

Implicit Bias of Gradient Descent under Feature-Mediated Spurious Correlations

Anonymous Author(s)

Affiliation

Address

email

Abstract

We study how *feature-mediated* spurious correlations affect the learning dynamics of gradient descent on linearly separable data. Unlike the widely studied *label-mediated* setting, where the spurious feature depends only on the class label and is conditionally independent of the causal features, feature-mediated correlations entangle spurious and causal signals at the instance level through group-specific operators, making them both harder to detect in practice and more complex to analyse theoretically. We first generalise the implicit-bias result of Soudry et al. (2018) to continuous distributions. Building on this, we derive closed-form group-wise classification error rates. In the general setting, each group’s error decay at a rate governed by its hard-margin γ_g . In the *isotropic regime*, we identify a single exponent α that captures the competition between the minority’s geometric margin advantage and the coupling induced by the spurious alignment operators. A phase transition occurs at $\alpha = 1$: below it, both groups decay at rate $\kappa_g^t / (\varepsilon_g z_t)$ and group proportion directly controls learning speed; above it, the minority escapes the ε -dependence entirely and converges at the faster polynomial rate $\ln(z_t)^{\alpha-1} z_t^{-\alpha \xi (1+\delta_{\text{maj}}) + \delta_{\text{min}}}$. These results formalise in closed form the widely observed phenomenon that group imbalance implicitly biases optimisation dynamics toward the majority, while revealing that the effect of imbalance can be entirely superseded by geometric margin advantages. Finally we establish a non-vanishing error at test time, under spurious correlation shift.

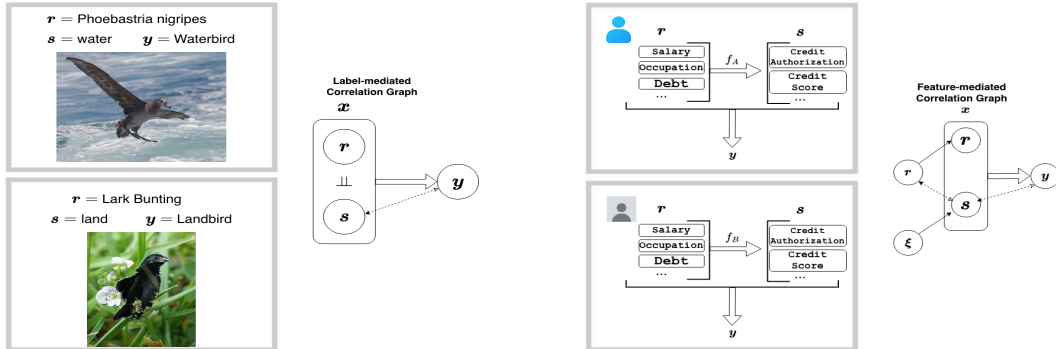


Figure 1: Two structural mechanisms behind spurious correlations. **Left (label-mediated):** the spurious feature s (background) depends on the label y but is conditionally independent of the causal features r given y and the group. **Right (feature-mediated):** the spurious features s (credit score etc.) depend directly on r through two group-specific operators f_A, f_B , entangling the two signals.

1 Introduction

Modern neural networks trained via Empirical Risk Minimisation (ERM) (Vapnik, 1999) are susceptible to *spurious correlations*: statistical associations between non-causal features and the target label that hold for most but not all of the training distribution. A classifier that relies on such associations performs well on the majority of the data but fails on minority subgroups where the correlation is reversed or absent, a failure mode documented across vision (Sagawa et al., 2020a; Idrissi et al., 2022), language and medical imaging applications.

These correlations can arise through different mechanisms (see Figure 1). In *label-mediated* spurious correlations, the spurious feature depends on the class label y but is conditionally independent of the causal features $\mathbf{r}$ given y and the group, as in the Waterbirds benchmark (Sagawa et al., 2020a) where the background scene correlates with the bird species through the labelling process. Because the conditional independence $\mathbf{s} \perp \mathbf{r} \mid (y, g)$ allows the spurious and causal features to be manipulated independently, label-mediated correlations are comparatively straightforward to study, and the vast majority of both empirical benchmarks and theoretical analyses have focused on this case.

In *feature-mediated* spurious correlations, by contrast, the spurious feature $\mathbf{s}$ depends on the *specific value* of $\mathbf{r}$, not merely on y . The causal and spurious signals are entangled at the instance level, making the spurious component harder to isolate experimentally and substantially more complex to analyse theoretically. To our knowledge, no prior work has specifically provided a theoretical analysis of learning dynamics under feature-mediated spurious correlations.

This work. We provide, to the best of our knowledge, the first closed-form characterisation of how features-mediated spurious correlations shape the *per-group* learning speed of gradient descent. Our analysis operates at the population level: we study gradient descent on the expected risk under a neural net learned representation distribution rather than on an empirical average over a finite training set. This yields closed-form dynamics free of finite-sample fluctuations, while remaining predictive of the behaviour observed in practice. Our contributions can be summarized this way:

- We generalise Soudry et al. (2018, *Theorem 9*) to continuous, compactly supported distributions (Proposition 5.4). The iterates still converge in direction to the max-margin classifier $\hat{w}$, but the residual grows as $-\hat{w} \log \log t + O(1)$, in contrast to the $O(1)$ residual of the finite-sample case.
- We investigate how features-mediated spurious correlation affects neural networks learned representation $\Phi(x)$ (Sections 4 and B).
- We derive explicit decay rates for the classification error on each group. In the general case, each group’s error decays as $\Theta(z_t^{-\gamma_g} (\ln z_t)^{\gamma_g + \delta_g + \tilde{\delta}_g - 1})$, governed by its hard-margin γ_g (Theorem 5.1). In the isotropic regime, we obtain sharp asymptotics controlled by a single exponent α that encodes the competition between geometric margins and spurious alignment operators (Theorem 5.3).
- We identify a phase transition at $\alpha = 1$. Below this threshold, the error of every group decays as $\kappa_g / (\varepsilon_g z_t)$, so the group proportion ε_g directly scales learning speed. Above it, the group with the larger $\mathbf{r}$ -margin escapes ε -dependence entirely and converges at the faster rate $z_t^{-\alpha}$ (Figure 4). We confirm the sharpness of this transition empirically (Figure 4).
- We prove a generalisation failure result (Corollary 7.1): under correlation shift at test time, the classifier may reach a non-vanishing error floor when the test margin $\gamma_{\text{test}} < 0$.

2 Problem Setup

A classical perspective in statistical learning theory frames nonlinear classification as a two-stage process: (i) a data representation $\Phi : \mathcal{X} \rightarrow \mathcal{H}$ lifts inputs into a reproducing kernel Hilbert space (RKHS) $\mathcal{H}$, in which (ii) a linear classifier finds a max-margin separator (SVM) (Schölkopf & Smola, 2002; Steinwart & Christmann, 2008; Chen et al., 2021). In this work we consider a data-generating distribution of pairs $(\mathbf{x}, y)$, where $\mathbf{x} \in \mathbb{R}^d$. The data representation $\Phi(\mathbf{x}) = (\mathbf{r}, \mathbf{s}) \in \mathbb{R}^{d_r + s_s}$ satisfies $(\mathbf{r}, \mathbf{s}) \sim \mathcal{D}_{\text{train}}$, where $\mathcal{D}_{\text{train}}$ is a continuous compactly supported distribution and where:

- $\mathbf{r} \in \mathbb{R}^{d_r}$ represents the *causal features*. Following the principle of Invariant Risk Minimization (IRM) (Arjovsky et al., 2019), these features are the invariant data representation that elicits an unbiased optimal classifier $f(\mathbf{r}) = y$ stable across all groups. Following the same formalism of IRM, $\mathbf{r}$ regroups the parents of y in the Structural Equation Model (SEM) of the problem.

• $\mathbf{s} \in \mathbb{R}^{d_s}$ represents the *spurious features*. These features are predictive of y only through group-specific correlations. Consequently, a predictor based on $\mathbf{s}$ cannot be simultaneously optimal across different groups.

Since our analysis operates at the population level rather than on a finite dataset, we require a notion of max-margin separation that is defined directly over the learned representation distribution:

Definition 2.1 (*r*-separability). *We say that the relevant features are **r**-separable whenever the following minimum-norm problem admits a solution:*

$$\mathbf{v} = \underset{\theta \in \mathbb{R}^{d_r}}{\operatorname{argmin}} \|\theta\|^2 \quad \text{s.t.} \quad P_{\mathcal{D}_{\text{train}}}(y \theta^\top \mathbf{r} \geq 1) = 1. \quad (1)$$

Throughout this work we assume $\|\mathbf{v}\| = 1$ to facilitate the analysis. Similarly, we have:

Definition 2.2 (full-separability). *We say that the full feature representation $\Phi(\mathbf{x}) = (\mathbf{r}, \mathbf{s})$ is **(r, s)**-separable whenever the following minimum-norm problem admits a solution:*

$$\hat{w} = \underset{w \in \mathbb{R}^{d_r + d_s}}{\operatorname{argmin}} \|w\|^2 \quad \text{s.t.} \quad P_{\mathcal{D}_{\text{train}}}(y w \cdot \Phi(\mathbf{x}) \geq 1) = 1. \quad (2)$$

We are not aware of prior work that explicitly formulates and uses these population-level minimum-norm separators; we therefore state the definitions here for clarity and assume that $\mathbf{v}$ and $\hat{w}$ exist.

Definition 2.3 (Features-mediated linear spurious correlation). *We say that the feature representation $\Phi(\mathbf{x})$ exhibits features-mediated linear spurious correlation when the conditional law $\mathbf{s} \mid \mathbf{r}$ satisfies:*

$$\mathbf{s} \mid \mathbf{r} \sim \begin{cases} A\mathbf{r} + \xi, & \text{with probability } 1 - \varepsilon, \\ B\mathbf{r} + \xi, & \text{with probability } \varepsilon, \end{cases} \quad (3)$$

where $\varepsilon \in [0, 1]$. The mappings $A, B \in \mathbb{R}^{d_s \times d_r}$ are spurious alignment operators encoding how the spurious signal linearly correlates with the relevant features in each group. ξ represents the residual and bounded random noise, dependent only on the **r**-margin variable:

$$\xi \perp (\mathbf{r}_\perp, y) \mid \mathbf{v} \cdot \mathbf{r}, \quad \text{with} \quad \mathbf{r}_\perp := \mathbf{r} - (\mathbf{v} \cdot \mathbf{r})\mathbf{v}. \quad (4)$$

Under this definition, the population naturally decomposes into two disjoint groups:

$$\mathcal{G}_1: (\mathbf{r}, \mathbf{s}, y) \text{ s.t. } \mathbf{s} = A\mathbf{r} + \xi, \quad \mathcal{G}_2: (\mathbf{r}, \mathbf{s}, y) \text{ s.t. } \mathbf{s} = B\mathbf{r} + \xi, \quad \varepsilon = P_{\mathcal{D}_{\text{train}}}(\mathcal{G}_2). \quad (5)$$

The spuriousness of $\mathbf{s}$ is defined by the condition $A \neq B$. This ensures that the optimal linear classifier from $\mathbf{s}$ to y varies across groups, thereby satisfying the definition of Arjovsky et al. (2019). Throughout all this work we use the convention $\mathcal{G}_{\text{maj}} = \mathcal{G}_1$, $\mathcal{G}_{\text{min}} = \mathcal{G}_2$, and $\Phi(\mathbf{x}) \leftarrow \mathbf{x}$ for simplicity.

3 Related Works

Theoretical analyses of spurious correlations. Prior theoretical explanations for why ERM relies on spurious correlations broadly fall into two classes. The first posits that both invariant and spurious features are only *partially* predictive of the label, so that the accuracy-maximising classifier cannot ignore the spurious component (Tsipras et al., 2018; Sagawa et al., 2020b; Arjovsky et al., 2019; Ilyas et al., 2019; Khani & Liang, 2020). The second invokes the *simplicity bias* of gradient-based training Rahaman et al. (2019); Kalimeris et al. (2019); Shah et al. (2020): even when both features are fully predictive, this model posits that the spurious feature is inherently “simpler” to learn than the invariant one, so that gradient descent preferentially relies on it (Shah et al., 2020; Hermann & Lampinen, 2020). Both frameworks have been developed primarily in the context of *label-mediated* spurious correlations, and therefore structurally differ from our setting.

At the model level, Sagawa et al. (2020b) show that overparameterisation creates a non-vanishing worst-group error floor in a linear model with spurious features. Nagarajan et al. (2020) identify geometric and statistical failure modes of ERM and show that the spurious weight component decays slowly under gradient descent. In the high-dimensional regime, Bombari & Mondelli (2025) characterises the spurious correlation learned by ridge regression through the covariance spectrum. Closest to our setting, Jain et al. (2024) derive closed-form ODE solutions for the SGD dynamics of a linear classifier on a Gaussian mixture with heterogeneous subpopulations, uncovering a three-phase

transient behaviour where variance, representation, and spurious shift dominate at different timescales. Our work differs in loss function (logistic vs. squared), data model (compact support vs. Gaussian), and the nature of the results: we obtain *asymptotic* per-group error exponents and a sharp phase transition at $\alpha = 1$, whereas their analysis characterises the full transient trajectory without producing such exponents. Furthermore, Yang et al. (2024) prove that the spurious contribution grows linearly with correlation strength during early training, and Bachoc et al. (2025) provide lower bounds on the additional training time needed to escape stereotypical predictors that neglect minority features. These works are complementary to our population-level, linear, asymptotic analysis.

Group robustness and mitigation strategies. Group robustness (Sagawa et al., 2020a) measures a model’s ability to maintain consistent performance across subgroups. Because the global empirical loss can mask poor minority-group accuracy, it is common to evaluate via the Worst-Group Accuracy (WGA). Distributionally Robust Optimisation (DRO) (Sagawa et al., 2020a) directly optimises WGA but requires group labels at training time. A wide range of methods instead attempt to achieve robustness without group annotations, including data augmentation (Yao et al., 2022), two-stage retraining (Liu et al., 2021; Kirichenko et al., 2023), balanced sampling (Idrissi et al., 2022), committee-based debiasing (Kim et al., 2022), shortcut probing (Zheng et al., 2025), and gradient-based proxy selection (Kenfack et al., 2025). A recurrent empirical finding across these works is that *group balancing*, is a key ingredient. Our theoretical results confirm that group proportion ε_g is indeed a first-order determinant of learning speed in the regime $\alpha < 1$, but reveal that it becomes irrelevant once the margin advantage exceeds the critical threshold ($\alpha \geq 1$).

4 Feature-mediated Spurious Correlations in Practice



Figure 2: Colored-MNIST dataset. **Top row:** majority group (maroon background), digits 0–9. Background intensity increases with digit class. **Bottom row:** minority group (teal background), digits 0–9. Background intensity *decreases* with digit class. The spurious feature (background intensity) depends on the causal feature (foreground image) through the digit class, a nonlinear function of the causal feature. The ground truth label $y = \mathbf{1}_{\text{digit} \geq 5}$, making this a feature-mediated spurious correlation.

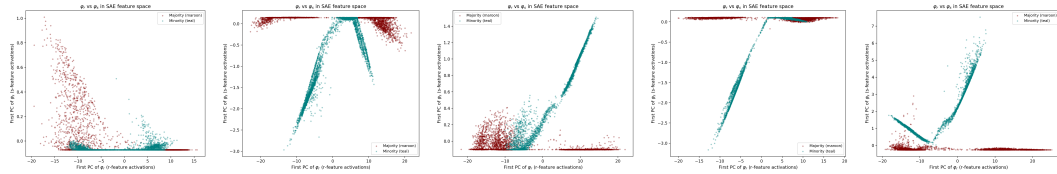


Figure 3: First principal component of ϕ_r (horizontal) vs. first principal component of ϕ_s (vertical) at epochs 1–5 (left to right). **Teal:** minority group; **maroon:** majority group. The minority points form a tight linear cloud whose slope becomes sharper as training progresses. The majority points are concentrated near $\phi_s \approx 0$.

While feature-mediated spurious correlation may appear restrictive, it arises naturally in several important settings. The most direct application is the *last-layer retraining* paradigm (Kirichenko et al., 2023), where a feature extractor Φ is frozen and only a linear head is retrained on balanced data. Motivated by the linear representation hypothesis (Park et al., 2023), and by the analysis of feature superposition (Elhage et al., 2022), we experimentally confirmed (Figures 2-3) that a correlation (even non linear) mediated by the raw data features becomes approximately linear in the model’s learned representation (Section B). Beyond learned representations, linear correlation structure appears in sensor fusion (e.g. radar and lidar reflectance measurements of the same scene

are related by a linear calibration model that varies with environmental conditions), in tabular data (where a protected attribute linearly predicts a correlated feature such as credit score, with coefficients that differ by demographic group or by provider – see Figure 1), in signal processing, where spatial filtering operations such as convolutions and blurs are inherently linear, or in Language Model alignment (PPO, DPO), where response length acts as a linearly correlated proxy for underlying semantic content (length bias Park et al. (2024)).

5 Main Results

Consider a linear model $f(x, w)$ with $w \in \mathbb{R}^{d_r + d_s}$ trained for a binary classification task: $y \in \{-1, 1\}$, with logistic loss function. We are now ready to state our main results. Our main theoretical results characterizes how features-mediated spurious correlation govern the asymptotic learning pace of each subgroup.

