
Promoting User Data Autonomy During the Dissolution of a Monopolistic Firm

Rushabh Solanki

University of Waterloo, Vector Institute
r7solanki@uwaterloo.ca

Elliot Creager

University of Waterloo, Vector Institute
creager@uwaterloo.ca

Abstract

The deployment of AI in consumer products is currently focused on the use of so-called foundation models, large neural networks pre-trained on massive corpora of digital records. This emphasis on scaling up datasets and pre-training computation raises the risk of further consolidating the industry, and enabling monopolistic (or oligopolistic) behavior. Judges and regulators seeking to improve market competition may employ various remedies. This paper explores *dissolution*—the breaking up of a monopolistic entity into smaller firms—as one such remedy, focusing in particular on the technical challenges and opportunities involved in the breaking up of large models and datasets. We show how the framework of Conscious Data Contribution can enable user autonomy during under dissolution. Through a simulation study, we explore how fine-tuning and the phenomenon of “catastrophic forgetting” could actually prove beneficial as a type of machine unlearning that allows users to specify which data they want used for what purposes.

1 Introduction

In August 2024, the D.C. District Court delivered a landmark ruling, finding that Google had engaged in anti-competitive practices in its search engine business, in violation of the Sherman Act [Goldstein, 2024]. This ruling signals an increasing willingness by the US Department of Justice to challenge the concentration of power in the technology sector. It also coincides with the waxing influence of AI-powered software on consumer markets, where the increasing reliance on so-called foundation models could further exacerbate competition issues [Bommasani et al., 2022, Brynjolfsson, 2022].

The Google anti-trust case is far from resolved: pending an ongoing appeal, it will conclude with judges suggesting practical “remedies,” such as fines or restrictions on business practices, to encourage market competition going forward. This paper explores the hypothetical use of a particularly stern remedy to address monopolistic behavior by an AI firm: the *dissolution* of a offending firm into smaller successor companies. While dissolution was employed against 20th century monopolists such as Standard Oil and AT&T, it has not been used since the advent of the digital economy.

The potential for dissolution of AI firms is muddled by mirky technical questions such as how best to “break up” the large datasets and models of a monopolist firm. We note that the dissolution of an AI firm (or any serious corporate restructuring imposed by a regulator) presents the chance to rethink standard data curation practices that emphasize scale over consent [Hanna and Park, 2020, Andreotta et al., 2022]. We advocate for the importance of data autonomy in the dissolution process, drawing on the framework of Conscious Data Contribution [Vincent and Hecht, 2021] to ensure that users can have say in how their data, previously owned by the monopolist, is used by the successor firms going forward. Using a simple and stylized simulation study, we explore the possibility of realizing dissolution that respects user data autonomy using standard fine-tuning practices.

2 Background

Data Autonomy Our explorations into regulating AI during dissolution emphasize user data autonomy, and are inspired by Data Leverage [Vincent et al., 2021]. This framework was originally introduced to broadly describe how users can exert indirect influence over AI systems by modifying their data. We focus in particular on Conscious Data Contribution (CDC) Vincent and Hecht [2021]—a particular data lever whereby users choose to share their data with smaller firms—as a way to ensure users’ data consent is respected when dissolution is carried out by regulators.

Model Adaptation Prior work on CDC focused on how users’ data contributions have a disproportionately positive impact for smaller firms who train their AI models from scratch [Vincent and Hecht, 2021]. Because we focus instead on applying the CDC principle to *pre-trained* foundation models owned by the monopolistic firm, our work also relates to other approaches seeking to adapt large models. Model merging seeks to combine multiple pre-trained models directly in their parameter space, thus avoiding the need to retrain from scratch when the source dataset changes [Ainsworth et al., 2023, Matena and Raffel, 2022]. By analogy, this might represent the type of technique suitable for the merger of two firms, rather than the dissolution of a single firm.

