

REFERENCES

- [1] X. Wang, Y. Lu, Y. Xuehu, and L. Yu, “Wet paper coding-based deep neural network watermarking,” *Sensors*, vol. 22, no. 9, p. 3489, May 2022.
- [2] B. Joe, Y. Park, J. Hamm, I. Shin, and J. Lee, “Exploiting missing value patterns for a backdoor attack on machine learning models of electronic health records: Development and validation study,” *JMIR Med Inform*, vol. 10, no. 8, p. e38440, Aug. 19 2022.
- [3] J. Souza, L. Oliveira, Y. Gumiel, D. Carvalho, and C. Moro, *Exploiting Siamese Neural Networks on Short Text Similarity Tasks for Multiple Domains and Languages*, 02 2020, pp. 357–367.
- [4] M. Goldblum, D. Tsipras, C. Xie, X. Chen, A. Schwarzschild, D. Song, A. Madry, B. Li, and T. Goldstein, “Dataset security for machine learning: Data poisoning, backdoor attacks, and defenses,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 45, no. 2, pp. 1563–1580, 2023.
- [5] U. Yasuke, Y. Nagai, S. Sakazawa, and S. Satoh, “Embedding watermarks into deep neural networks,” in *Proceedings of the 2017 ACM on International Conference on Multimedia Retrieval*. New York, NY, USA: ACM, Jun. 2017, pp. 269–277.
- [6] B. D. Rouhani, H. Chen, and F. Koushanfar, “Deepsigns: An end-to-end watermarking framework for protecting the ownership of deep neural networks,” in *Proceedings of the Twenty-Fourth International Conference on Architectural Support for Programming Languages and Operating Systems*, New York, Apr. 2019.
- [7] L. Fan, K. W. Ng, and C. C. Seng, “Rethinking deep neural network ownership verification: Embedding passports to defeat ambiguity attacks,” in *Annual Conference on Neural Information Processing Systems*, Canada, Dec. 2019.

- [8] L. Feng and X. Zhang, “Watermarking neural network with compensation mechanism,” 2020, pp. 363–375.
- [9] J. Zhang, D. Chen, J. Liao, H. Fang, W. Zhang, W. Zhou, H. Cui, and N. Yu, “Model watermarking for image processing networks,” *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 34, no. 07, pp. 12 805–12 812, Apr. 2020.
- [10] F. Cai, “‘hack’ the neural network: Tencent reveals new ai attack methods,” Apr. 17 2023, [Online]. Available: <https://www.infoq.cn/article/9X9srGHSZpG9hC1MF06s>.
- [11] Z. Wang, C. Liu, and X. Cui, “Evilmodel: Hiding malware inside of neural network models,” in *2021 IEEE Symposium on Computers and Communications (ISCC)*. IEEE, Sep. 2021, pp. 1–7.
- [12] A. Cohen, A. Cohen, and N. Nissim, “Assaf: Advanced and slim steganalysis detection framework for jpeg images based on deep convolutional denoising autoencoder and siamese networks,” *Neural Networks*, vol. 131, pp. 64–77, Nov. 2020.
- [13] H. Li, J. Wang, N. Xiong, Y. Zhang, A. V. Vasilakos, and X. Luo, “A siamese inverted residuals network image steganalysis scheme based on deep learning,” *ACM Transactions on Multimedia Computing, Communications, and Applications*, Jul. 2023.
- [14] H. Tao, L. Chongmin, J. M. Zain, and A. N. Abdalla, “Robust image watermarking theories and techniques: A review,” *Journal of Applied Research and Technology*, vol. 12, no. 1, pp. 122–138, Feb. 2014.
- [15] A. Tirkel and T. Hall, “A unique watermark for every image,” *IEEE MultiMedia*, vol. 8, no. 4, pp. 30–37, Oct.-Dec. 2001.
- [16] I. H. Sharker, “Deep learning: A comprehensive overview on techniques, taxonomy, applications, and research directions,” vol. 2, no. 420. Springer Nature Singapore Pte Ltd, August 2021.
- [17] C. Gupta and P. Thakral, “Ascii conversion based two keys v4s scheme for encryption and decryption — a four step approach,” in *2017 8th International Conference on Computing, Communication and Networking Technologies (ICCCNT)*. Kurukshetra, India: IEEE, Dec 2019, pp. 165–168.

- [18] S. Boujerfaoui, R. Riad, H. Douzi, F. Ros, and R. Harba, “Image watermarking between conventional and learning-based techniques: A literature review,” vol. 12, 2023, [Online]. Available: <https://www.mdpi.com/2079-9292/12/1/74f>.