Theorem 5.1 (General Decay Rates). *Assume the learning rates sequence $(h_t)_{t \geq 0}$ satisfies the condition (21). Assume the infimums in (25) are achieved at the boundaries $y \mathbf{v} \cdot \mathbf{r} = \tilde{\gamma}_g$, then:*

$$\mathbb{E}_{\mathcal{G}_{maj}}[1 - p_y(\mathbf{x}, w^t)] = \Theta\left(\frac{(\ln z_t)^{\gamma_{maj} + \delta_{maj} + \tilde{\delta}_{maj} - 1}}{z_t^{\gamma_{maj} + o(1)}}\right), \quad (6)$$

$$\mathbb{E}_{\mathcal{G}_{min}}[1 - p_y(\mathbf{x}, w^t)] = \Theta\left(\frac{(\ln z_t)^{\gamma_{min} + \delta_{min} + \tilde{\delta}_{min} - 1}}{z_t^{\gamma_{min} + o(1)}}\right), \quad (7)$$

where γ_g are the group-specific hard-margins defined in Eq (25), $\tilde{\gamma}_g$ are the group-specific $\mathbf{r}$ -hard-margins Eq (26), $(\delta_g, \tilde{\delta}_g)$ are defined in Eqs (118)-(175), the $o(1)$ are uniformly noise-dependent, and $z_t := \sum_{k=0}^{t-1} h_k$ is the cumulative learning steps.

151

We also focus on this following special spurious correlation regime:

Definition 5.2 (Isotropic Regime). *A spurious correlation setting is isotropic if $\mathbf{v}$ is an eigenvector of $A^\top A$, $A^\top B$, $B^\top B$ and $B^\top A$. In plain terms, there exist $\mu_A, \mu_B \geq 0$, and μ such that:*

154

$$A^\top A \mathbf{v} = \mu_A \mathbf{v}, \quad A^\top B \mathbf{v} = \mu \mathbf{v} \quad (8)$$

$$B^\top B \mathbf{v} = \mu_B \mathbf{v}, \quad B^\top A \mathbf{v} = \mu \mathbf{v}. \quad (9)$$

In this regime, we assume

$$-1 \leq \mu < \min(\mu_A, \mu_B). \quad (10)$$

We now present our second result.

156

Theorem 5.3 (Isotropic regime). *Assume the spurious correlation is isotropic 5.2, and assume the noise ξ amplitude is small enough (see Eq (106)). Without loss of generality, assume $\tilde{\gamma}_{maj} = 1$ and $\tilde{\gamma}_{min} \geq 1$ and define*

$$\alpha := \frac{\tilde{\gamma}_{min}(1 + \mu)}{1 + \mu_A}, \quad (11)$$

where $\tilde{\gamma}_g$ are the group-specific $\mathbf{r}$ -hard-margins Eq (26). Then, under suitable regularity conditions on the data distribution (see Theorem D.4 for the precise form)

(i) If $\alpha < 1$, there exist two bounded factors $\kappa_{maj}^t(\tilde{\gamma}_{min})$ and $\kappa_{min}^t(\tilde{\gamma}_{min})$ independent of ε and independent of t when the noise $\xi = 0$, such that:

$$\mathbb{E}_{\mathcal{G}_{maj}}[1 - p_y^t] \underset{t \rightarrow \infty}{\sim} \frac{\kappa_{maj}^t(\tilde{\gamma}_{min})}{(1 - \varepsilon) z_t}, \quad \mathbb{E}_{\mathcal{G}_{min}}[1 - p_y^t] \underset{t \rightarrow \infty}{\sim} \frac{\kappa_{min}^t(\tilde{\gamma}_{min})}{\varepsilon z_t}, \quad (12)$$

(ii) If $\alpha > 1$, the classification error decay at different rates: there exists a positive bounded factor κ_{maj}^t independent on ε and $\tilde{\gamma}_{min}$, and independent on t when the noise $\xi = 0$, such that:

$$\mathbb{E}_{\mathcal{G}_{maj}}[1 - p_y^t] \underset{t \rightarrow \infty}{\sim} \frac{\kappa_{maj}^t}{(1 - \varepsilon) z_t}, \quad \mathbb{E}_{\mathcal{G}_{min}}[1 - p_y^t] = \Theta\left(\frac{(\ln z_t)^{\alpha_\xi - 1}}{z_t^{\alpha_\xi(1 + \delta_{maj}) - \delta_{min} + o(1)}}\right), \quad (13)$$

157

with the $o(1)$ uniformly noise-dependent, δ_g defined in Eq (118) noise-dependent, and $\alpha_\xi = \alpha + O(\|\xi\|_\infty) \xrightarrow{\|\xi\| \rightarrow 0} \alpha$ independently on t .

(iii) Furthermore, these different rate decay are continuous with respect to $\tilde{\gamma}_{\min}$, and ξ at the phase transition point $\alpha = 1$.

158

159 Note that the part (ii) condition $\alpha > 1$ is more precisely $\alpha_\xi > \frac{1+\delta_{\min}}{1+\delta_{\max}} \geq 1$ as we showed in Proposition F.7.

160 Theorem 5.3 naturally extends analogously to the mirror case where $\tilde{\gamma}_{\min}=1$ and $\tilde{\gamma}_{\max} \geq 1$, by
 161 substituting $(\mu_A, \tilde{\gamma}_{\min}) \leftarrow (\mu_B, \tilde{\gamma}_{\max})$. In the Appendix E-F, we present a detailed proof for a
 162 simplified setting of Theorem 5.3 and in Appendix G.1 we present a proof of Theorem 5.1.

163 Our main results rely on the following generalisation of Soudry et al. (2018, Theorem 9) for our
 164 continuous-distribution setting.

Proposition 5.4 (Implicit bias of GD - continuous-distribution setting). *Under Assumptions I.1-I.4, let w^t denote the iterates of population gradient descent (210) with any initial point w^0 and learning rate sequence $(\eta_t)_t$ satisfying condition (21). Then, as $t \rightarrow \infty$:*

$$w^t = \hat{w} \log z_t - \hat{w} \log \log z_t + O(1), \quad (14)$$

where $z_t := \sum_{k=0}^{t-1} h_k$ is the cumulative learning steps. In particular, the direction of w^t still converges to the max-margin direction:

$$\lim_{t \rightarrow \infty} \frac{w^t}{\|w^t\|} = \frac{\hat{w}}{\|\hat{w}\|}.$$

165

166 A proof of this result is provided in Section I with two versions: variable and fixed step-sizes. This
 167 result shows that unlike in the finite dataset case, the residual term $\rho^t := w^t - \hat{w} \log t$ is no longer
 168 bounded. Section 7 discusses this difference.

169 6 Experiments

170 We validate our theoretical predictions on a synthetic data pipeline that instantiates the isotropic
 171 regime (Definition 5.2) in a realistic tabular setting. All experiments use full-batch gradient descent
 172 with logistic loss, constant learning rate $h = 0.01$. Code is provided in the supplementary material.

173 **Experimental setup.** In many real-world tabular pipelines, practitioners concatenate raw, verified
 174 measurements (e.g. lab results, income records) that get encoded in the feature space as $\mathbf{r} \in \mathbb{R}^{d_r}$,
 175 with derived risk scores produced by external vendors (see Figure 1), encoded as $\mathbf{s} \in \mathbb{R}^{d_s}$. As
 176 demonstrated empirically in Section 4, the dependence of these derived risk scores on the raw
 177 measurements linearises to take this form:

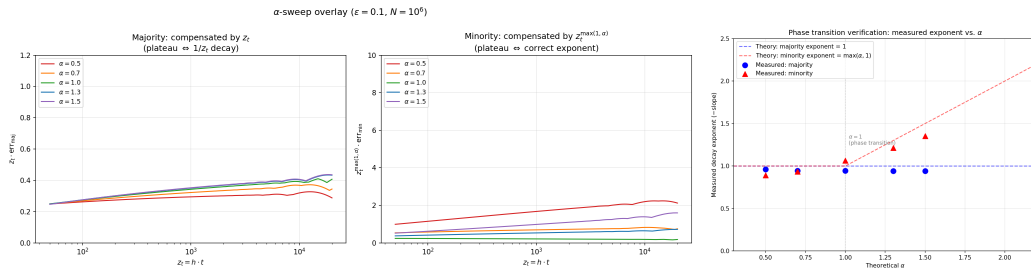


Figure 4: α -sweep overlay ($\varepsilon = 0.1$, $N = 10^6$, $T = 2 \cdot 10^6$). **Left:** majority errors compensated by z_t plateau for all α , confirming the universal $1/z_t$ rate. **Middle:** minority errors compensated by $z_t^{\max(1, \alpha)}$ also plateau, confirming that the theoretical exponent is correct across the full sweep. **Right:** empirical decay exponent β vs. theoretical α . Dashed lines: theoretical predictions (flat at 1 for majority; $\max(\alpha, 1)$ for minority). This provides evidence for the phase transition Theorem 5.3.

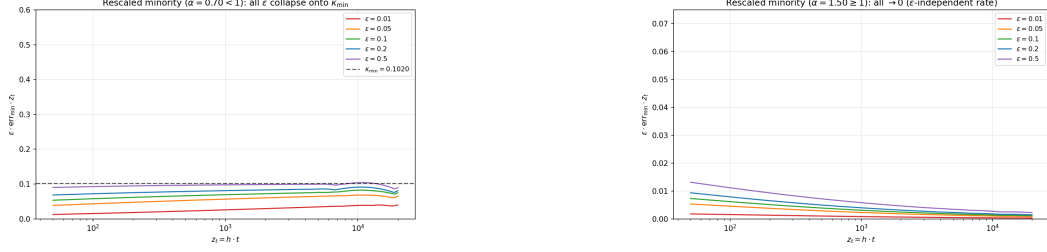


Figure 5: Rescaled minority error $\varepsilon \cdot \text{err}_{\min} \cdot z_t$ for various ε . **Left** ($\alpha = 0.7$): all curves collapse onto the theoretical constant $\kappa_{\min}$ (dashed line), confirming the exact prefactor predicted by Theorem 5.3-(i). **Right** ($\alpha = 1.5$): all curves tend to zero together, confirming ε -independent decay.

$$\mathbf{s} = A \mathbf{r} + \xi \quad (\text{for vendor 0}), \quad \mathbf{s} = B \mathbf{r} + \xi \quad (\text{for vendor 1}), \quad (15)$$

Because both vendors attempt to capture the same underlying phenomenon, we can naturally assume the isotropic regime by construction. We fix $\mu_A = \mu_B = 1$, $\mu = 0$, $d_r = d_s = 8$, and vary $\tilde{\gamma}_{\min}$ through α . All runs last $T = 2 \cdot 10^6$ training steps, and use $N = 10^6$ samples. We used a truncated Gaussian noise according to condition (28).

Asymptotic rates and phase transition. We fix $\varepsilon = 0.1$ and sweep $\alpha \in \{0.5, 0.7, 1.0, 1.3, 1.5\}$. Figure 4 displays the compensated errors for all α values overlaid. The left panel shows the majority error decays at $1/z_t$ rate since all curves plateau at $\kappa_{\text{maj}}(\alpha)/(1-\varepsilon)$. The middle panel shows the minority error decays at $1/z_t^{\max(1, \alpha)}$ rate since all curves plateau, matching Theorem 5.3 predictions. The right panel illustrates the phase transition by plotting the empirical decay exponent β (obtained via log-log regression) against the theoretical α .

Impact of group proportion.

We fix two representative α values (one below and one above the phase transition) and sweep $\varepsilon \in \{0.01, 0.05, 0.1, 0.2, 0.5\}$. Figure 5 (left) shows the rescaled quantity $\varepsilon \cdot \text{err}_{\min} \cdot z_t^{\max(1, \alpha)}$ for both α . At $\alpha < 1$ (left panel), all five curves collapse onto the single constant $\kappa_{\min}$, confirming the $1/(\varepsilon \cdot z_t^{\max(1, \alpha)})$ rate decay of the classification error. At $\alpha \geq 1$ (right panel), all curves tend toward zero together regardless of ε , confirming that the minority rate $\Theta(z_t^{-\alpha} (\ln z_t)^{\alpha-1})$ is ε -independent once $\alpha \geq 1$.

7 Discussion

Scope of the Isotropic Regime. The isotropic regime (Definition 5.2) is less restrictive than it may first appear: we show in Section C that any operators (A, B) with $d_s \geq 2d_r$ admits an isotropic reparametrisation (A', B') provided the minimum displacement d^* (Eq (40)) is small enough for the noise condition (106) to hold. In the overparameterised setting $d_s \gg d_r$ that is typical of learned representations (e.g. the last-layer retraining paradigm), this dimensional requirement is easily met.

Table 1: Summary of differences from Soudry et al. (2018). γ_g is the group-specific margin, ε_g the group proportion, $\tilde{\gamma}_g$ the group-specific $\mathbf{r}$ -only margin.

	Finite dataset (Soudry)	Continuous dist. (ours)
Gradient structure	finite sum	Laplace integral
Support set	finite $\mathcal{S}$ with $ \mathcal{S} \leq d$	$\{\mathbf{x} \cdot \hat{\mathbf{w}} = 1\}$, measure zero
Residual $\rho(t)$	$O(1)$ (Thm. 9)	$-\hat{\mathbf{w}} \log \log z_t + O(1)$
Direction convergence	$O(\log \log t / \log t)$	$O(\log \log z_t / \log z_t)$
Classification error	$\Theta(t^{-1})$	$\kappa_g^t(\tilde{\gamma}_g)/(z_t \varepsilon_g)^*$ $\kappa_g^t(\tilde{\gamma}_g)/(z_t \varepsilon_g)$ and $\Theta(z_t^{-\alpha \xi(1+\delta_{\tilde{g}})+\delta_g} (\log z_t)^{\alpha \xi-1})^{**}$ $\Theta(z_t^{-\gamma_g} (\log z_t)^{\gamma_g+\delta_g+\delta_g-1})^{***}$

* Decay rate in the isotropic regime ($\alpha < 1$).

** Decay rate in the isotropic regime ($\alpha > 1$), $g = \min$

*** General coarse decay rate.

201 **When is group balancing sufficient?** A widely adopted conclusion in the group-robustness liter-
 202 ature (Liu et al., 2021; Yao et al., 2022; Idrissi et al., 2022; Kim et al., 2022) is that balancing the
 203 group proportions, whether by resampling, reweighting, or data augmentation, is the primary lever
 204 for improving worst-group performance. Our isotropic regime results provide a precise theoretical
 205 grounding for this claim, and at the same time delineate its limits.

206 When $\alpha < 1$, the error on each group scales as $\kappa_g^t / (\varepsilon_g z_t)$. In this regime, the group proportion ε_g
 207 acts as a direct multiplier on the learning speed: a group that is ten times smaller learns ten times
 208 slower. Balancing ($\varepsilon = 1/2$) equalises the denominators and is therefore an effective intervention.

209 However, when $\alpha > 1$, the minority error decays as $\Theta\left(z_t^{\alpha_\xi(1+\delta_{\text{maj}})-\delta_{\text{min}}+o(1)}(\ln z_t)^{\alpha_\xi-1}\right)$, a rate
 210 that *does not depend on ε* at all: the minority’s geometric margin advantage is large enough that it
 211 learns faster than the majority regardless of how small it is. In this regime, the majority group is the
 212 bottleneck, and its rate $1/((1-\varepsilon)z_t)$ is the one that benefits from balancing.

213 This reveals a range of $\tilde{\gamma}_{\text{min}}$ values, specifically $1 \leq \tilde{\gamma}_{\text{min}} < (1+\mu_A)/(1+\mu)$, where the spurious
 214 correlation still influences the learning dynamics through ε , and balancing is a meaningful intervention.
 215 Beyond this threshold, balancing the groups no longer accelerates the slower group; the margin
 216 geometry has become the binding constraint.

217 **Interpretation of the main results.** (i) *The role of α .* Theorems 5.1 and 5.3 together paint a coherent
 218 picture of how spurious correlations govern optimisation dynamics. Theorem 5.1 generalises the
 219 classification error analysis of Soudry et al. (2018) to a multi-group population and shows that, in
 220 full generality, each group’s error decays at a rate controlled by its hard-margin γ_g . Theorem 5.3
 221 then zooms into the isotropic regime and reveals that the *effective* quantity controlling the decay
 222 rate is the exponent $\alpha = \tilde{\gamma}_{\text{min}}(1+\mu)/(1+\mu_A)$, not the **r**-margin $\tilde{\gamma}_{\text{min}}$ alone. The exponent α captures
 223 a competition: the minority group’s geometric margin advantage ($\tilde{\gamma}_{\text{min}} \geq 1$) is discounted by the
 224 coupling ratio $(1+\mu_A)/(1+\mu) \geq 1$, which measures how efficiently the majority’s gradient drift
 225 propagates to the minority logits. In other words, the spurious features “absorb” part of the minority’s
 226 margin advantage, and α quantifies exactly how much survives.

227 Aligning the two theorems in the isotropic regime yields the identity $\gamma_{\text{min}} = \max(1, \alpha_\xi(1 + \delta_{\text{maj}}) -$
 228 $\delta_{\text{min}}) \approx \max(1, \alpha)$, providing an explicit relationship between the global hard-margin γ_{min} (which
 229 determines the coarse rate in Theorem 5.1) and the **r**-margin $\tilde{\gamma}_{\text{min}}$ modulated by the spurious alignment
 230 eigenvalues. The critical threshold $\gamma_{\text{crit}} = (1+\mu_A)/(1+\mu)$ depends on the coupling μ between the
 231 spurious alignment operators and the “diagonal” eigenvalue μ_A .