Removing the influence of specific data points from a model is the focus of Machine Unlearning [Bourtoule et al., 2021], motivated in part by the “right to be forgotten” principle within the EU GDPR [Wachter et al., 2018]. Unlearning is difficult to implement due to an inherent tension between forgetting the target data point and preserving the model’s fit to the remaining distribution [Kurmanji et al., 2023]; accordingly, contemporary unlearning schemes often still rely on retraining from scratch. On the other hand, there is evidence that fine-tuning a model on a new task induces *catastrophic forgetting*, a type of natural unlearning resulting from the collapse of the model’s representation to retain only features relevant to the new task [French, 1999]. Below we discuss how this phenomenon makes fine-tuning a simple and appealing approach for implementing CDC under dissolution. Moreover, existing machine unlearning methods are primarily aimed at removing specified objects, patterns or stylistic elements, rather than addressing varied diverse data as in our case – making fine-tuning an even more attractive option.

AI and Market Competition The heavy emphasis on scaling up data and compute power to build modern AI products raises concerns over the ability for smaller firms to compete [Brynjolfsson, 2022]. Luitse and Denkena [2021] analyze the political economy of the transformer architecture, including the tendency of transformer-based LLMs to consolidate power within the tech industry. Burkhardt and Rieder [2024] discuss how the focus of prompting as an interface for foundation models further supports the recent *platformization* of the software industry [Helmond, 2015], possibly strengthening the market hold of larger firms. Recently, simulation studies have suggested that automated AI decision making in the domain of pricing algorithms can unexpectedly lead to collusive behavior, both for models based on Q-learning [Calvano et al., 2020] and LLMs [Fish et al., 2024].

3 Conscious Data Contribution Under Firm Dissolution

We consider a scenario in which a monopolistic firm F is dissolved into k successor companies $\{S_1, \dots, S_k\}$. Before dissolution, F owns a dataset $U \in \mathbb{R}^{N \times d}$ of d -dimensional data records from each of its N users. We denote $u_i \subset U$ as the data about the i -th user within the dataset, and further assume that this user data is a vector $u_i \in \mathbb{R}^d$ built by concatenating J task-specific user data vectors:

$$u_i = \begin{pmatrix} u_i^{(1)} \\ u_i^{(2)} \\ \vdots \\ u_i^{(J)} \end{pmatrix},$$

where $u_i^{(j)} \in \mathbb{R}^{d_j}$ and $\sum_j d_j = d$. For example, $u_i^{(1)}$ might correspond to location data collected about user u_i , while $u_i^{(2)}$ might correspond to that user’s music recommendation data.

As the monopolistic firm F is dissolved into successor firms $\{S_1, \dots, S_k\}$, the key regulatory question is what will happen to the dataset U , and any models derived from it. To allow successor

firms to build expressive models in a sample efficient manor, while preserving data autonomy, we posit an avenue for Conscious Data Contribution [Vincent and Hecht, 2021] based on fine-tuning. We introduce a binary matrix for each user $C(u_i) \in \{0, 1\}^{K \times J}$, whose entries correspond to that user consciously contributing task-specific data to particular successor firms. In other words, we have

$$[C(u_i)]_{j,k} = \begin{cases} 1 & \text{if user } i \text{ contributes their task-}j \text{ data to successor } k, \\ 0 & \text{else.} \end{cases}$$

For example, if S_k represents the successor company taking on the music streaming business from F , then user U_i may wish to set $[C(u_i)]_{1,k} = 1$ in order to contribute their music listening history to that successor company, while setting $[C(u_i)]_{2,k} = 0$ to preclude the use of location data by S_k .¹ Thus the i -th user consciously contributes the following data vector to the k -th successor firm:

$$u_i^{\text{CDC}-S_k} = \begin{pmatrix} u_i^{(1)} \cdot [C(u_i)]_{1,k} \\ u_i^{(2)} \cdot [C(u_i)]_{2,k} \\ \vdots \\ u_i^{(J)} \cdot [C(u_i)]_{J,k} \end{pmatrix}.$$

Accordingly, S_k receives $U^{\text{CDC}-S_k} = [u_1^{\text{CDC}-S_k}, \dots, u_N^{\text{CDC}-S_k}]^T$, a dataset of consciously contributed data for the purposes of training their models, with non-consented entries zeroed out.

