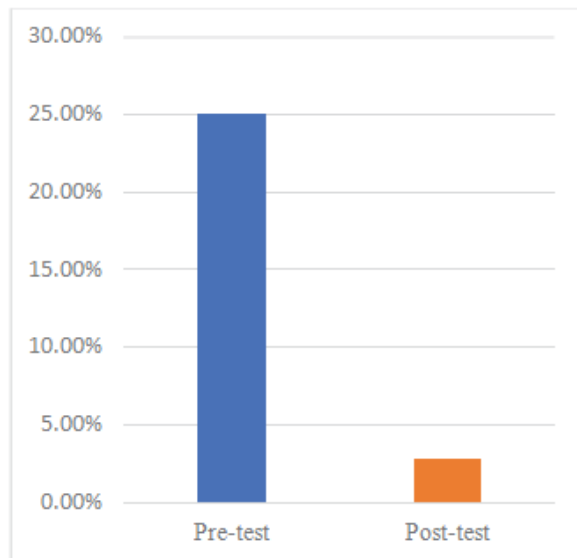


Fig.6.PercentageofFrontEndBudget

The rate of legitimate photographs created each day is determined by selecting just the highest quality images from all those generated. There was a backlog in the pre-test since photos sometimes take days to supply and were occasionally rejected. A pace of advancement of fewer than one picture each day is caused by this condition. The end result was an average of 0.24 photographs, being on ex day work. There were a total of 21 photos collected throughout the two-week production run in the post-test, with an average of 1.5 photographs each day. This is a speed boost. The implementation results in a 609.76% decrease.

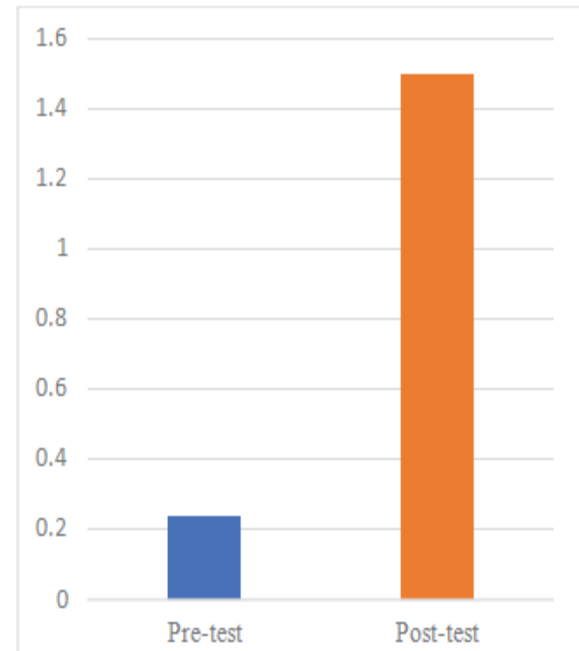


Fig.7.Rateofvalidimagesmadexday

Percentage of photographs utilized: According to the preceding indication, a small number of images were discarded during the pre-test. On average, only 33.01% of the images that were created were actually used, which is a significant amount of wasted resources. While the post-test uses a much less percentage of photos (7.47 percent), the creation speed is so high that the waste of images does not result in any resource waste. This statistic shows a rise of 25.54% due to the implementation.

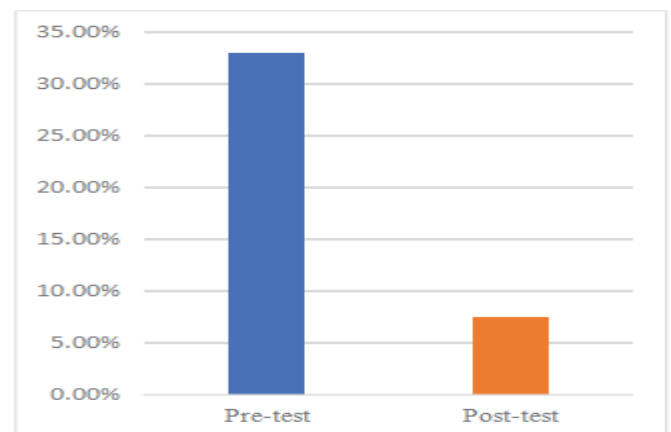


Fig.8.Percentage of imagesused

Conclusions

Finally, the data acquired from the use of text-to-image generators shows that the efficiency of user interface creation has been much enhanced. The development team really benefited from the creation of the requisite pictures utilizing Stable Diffusion technology.