232 (ii) *Phase transition at $\alpha = 1$.* Below the threshold ($\alpha < 1$), both groups’ errors decay at the same
 233 $1/z_t$ rate, but with group-specific prefactors κ_g/ε_g . Above it ($\alpha > 1$), the two groups decouple: the
 234 majority still decays as $1/((1-\varepsilon)(1+\mu_A)z_t)$, while the minority accelerates to $z_t^{-\alpha}(\ln z_t)^{\alpha-1}$ in the
 235 noise-free case. When feature noise is present, the minority exponent becomes $\alpha_\xi(1+\delta_{\text{maj}})-\delta_{\text{min}} \geq 1$,
 236 a noise-dependent perturbation of α that converges back to α as the noise amplitude vanishes (see
 237 Eq (140)). The phase transition is continuous: at $\alpha = 1$, both the prefactors and the exponents match
 238 smoothly across the two regimes.

239 **Features-mediated spuriousness Generalisation Failure.** Most works on label-mediated spurious
 240 correlation observe a generalisation failure of the *worst-group accuracy* when the model is evaluated
 241 on minority-group samples underrepresented during training (Kirichenko et al., 2023; Sagawa et al.,
 242 2020a; Nagarajan et al., 2020). Under ERM, the classifier places excessive weight on the spurious
 243 component, making it vulnerable to any shift in the spurious correlation structure between training
 244 and test. This naturally extends to our *Expected Risk Minimization* setup by defining a continuous
 245 test distribution $\mathcal{D}_{\text{test}}$ with a shifted spurious correlation:

$$\tilde{\gamma}_{\text{test}} := \text{ess inf}_{\mathbf{r} | \mathcal{D}_{\text{test}}} y \mathbf{v} \cdot \mathbf{r}, \quad \mathbf{s} | \mathbf{r} \sim C\mathbf{r} + \xi, \quad \text{under } \mathcal{D}_{\text{test}}, \quad (16)$$

246 where $C \in \mathbb{R}^{d_s \times d_r}$ is an arbitrary spurious alignment operator. Define the *projected test weight* and
 247 test margins: writing $\hat{\mathbf{w}} = [\hat{\mathbf{w}}_r \quad \hat{\mathbf{w}}_s]$ (see Eq (2)),

$$\hat{\chi} := \hat{\mathbf{w}}_r + C^\top \hat{\mathbf{w}}_s, \quad \gamma_{\text{test}}(C) := \text{ess inf}_{(\mathbf{r}, \xi) \sim \mathcal{D}_{\text{test}}} (\hat{\chi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi), \quad \delta_{\text{test}} := \delta_+(\tilde{\gamma}_{\text{test}}), \quad \tilde{\delta}_{\text{test}} := \tilde{\delta}_+(\tilde{\gamma}_{\text{test}}), \quad (17)$$

248 with $\delta_+, \tilde{\delta}_+$ defined in Eqs (117) and (168). We formulate these following generalisation results:

Corollary 7.1 (Accuracy ceiling under correlation shift). *Assume the infimum in (17) is achieved at the boundary $y \mathbf{v} \cdot \mathbf{r} = \tilde{\gamma}_{\text{test}} : \gamma_{\text{test}} = \gamma_{\text{test}}^{(r)} - \delta_{\text{test}}$, with $\gamma_{\text{test}}^{(r)} := \text{ess inf}_{\mathbf{r} | y \mathbf{v} \cdot \mathbf{r} = \tilde{\gamma}_{\text{test}}} \{\hat{\chi} \cdot \mathbf{r}\}$. We have the following cases:*

(i) *If $\gamma_{\text{test}}(C) > 0$ (test group still separated by $\hat{w}$), then*

$$\mathbb{E}_{\mathcal{D}_{\text{test}}}[1 - p_y(\mathbf{x}, w^t)] = \Theta\left(\frac{(\log z_t)^{\gamma_{\text{test}} + \delta_{\text{test}} + \tilde{\delta}_{\text{test}} - 1}}{z_t^{\gamma_{\text{test}} + o(1)}}\right), \quad (18)$$

where the $o(1)$ is uniformly noise dependent. The test group error still vanishes, but at a rate governed by $\gamma_{\text{test}}(C)$.

(ii) *If $\gamma_{\text{test}}(C) < 0$ (test group not separated by $\hat{w}$):*

$$\lim_{t \rightarrow \infty} \mathbb{E}_{\mathcal{D}_{\text{test}}}[1 - p_y(\mathbf{x}, w^t)] = P_{\mathcal{D}_{\text{test}}}(\hat{\chi} \cdot \mathbf{r} + \hat{w}_s \cdot \xi < 0) > 0. \quad (19)$$

249

250 Part (ii) reveals a non-vanishing error floor caused by the correlation shift (proof in Section G).

251 **Differences between the finite and the continuous settings.** Table 1 summarises the structural
252 differences between *Expected* and *Empirical* Risk Minimization settings Soudry et al. (2018).

253 (i) *Origin of the $\log \log t$ term in Eq (14).* In the finite-sample setting, the gradient at time t is a finite
254 sum over support vectors, and its leading term matches $\hat{w}/t$ exactly via Soudry et al. (2018, Eq 21),
255 since the residual gradient $\dot{r}$ (reusing their notations) is strictly dominated because it is integrable.
256 In the continuous-distribution setting, the gradient is a Laplace integral (via the disintegration
257 theorem, Eq 236). By Watson’s Lemma the leading term is $t^{-1}/\log t$. The additional $1/\log t$ factor
258 creates a systematic mismatch between the gradient and $\hat{w}/t$, which accumulates over time into the
259 $-\hat{w} \log \log t$ correction. This mechanism is entirely distinct from the degeneracy-driven $O(\log \log t)$
260 growth in Soudry et al. (2018, Theorem 13).

261 (ii) *Absence of the zero-activation degeneracy.* In the finite-sample setting, when certain support
262 vectors have exactly zero activation (i.e. $\hat{w} \cdot x_i = 0$ for some support vector x_i), the residual $\rho(t)$ can
263 grow as $O(\log \log t)$ rather than remaining bounded – the degenerate case treated in Soudry et al.
264 (2018, Theorem 13). Our continuous-distribution framework sidesteps this issue entirely: the margin
265 density $\lambda(z)$ is assumed strictly positive and continuous at the boundary $z = 1$, so the gradient
266 contribution is spread over a continuum of margin values rather than concentrated on isolated support
267 vectors. This is at once a theoretical strength – the analysis is cleaner and the resulting rates are
268 sharper – and a modelling limitation, since finite datasets can exhibit such degeneracies in practice.
269 The finite-sample concentration analysis that would bridge the two regimes is left to future work.

270 8 Conclusion and Limitations

271 We have presented a mathematical analysis of how features-mediated spurious correlations shape
272 the per-group learning dynamics of gradient descent. By working at the population level on com-
273 pactly supported distributions, we obtained closed-form group-wise classification error rates that are
274 governed by two quantities: the group proportion ε_g and a composite exponent α that encodes the
275 competition between geometric margins and spurious alignment operators. The phase transition at
276 $\alpha = 1$ cleanly separates a regime where group proportion is the dominant factor from one where
277 margin geometry supersedes it. These results formalise the empirical intuition that group imbalance
278 slows minority learning, while clarifying the precise conditions under which this effect can be entirely
279 overcome by favourable geometry. Beyond training-time dynamics, our analysis reveals that when
280 the testing group is not separated by the learned weights: $\gamma_{\text{test}}(C) < 0$, a *non-vanishing error floor*
281 emerges, irrespective of how long training continues (Corollary 7.1).

282 **Limitations.** Our analysis rests on several simplifying assumptions. First, the population-level
283 analysis avoids finite-sample effects; bridging the gap to empirical risk via concentration inequalities
284 remains open. Second, we treat *binary* classification only. These limitations pave the ground to
285 several potential future works. The *multi-categories* generalisation, where the linear relationship
286 takes more than two forms, would significantly broaden the scope of the theory, and a *finite-sample*
287 version of our results, connecting the population-level rates to empirical convergence via uniform
288 concentration bounds, would complete the bridge between theory and practice.

References

- Luigi Ambrosio, Nicola Gigli, and Giuseppe Savaré. *Gradient flows: in metric spaces and in the space of probability measures*. Springer, 2005.
- Martin Arjovsky, Léon Bottou, Ishaan Gulrajani, and David Lopez-Paz. Invariant risk minimization. *arXiv preprint arXiv:1907.02893*, 2019.
- François Bachoc, Jérôme Bolte, Ryan Boustany, and Jean-Michel Loubes. When majority rules, minority loses: bias amplification of gradient descent. *arXiv preprint arXiv:2505.13122*, 2025.
- Simone Bombari and Marco Mondelli. Spurious correlations in high dimensional regression: The roles of regularization, simplicity bias and over-parameterization. In Aarti Singh, Maryam Fazel, Daniel Hsu, Simon Lacoste-Julien, Felix Berkenkamp, Tegan Maharaj, Kiri Wagstaff, and Jerry Zhu (eds.), *Proceedings of the 42nd International Conference on Machine Learning*, volume 267 of *Proceedings of Machine Learning Research*, pp. 4839–4873. PMLR, 13–19 Jul 2025. URL <https://proceedings.mlr.press/v267/bombari25a.html>.
- Yilan Chen, Wei Huang, Lam Nguyen, and Tsui-Wei Weng. On the equivalence between neural network and support vector machine. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan (eds.), *Advances in Neural Information Processing Systems*, volume 34, pp. 23478–23490. Curran Associates, Inc., 2021. URL https://proceedings.neurips.cc/paper_files/paper/2021/file/c559da2ba967eb820766939a658022c8-Paper.pdf.
- Nelson Elhage, Tristan Hume, Catherine Olsson, Nicholas Schiefer, Tom Henighan, Shauna Kravec, Zac Hatfield-Dodds, Robert Lasenby, Dawn Drain, Carol Chen, et al. Toy models of superposition. *arXiv preprint arXiv:2209.10652*, 2022.
- William Feller et al. *An introduction to probability theory and its applications*. Wiley New York, 1971.
- Katherine Hermann and Andrew Lampinen. What shapes feature representations? exploring datasets, architectures, and training. *Advances in Neural Information Processing Systems*, 33:9995–10006, 2020.
- Robert Huben, Hoagy Cunningham, Logan Smith, Aidan Ewart, and Lee Sharkey. Sparse autoencoders find highly interpretable features in language models. In *International Conference on Learning Representations*, volume 2024, pp. 7827–7845, 2024.
- Badr Youbi Idrissi, Martin Arjovsky, Mohammad Pezeshki, and David Lopez-Paz. Simple data balancing achieves competitive worst-group-accuracy. In Bernhard Schölkopf, Caroline Uhler, and Kun Zhang (eds.), *Proceedings of the First Conference on Causal Learning and Reasoning*, volume 177 of *Proceedings of Machine Learning Research*, pp. 336–351. PMLR, 11–13 Apr 2022. URL <https://proceedings.mlr.press/v177/idrissi22a.html>.
- Andrew Ilyas, Shibani Santurkar, Dimitris Tsipras, Logan Engstrom, Brandon Tran, and Aleksander Madry. Adversarial examples are not bugs, they are features. *Advances in neural information processing systems*, 32, 2019.
- Anchit Jain, Rozhin Nobahari, Aristide Baratin, and Stefano Sarao Mannelli. Bias in motion: Theoretical insights into the dynamics of bias in sgd training. *Advances in Neural Information Processing Systems*, 37:24435–24471, 2024.
- Eric Jourdain. Intégrales généralisées : Théorème d’intégration des équivalents. <https://perso.univ-rennes1.fr/eric.jourdain/AP3/Chapitre4.pdf>, 2018. Théorème 4.11 — *Intégration des équivalents*.
- Dimitris Kalimeris, Gal Kaplun, Preetum Nakkiran, Benjamin Edelman, Tristan Yang, Boaz Barak, and Haofeng Zhang. Sgd on neural networks learns functions of increasing complexity. *Advances in neural information processing systems*, 32, 2019.
- Patrik Kenfack, Ulrich Aïvodji, and Samira Ebrahimi Kahou. Gradtune: Last-layer fine-tuning for group robustness without group annotation. In *Workshop on Spurious Correlation and Shortcut Learning: Foundations and Solutions*, 2025.

338 Hassan K. Khalil. *Nonlinear Systems*. Prentice Hall, 3rd edition, 2002.

339 Fereshte Khani and Percy Liang. Feature noise induces loss discrepancy across groups. In *International conference on machine learning*, pp. 5209–5219. PMLR, 2020.

341 Nayeong Kim, Sehyun Hwang, Sungsoo Ahn, Jaesik Park, and Suha Kwak. Learning debiased
342 classifier with biased committee. In *ICML 2022: Workshop on Spurious Correlations, Invariance
343 and Stability*, 2022. URL <https://openreview.net/forum?id=bcxmUnTPwny>.

344 Polina Kirichenko, Pavel Izmailov, and Andrew Gordon Wilson. Last layer re-training is sufficient
345 for robustness to spurious correlations. In *The Eleventh International Conference on Learning
346 Representations*, 2023. URL <https://openreview.net/forum?id=Zb6c8A-Fghk>.

347 Evan Z Liu, Behzad Haghighi, Annie S Chen, Aditi Raghunathan, Pang Wei Koh, Shiori Sagawa,
348 Percy Liang, and Chelsea Finn. Just train twice: Improving group robustness without training
349 group information. In *International Conference on Machine Learning*, pp. 6781–6792. PMLR,
350 2021.

351 Vaishnavh Nagarajan, Anders Andreassen, and Behnam Neyshabur. Understanding the failure modes
352 of out-of-distribution generalization. *arXiv preprint arXiv:2010.15775*, 2020.

353 Kiho Park, Yo Joong Choe, and Victor Veitch. The linear representation hypothesis and the geometry
354 of large language models. *arXiv preprint arXiv:2311.03658*, 2023.

355 Ryan Park, Rafael Rafailov, Stefano Ermon, and Chelsea Finn. Disentangling length from quality in
356 direct preference optimization. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar (eds.), *Findings
357 of the Association for Computational Linguistics: ACL 2024*, pp. 4998–5017, Bangkok, Thailand,
358 August 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.findings-acl.297.
359 URL <https://aclanthology.org/2024.findings-acl.297/>.

360 Nasim Rahaman, Aristide Baratin, Devansh Arpit, Felix Draxler, Min Lin, Fred Hamprecht, Yoshua
361 Bengio, and Aaron Courville. On the spectral bias of neural networks. In *International conference
362 on machine learning*, pp. 5301–5310. PMLR, 2019.

363 Shiori Sagawa, Pang Wei Koh, Tatsunori B Hashimoto, and Percy Liang. Distributionally robust
364 neural networks. In *International Conference on Learning Representations*, 2020a.

365 Shiori Sagawa, Aditi Raghunathan, Pang Wei Koh, and Percy Liang. An investigation of why
366 overparameterization exacerbates spurious correlations. In Hal Daumé III and Aarti Singh
367 (eds.), *Proceedings of the 37th International Conference on Machine Learning*, volume 119
368 of *Proceedings of Machine Learning Research*, pp. 8346–8356. PMLR, 13–18 Jul 2020b. URL
369 <https://proceedings.mlr.press/v119/sagawa20a.html>.

370 Bernhard Schölkopf and Alexander J Smola. *Learning with kernels: support vector machines,
371 regularization, optimization, and beyond*. MIT press, 2002.

372 Harshay Shah, Kaustav Tamuly, Aditi Raghunathan, Prateek Jain, and Praneeth Netrapalli. The
373 pitfalls of simplicity bias in neural networks. *Advances in Neural Information Processing Systems*,
374 33:9573–9585, 2020.

375 Daniel Soudry, Elad Hoffer, Mor Shpigel Nacson, Suriya Gunasekar, and Nathan Srebro. The implicit
376 bias of gradient descent on separable data. *Journal of Machine Learning Research*, 19(70):1–57,
377 2018.

378 Ingo Steinwart and Andreas Christmann. *Support vector machines*. Springer Science & Business
379 Media, 2008.

380 Dimitris Tsipras, Shibani Santurkar, Logan Engstrom, Alexander Turner, and Aleksander Madry.
381 Robustness may be at odds with accuracy. *arXiv preprint arXiv:1805.12152*, 2018.

382 Vladimir N Vapnik. An overview of statistical learning theory. *IEEE transactions on neural networks*,
383 10(5):988–999, 1999.

384 Matt Visser. Primes and the lambert w function. *Mathematics*, 6(4):56, 2018.

- 385 Roderick Wong. *Asymptotic approximations of integrals*. SIAM, 2001.
- 386 Yu Yang, Eric Gan, Gintare Karolina Dziugaite, and Baharan Mirzasoleiman. Identifying spurious
 387 biases early in training through the lens of simplicity bias. In Sanjoy Dasgupta, Stephan Mandt,
 388 and Yingzhen Li (eds.), *Proceedings of The 27th International Conference on Artificial Intelligence
 389 and Statistics*, volume 238 of *Proceedings of Machine Learning Research*, pp. 2953–2961. PMLR,
 390 02–04 May 2024. URL <https://proceedings.mlr.press/v238/yang24c.html>.
- 391 Huaxiu Yao, Yu Wang, Sai Li, Linjun Zhang, Weixin Liang, James Zou, and Chelsea Finn. Improving
 392 out-of-distribution robustness via selective augmentation. In Kamalika Chaudhuri, Stefanie
 393 Jegelka, Le Song, Csaba Szepesvari, Gang Niu, and Sivan Sabato (eds.), *Proceedings of the 39th
 394 International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning
 395 Research*, pp. 25407–25437. PMLR, 17–23 Jul 2022. URL [https://proceedings.mlr.press/
 396 v162/yao22b.html](https://proceedings.mlr.press/v162/yao22b.html).
- 397 Guangtao Zheng, Wenqian Ye, and Aidong Zhang. Shortcutprobe: Probing prediction shortcuts for
 398 learning robust models. In James Kwok (ed.), *Proceedings of the Thirty-Fourth International Joint
 399 Conference on Artificial Intelligence, IJCAI-25*, pp. 7146–7154. International Joint Conferences
 400 on Artificial Intelligence Organization, 8 2025. doi: 10.24963/ijcai.2025/795. URL [https:
 401 //doi.org/10.24963/ijcai.2025/795](https://doi.org/10.24963/ijcai.2025/795). Main Track.