Through $C(u_i)$, each user has the opportunity to consciously contribute portions of their data (organized by task-appropriateness) to each successor company. In order for this scheme to promote user data autonomy within a competitive market, regulators would need to establish a framework for ensuring that users' wishes are respected by the firms. In particular, this introduces the following two regulatory challenges:

1. Firms should only gain access to data permitted under $C(u_i)$;
2. When $[C(u_i)]_{j,k} = 0$, reasonable steps should be taken to ensure that information about user i is not used by firm k to solve task j .

Satisfying these requirements is complicated for firms that rely on foundation models: the successor companies may not have sufficient data within $U^{\text{CDC}-S_k}$ to train a performant model from scratch (impeding (1)), and the original foundation model of F may not be directly interpretable or explainable (impeding (2)). To take a first step towards realizing a CDC approach for foundation models, we next examine a stylized setting focused on the dynamics of fine-tuning smaller neural networks, where we find the common phenomenon of catastrophic forgetting promotes user data autonomy while allowing the firms a sample-efficient way to comply with these regulatory challenges.

4 Simulation Studies

We conduct proof-of-concept simulations across two data modalities. Upon dissolution, we suppose a successor firm S_k can preserve only a portion of the original data obtained through Conscious Data Contribution, that is $U^{\text{CDC}-S_k}$, but may fine tune some model previously trained on U . For simplicity we assume $k = 1$, then denote U^{CDC} as the available data to the single successor firm and $\neg U^{\text{CDC}} := U - U^{\text{CDC}}$ as the data the successor needs to unlearn.

4.1 Datasets

Image generation We begin with a simple setting where the firm's data U was used to train a class-conditional diffusion model. We carry out the same experimental procedure first on MNIST and then on CIFAR-10. To keep it simple, we suppose the sole successor S receives the data U^{CDC} comprising three of the ten classes. While the original diffusion model can generate images that match its conditioned class, if user autonomy is respected then the model fine-tuned with U^{CDC} will produce samples according to the new distribution, even if conditioned on classes outside of U^{CDC} .

¹Apple's "ask app not to track" iOS feature [Chen, 2021] can be seen as a realization of this type of CDC where $\{S_1, \dots, S_K\}$ represent (possibly competing) apps rather than successor companies.

Table 1: Main Results for Diffusion Experiments

	MNIST @ 100 epochs		CIFAR10 @ 100 epochs		
	Forget Rate on $\neg U^{\text{CDC}}$ ($\uparrow$)	Retain Rate on U^{CDC} ($\uparrow$)	Forget Rate on $\neg U^{\text{CDC}}$ ($\uparrow$)	Retain Rate on U^{CDC} ($\uparrow$)	FID scores ($\downarrow$)
Original	0.016	0.97	0.381	0.544	71.013
Retraining	0.962	0.955	0.918	0.266	132.00
Fine-tuning	0.955	0.914	0.844	0.536	57.550

Text classification We train a DistillBERT model to predict which sets of skills are described in a dataset of 30,000 resumes from a major gig labor platform [Jiechiew and Tsopze, 2021]. Importantly, this is multi-label classification as skills may not be unique. We consider a single employer F (hiring from all skills) that is dissolved into one successor S hiring from a single skill set, so U^{CDC} as all resumes with “administrative” skills and $\neg U^{\text{CDC}}$ as all resumes with “developer” skills. As before, S may optionally fine-tune on U^{CDC} , starting from F ’s model previously trained on U .

Table 2: Main Results for Resume Dataset (@ 25 epochs)

	Forget Rate (F1) on $\neg U^{\text{CDC}}$ ($\uparrow$)	Retain Rate (F1) on U^{CDC} ($\uparrow$)
Original	-	0.898
Retraining (pre-trained)	0.937	0.901
Retraining (scratch)	0.946	0.805
Fine-tuning	0.917	0.964

4.2 Results

The successor S is satisfied if they have a sample efficient way to produce a model that generalizes well to unseen data from $\mathbb{P}(U^{\text{CDC}})$, while the users and regulator are satisfied if this model generalizes poorly to $\mathbb{P}(\neg U^{\text{CDC}})$. Thus we measure the “retain rate” on held out data from U^{CDC} and “forget rate” on data from $\neg U^{\text{CDC}}$, where both metrics are “the higher the better”. For the image generation task, given a snapshot of a trained diffusion model, 10K images were sampled comprising 1000 samples per class. The forget rate was measured for 7K samples corresponding to $\neg U^{\text{CDC}}$ classes using 1 minus accuracy, and we used accuracy for the remaining 3K samples. For the resume dataset, we instead use $F1$ score of the DistillBERT classifier evaluated on successor data for retain rate, and 1 minus $F1$ score on non-successor data for forget rate. See Appendix A for details.