The data analysis showed that the time and money saved were a direct result of the rapidity with which images could be generated and the fact that any team member could carry out this procedure. In most cases, a text-to-image converter that can produce pictures with the necessary styles makes the process of developing a user interface for a website much faster. What's more, you should think about how fast picture creation technology is progressing; in the future, it could be more efficient to use a service than of making your own text-to-image generator. HyperDreamBooth[14] and other modern technologies may help us improve our diffusion models; we can also use metrics like FID SCORE [8] to evaluate the quality of our images.

References

- [1] Beaird, J., Walker, A., George, J., O'Reilly for Higher Education (Firm), & Safari, an O.M. Company. (n.d.). *The Principles of Beautiful Web Design, 4th Edition*. 282. Retrieved October 21, 2023, from https://books.google.com/books/about/The_Principles_of_Beautiful_Web_Design.html?hl=es&id=BczDEAAAQBAJ
- [2] Black, K., Janner, M., Du, Y., Kostrikov, I., & Levine, S. (2023). *Training Diffusion Models with Reinforcement Learning*. <https://arxiv.org/abs/2305.13301v3>
- [3] Decreto Legislativo N.º 822 - Normas y documentos legales - Presidencia del Consejo de Ministros - Plataforma del Estado Peruano. (n.d.). Retrieved October 21, 2023, from <https://www.gob.pe/institucion/pcm/normas-legales/1670023-822>
- [4] Dhariwal, P., Openai, & Nichol, A. (2021). Diffusion Models Beat GANs on Image Synthesis. *Advances in Neural Information Processing Systems*, 34, 8780–8794.
- [5] Gal, R., Arar, M., Atzmon, Y., Bermano, A. H., Chechik, G., & Cohen-Or, D. (2023). Encoder-based Domain Tuning for Fast Personalization of Text-to-Image Models. *ACM Transactions on Graphics*, 42(4). <https://doi.org/10.1145/3592133>
- [6] Ghazanfari, S., Garg, S., Krishnamurthy, P., Khorrami, F., & Araujo, A. (2023). *R-LPIPS: An Adversarially Robust Perceptual Similarity Metric*. <https://arxiv.org/abs/2307.15157v2>
- [7] Heusel, M., Ramsauer, H., Unterthiner, T., Nessler, B., & Hochreiter, S. (2017). GANs Trained by a Two Time-Scale Update Rule Converge to a Local Nash Equilibrium. *Advances in Neural Information Processing Systems*, 30.
- [8] Rao, S. Govinda, R. Ram Babu, BS Anil Kumar, V. Srinivas, and P. Varaprasada Rao. "Detection of traffic congestion from surveillance videos using machine learning techniques." In *2022 Sixth International Conference on I-SMAC (IoT in Social, Mobile, Analytics and Cloud) (I-SMAC)*, pp. 572–579. IEEE, 2022.
- [9] Karras, T., Aila, T., Laine, S., & Lehtinen, J. (2017). Progressive Growing of GANs for Improved Quality, Stability, and Variation. *6th International Conference on Learning Representations, ICLR 2018 - Conference Track Proceedings*. <https://arxiv.org/abs/1710.10196v3>
- [10] Nightingale, S. J., & Farid, H. (2022). AI-synthesized faces are indistinguishable from real faces and more trustworthy. *Proceedings of the National Academy of Sciences*, 119(8), e2120481119. <https://doi.org/10.1073/PNAS.2120481119>
- [11] Oppenlaender, J. (2022a). *A Taxonomy of Prompt Modifiers for Text-To-Image Generation*. <https://arxiv.org/abs/2204.13988v3>
- [12] Oppenlaender, J. (2022b). The Creativity of Text-to-Image Generation. *ACM International Conference Proceeding Series*, 11, 192–202. <https://doi.org/10.1145/3569219.3569352>
- [13] Ruiz, N., Li, Y., Jampani, V., Pritch, Y., Rubinstein, M., & Aberman, K. (2023). *DreamBooth: Fine Tuning Text-to-Image Diffusion Models for Subject-Driven Generation* (pp. 22500–22510). <https://dreambooth.github.io/>
- [14] Ruiz, N., Li, Y., Jampani, V., Wei, W., Hou, T., Pritch, Y., Wadhwa, N., Rubinstein, M., & Aberman, K. (2023). *HyperDreamBooth: HyperNetworks for Fast Personalization of Text-to-Image Models*. <https://arxiv.org/abs/2307.06949v1>
- [15] *The future of jobs report 2020 | VOCEDplus, the international tertiary education and research database*. (n.d.). Retrieved October 25, 2023, from <https://www.voced.edu.au/content/ngv:88417>
- [16] Zhang, C., Zhang, C., Zhang, M., & Kweon, I. S. (2023). *Text-to-image Diffusion Models in Generative AI: A Survey*. <https://arxiv.org/abs/2303.07909v2>