402 **Appendix**

403 **Table of Contents**

404	A	Notations	14
405	B	Features-mediated Spurious Correlation Linearises in Learned Representation	15
406	B.1	Experimental Methodology	15
407	B.2	Results and interpretation	16
408	C	Scope of the Isotropic Regime.	17
409	D	Relaxed Version of Theorem 5.3	19
410	E	Proof of Theorem D.4	20
411	E.1	Auxiliary Results	20
412	E.2	Discrete and Continuous Error Dynamics	21
413	E.3	ODE Synchronisation	22
414	E.4	Conclusion of the Proof when $\alpha < 1$	26
415	E.5	Conclusion of the Proof when $\alpha \geq 1$	26
416	F	Generalizing to the proof of Theorem 5.3	28
417	F.1	Relaxing Assumptions D.1 and D.2	28
418	F.2	Relaxing Assumption D.3 (feature noise)	28
419	G	Proof of Corollary 7.1	36
420	G.1	Application to the Proof of Theorem 5.1	37
421	H	Toolbox	39
422	H.1	Asymptotic Behavior of $\mathcal{J}$ and $\mathcal{I}$	39
423	H.2	Asymptotic Behavior of the Solution of 1D ODEs	40
424	H.3	Euler Scheme Error Bound	43
425	H.4	Synchronisation Method Lemmas	44
426	H.5	Discrete Synchronisation	46
427	I	Generalizing Theorem 9 of Soudry et al.	48
428	I.1	Setup and assumptions	48
429	I.2	Proof of Proposition 5.4	49
430	I.3	Generalisation to variable step-size	53

A Notations

The model parameters after t training steps are denoted $w^t = [w_r^t \ w_s^t] \in \mathbb{R}^{d_r+d_s}$.

Denote $p_c(x, w) = \sigma(c w \cdot x)$ the predicted probability of input x being in class $c \in \{-1, 1\}$, where σ is the sigmoid function $1/(1 + e^{-u})$.

We use the logistic loss function $\ell(f(x, w), y) = -\log p_y(x, w)$, and the objective function is $\mathcal{L}(\mathcal{D}_{\text{train}}, w) = \mathbb{E}_{\mathcal{D}_{\text{train}}}[\ell(f(\mathbf{x}, w), y)]$. The parameters are updated via population gradient descent with step size h_t :

$$w^{t+1} = w^t - h_t \nabla \mathcal{L}(\mathcal{D}_{\text{train}}, w^t). \quad (20)$$

We assume that the sequence $(h_t)_{t \geq 0}$ satisfies

$$\sum_{t=0}^{\infty} h_t = +\infty \quad \text{and} \quad \sum_{t=0}^{\infty} h_t^2 < \infty, \quad h_t > 0, \quad (21)$$

and we denote $z_t := \sum_{k=0}^{t-1} h_k$ is the cumulative learning steps.

For any vector $u \in \mathbb{R}^{d_r} \setminus \{0\}$, we denote the orthogonal projections

$$\Pi_u : x \in \mathbb{R}^{d_r} \mapsto \frac{u \cdot x}{\|u\|} \frac{u}{\|u\|} \in \text{Span}\{u\}, \quad \bar{\Pi}_u : x \in \mathbb{R}^{d_r} \mapsto x - \frac{u \cdot x}{\|u\|} \frac{u}{\|u\|} \in \text{Span}\{u\}^\perp, \quad (22)$$

so that $\Pi_u + \bar{\Pi}_u = I_{d_r}$.

$\mathbf{v}$ is the solution of the minimum-norm hard-margin problem:

$$\mathbf{v} = \underset{\theta \in \mathbb{R}^{d_r}}{\text{argmin}} \|\theta\|^2 \quad \text{s.t.} \quad P_{\mathcal{D}_{\text{train}}}(y \theta^\top \mathbf{r} \geq 1) = 1, \quad (23)$$

and throughout this work we assume $\|\mathbf{v}\| = 1$ to facilitate the analysis.

$\hat{w}$ is the solution of the minimum-norm hard-margin problem:

$$\hat{w} = \underset{w \in \mathbb{R}^{d_r+d_s}}{\text{argmin}} \|w\|^2 \quad \text{s.t.} \quad P_{\mathcal{D}_{\text{train}}}(y w \cdot \mathbf{x} \geq 1) = 1 \quad (24)$$

Both $\mathbf{v}$ and $\hat{w}$ are assumed to exist in our setting.

Define the group-specific hard-margins for the two groups as:

$$\gamma_{\text{maj}} = \text{ess inf}_{(\mathbf{x}, y) \in \mathcal{G}_{\text{maj}}} y \hat{w} \cdot \mathbf{x}, \quad \gamma_{\text{min}} = \text{ess inf}_{(\mathbf{x}, y) \in \mathcal{G}_{\text{min}}} y \hat{w} \cdot \mathbf{x}, \quad \min(\gamma_{\text{maj}}, \gamma_{\text{min}}) = 1, \quad (25)$$

and the group-specific $\mathbf{r}$ -hard-margins for the two groups as

$$\tilde{\gamma}_{\text{maj}} = \text{ess inf}_{(\mathbf{r}, y) \in \mathcal{G}_{\text{maj}}} y \mathbf{v} \cdot \mathbf{r}, \quad \tilde{\gamma}_{\text{min}} = \text{ess inf}_{(\mathbf{r}, y) \in \mathcal{G}_{\text{min}}} y \mathbf{v} \cdot \mathbf{r}, \quad \min(\tilde{\gamma}_{\text{maj}}, \tilde{\gamma}_{\text{min}}) = 1. \quad (26)$$

We work with compactly supported distributions, therefore we define the constants K , and R :

$$\|\mathbf{r}\| \leq K \quad \text{almost surely}, \quad (27)$$

$$\|\xi\| \leq R \quad \text{almost surely}. \quad (28)$$

In the isotropic regime Definition 5.2, $\mathbf{v}$ is a common eigenvector for $A^\top A$, $A^\top B$, $B^\top B$ and $B^\top A$.

We therefore denote μ_A , μ_B , and μ their corresponding positive eigenvalues:

$$A^\top A \mathbf{v} = \mu_A \mathbf{v}, \quad A^\top B \mathbf{v} = \mu \mathbf{v} \quad (29)$$

$$B^\top B \mathbf{v} = \mu_B \mathbf{v}, \quad B^\top A \mathbf{v} = \mu \mathbf{v}. \quad (30)$$

Define the quantity

$$\Sigma := (1 + \mu_A)(1 + \mu_B) - (1 + \mu)^2 > 0, \quad (31)$$

which is strictly positive under condition 10.

Write $Z := \mathbf{v} \cdot \mathbf{r}$, the $\mathbf{r}$ -margin variable. We denote λ_{maj} and λ_{min} the continuous probability density functions, such that $\lambda_{\text{maj}}(\tilde{\gamma}_{\text{maj}}) > 0$, $\lambda_{\text{min}}(\tilde{\gamma}_{\text{min}}) > 0$, and

$$p_Z(z | g) = \begin{cases} \lambda_{\text{maj}}(z) \mathbb{1}_{[\tilde{\gamma}_{\text{maj}}, K]}(z) & \text{if } g = \mathcal{G}_{\text{maj}}, \\ \lambda_{\text{min}}(z) \mathbb{1}_{[\tilde{\gamma}_{\text{min}}, K]}(z) & \text{if } g = \mathcal{G}_{\text{min}}, \end{cases} \quad (32)$$

We absorb the label into the feature: redefine $\mathbf{x} \leftarrow y\mathbf{x}$ so that we may assume $y = +1$ almost surely.

In this convention also note $q^t := 1 - p_y^t = 1 - p_1^t$.

B Features-mediated Spurious Correlation Linearises in Learned Representation

A natural question to ask is: what does the feature representation $\Phi(\mathbf{x})$ look like when the raw-data $\mathbf{x}$ present correlation between the causal and spurious features? Does the learned representation preserves this correlation? If yes, does that correlation linearises or does it take another form?

We construct a controlled experiment on a variant of colored MNIST that introduces a nonlinear, feature-mediated spurious correlation by design, train a nonlinear classifier (MLP), and use sparse autoencoders (SAEs) (Huben et al., 2024) to decompose the learned representation into causal and spurious components. The main finding is that even though the input-space correlation is mediated by a highly nonlinear function (digit class, which is a nonlinear function of raw pixels), the representation-space relationship between the causal and spurious SAE features is approximately linear within each group, therefore aligning with the linear feature-mediated model assumed in our theory.

All the codes are available in the supplementary files.

B.1 Experimental Methodology

Dataset construction. We build on the MNIST dataset ($N = 70,000$ images) and introduce a feature-mediated spurious correlation through the background color. Each image is assigned a binary label $y = \mathbf{1}_{\text{digit} \geq 5}$ and, independently of y , is placed into one of two groups:

$$g = \begin{cases} 0 \text{ (majority)} & \text{with probability } 1 - \varepsilon, \\ 1 \text{ (minority)} & \text{with probability } \varepsilon, \end{cases} \quad (33)$$

with $\varepsilon = 0.1$. The digit foreground is rendered in grayscale. The background color and intensity are then set as follows:

- *Majority (maroon background):* intensity = $\delta + (1 - \delta) \cdot (\text{digit_class} + 1) / 10$, increasing with digit class.
- *Minority (teal background):* intensity = $\delta + (1 - \delta) \cdot (1 - (\text{digit_class} + 1) / 10)$, decreasing with digit class.

Here $\delta = 0.15$ is a floor that prevents pure-black backgrounds. The intensity is monotonically determined by the digit class, but the digit class is a highly nonlinear function of the raw pixel values $\mathbf{r}$. Figure 2 illustrates the resulting images.

This construction produces a feature-mediated spurious correlation: the spurious feature $\mathbf{s}$ (background intensity) correlates with the label y through the causal feature $\mathbf{r}$ (through the digit class). That correlations takes two forms depending on the group: in the majority, higher digits produce brighter backgrounds; in the minority, higher digits produce darker backgrounds. This mirrors the $A \neq B$ structure of Definition 2.3, but with a nonlinear input-space map because the map from raw pixels to digit class is highly nonlinear.

Model and representation extraction. We train a 4-layer MLP, binary cross-entropy loss, Adam optimiser, learning rate 10^{-3} and extract the penultimate-layer activations $\phi(x) \in \mathbb{R}^{128}$ at training checkpoints (epochs 1–5). Early stopping stopped the training after 5 epochs, preventing overfitting.

Sparse autoencoder decomposition. To decompose $\phi(x)$ into causal and spurious components, we train a sparse autoencoder (SAE) on the extracted representations at each epoch. Following Huben et al. (2024), we learned a dictionary matrix $M \in \mathbb{R}^{d_h \times 128}$ with $d_h = 512$, four times the feature space dimension of $\phi(x)$ and with an L_1 regularisation coefficient $\alpha = 0.03$. All SAE achieves $> 99.9\%$ variance explained with an average of ~ 20 active features per sample.

Furthermore, writing

$$c = \text{ReLU}(M \phi(x) + b), \quad \hat{\phi} = M^\top c, \quad (34)$$

then c_k represents the feature f_k (whose direction is the k -th row of M) activation on instance $\phi(x)$.

499 **Feature classification via group-conditional correlations.** Each SAE feature f_k is classified as
500 causal, spurious, or weak/dead using group-conditional correlations with the digit class $\mathbf{d}$:

$$\rho_0(k) = \text{corr}(c_k, \mathbf{d} \mid g=0), \quad \rho_1(k) = \text{corr}(c_k, \mathbf{d} \mid g=1). \quad (35)$$

501 Within the majority ($g=0$), background intensity increases with $\mathbf{d}$, so both **r**-features and **s**-features
502 have $\rho_0 > 0$. Within the minority ($g=1$), background intensity *decreases* with $\mathbf{d}$, so an **r**-feature
503 (tracking digit shape, invariant to color) retains $\rho_1 > 0$, while an **s**-feature (tracking background
504 color) has $\rho_1 < 0$. This yields the classification rule:

$$f_k \text{ is } \begin{cases} \text{r-type} & \text{if } \rho_0(k) \rho_1(k) > 0 \text{ and } \min(|\rho_0(k)|, |\rho_1(k)|) \geq \tau, \\ \text{s-type} & \text{if } \rho_0(k) \rho_1(k) < 0 \text{ and } \min(|\rho_0(k)|, |\rho_1(k)|) \geq \tau, \\ \text{weak} & \text{otherwise,} \end{cases} \quad (36)$$

505 with threshold $\tau = 0.2$. Figure 6 shows the (ρ_0, ρ_1) scatter for a representative epoch. The **r**-features
506 (blue) cluster along the diagonal $\rho_0 = \rho_1$, while the **s**-features (red) appear in the off-diagonal
507 quadrants ($\rho_0 = -\rho_1$), confirming that the classification cleanly separates the two types.

508 **Linearity of the $\phi_r \rightarrow \phi_s$ map.** Let $\phi_r = (c_k)_{k \in \mathcal{R}}$ and $\phi_s = (c_k)_{k \in \mathcal{S}}$ denote the subvectors of
509 SAE activations c corresponding to **r**-type and **s**-type features respectively. We fit a Ridge regression
510 $\phi_r \rightarrow \phi_s$ within each group and measure R^2 (variance-weighted, multi-output). Table 2 reports the
511 results across all five checkpoints.

512 **Null-distribution control.** To verify that the R^2 values reflect genuine structure rather than dimen-
513 sional artefacts, we perform a permutation test: the rows of ϕ_s are shuffled (as a block, preserving the
514 internal correlation structure among **s**-features) within each group, and the Ridge regression is refit.
515 This is repeated 1000 times to obtain a null baseline. The null R^2 values are three to four orders of
516 magnitude below the observed values (Table 2), with all empirical p -values below 0.001.

Table 2: SAE decomposition and linearity across training epochs ($\alpha_{\text{SAE}} = 0.03$, $\tau = 0.2$). n_r ,
 n_s : number of classified **r**- and **s**-features. R^2_{maj} , R^2_{min} : coefficient of determination of the Ridge
regression $\phi_r \rightarrow \phi_s$ within each group. R^2_{null} : mean R^2 under 1000 random permutations of ϕ_s
(within-group), establishing the null baseline.

Epoch	n_r	n_s	R^2_{maj}	$R^2_{\text{null, maj}}$	R^2_{min}	$R^2_{\text{null, min}}$	Avg. active
1	8	1	0.669	0.0001	0.417	0.001	16.4
2	19	13	0.792	0.0002	0.977	0.002	26.7
3	8	5	0.415	0.0002	0.998	0.001	20.2
4	16	6	0.382	0.0005	0.998	0.001	20.7
5	14	14	0.777	0.0003	0.997	0.001	21.8

517 B.2 Results and interpretation

518 **Minority group: near-perfect linearity.** From epoch 2 onward, $R^2_{\text{min}} \geq 0.977$, indicating that
519 the linear map $\phi_r \rightarrow \phi_s$ explains essentially all the variance of the spurious features within the
520 minority group. Furthermore we observe that even though the input-space relationship between digit
521 identity and background intensity is highly nonlinear, the MLP’s learned representation linearises
522 this relationship almost perfectly.

523 **Majority group: weaker but significant coupling.** The majority R^2_{maj} is lower and more variable
524 across epochs (0.38–0.79), but remains orders of magnitude above the null baseline. This still
525 supports the assumption of correlation linearisation in the learned representation feature space.

526 Figure 3 displays the first principal component of ϕ_r against the first principal component of ϕ_s for
527 all five epochs, with points colored by group. Two consistent patterns are visible across all epochs:
528 (i) The *minority* (teal) points form a tight, approximately linear cloud with a clear slope. (ii) The
529 *majority* (maroon) points are concentrated near $\phi_s \approx 0$ (flat), indicating that the **s**-features carry little
530 additional information beyond what is already captured by the **r**-features.

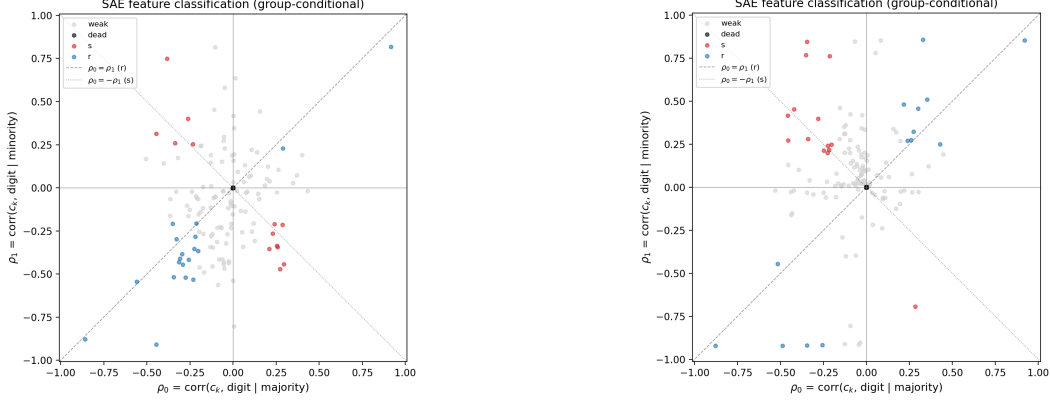


Figure 6: SAE feature classification at epoch 2 (left) and epoch 5 (right). Each point is one SAE feature; axes are group-conditional correlations ρ_0 (majority) and ρ_1 (minority) with digit class. Blue ($\mathbf{r}$ -features) cluster along $\rho_0 = \rho_1$; red ($\mathbf{s}$ -features) in the off-diagonal quadrants.