Our proposed approach is to fine-tune on U^{CDC} starting from a model pre-trained on U , which we compare with two baselines: training on U^{CDC} from scratch, and fine-tuning from a domain-agnostic pre-training initialization (e.g. huggingface weights). While retraining from scratch works well for small datasets like MNIST, it suffers from poor utility on larger datasets, where we find that fine-tuning more favorably balances successor utility with user data autonomy (Table 1 and Table 2). We also attempted gradient ascent-based unlearning, but it became unstable after a few epochs when applied to unlearn large data. These experiments are performed under a simple setting, whereas the process of real-world dissolution will introduce multiple complexities in the distribution of U^{CDC} . While limited in scope, these results show that fine-tuning provides a natural unlearning effect for users’ data that was not consciously contributed to the successor, suggesting that promoting users’ data autonomy within a complex regulatory schema may be feasible through simple means.

5 Conclusion

Dissolving a large firm would be a substantial undertaking, which would almost certainly involve political and legal friction. For example, in the 1990s monopoly case brought against Microsoft, a dissolution remedy initially explored by the judiciary was ultimately watered down, and the actual remedies realized by the courts were less disruptive to Microsoft’s business [Goldstein, 2024]. Nevertheless, we feel that the emerging field of AI regulation should explore frameworks that account for all possible remedies available to counteract monopolistic behavior. In this paper we have taken one step towards exploring a regulatory framework for dissolution through the lens of Conscious Data Contribution [Vincent and Hecht, 2021], with the aims of ensuring that when a monopolistic firm is broken up, its successor firms only use data and models with explicit user consent. Through a simple proof-of-concept simulation study we demonstrated that the phenomenon of catastrophic forgetting [French, 1999] during fine-tuning promotes user autonomy in this case.

Acknowledgments

We are grateful to Nicholas Vincent and the anonymous reviewers for their constructive feedback. Resources used in preparing this research were provided, in part, by the Province of Ontario, the Government of Canada through CIFAR, and companies sponsoring the Vector Institute www.vectorinstitute.ai/partnerships/.

References

- Samuel K. Ainsworth, Jonathan Hayase, and Siddhartha Srinivasa. Git Re-Basin: Merging Models modulo Permutation Symmetries, March 2023. URL <http://arxiv.org/abs/2209.04836>. arXiv:2209.04836 [cs].
- Adam J. Andreotta, Nin Kirkham, and Marco Rizzi. AI, big data, and the future of consent. *AI & SOCIETY*, 37(4):1715–1728, December 2022. ISSN 0951-5666, 1435-5655. doi: 10.1007/s00146-021-01262-5. URL <https://link.springer.com/10.1007/s00146-021-01262-5>.
- Rishi Bommasani, Drew A. Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S. Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, Erik Brynjolfsson, Shyamal Buch, Dallas Card, Rodrigo Castellon, Niladri Chatterji, Annie Chen, Kathleen Creel, Jared Quincy Davis, Dora Demszky, Chris Donahue, Moussa Doumbouya, Esin Durmus, Stefano Ermon, John Etchemendy, Kavin Ethayarajh, Li Fei-Fei, Chelsea Finn, Trevor Gale, Lauren Gillespie, Karan Goel, Noah Goodman, Shelby Grossman, Neel Guha, Tatsunori Hashimoto, Peter Henderson, John Hewitt, Daniel E. Ho, Jenny Hong, Kyle Hsu, Jing Huang, Thomas Icard, Saahil Jain, Dan Jurafsky, Pratyusha Kalluri, Siddharth Karamcheti, Geoff Keeling, Fereshte Khani, Omar Khattab, Pang Wei Koh, Mark Krass, Ranjay Krishna, Rohith Kuditipudi, Ananya Kumar, Faisal Ladhak, Mina Lee, Tony Lee, Jure Leskovec, Isabelle Levent, Xiang Lisa Li, Xuechen Li, Tengyu Ma, Ali Malik, Christopher D. Manning, Suvir Mirchandani, Eric Mitchell, Zanele Munyikwa, Suraj Nair, Avani Narayan, Deepak Narayanan, Ben Newman, Allen Nie, Juan Carlos Niebles, Hamed Nilforoshan, Julian Nyarko, Giray Ogut, Laurel Orr, Isabel Papadimitriou, Joon Sung Park, Chris Piech, Eva Portelance, Christopher Potts, Aditi Raghunathan, Rob Reich, Hongyu Ren, Frieda Rong, Yusuf Roohani, Camilo Ruiz, Jack Ryan, Christopher Ré, Dorsa Sadigh, Shiori Sagawa, Keshav Santhanam, Andy Shih, Krishnan Srinivasan, Alex Tamkin, Rohan Taori, Armin W. Thomas, Florian Tramèr, Rose E. Wang, William Wang, Bohan Wu, Jiajun Wu, Yuhuai Wu, Sang Michael Xie, Michihiro Yasunaga, Jiaxuan You, Matei Zaharia, Michael Zhang, Tianyi Zhang, Xikun Zhang, Yuhui Zhang, Lucia Zheng, Kaitlyn Zhou, and Percy Liang. On the Opportunities and Risks of Foundation Models, July 2022. URL <http://arxiv.org/abs/2108.07258>. arXiv:2108.07258 [cs].
- Lucas Bourtole, Varun Chandrasekaran, Christopher A. Choquette-Choo, Hengrui Jia, Adelin Travers, Baiwu Zhang, David Lie, and Nicolas Papernot. Machine Unlearning. In *2021 IEEE Symposium on Security and Privacy (SP)*, pages 141–159, May 2021. doi: 10.1109/SP40001.2021.00019. URL <https://ieeexplore.ieee.org/document/9519428/?arnumber=9519428>. ISSN: 2375-1207.
- Erik Brynjolfsson. The Turing Trap: The Promise & Peril of Human-Like Artificial Intelligence. *Daedalus*, 151(2):272–287, May 2022. ISSN 0011-5266, 1548-6192. doi: 10.1162/daed_a_01915. URL <https://direct.mit.edu/daed/article/151/2/272/110622/The-Turing-Trap-The-Promise-and-Peril-of-Human>.
- Sarah Burkhardt and Bernhard Rieder. Foundation models are platform models: Prompting and the political economy of AI. *Big Data & Society*, 11(2):20539517241247839, 2024. doi: 10.1177/20539517241247839. URL <https://doi.org/10.1177/20539517241247839>. _eprint: <https://doi.org/10.1177/20539517241247839>.
- Emilio Calvano, Giacomo Calzolari, Vincenzo Denicolò, and Sergio Pastorello. Artificial Intelligence, Algorithmic Pricing, and Collusion. *The American Economic Review*, 110(10):3267–3297, 2020. ISSN 0002-8282. URL <https://www.jstor.org/stable/26966472>. Publisher: American Economic Association.
- Brian X. Chen. To Be Tracked or Not? Apple Is Now Giving Us the Choice. *The New York Times*, April 2021. URL <https://www.nytimes.com/2021/04/26/technology/personaltech/apple-app-tracking-transparency.html>. Publication Title: The New York Times.
- Sara Fish, Yannai A. Gonczarowski, and Ran I. Shorrer. Algorithmic Collusion by Large Language Models, March 2024. URL <http://arxiv.org/abs/2404.00806>. arXiv:2404.00806 [cs, econ, q-fin].
- Robert M. French. Catastrophic forgetting in connectionist networks. *Trends in Cognitive Sciences*, 3(4):128–135, April 1999. ISSN 1364-6613. doi: 10.1016/S1364-6613(99)01294-2. URL <https://www.sciencedirect.com/science/article/pii/S1364661399012942>.