Conclusion. The colored-MNIST experiment demonstrates that the linear feature-mediated model $\mathbf{s} = A_g \mathbf{r} + \xi$ can emerge in a learned representation even when the input-space relationship is nonlinear. The key mechanism is that the MLP, through gradient descent, learns a representation in which the nonlinear digit-class information is encoded as a set of approximately linear features. This finding goes further than our initial assumption: rather than merely *preserving* a pre-existing linear relationship, the representation actively *linearises* a nonlinear one.

C Scope of the Isotropic Regime.

Our sharpest rates (Theorem 5.3) require the *isotropic regime* (Definition 5.2), in which $\mathbf{v}$ is a common eigenvector of the four alignment products $A^\top A$, $A^\top B$, $B^\top B$, $B^\top A$. A natural question is whether this condition is genuinely restrictive. We show that the class of operator pairs (A, B) to which Theorem 5.3 applies is strictly larger than those that are already isotropic.

Because the noise ξ depends on the margin variable $Z = \mathbf{v} \cdot \mathbf{r}$ (see Eq (4)), one may decompose any pair (A, B) as

$$A = A' + \mathbf{u}_A \mathbf{v}^\top, \quad B = B' + \mathbf{u}_B \mathbf{v}^\top, \quad (37)$$

for vectors $\mathbf{u}_A, \mathbf{u}_B \in \mathbb{R}^{d_s}$, and treat the residuals $\xi'_g := Z \mathbf{u}_g$ as additional noise. By construction ξ'_g depends only on Z , so the conditional independence $\xi' \perp\!\!\!\perp (\mathbf{r}_\perp, y) \mid Z$ holds automatically, and the total noise $\xi + \xi'$ remains compatible with our setup.

For (A', B') to be isotropic, one needs $A'^\top C' \mathbf{v} \propto \mathbf{v}$ for each $C' \in \{A', B'\}$ and similarly with $B'^\top$:

$$A' \mathbf{v}, B' \mathbf{v} \in \mathcal{K} := \ker(\bar{\Pi}_{\mathbf{v}} \circ A^\top) \cap \ker(\bar{\Pi}_{\mathbf{v}} \circ B^\top), \quad (38)$$

where $\bar{\Pi}_{\mathbf{v}}$ is a notation introduced in Eq (22). Each map $\bar{\Pi}_{\mathbf{v}} \circ A^\top, \bar{\Pi}_{\mathbf{v}} \circ B^\top : \mathbb{R}^{d_s} \rightarrow \text{Span}\{\mathbf{v}\}^\perp$ has rank at most $d_r - 1$, so $\dim(\mathcal{K}) \geq d_s - 2(d_r - 1)$. The subspace $\mathcal{K}$ is therefore nontrivial whenever $d_s \geq 2d_r - 1$. Finding the closest isotropic pair amounts to minimising

$$\begin{aligned} d_S((A', B'), (A, B)) &= \sup_{r \in \text{supp}(\mathbf{r})} \|(A' - A)r\| + \sup_{r \in \text{supp}(\mathbf{r})} \|(B' - B)r\| \\ &\stackrel{(1)}{=} K(\|\mathbf{u}_A\| + \|\mathbf{u}_B\|) \end{aligned} \quad (39)$$

$$= K(\|A' \mathbf{v} - A \mathbf{v}\| + \|B' \mathbf{v} - B \mathbf{v}\|) \quad \text{such that,}$$

$$\begin{aligned} \inf_{A', B'} d_S((A', B'), (A, B)) &\stackrel{(2)}{\geq} K(\|\Pi_{\mathcal{K}^\perp}(A \mathbf{v})\| + \|\Pi_{\mathcal{K}^\perp}(B \mathbf{v})\|) \\ &=: d^* \quad (= 0 \text{ iff } (A, B) \text{ is already isotropic}), \end{aligned} \quad (40)$$

where (1) uses Eq (37) and (2) uses the condition (38). The resulting eigenvalues are $\mu_{A'} = \|A' \mathbf{v}\|^2$, $\mu_{B'} = \|B' \mathbf{v}\|^2$, and $\mu' = A' \mathbf{v} \cdot B' \mathbf{v}$, and the attractive regime condition (10) is satisfiable whenever

553 $\dim(\mathcal{K}) \geq 2$ (i.e. $d_s \geq 2d_r$), since one can break positive collinearity of $A'\mathbf{v}$ and $B'\mathbf{v}$ by perturbing
 554 within $\mathcal{K}$.
 555 When all these conditions are met, the absorbed noise inflates the amplitude bound to $R' = R +$
 556 $K \max(\|\mathbf{u}_A\|, \|\mathbf{u}_B\|)$ and simultaneously reduces the signal strength $\hat{\psi}'^g \cdot \mathbf{v} = \hat{\psi}^g \cdot \mathbf{v} - \mathbf{u}_g \cdot \hat{\mathbf{w}}_s$
 557 (see Eq (103)), so the noise condition (106) may fail for large $\|\mathbf{u}_g\|$.
 558 In summary, the isotropic regime is less restrictive than it may first appear: any (A, B) with $d_s \geq 2d_r$
 559 admits an isotropic reparametrisation provided the minimum displacement d^* is small enough for
 560 the noise condition to hold. In the overparameterised setting $d_s \gg d_r$ that is typical of learned
 561 representations (such as the last-layer retraining paradigm), this dimensional requirement is easily
 562 met.

D Relaxed Version of Theorem 5.3

We make the following very natural extra assumptions to make the proof simpler and we explain in Section F how to lift them.

Assumption D.1 (Zero-mean Component). Write $\mathbf{r}_\perp := \mathbf{r} - Z\mathbf{v} \in \text{Span}\{\mathbf{v}\}^\perp$, the component of $\mathbf{r}$ perpendicular to $\mathbf{v}$. Then,

$$\mathbb{E}[\mathbf{r}_\perp \mid Z, g] = 0 \quad \text{for all } g \in \{\mathcal{G}_{\text{maj}}, \mathcal{G}_{\text{min}}\}. \quad (41)$$

Assumption D.2 (Initial condition). Writing the projected parameters

$$\psi_t := w_r^t + A^\top w_s^t \in \mathbb{R}^{d_r}, \quad \phi_t := w_r^t + B^\top w_s^t \in \mathbb{R}^{d_r}, \quad (42)$$

the model initialization satisfies

$$\mathbf{v} \cdot \psi_0 \geq 1, \quad \mathbf{v} \cdot \phi_0 \geq 1, \quad \bar{\Pi}_\mathbf{v} \psi_0 = 0, \quad \bar{\Pi}_\mathbf{v} \phi_0 = 0. \quad (43)$$

Special extra-condition $\star$ In Section E.4, we further assume $(\mathbf{v} \cdot \psi_0, \mathbf{v} \cdot \phi_0) \in \mathcal{T}$, where $\mathcal{T}$ is introduced in Eq (74).

Assumption D.3 (Noise-free case). For clarity, we analyze the dynamics without perturbation, setting $\xi = 0$.

Assumption D.1 is simply made to make the perpendicular gradient vanish (Lemma E.1). In Section F.1 we discuss how to lift this assumption.

Assumption D.2 is very reasonable in light of Remark E.4 and guarantees numerical stability as required in the proof of Propositions H.3 and H.4. The orthogonality condition $\bar{\Pi}_\mathbf{v} \psi_0 = \bar{\Pi}_\mathbf{v} \phi_0 = 0$ goes hand in hand with condition (41). A discussion on how to lift this assumption is presented in Section F.1. The special condition $\star$ is very mild as well as explained in Remark E.6.

Assumption D.3 is introduced solely to simplify the computations in the proof. This assumption might seem to overly simplify the problem, but it is still aligned with formalism of spurious correlation Section 2. In fact, although it makes $\mathbf{s}$ a deterministic function of $\mathbf{r}$, the function itself is not linear ($A \neq B$). Therefore, a linear model that learns to rely on $\mathbf{s}$ would fall in a spurious pattern trap. A discussion on how to remove this assumption is presented in Section F.2.

We now restate the theorem in the isotropic regime.

Theorem D.4 (Isotropic regime). Assume the spurious correlation is isotropic 5.2, and that assumptions D.1-D.3 are verified. Without loss of generality, assume $\tilde{\gamma}_{\text{maj}} = 1$ and $\tilde{\gamma}_{\text{min}} \geq 1$ and define

$$\alpha := \frac{\tilde{\gamma}_{\text{min}}(1 + \mu)}{1 + \mu_A}. \quad (44)$$

Then

(i) If $\alpha < 1$, there exist two positive constants κ_{maj} and κ_{min} independent on ε such that:

$$\mathbb{E}_{\mathcal{G}_{\text{maj}}}[1 - p_y^t] \underset{t \rightarrow \infty}{\sim} \frac{\kappa_{\text{maj}}}{(1 - \varepsilon) z_t}, \quad \kappa_{\text{maj}} = \frac{\tilde{\gamma}_{\text{min}}(1 + \mu_B) - (1 + \mu)}{\tilde{\gamma}_{\text{min}} \Sigma}, \quad (45)$$

$$\mathbb{E}_{\mathcal{G}_{\text{min}}}[1 - p_y^t] \underset{t \rightarrow \infty}{\sim} \frac{\kappa_{\text{min}}}{\varepsilon z_t}, \quad \kappa_{\text{min}} = \frac{(1 + \mu_A) - \tilde{\gamma}_{\text{min}}(1 + \mu)}{\tilde{\gamma}_{\text{min}}^2 \Sigma}, \quad (46)$$

with Σ defined in Eq (31).

(ii) If $\alpha \geq 1$, the classification error decay at different rates:

$$\mathbb{E}_{\mathcal{G}_{\text{maj}}}[1 - p_y^t] \underset{t \rightarrow \infty}{\sim} \frac{1}{(1 - \varepsilon)(1 + \mu_A) z_t}, \quad (47)$$

$$\mathbb{E}_{\mathcal{G}_{\text{min}}}[1 - p_y^t] = \Theta\left(\frac{(\ln z_t)^{\alpha-1}}{z_t^\alpha}\right). \quad (48)$$

(iii) Furthermore, these different rate decay are continuous with respect to $\tilde{\gamma}_{\text{min}}$ at the phase transition point $\alpha = 1$.

587 **Continuity at the phase transition** At the threshold $\alpha = 1$ (i.e. $\tilde{\gamma}_{\min} = (1+\mu_A)/(1+\mu)$):

- 588 • $\kappa_{\min}(\tilde{\gamma}_{\min}) \rightarrow 0$: the factor of the minority's $1/z_t$ rate vanishes continuously.
- 589 • $\kappa_{\text{maj}}(\tilde{\gamma}_{\min}) \rightarrow 1/(1+\mu_A)$: the majority factor continuously matches the Regime (ii) rate (47).
- 590 • The minority rate $(\ln z_t)^{\alpha-1}/z_t^\alpha$ with $\alpha \rightarrow 1^+$ approaches $1/z_t$, matching Regime (i) rate (46).

591 E Proof of Theorem D.4

592 The proof proceeds in four stages. First, we establish geometric properties of the gradient updates
 593 (Lemma E.1 and Proposition E.2): the weight symmetry inherited from the loss and the conditional
 594 zero-mean structure of $\mathbf{r}_\perp$ together confine the parameter trajectory to the one-dimensional subspace
 595 $\text{Span}\{\mathbf{v}\}$, reducing the high-dimensional dynamics to the scalar projections $a_t := \mathbf{v} \cdot \psi_t$ and $b_t := \mathbf{v} \cdot \phi_t$.
 596 Second, we derive the coupled recurrences satisfied by (a_t, b_t) (Proposition E.3), identify them as an
 597 Euler discretisation of a 2D autonomous ODE, and control the approximation error (Lemma H.6).
 598 Third, an asymptotic synchronisation argument shows that the gap $a(z) - b(z)$ converges to an
 599 explicit constant d determined by the eigenvalues μ_A, μ_B, μ and the group proportions, effectively
 600 reducing the 2D system to a single perturbed 1D equation. Finally, Laplace-method asymptotics and
 601 the solution of the resulting ODE (Proposition H.4) combine to yield the group-wise rates κ_{maj} and
 602 $\kappa_{\min}$.

603 E.1 Auxiliary Results

Lemma E.1 (Perpendicular Gradient Vanishes). *For all $t \geq 0$ and $g \in \{\mathcal{G}_{\text{maj}}, \mathcal{G}_{\min}\}$,*

$$\mathbb{E}_g[q^t \mathbf{r}_\perp] = 0, \quad \text{where } \mathbf{r}_\perp := \mathbf{r} - \Pi_{\mathbf{v}} \mathbf{r} = \bar{\Pi}_{\mathbf{v}} \mathbf{r}.$$

605 *Proof.* We proceed step by step.

606 *Joint induction.* We prove the following two facts simultaneously by induction on $t \geq 0$:

607 (I) $\psi_t \cdot \mathbf{r}_\perp = 0$ and $\phi_t \cdot \mathbf{r}_\perp = 0$ almost surely.

608 (II) $\mathbb{E}_g[q^t \mathbf{r}_\perp] = 0$ for $g \in \{\mathcal{G}_{\text{maj}}, \mathcal{G}_{\min}\}$.

609 *Base case ($t = 0$):* (I) holds by Assumption D.2. (II) then follows from the cancellation argument
 610 below, applied with $\bar{\Pi}_{\mathbf{v}} \psi_0 = \bar{\Pi}_{\mathbf{v}} \phi_0 = 0$.

611 *Inductive step:* Assuming (I) at time t , the argument below establishes (II) at time t . Proposition E.2
 612 then uses (II) at time t to establish (I) at time $t + 1$.

613 The model's predicted logit is

$$w^t \cdot \mathbf{x} = (w_{\mathbf{r}}^t) \cdot \mathbf{r} + (w_{\mathbf{s}}^t) \cdot \mathbf{s} = \begin{cases} (w_{\mathbf{r}}^t + A^\top w_{\mathbf{s}}^t) \cdot \mathbf{r} = \psi_t \cdot \mathbf{r} & \text{on } \mathcal{G}_{\text{maj}}, \\ (w_{\mathbf{r}}^t + B^\top w_{\mathbf{s}}^t) \cdot \mathbf{r} = \phi_t \cdot \mathbf{r} & \text{on } \mathcal{G}_{\min}, \end{cases} \quad (49)$$

614 so that $q^t = \sigma(-\psi_t \cdot \mathbf{r})$ on $\mathcal{G}_{\text{maj}}$, and $q^t = \sigma(-\phi_t \cdot \mathbf{r})$ on $\mathcal{G}_{\min}$. Expanding $\mathbf{r}$ with $\mathbf{r} = Z\mathbf{v} + \mathbf{r}_\perp$ (see
 615 Assumption D.1):

$$\text{On } \mathcal{G}_{\text{maj}} : q^t = \sigma(-\psi_t \cdot \mathbf{r}) = \sigma(-\mathbf{v} \cdot \psi_t Z - \mathbf{r}_\perp \cdot \psi_t) = \sigma(-\mathbf{v} \cdot \psi_t Z), \quad (50)$$

$$\text{On } \mathcal{G}_{\min} : q^t = \sigma(-\phi_t \cdot \mathbf{r}) = \sigma(-\mathbf{v} \cdot \phi_t Z - \mathbf{r}_\perp \cdot \phi_t) = \sigma(-\mathbf{v} \cdot \phi_t Z), \quad (51)$$

616 where we used the inductive hypothesis (I), $\mathbf{r}_\perp \cdot \psi_t = 0$ and $\mathbf{r}_\perp \cdot \phi_t = 0$ almost surely to remove the
 617 $\mathbf{r}_\perp$ -terms. Therefore, q^t becomes a function of Z and by the total expectation property we have:

$$\mathbb{E}_{\mathcal{G}_{\min}}[q^t \mathbf{r}_\perp] = \mathbb{E}_{\mathcal{G}_{\min}} \left[\sigma(-\mathbf{v} \cdot \phi_t Z) \cdot \underbrace{\mathbb{E}[\mathbf{r}_\perp \mid Z, \mathcal{G}_{\min}]}_{= 0 \text{ by Eq (41)} } \right] = 0.$$

618 The same argument with ψ_t instead of ϕ_t yields $\mathbb{E}_{\mathcal{G}_{\text{maj}}}[q^t \mathbf{r}_\perp] = 0$. □

Proposition E.2 (Parameter Updates Stay in $\text{Span}\{\mathbf{v}\}$). *For all $t \geq 0$,*

$$\psi_t \cdot \mathbf{r}_\perp = 0 \quad \text{and} \quad \phi_t \cdot \mathbf{r}_\perp = 0 \quad \text{almost surely.} \quad (52)$$

619

620 *Proof.* Let's write $\psi_{t+1} - \psi_t = -h_t(\nabla_{w_r} \mathcal{L}^t + A^\top \nabla_{w_s} \mathcal{L}^t)$.

621 On one hand, using Lemma E.1 and $\mathbf{r} = Z\mathbf{v} + \mathbf{r}_\perp$:

$$\nabla_{w_r} \mathcal{L}^t = -\mathbb{E}[q^t \mathbf{r}] = -\mathbb{E}[Z q^t] \mathbf{v} - \underbrace{\mathbb{E}[q^t \mathbf{r}_\perp]}_{=0} = -\mathbb{E}[Z q^t] \mathbf{v}. \quad (53)$$

622 On the other hand, by Eq (5):

$$\begin{aligned} \nabla_{w_s} \mathcal{L}^t &= -\mathbb{E}[q^t \mathbf{s}] \\ &= -\left((1-\varepsilon) A \mathbb{E}_{\mathcal{G}_{\text{maj}}}[q^t \mathbf{r}] + \varepsilon B \mathbb{E}_{\mathcal{G}_{\text{min}}}[q^t \mathbf{r}] \right) \\ &= -\left((1-\varepsilon) \mathbb{E}_{\mathcal{G}_{\text{maj}}}[Z q^t] A \mathbf{v} + \varepsilon \mathbb{E}_{\mathcal{G}_{\text{min}}}[Z q^t] B \mathbf{v} \right) \quad (\text{by Lemma E.1}) \end{aligned} \quad (54)$$

623 Since $\bar{\Pi}_\mathbf{v} \psi_0 = 0$ (Assumption D.2), and the update $\psi_{t+1} - \psi_t = -h_t(\nabla_{w_r} \mathcal{L}^t + A^\top \nabla_{w_s} \mathcal{L}^t)$ lies
624 in $\text{Span}\{\mathbf{v}\} + A^\top \text{Span}\{A\mathbf{v}\} + A^\top \text{Span}\{B\mathbf{v}\} = \text{Span}\{\mathbf{v}\}$ (under isotropic regime 5.2), it follows
625 by induction that $\psi_t \cdot \mathbf{r}_\perp = 0$ almost surely for all t . The same argument gives $\phi_t \cdot \mathbf{r}_\perp = 0$ almost
626 surely. $\square$

627 E.2 Discrete and Continuous Error Dynamics

628 Now we can write the scalar equations that characterize the group-wise classification errors. We show
629 that these discrete recurrence equation can be seen as approximations of a 2D-ODE.

630 Define the group-specific drift functions

$$\mathcal{I}_{\text{maj}}(x) := \int_1^K z \sigma(-xz) \lambda_{\text{maj}}(z) dz, \quad \mathcal{I}_{\text{min}}(x) := \int_{\tilde{\gamma}_{\text{min}}}^K z \sigma(-xz) \lambda_{\text{min}}(z) dz, \quad (55)$$

631 where λ_{maj} and λ_{min} are the group densities from Assumption D.1. Both functions are decreasing.

632 Define the scalar projections

$$a_t := \mathbf{v} \cdot \psi_t, \quad b_t := \mathbf{v} \cdot \phi_t. \quad (56)$$

Proposition E.3 (Coupled Scalar Dynamics). *Under Assumptions D.1-D.3, the scalars a_t and b_t evolve according to the coupled recurrences*

$$a_{t+1} - a_t = h_t(1-\varepsilon)(1+\mu_A) \mathcal{I}_{\text{maj}}(a_t) + h_t \varepsilon (1+\mu) \mathcal{I}_{\text{min}}(b_t), \quad (57)$$

$$b_{t+1} - b_t = h_t \varepsilon (1+\mu_B) \mathcal{I}_{\text{min}}(b_t) + h_t(1-\varepsilon)(1+\mu) \mathcal{I}_{\text{maj}}(a_t). \quad (58)$$

Moreover, under condition 10 the two sequences are strictly increasing and for all $t \geq 0$,

$$a_t \geq 1 \quad \text{and} \quad b_t \geq 1. \quad (59)$$

633

634 *Proof.* By combining Eqs (53)-(54),

$$\begin{aligned} \psi_{t+1} - \psi_t &= -h_t(\nabla_{w_r} \mathcal{L}^t + A^\top \nabla_{w_s} \mathcal{L}^t) \\ &= h_t \mathbb{E}[Z q^t] \mathbf{v} + (1-\varepsilon) h_t \mathbb{E}_{\mathcal{G}_{\text{maj}}}[Z q^t] A^\top A \mathbf{v} + \varepsilon h_t \mathbb{E}_{\mathcal{G}_{\text{min}}}[Z q^t] A^\top B \mathbf{v} \\ &= h_t(1-\varepsilon) \mathbb{E}_{\mathcal{G}_{\text{maj}}}[Z q^t] (\mathbf{v} + A^\top A \mathbf{v}) + \varepsilon h_t \mathbb{E}_{\mathcal{G}_{\text{min}}}[Z q^t] (\mathbf{v} + A^\top B \mathbf{v}) \\ &= h_t(1-\varepsilon)(1+\mu_A) \mathbb{E}_{\mathcal{G}_{\text{maj}}}[Z q^t] \mathbf{v} + h_t \varepsilon (1+\mu) \mathbb{E}_{\mathcal{G}_{\text{min}}}[Z q^t] \mathbf{v}, \end{aligned} \quad (60)$$

635 using Eq (29) in the last line. By an identical argument we also get:

$$\phi_{t+1} - \phi_t = h_t \varepsilon (1+\mu_B) \mathbb{E}_{\mathcal{G}_{\text{min}}}[Z q^t] \mathbf{v} + h_t(1-\varepsilon)(1+\mu) \mathbb{E}_{\mathcal{G}_{\text{maj}}}[Z q^t] \mathbf{v}. \quad (61)$$

636 Projecting Eq. (60) onto $\mathbf{v}$:

$$a_{t+1} - a_t = h_t(1 - \varepsilon)(1 + \mu_A) \mathbb{E}_{\mathcal{G}_{\text{maj}}}[Z q^t] + h_t \varepsilon(1 + \mu) \mathbb{E}_{\mathcal{G}_{\text{min}}}[Z q^t]. \quad (62)$$

637 Using Eq (50)-(51), $q^t = \sigma(-a_t Z)$ on $\mathcal{G}_{\text{maj}}$ and $\sigma(-b_t Z)$ on $\mathcal{G}_{\text{min}}$, and since $Z \mid \mathcal{G}_g$ has density λ_g
 638 for $g \in \{\text{min}, \text{maj}\}$ (Assumption D.1) we then have:

$$\begin{aligned} \mathbb{E}_{\mathcal{G}_{\text{maj}}}[Z q^t] &= \mathbb{E}_{\mathcal{G}_{\text{maj}}}[Z \sigma(-a_t Z)] = \int_1^K z \sigma(-a_t z) \lambda_{\text{maj}}(z) dz = \mathcal{I}_{\text{maj}}(a_t) \\ \mathbb{E}_{\mathcal{G}_{\text{min}}}[Z q^t] &= \mathbb{E}_{\mathcal{G}_{\text{min}}}[Z \sigma(-b_t Z)] = \int_{\tilde{\gamma}_{\text{min}}}^K z \sigma(-b_t z) \lambda_{\text{min}}(z) dz = \mathcal{I}_{\text{min}}(b_t). \end{aligned}$$

639 Substituting back:

$$a_{t+1} - a_t = h_t(1 - \varepsilon)(1 + \mu_A) \mathcal{I}_{\text{maj}}(a_t) + h_t \varepsilon(1 + \mu) \mathcal{I}_{\text{min}}(b_t).$$

640 Because $\mathcal{I}_{\text{maj}}(x)$, $\mathcal{I}_{\text{min}}(x) > 0$ for all x , and given that $1 + \mu \geq 0$ by condition 10, and $(a_0, b_0) =$
 641 $(\mathbf{v} \cdot \psi_0, \mathbf{v} \cdot \phi_0) \geq 1$ by Assumption D.2, the sequences (a_t) and (b_t) are strictly increasing and
 642 $a_t, b_t \geq 1$ for all t by induction. $\square$

Remark E.4. We highlight two useful observations:

(i) The sequences $(a_t)_{t \geq 0}$ and $(b_t)_{t \geq 0}$ are strictly increasing and diverge to $+\infty$. In fact, if for instance $(a_t)_{t \geq 0}$ was upper bounded, and if we denote $a_\infty, b_\infty \in \mathbb{R} \cup \{+\infty\}$ their limits, one would get

$$\left((1 - \varepsilon)(1 + \mu_A) \mathcal{I}_{\text{maj}}(a_\infty) + \varepsilon(1 + \mu_B) \mathcal{I}_{\text{min}}(b_\infty) \right) \sum_{t \geq 0} h_t \leq \sum_{t \geq 0} (a_{t+1} - a_t) < \infty.$$

Under condition (21), this is only possible if $a_\infty = b_\infty = +\infty$, which is impossible since we assumed (a_t) was bounded.

(ii) Equations (57)–(58) are explicit Euler schemes for the 2D autonomous ODE

$$(\dot{u}(z), \dot{v}(z)) = F(u, v), \quad (u(0), v(0)) = (a_0, b_0) \quad (63)$$

where the 2D vector field F is defined as

$$F(a, b) := \begin{pmatrix} (1 - \varepsilon)(1 + \mu_A) \mathcal{I}_{\text{maj}}(a) + \varepsilon(1 + \mu) \mathcal{I}_{\text{min}}(b) \\ \varepsilon(1 + \mu_B) \mathcal{I}_{\text{min}}(b) + (1 - \varepsilon)(1 + \mu) \mathcal{I}_{\text{maj}}(a) \end{pmatrix}, \quad (64)$$

Defining $z_t := \sum_{k=0}^{t-1} h_k$, we will approximate $a_t \approx u(z_t)$ and $b_t \approx v(z_t)$ in Lemma H.6.

643

644 From now on we work with the continuous trajectory $(a(\cdot), b(\cdot))$ for the synchronisation analysis;
 645 the discrete dynamics is recovered *a posteriori* via Lemma H.6 and Lemma H.9 in the conclusion
 646 (Section E.4).

647 E.3 ODE Synchronisation

648 Let $\Delta(z) := a(z) - \tilde{\gamma}_{\text{min}} b(z)$ denote the gap between the two trajectories. Subtracting the two
 649 coordinates of F (64) and using α introduced in Eq (44),

$$\begin{aligned} \dot{\Delta} &= (1 - \varepsilon)((1 + \mu_A) - \tilde{\gamma}_{\text{min}}(1 + \mu)) \mathcal{I}_{\text{maj}}(a) - \varepsilon(\tilde{\gamma}_{\text{min}}(1 + \mu_B) - (1 + \mu)) \mathcal{I}_{\text{min}}(b) \\ &= (1 - \varepsilon)(1 + \mu_A)(1 - \alpha) \mathcal{I}_{\text{maj}}(a) - \varepsilon(\tilde{\gamma}_{\text{min}}(1 + \mu_B) - (1 + \mu)) \mathcal{I}_{\text{min}}(b). \end{aligned} \quad (65)$$

650 We seek a constant gap $d \in \mathbb{R}$ for which the leading-order forcing of the gap dynamics cancels on the
 651 manifold $\mathcal{M}_d := \{(a, b) \in \mathbb{R}^2 : a = \tilde{\gamma}_{\text{min}} b + d\}$. $\mathcal{M}_d$ is never invariant at finite time (i.e a solution
 652 initialised in $\mathcal{M}_d$ does not stay inside), however we'd like to keep the trajectory "close" enough,
 653 which we'll rigorously formulate below.

654 As $b \rightarrow +\infty$, and for any scalar d , Lemma H.2 gives

$$\begin{aligned}
\mathcal{I}_{\text{maj}}(\tilde{\gamma}_{\min} b + d) &= \frac{\lambda_{\text{maj}}(1)}{\tilde{\gamma}_{\min} \lambda_{\min}(\tilde{\gamma}_{\min})} e^{-d} \frac{b}{\tilde{\gamma}_{\min} b + d} \mathcal{I}_{\min}(b) (1 + O(1/b)) \\
&= \frac{\lambda_{\text{maj}}(1)}{\tilde{\gamma}_{\min}^2 \lambda_{\min}(\tilde{\gamma}_{\min})} e^{-d} \mathcal{I}_{\min}(b) (1 + O(1/b)),
\end{aligned} \tag{66}$$

so that, when $(a(z), b(z)) \in \mathcal{M}_d$ i.e. $a(z) = \tilde{\gamma}_{\min} b(z) + d$,

$$\begin{aligned}
\dot{\Delta}(z) &= \mathcal{I}_{\min}(b(z)) \left[(1 - \varepsilon)(1 + \mu_A)(1 - \alpha) \frac{\lambda_{\text{maj}}(1)}{\tilde{\gamma}_{\min}^2 \lambda_{\min}(\tilde{\gamma}_{\min})} e^{-d} \right. \\
&\quad \left. - \varepsilon (\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu)) \right] + O(\mathcal{I}_{\min}(b(z))/|b(z)|).
\end{aligned} \tag{67}$$

We choose d so that the bracket in (67) vanishes, making $\dot{\Delta}(z) = O(\mathcal{I}_{\min}(b(z))/|b(z)|)$. The resulting equation for d admits a unique real solution only when $\alpha < 1$:

$$\begin{aligned}
e^{-d} &= \frac{\varepsilon}{1 - \varepsilon} \frac{\tilde{\gamma}_{\min}^2 \lambda_{\min}(\tilde{\gamma}_{\min})}{\lambda_{\text{maj}}(1)} \frac{\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu)}{(1 + \mu_A)(1 - \alpha)}, \\
d &= -\ln \left(\frac{\varepsilon}{1 - \varepsilon} \frac{\tilde{\gamma}_{\min}^2 \lambda_{\min}(\tilde{\gamma}_{\min})}{\lambda_{\text{maj}}(1)} \frac{\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu)}{(1 + \mu_A)(1 - \alpha)} \right).
\end{aligned} \tag{68}$$

The complementary case $\alpha \geq 1$ will be covered in Section E.5.

Having fixed the synchronisation constant d , we define the scalar gap $\eta(z) := a(z) - \tilde{\gamma}_{\min} b(z) - d$, which satisfies

$$\dot{\eta} = (1 - \varepsilon)(1 + \mu_A)(1 - \alpha) \mathcal{I}_{\text{maj}}(\tilde{\gamma}_{\min} b + d + \eta) - \varepsilon (\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu)) \mathcal{I}_{\min}(b), \tag{69}$$

by substituting $a = \tilde{\gamma}_{\min} b + d + \eta$ into (65). Thus the 2D ODE reduces to the scalar non-autonomous ODE (69) for η driven by the exogenous signal $b(z)$.

Lemma E.5 (Asymptotic synchronisation). *Under condition 10, there exist $\eta_\star > 0$ and $b_\star > 0$ such that every trajectory (a, b) of (63) satisfying $|a(z_0) - \tilde{\gamma}_{\min} b(z_0) - d| \leq \eta_\star$ and $b(z_0) \geq b_\star$ for some $z_0 \geq 0$ verifies*

$$|\Delta(z) - d| = O\left(\frac{1}{\ln z}\right) \quad \text{as } z \rightarrow +\infty.$$

In particular, $\Delta(z) \rightarrow d$.

Proof. Recall $\eta := \Delta - d = a - \tilde{\gamma}_{\min} b - d$. Writing

$$\begin{aligned}
\mathcal{I}_{\text{maj}}(\tilde{\gamma}_{\min} b + d + \eta) &= \frac{\lambda_{\text{maj}}(1)}{\tilde{\gamma}_{\min}^2 \lambda_{\min}(\tilde{\gamma}_{\min})} e^{-d-\eta} (1 + s(z, \eta)) \mathcal{I}_{\min}(b), \\
s(z, \eta) &:= \frac{\tilde{\gamma}_{\min} b}{\tilde{\gamma}_{\min} b + d + \eta} \cdot \frac{1 + r_{\text{maj}}(\tilde{\gamma}_{\min} b + d + \eta)}{1 + r_{\min}(b)} - 1,
\end{aligned} \tag{70}$$

with $r_g(x) := \mathcal{I}_g(x)/(\tilde{\gamma}_g \lambda_g(\tilde{\gamma}_g) e^{-\tilde{\gamma}_g x}/x) - 1$, and using Eq (68), Eq (69) becomes

$$\dot{\eta} = \varepsilon (\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu)) \mathcal{I}_{\min}(b) [e^{-\eta} (1 + s(z, \eta)) - 1]. \tag{71}$$

Proposition H.2 gives $r_g(x) = \nu_g/x + O(1/x^2)$, so that we have the expansions

$$\begin{aligned}
\frac{\tilde{\gamma}_{\min} b}{\tilde{\gamma}_{\min} b + d + \eta} &= 1 - \frac{d + \eta}{\tilde{\gamma}_{\min} b} + O\left(\frac{1}{b^2}\right) \\
\frac{1 + r_{\text{maj}}(\tilde{\gamma}_{\min} b + d + \eta)}{1 + r_{\min}(b)} &= 1 + \frac{\nu_{\text{maj}}/\tilde{\gamma}_{\min} - \nu_{\min}}{b} + O\left(\frac{1}{b^2}\right),
\end{aligned}$$

uniformly for $|\eta| \leq 1$. Therefore,

$$s(z, \eta) = \frac{1}{b(z)} \left(\nu_{\text{maj}} / \tilde{\gamma}_{\min} - \nu_{\min} - (d + \eta) / \tilde{\gamma}_{\min} \right) + O\left(\frac{1}{b(z)^2}\right), \quad |s(z, \eta)| \leq \frac{K_s}{b(z)} \quad (72)$$

668 for some constant $K_s > 0$ and all z with $b(z)$ large enough and $|\eta| \leq 1$. Moreover

$$|s(z, \eta) - s(z, 0)| \leq K'_s |\eta| / b(z) \quad \text{for some } K'_s > 0. \quad (73)$$

669 From this point, the proof proceeds in three steps: (i) forward invariance of a tube $\{|\eta| \leq \eta_*, b \geq b_*\}$,
 670 (ii) the two-sided asymptotic for b via Proposition H.3, and (iii) the Comparison-Lemma argument
 671 yielding the $1/\ln z$ rate.