- Luke Goldstein. Will Google’s Monopoly Be Vanquished? - The American Prospect. *The American Prospect*, August 2024. URL <https://prospect.org/justice/2024-08-09-will-googles-monopoly-be-vanquished/>.
- Alex Hanna and Tina M. Park. Against Scale: Provocations and Resistances to Scale Thinking, November 2020. URL <http://arxiv.org/abs/2010.08850>. arXiv:2010.08850 [cs].
- Anne Helmond. The Platformization of the Web: Making Web Data Platform Ready. *Social Media + Society*, 1(2):2056305115603080, July 2015. ISSN 2056-3051. doi: 10.1177/2056305115603080. URL <https://doi.org/10.1177/2056305115603080>. Publisher: SAGE Publications Ltd.
- Kameni Florentin Flambeau Jiechieu and Norbert Tsopze. Skills prediction based on multi-label resume classification using cnn with model predictions explanation. *Neural Computing and Applications*, 33(10):5069–5087, 2021.
- Meghdad Kurmanji, Peter Triantafillou, Jamie Hayes, and Eleni Triantafillou. Towards Unbounded Machine Unlearning, September 2023. URL <http://arxiv.org/abs/2302.09880>. arXiv:2302.09880 [cs].
- Dieuwertje Luitse and Wiebke Denkena. The great Transformer: Examining the role of large language models in the political economy of AI. *Big Data & Society*, 8(2):205395172110477, July 2021. ISSN 2053-9517, 2053-9517. doi: 10.1177/20539517211047734. URL <http://journals.sagepub.com/doi/10.1177/20539517211047734>.
- Michael Matena and Colin Raffel. Merging Models with Fisher-Weighted Averaging. In *36th Conference on Neural Information Processing Systems (NeurIPS 2022)*, 2022.
- Victor Sanh, Lysandre Debut, Julien Chaumond, and Thomas Wolf. Distilbert, a distilled version of bert: smaller, faster, cheaper and lighter, 2020. URL <https://arxiv.org/abs/1910.01108>.
- Nicholas Vincent and Brent Hecht. Can "Conscious Data Contribution" Help Users to Exert "Data Leverage" Against Technology Companies? *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW1):1–23, April 2021. ISSN 2573-0142. doi: 10.1145/3449177. URL <https://dl.acm.org/doi/10.1145/3449177>.
- Nicholas Vincent, Hanlin Li, Nicole Tilly, Stevie Chancellor, and Brent Hecht. Data Leverage: A Framework for Empowering the Public in its Relationship with Technology Companies, February 2021. URL <http://arxiv.org/abs/2012.09995>. arXiv:2012.09995 [cs].
- Sandra Wachter, Brent Mittelstadt, and Chris Russell. Counterfactual Explanations without Opening the Black Box: Automated Decisions and the GDPR, March 2018. URL <http://arxiv.org/abs/1711.00399>. arXiv:1711.00399 [cs].

A Experimental Details

Image generation Our conditional diffusion model was implemented using the Hugging Face’s Diffusers library. We utilized the standard configuration, which uses the DDPM scheduler with 1000 time steps and a cosine variance schedule. In all the experiments, we have used the batch size of 128. For MNIST data, we used the UNet architecture, with approximately 1.71M trainable parameters, which was trained for 46.9K parameter updates. Whereas the model trained on the CIFAR10 dataset has around 6.34M parameters and was trained for 195.5K parameter updates. We chose the ADAM optimizer with a learning rate 0.001 and used Mean Squared Error (MSE) as the loss function for both datasets.

The MNIST classifier used to compute the forget rate was trained from scratch with PyTorch’s ResNet18 implementation for roughly 164K parameter updates. Meanwhile, the CIFAR10 classifier was fine-tuned from pre-trained ImageNet weights using a ResNet50 model, with around 39K parameter updates. Test set accuracies for both datasets are provided in Table 3.