672 (i) *Forward invariance of the tube and asymptotic of b .* Choose $\eta_* \in (0, \frac{1}{2})$. Using $|s| \leq K_s/b$
 673 from (72), pick $b_* > 0$ large enough that $\frac{K_s}{b_*} \leq \sigma$ where $\sigma := \frac{1-e^{-\eta_*}}{1+e^{-\eta_*}} \in (0, 1)$. Define the tube

$$\mathcal{T} := \{(a, b) \in \mathbb{R}^2 : |a - \tilde{\gamma}_{\min} b - d| \leq \eta_*, b \geq b_*\}. \quad (74)$$

674 On the boundary $\eta = \eta_*, b \geq b_*$, $e^{-\eta_*}(1+s) - 1 \leq e^{-\eta_*}(1+\sigma) - 1 = \frac{2e^{-\eta_*}}{1+e^{-\eta_*}} - 1 = -\sigma < 0$.

675 Since $\mathcal{I}_{\min}(b) > 0$ and $\varepsilon(\tilde{\gamma}_{\min}(1+\mu_B) - (1+\mu)) > 0$, (71) yields $\dot{\eta} < 0$ on this boundary.

676 Symmetrically, on $\eta = -\eta_*, b \geq b_*$, $e^{\eta_*}(1+s) - 1 \geq e^{\eta_*}(1-\sigma) - 1 = \frac{2}{1+e^{-\eta_*}} - 1 = \sigma > 0$,
 677 and $\dot{\eta} > 0$.

678 Both pushes are inward, so the tube $\mathcal{T}$ is forward invariant (i.e when a trajectory starts inside,
 679 it remains inside) as long as $b(z)$ stays over b_* ; the latter is guaranteed because $b(z) \rightarrow +\infty$
 680 monotonically (by the same argument as in Remark E.4).

681 (ii) *Proposition H.3 requirements.* The second line of Eq (63) can be rewritten,

$$\dot{b}(z) = c(z) \mathcal{I}_{\min}(b(z)), \quad c(z) := \varepsilon(1+\mu_B) + (1-\varepsilon)(1+\mu) \frac{\mathcal{I}_{\text{maj}}(a(z))}{\mathcal{I}_{\min}(b(z))}. \quad (75)$$

682 We show that for every sufficiently small $\eta_* > 0$, there exist constants $0 < c_-(\eta_*) \leq c_+(\eta_*) < \infty$
 683 (independent of z) with $c(z) \in [c_-(\eta_*), c_+(\eta_*)]$ for all z large enough.

684 Inside $\mathcal{T}$ we have $a = \tilde{\gamma}_{\min} b + d + \eta$ with $\eta = \eta(z) \in [-\eta_*, \eta_*]$. Lemma H.2 combined with Eq (68)
 685 yields

$$\begin{aligned} \frac{\mathcal{I}_{\text{maj}}(a(z))}{\mathcal{I}_{\min}(b(z))} &\sim \frac{\lambda_{\text{maj}}(1)}{\tilde{\gamma}_{\min} \lambda_{\min}(\tilde{\gamma}_{\min})} e^{-d-\eta(z)} \frac{b(z)}{\tilde{\gamma}_{\min} b(z) + d + \eta(z)} \\ &\sim \underbrace{\frac{\varepsilon(\tilde{\gamma}_{\min}(1+\mu_B) - (1+\mu))}{(1-\varepsilon)(1+\mu_A)(1-\alpha)}}_{=: R_0} e^{-\eta(z)} \end{aligned} \quad (76)$$

686 since $|\eta(z)| \leq \eta_*$ and $b(z) \rightarrow +\infty$. Hence there exist $z_a(\eta_*) \geq z_0$ and constants

$$\begin{aligned} R_{\pm}(\eta_*) &:= R_0 e^{\pm \eta_*} (1 \pm o_{\eta_*}(1)), \quad \text{s.t.} \quad R_-(\eta_*) \leq \frac{\mathcal{I}_{\text{maj}}(a(z))}{\mathcal{I}_{\min}(b(z))} \leq R_+(\eta_*) \\ c_{\pm}(\eta_*) &:= \varepsilon(1+\mu_B) + (1-\varepsilon)(1+\mu) R_{\pm}(\eta_*), \quad c_-(\eta_*) \leq c(z) \leq c_+(\eta_*) \quad \forall z \geq z_a. \end{aligned} \quad (77)$$

687 Both $R_{\pm}(\eta_*) \rightarrow R_0$ as $\eta_* \rightarrow 0$, so $c_{\pm}(\eta_*)$ share the common limit

$$\begin{aligned} c_0 &:= \varepsilon(1+\mu_B) + (1-\varepsilon)(1+\mu) R_0 = \varepsilon \left((1+\mu_B) + (1+\mu) \frac{\tilde{\gamma}_{\min}(1+\mu_B) - (1+\mu)}{(1+\mu_A)(1-\alpha)} \right) \\ &= \frac{\varepsilon}{(1+\mu_A)(1-\alpha)} \left((1-\alpha)(1+\mu_A)(1+\mu_B) + \alpha(1+\mu_A)(1+\mu_B) - (1+\mu)^2 \right) \\ &= \frac{\varepsilon \Sigma}{(1+\mu_A)(1-\alpha)}, \end{aligned} \quad (78)$$

with Σ defined in Eq (31). The quantity c_0 is strictly positive with condition 10. By continuity in η_* of $c_{\pm}(\eta_*)$, fix $\eta_* > 0$ small enough that $c_{-}(\eta_*) \geq c_0/2 > 0$. With this fixed η_* , Eq (77) gives $c(z) \in [c_{-}(\eta_*), c_{+}(\eta_*)] \subset (0, \infty)$ for all $z \geq z_a$, which is the hypothesis required by Proposition H.3.

(iii) *Comparison-lemma.* Proposition H.3 applied to b then gives

$$b(z) = \frac{1}{\tilde{\gamma}_{\min}} \left(\ln z - \ln \ln z + O(1) \right), \quad \frac{M_-}{z} \leq \mathcal{I}_{\min}(b(z)) \leq \frac{M_+}{z} \quad \text{for all } z \geq z_b, \quad (79)$$

for some constants $0 < M_- \leq M_+$ and $z_b \geq z_a$.

Taylor-expanding $e^{-\eta} = 1 - \eta + \rho(\eta)$ with $|\rho(\eta)| \leq \eta^2 e^{\eta} / 2$ for $|\eta| \leq \eta_*$, and separating the s -contribution, Eq (71) rewrites as

$$\dot{\eta} = -\gamma(z) \eta + \beta(z) + R(z, \eta), \quad (80)$$

with

$$\gamma(z) := \varepsilon \left(\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu) \right) \mathcal{I}_{\min}(b(z)), \quad (81)$$

$$\beta(z) := \varepsilon \left(\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu) \right) \mathcal{I}_{\min}(b(z)) s(z, 0), \quad (82)$$

and a remainder $R(z, \eta)$ gathering the remaining $O(\rho(\eta))$, $O(s(z, \eta) - s(z, 0))$, $O(\eta s(z, 0))$ and $O(\eta(s(z, \eta) - s(z, 0)))$ each multiplied by $\varepsilon \left(\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu) \right) \mathcal{I}_{\min}(b)$. By Eqs (79), (72), and (73),

$$\frac{\gamma_-}{z} \leq \gamma(z) \leq \frac{\gamma_+}{z}, \quad |\beta(z)| \leq \frac{B}{z \ln z}, \quad |R(z, \eta)| \leq \frac{K}{z} \left(\eta^2 + \frac{|\eta|}{\ln z} \right), \quad (83)$$

for $z \geq z_b$ (possibly enlarging z_b), with $\gamma_{\pm} := \varepsilon \left(\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu) \right) M_{\pm} > 0$, and explicit constants $B, K > 0$.

Choose now η_* small enough and $z_c \geq z_b$ large enough that $K \eta_* \leq \frac{\gamma_-}{4}$ and $\frac{K}{\ln z_c} \leq \frac{\gamma_-}{4}$.

Then for $z \geq z_c$ and $|\eta| \leq \eta_*$,

$$\eta(-\gamma(z)\eta + R(z, \eta)) \leq -\gamma(z)\eta^2 + \frac{K}{z} \left(|\eta|^3 + \frac{\eta^2}{\ln z} \right) \leq -\frac{\gamma_-}{2z} \eta^2, \quad (84)$$

so multiplying (80) by $\eta/|\eta|$ and using $\eta \dot{\eta} = \frac{d}{dz}(\eta^2/2) = |\eta| D^+ |\eta|$ (Dini derivative Khalil (2002, Appendix C2)), one obtains the scalar differential inequality

$$D^+ |\eta(z)| \leq -\frac{\gamma_-}{2z} |\eta(z)| + \frac{B}{z \ln z}, \quad z \geq z_c. \quad (85)$$

Khalil's Comparison Lemma (Khalil, 2002, Lemma 3.4) applied to (85) yields

$$|\eta(z)| \leq |\eta(z_c)| \exp\left(-\int_{z_c}^z \frac{\gamma_-}{2s} ds\right) + \int_{z_c}^z \exp\left(-\int_s^z \frac{\gamma_-}{2r} dr\right) \frac{B}{s \ln s} ds. \quad (86)$$

The homogeneous factor evaluates to $\exp(-\frac{\gamma_-}{2} \ln(z/z_c)) = (z_c/z)^{\gamma_-/2}$, hence the first term of (86) is $O(z^{-\gamma_-/2})$, i.e. polynomial decay. Lemma H.7, the second term equals $O(1/\ln z)$.

Combining, $|\eta(z)| = O(z^{-\gamma_-/2}) + O(1/\ln z) = O(1/\ln z)$, which proves the lemma. $\square$

This synchronisation Lemma naturally generalises to the discrete dynamics as shown in Lemma H.9.

Remark E.6 (On entering the tube). The Lemmas E.5 and H.9 both show that $\mathcal{T}$ is forward invariant. So we just need to prove that both continuous and discrete trajectories eventually fall inside. Noticing that for any initialisation with $a_0, b_0 \geq 1$, the sign of the drift outside the tube pushes η_t inward, the trajectory is expected to reach the tube in finite time. This can be made rigorous at the expense of a longer argument (quantifying the discrete drift rate against the decay

of $\mathcal{I}_{\min}(b_t)$), but we opt for the cleaner approach of assuming it directly (*special condition $\star$ in Assumption D.2*, in the same spirit as the magnitude condition $\mathbf{v} \cdot \psi_0 \geq 1$).

712

713 **E.4 Conclusion of the Proof when $\alpha < 1$**

714 *Proof.* We now complete the proof of Theorem D.4 when $\alpha < 1$.

715 By Lemma H.9 with Assumption D.2 (special condition $\star$), we have $\eta_t = a_t - \tilde{\gamma}_{\min} b_t - d \rightarrow 0$.
716 Therefore the coefficient

$$\tilde{c}_t := \varepsilon(1 + \mu_B) + (1 - \varepsilon)(1 + \mu) \frac{\mathcal{I}_{\text{maj}}(a_t)}{\mathcal{I}_{\min}(b_t)} = \tilde{c}_t = c_0 + \tilde{\delta}_t, \quad (87)$$

717 with $\tilde{\delta}_t \rightarrow 0$, by exactly the same computation as (76) (with $a_t = \tilde{\gamma}_{\min} b_t + d + \eta_t$ replacing
718 $a(z) = \tilde{\gamma}_{\min} b(z) + d + \eta(z)$). The recurrence (58) rewrites as $b_{t+1} = b_t + h_t \tilde{c}_t \mathcal{I}_{\min}(b_t)$, which is
719 exactly the setup of Lemma H.6 with $c_t = \tilde{c}_t$. Part (iii) of that lemma gives

$$b_t = b(z_t) + o(1) \quad \text{as } t \rightarrow \infty. \quad (88)$$

720 By symmetry (exchanging the roles of a and b), the same argument gives $a_t = a(z_t) + o(1)$.

721 *Minority group error.* Proposition H.4 with c_0 defined in Eq (78), $\lambda_0 = \lambda_{\min}(\gamma)$, and $\gamma = \tilde{\gamma}_{\min}$ gives:

$$\tilde{\gamma}_{\min} b(z) = \ln z - \ln \ln z + \ln(\tilde{\gamma}_{\min}^3 c_0 \lambda_{\min}(\tilde{\gamma}_{\min})) + o(1),$$

722 so $\tilde{\gamma}_{\min} b(z_t) = \ln(z_t) - \ln \ln(z_t) + \ln(\tilde{\gamma}_{\min}^3 c_0 \lambda_{\min}(\tilde{\gamma}_{\min})) + o(1)$ since $z_t \rightarrow +\infty$ by Condition (21).

723 Using Eq (88), we have:

$$\begin{aligned} \frac{e^{-\tilde{\gamma}_{\min} b_t}}{b_t} &= \frac{e^{-\tilde{\gamma}_{\min} b(z_t) + o(1)}}{b(z_t) + o(1)} = \frac{e^{-b(z_t)}}{b(z_t)} \cdot (1 + o(1)) \\ &= \tilde{\gamma}_{\min} \frac{e^{-\ln z_t + \ln \ln z_t - \ln(\tilde{\gamma}_{\min}^3 c_0 \lambda_{\min}(\tilde{\gamma}_{\min})) + o(1)}}{\ln z_t + o(\ln z_t)} \\ &= \frac{1}{\tilde{\gamma}_{\min}^2 c_0 \lambda_{\min}(\tilde{\gamma}_{\min}) z_t} \cdot (1 + o(1)). \end{aligned}$$

724 Therefore with $\mathbb{E}_{\mathcal{G}_{\min}}[q^t] = \int_{\tilde{\gamma}_{\min}}^K \sigma(-b_t z) \lambda_{\min}(z) dz$, Lemma H.1 gives

$$\mathbb{E}_{\mathcal{G}_{\min}}[q^t] \sim \frac{\lambda_{\min}(\tilde{\gamma}_{\min}) e^{-\tilde{\gamma}_{\min} b_t}}{b_t} \sim \frac{1}{\tilde{\gamma}_{\min}^2 c_0 z_t}.$$

725 We retrieve Eq (46) by expanding c_0 with Eq (78) and using Eq (44).

726 *Majority group error.* The identical argument applied to (a_t) with the analogue value for c_0 , λ_0 gives
727 Eq (45). $\square$

728 **E.5 Conclusion of the Proof when $\alpha \geq 1$**

729 *Proof.* When $\alpha \geq 1$, the synchronisation observed in Section E.3 does no longer happen. The
730 proof proceeds in three steps: (i) the drift ratio $\mathcal{I}_{\min}(b)/\mathcal{I}_{\text{maj}}(a) \rightarrow 0$ is deduced, (ii) the decoupled
731 asymptotics for a and b are derived, and (iii) the classification errors are computed.

732 (i) *The drift ratio vanishes.* Recall that by Eq (65),

$$\dot{\Delta} = \underbrace{(1-\varepsilon)(1+\mu_A)(1-\alpha)}_{\leq 0} \mathcal{I}_{\text{maj}}(a) + \underbrace{\varepsilon((1+\mu) - \tilde{\gamma}_{\min}(1+\mu_B))}_{\leq 0} \mathcal{I}_{\min}(b). \quad (89)$$

733 Since both $\mathcal{I}_{\text{maj}}(a)$ and $\mathcal{I}_{\min}(b)$ are strictly positive and since both factors $(1-\varepsilon)(1+\mu_A)(1-\alpha)$ and
734 $\varepsilon((1+\mu) - \tilde{\gamma}_{\min}(1+\mu_B))$ are not zeros simultaneously, $\Delta \rightarrow -\infty$ as soon as one of $\int_0^\infty \mathcal{I}_{\text{maj}}(a) ds$
735 or $\int_0^\infty \mathcal{I}_{\min}(b) ds$ is ∞ . From the first line of the 2D-ODE drift function (64),

$$a(z) - a(0) = (1-\varepsilon)(1+\mu_A) \int_0^z \mathcal{I}_{\text{maj}}(a(s)) \, ds + \varepsilon(1+\mu) \int_0^z \mathcal{I}_{\text{min}}(b(s)) \, ds.$$

736 Since $a(z) \rightarrow +\infty$, their sum diverges:

$$(1-\varepsilon)(1+\mu_A) \int_0^\infty \mathcal{I}_{\text{maj}}(a) \, ds + \varepsilon(1+\mu) \int_0^\infty \mathcal{I}_{\text{min}}(b) \, ds = +\infty,$$

737 and $\Delta(z) \rightarrow -\infty$. By Lemma H.2, for large enough a and b ,

$$\frac{\mathcal{I}_{\text{min}}(b)}{\mathcal{I}_{\text{maj}}(a)} = \frac{\tilde{\gamma}_{\text{min}} \lambda_{\text{min}}(\tilde{\gamma}_{\text{min}})}{\lambda_{\text{maj}}(1)} \frac{a}{b} e^\Delta (1+o(1)) \leq \frac{\tilde{\gamma}_{\text{min}}^2 \lambda_{\text{min}}(\tilde{\gamma}_{\text{min}})}{\lambda_{\text{maj}}(1)} e^\Delta (1+o(1)) \xrightarrow{z \rightarrow \infty} 0. \quad (90)$$

738 where the last inequality uses $\Delta = a - \tilde{\gamma}_{\text{min}} b \leq 0$ since $\Delta \rightarrow -\infty$.