Table 3: The test performance of both the classifiers

Dataset	Test Accuracy
MNIST	0.997
CIFAR10	0.969

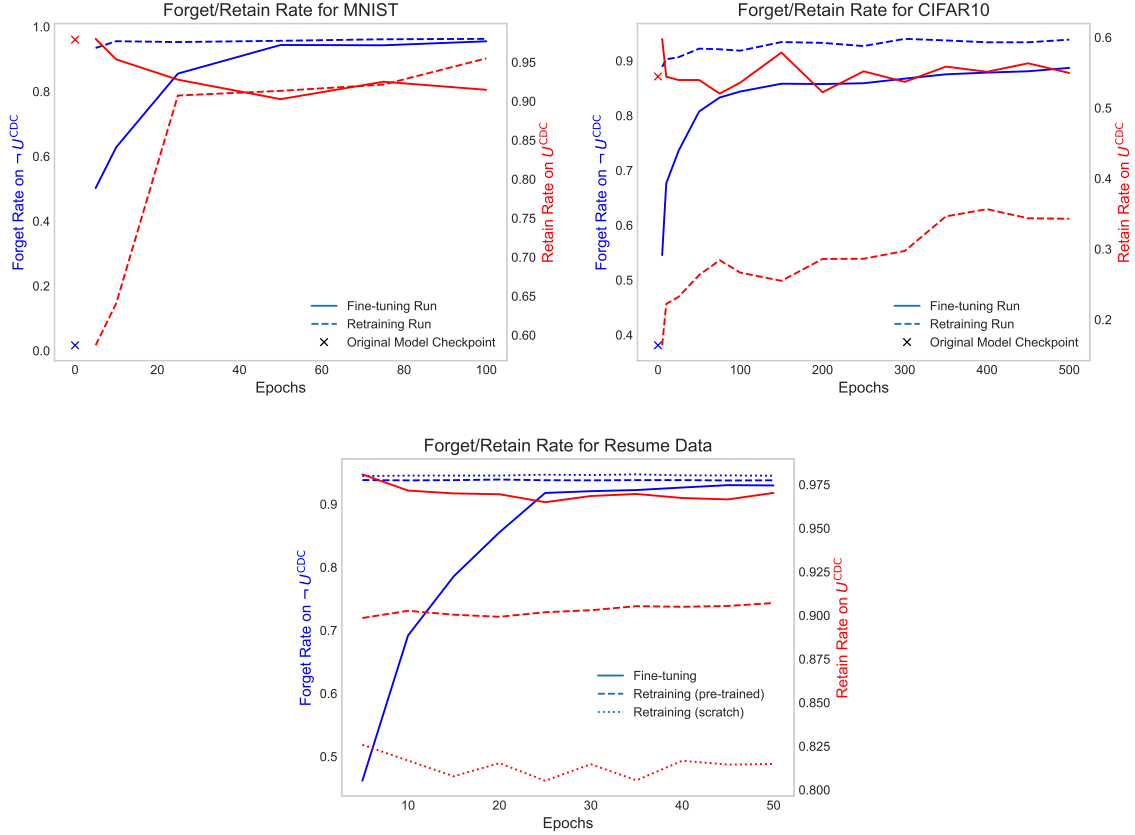


Figure 1: The plots illustrates the changes in forget and retain rates during the process of fine-tuning and retraining with consciously contributed data U^{CDC} to a successor S of a firm upon *dissolution*. Each plot displays the results for different datasets, with the top two for image generation and bottom one for text classification.

Text classification We used pre-trained `distilbert-base-uncased` from HuggingFace transformers library corresponding to the DistilBERT transformer [Sanh et al., 2020]. After preprocessing, the dataset contains a total of 29020 samples, split into 25000 for training and 4020 for testing. Following dissolution, U^{CDC} holds 14009 samples, divided into a training set of 11907 and a test set of 2102, while the remaining samples are treated as $\sim U^{CDC}$.

Figure 1 shows the increase in the forget rate as the fine-tuning progresses with almost 0% at the start of fine-tuning representing the base model. For MNIST data, at only the 25th epoch of fine-tuning, we get the forget rate of 91.6%. Attaining a comparable forget rate for the CIFAR10 involved a considerable number of epochs and plateaus of around 100 epochs. The forget rate for the resume data converges at around 25 epochs and overall fine-tuning generalizes well when comparing the retain rates.