739 (ii): *Decoupled asymptotics for a and b .* Denote $R(z) := \frac{\mathcal{I}_{\text{min}}(b(z))}{\mathcal{I}_{\text{maj}}(a(z))}$. From the first line of the 2D-ODE
740 drift function (64),

$$\dot{a}(z) = (1-\varepsilon)(1+\mu_A) \mathcal{I}_{\text{maj}}(a) \left(1 + \frac{\varepsilon(1+\mu)}{(1-\varepsilon)(1+\mu_A)} R(z) \right) = (1-\varepsilon)(1+\mu_A) \mathcal{I}_{\text{maj}}(a) (1+o(1)). \quad (91)$$

741 This is a single scalar ODE of the form treated in Propositions H.3 and H.4 with $c_0 = (1-\varepsilon)(1+\mu_A)$,
742 $\gamma = 1$ and $\lambda_0 = \lambda_{\text{maj}}(1)$. Hence

$$a(z) = \ln z - \ln \ln z + \ln \left((1-\varepsilon)(1+\mu_A) \lambda_{\text{maj}}(1) \right) + o(1) \quad (92)$$

$$\mathcal{I}_{\text{maj}}(a(z)) = \Theta(1/z). \quad (93)$$

743 On another hand, dividing the two lines of the 2D-ODE drift function (64) and using $R(z) \rightarrow 0$:

$$\begin{aligned} \frac{\dot{b}}{\dot{a}} &= \frac{\varepsilon(1+\mu_B) R(z) + (1-\varepsilon)(1+\mu)}{(1-\varepsilon)(1+\mu_A) + \varepsilon(1+\mu) R(z)} = \frac{1+\mu}{1+\mu_A} + O(R(z)), \\ \dot{b} &= \frac{1+\mu}{1+\mu_A} \dot{a} + O(\dot{a} R(z)). \end{aligned} \quad (94)$$

744 Combining Eqs (91) and (93) gives $\dot{a} = \Theta(1/z)$, while combining combining Eqs (89) and (93)
745 gives

$$\begin{aligned} \dot{\Delta} &\leq (1-\varepsilon)(1+\mu_A)(1-\alpha) \mathcal{I}_{\text{maj}}(a) \leq -\frac{C}{z} \\ \Delta(z) &\leq -C \ln z + O(1), \quad \text{s.t.} \\ R(z) &\leq \frac{\tilde{\gamma}_{\text{min}}^2 \lambda_{\text{min}}(\tilde{\gamma}_{\text{min}})}{\lambda_{\text{maj}}(1)} e^{\Delta(z)} (1+o(1)) \leq M/z^C \quad (\text{by Eq (90)}) \end{aligned} \quad (95)$$

746 for some positive constants C and M . Therefore $\dot{a} R(z) = O(1/z^{C+1})$ and integrating (94):

$$\begin{aligned} b(z) - b(z_0) &= \frac{1+\mu}{1+\mu_A} (a(z) - a(z_0)) + O(1), \quad \text{i.e.} \\ \tilde{\gamma}_{\text{min}} b(z) &= \alpha a(z) + O(1) \quad \text{using Eq (44)} \\ &= \alpha \ln z - \alpha \ln \ln z + O(1) \quad \text{by Eq (92)}. \end{aligned} \quad (96)$$

747 (iii) *Classification errors.* We conclude as in Section E.4. From Eq (92) and Lemma H.1 :

$$\mathbb{E}_{\mathcal{G}_{\text{maj}}} [1-p_y^t] \sim \frac{\lambda_{\text{maj}}(1) e^{-a_t}}{a_t} \sim \frac{\lambda_{\text{maj}}(1) e^{-a(z_t)}}{a(z_t)} \sim \frac{1}{(1-\varepsilon)(1+\mu_A) z_t}.$$

748 Likewise, Eq (96) and Lemma H.1 give:

$$\mathbb{E}_{\mathcal{G}_{\text{min}}} [1-p_y^t] \sim \frac{\lambda_{\text{min}}(\tilde{\gamma}_{\text{min}}) e^{-\tilde{\gamma}_{\text{min}} b_t}}{b_t} \sim \frac{\lambda_{\text{min}}(\tilde{\gamma}_{\text{min}}) e^{-\tilde{\gamma}_{\text{min}} b(z_t)}}{b(z_t)} = \Theta \left(\frac{(\ln z_t)^{\alpha-1}}{z_t^\alpha} \right).$$

749

□

F Generalizing to the proof of Theorem 5.3

The proof of Theorem D.4 rests on Assumptions D.1–D.3. This section discusses these assumptions: their necessity, what the proof can still say when they are relaxed, and which generalisations remain open.

F.1 Relaxing Assumptions D.1 and D.2

Assumptions D.1 and D.2 can be weakened by assuming simply that:

Assumption F.1 (Non-degeneracy of $\mathbf{r}_\perp$ at the margin boundary). For every unit direction $\hat{\delta} \in \mathbb{R}^{d_r}$ with $\mathbf{v}^\top \hat{\delta} = 0$, the conditional distribution of $\zeta := \hat{\delta}^\top \mathbf{r}_\perp$ given $Z = 1$ and any group g satisfies

$$\Pr(\zeta < 0 \mid Z=1, g) > 0. \quad (97)$$

This non-degeneracy assumption simply requires that for every direction $\hat{\delta}$ such that $\hat{\delta}^\top \mathbf{v} = 0$, $\mathbf{r}_\perp \mid Z=1$ never puts all its mass on one single side of the hyperplane $(\hat{\delta})^\perp$. In geometric terms, the origin lies in the interior of the convex hull of the support of $\mathbf{r}_\perp \mid Z=1$ in the subspace $(\mathbf{v})^\perp$.

In this case, the parameters ψ_t and ϕ_t acquire a perpendicular component:

$$\psi_t = a_t \mathbf{v} + \delta_t^\psi, \quad \phi_t = b_t \mathbf{v} + \delta_t^\phi, \quad \delta_t^\psi, \delta_t^\phi \in (\text{Span}\{\mathbf{v}\})^\perp. \quad (98)$$

Furthermore $\|\delta_t^\psi\|$ and $\|\delta_t^\phi\|$ remain $O(1)$, and when $\tilde{\gamma}_{\text{maj}} = 1$, $\tilde{\gamma}_{\text{min}} \geq 1$ (with analogous result in the mirror case), the asymptotic rates are slightly changed:

(i) If $\alpha < 1$,

$$\kappa_{\text{maj}} = \frac{\tilde{\gamma}_{\text{min}}(1+\mu_B) - (1+\mu)}{\tilde{\gamma}_{\text{min}} \Sigma} \cdot \frac{1}{\Lambda_{\text{maj}}}, \quad (99)$$

$$\kappa_{\text{min}} = \frac{(1+\mu_A) - \tilde{\gamma}_{\text{min}}(1+\mu)}{\tilde{\gamma}_{\text{min}}^2 \Sigma} \cdot \frac{1}{\Lambda_{\text{min}}}, \quad (100)$$

(ii) If $\alpha \geq 1$,

$$\mathbb{E}_{\mathcal{G}_{\text{maj}}}[1 - p_y^t] \underset{t \rightarrow \infty}{\sim} \frac{1}{(1-\varepsilon)(1+\mu_A) \Lambda_{\text{maj}} z_t}, \quad (101)$$

$$\mathbb{E}_{\mathcal{G}_{\text{min}}}[1 - p_y^t] = \Theta\left(\frac{(\ln z_t)^{\alpha-1}}{z_t^\alpha}\right), \quad (102)$$

with Λ_{maj} and Λ_{min} only depending on the group-wise distributions of $\mathbf{r}_\perp \mid Z = \tilde{\gamma}_g$. More precisely, when $\text{Law}(\mathbf{r}_\perp \mid Z = \tilde{\gamma}_{\text{maj}}, \mathcal{G}_{\text{maj}}) = \text{Law}(\mathbf{r}_\perp \mid Z = \tilde{\gamma}_{\text{min}}, \mathcal{G}_{\text{min}})$, we get that $\Lambda_{\text{maj}} = \Lambda_{\text{min}}$ which falls back to Eqs (45)-(46).

However, this would require considerable more complex mathematical development, therefore we leave this generalisation opened.

F.2 Relaxing Assumption D.3 (feature noise)

We simplify the notations in this section this way: $a_t^g, A_g, \psi_t^g, \mu_g, \bar{g}$ represent in this order $a_t, A, \psi_t, \mu_A, \mathcal{G}_{\text{min}}$ when $g = \text{maj}$ and $b_t, B, \phi_t, \mu_B, \mathcal{G}_{\text{maj}}$ when $g = \text{min}$.

When $\xi \neq 0$, the logit by Eqs (49) becomes $w^t \cdot \mathbf{x} = \begin{cases} \psi_t \cdot \mathbf{r} + w_s^t \cdot \xi & \text{on } \mathcal{G}_{\text{maj}} \\ \phi_t \cdot \mathbf{r} + w_s^t \cdot \xi & \text{on } \mathcal{G}_{\text{min}} \end{cases}$, where we combined

Eq (42) and Proposition 5.4 to define

$$\psi_t^g = \hat{\psi}^g (\log z_t - \log \log z_t) + O(1), \quad \hat{\psi}^g := \hat{\mathbf{w}}_r + A_g^\top \hat{\mathbf{w}}_s. \quad (103)$$

Define the effective sigmoids:

$$\sigma_{\text{eff}}^t(u, z) := \mathbb{E}_{\xi \mid Z=z} [\sigma(u - w_s^t \cdot \xi)], \quad (104)$$

777 and the noise symmetrical moment-generating functions:

$$\Lambda_\xi(w_s, z) := \mathbb{E}_{\xi|Z=z} [e^{-w_s \cdot \xi}]. \quad (105)$$

778 We have the upper bound $\Lambda_\xi(w_s^t, z) \leq e^{\|w_s^t\|R} =: e^{M_t}$, where $M_t := \|w_s^t\|R$.

779 (i) *Relaxed assumption: Noise bound.* Assumption D.3 can be greatly weakened by assuming simply
780 that R (see Eq (28)) is sufficiently small in the sense of:

781 **Assumption F.2 (Noise bound).** *The noise bound satisfies for all $g \in \{\text{maj}, \text{min}\}$*

$$(\hat{\psi}^g \cdot \mathbf{v}) \tilde{\gamma}_g > \|\hat{\mathbf{w}}_s\| R \quad (106)$$

782 To understand the utility of this relaxed assumption, let's first show the following result.

Lemma F.3. *For any u, z with $u > M_t$, $z \geq 1$:*

$$\sigma_{\text{eff}}^t(-u, z) = e^{-u} \Lambda_\xi(w_s^t, z) (1 + \rho_t(u, z)), \quad (107)$$

with

$$|\rho_t(u, z)| \leq e^{-(u-M_t)}. \quad (108)$$

783

784 *Proof.* For $u > M_t$, $u - w_s^t \cdot \xi \geq u - M_t > 0$ almost surely. Using the identity $\sigma(-v) =$
785 $e^{-v}/(1 + e^{-v})$:

$$\sigma(-u - w_s^t \cdot \xi) = e^{-(u+w_s^t \cdot \xi)} \cdot \frac{1}{1 + e^{-(u+w_s^t \cdot \xi)}} = e^{-u} e^{-w_s^t \cdot \xi} h(u, \xi), \quad (109)$$

786 where $h(u, \xi) := 1/(1 + e^{-(u+w_s^t \cdot \xi)})$. Since $u + w_s^t \cdot \xi > 0$, $h \in (1/2, 1)$ a.s. More precisely:

$$1 - e^{-(u+w_s^t \cdot \xi)} \leq h(u, \xi) \leq 1 \quad \text{a.s.} \quad (110)$$

787 Taking expectations in (109) under $Z = z$:

$$\begin{aligned} \sigma_{\text{eff}}^t(-u, z) &= e^{-u} \mathbb{E}_{Z=z} [e^{-w_s^t \cdot \xi} h(u, \xi)] \quad \text{s.t.} \\ \sigma_{\text{eff}}^t(-u, z) &\leq e^{-u} \Lambda_\xi(w_s^t, z) \quad \text{since } h \leq 1, \\ \sigma_{\text{eff}}^t(-u, z) &\geq e^{-u} \left(\Lambda_\xi(w_s^t, z) - e^{-u} \mathbb{E}_{Z=z} [e^{-2w_s^t \cdot \xi}] \right) \quad \text{since } h \geq 1 - e^{-(u+w_s^t \cdot \xi)} \\ &= e^{-u} \Lambda_\xi(w_s^t, z) \left(1 - \underbrace{e^{-u} \Lambda_\xi(2w_s^t, z) / \Lambda_\xi(w_s^t, z)}_{:= \rho_t(u, z)} \right), \end{aligned} \quad (111)$$

788 with

$$|\rho_t(u, z)| \leq e^{-u} \frac{\Lambda_\xi(2w_s^t, z)}{\Lambda_\xi(w_s^t, z)} \leq e^{-(u-M_t)}. \quad (112)$$

789 The last inequality is due to the pointwise inequality $e^{-2w_s^t \cdot \xi} \leq e^{-w_s^t \cdot \xi} e^{M_t}$ a.s. $\square$

790 By Proposition I.5 and Eq (103), for all group g

$$\begin{aligned} \frac{a_t^g}{\|w^t\|} &= \frac{\psi_t^g \cdot \mathbf{v}}{\|w^t\|} = \frac{\hat{\psi}^g \cdot \mathbf{v} (\log z_t - \log \log z_t) + O(1)}{\|\hat{\mathbf{w}}\| (\log z_t - \log \log z_t) + O(1)} \longrightarrow (\hat{\psi}^g \cdot \mathbf{v}) / \|\hat{\mathbf{w}}\| \\ &\quad \frac{\|w_s^t\|}{\|w^t\|} \longrightarrow \|\hat{\mathbf{w}}_s\| / \|\hat{\mathbf{w}}\|. \end{aligned} \quad (113)$$

791 Therefore:

$$\frac{a_t \tilde{\gamma}_{\text{maj}}}{\|w_s^t\| R} \longrightarrow \frac{\tilde{\gamma}_{\text{maj}} (\hat{\psi} \cdot \mathbf{v})}{\|\hat{\mathbf{w}}_s\| R} > 1, \quad \frac{b_t \tilde{\gamma}_{\text{min}}}{\|w_s^t\| R} \longrightarrow \frac{\tilde{\gamma}_{\text{min}} (\hat{\phi} \cdot \mathbf{v})}{\|\hat{\mathbf{w}}_s\| R} > 1, \quad (114)$$

792 the last inequalities come from Eq (106), thereby emphasizing their utility.

793 Consequently, for all t large enough and all group g :

$$a_t^g \tilde{\gamma}_g - M_t = \left(\frac{a_t^g \tilde{\gamma}_g}{\|w_s^t\| R} - 1 \right) M_t \geq c_0 M_t \geq c'_0 \log z_t, \quad (115)$$

794 for constants $c_0, c'_0 > 0$. Combining Lemma F.3 and Eq (115) yields:

$$\sigma_{\text{eff}}^t(-a_t^g Z, Z) = e^{-a_t^g Z} \Lambda_\xi(w_s^t, Z) (1 + O(z_t^{-c'_0})) \quad \mathcal{G}_g - \text{almost surely}, \quad (116)$$

795 and the $O(\dots)$ is uniform in $Z \mid \mathcal{G}_g$.

796 (ii) Λ_ξ growth. For $z \geq 1$, define the noise exponent function:

$$\delta_+(z) := \sup_{\xi \in \text{supp}(\xi \mid Z=z)} -\hat{\mathbf{w}}_s \cdot \xi, \quad (117)$$

797 and the boundary noise exponents for each group:

$$\delta_g := \delta_+(\tilde{\gamma}_g), \quad g \in \{\text{maj}, \text{min}\}. \quad (118)$$

798 The function $\delta_+(z)$ is non-decreasing in z : as Z grows (farther from the margin boundary), the margin
799 condition $\hat{\psi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi \geq \gamma_{\text{maj}}$ becomes less restrictive, allowing a wider range of ξ realisations, and
800 hence a larger $-\hat{\mathbf{w}}_s \cdot \xi$ supremum. In particular, $\delta_g = \delta_+(\tilde{\gamma}_g)$ is the *smallest* value of δ_+ over the
801 support of $Z \mid \mathcal{G}_g$. We also have the upper bound $\delta_g \leq \|\hat{\mathbf{w}}_s\| R$ since $|\hat{\mathbf{w}}_s \cdot \xi| \leq \|\hat{\mathbf{w}}_s\| R$ a.s.

802 This gives us this following result.

Lemma F.4 (Λ_ξ growth rate). *For each $z \geq 1$ with $\text{supp}(\xi \mid Z = z) \neq \emptyset$,*

$$\Lambda_\xi(w_s^t, z) = z_t^{\delta_+(z) + o(1)} \quad \text{as } t \rightarrow \infty, \quad (119)$$

with $o(1)$ uniform in z .

803 *Proof.* Write $w_s^t = \hat{\mathbf{w}}_s \log z_t - \hat{\mathbf{w}}_s \log \log z_t + O(1)$ (Proposition 5.4), so:

$$w_s^t \cdot \xi = (\hat{\mathbf{w}}_s \cdot \xi) \log z_t + O(\log \log z_t) \quad \text{a.s. uniformly in } \xi$$

805 Then on one hand,

$$\Lambda_\xi(w_s^t, z) \leq \exp \left(\sup_{\xi \in \text{supp}(\xi \mid Z=z)} -w_s^t \cdot \xi \right) = \exp(\delta_+(z) \log z_t + O(\log \log z_t)) = z_t^{\delta_+(z)} \cdot O((\log z_t)^C)$$

806 giving the upper bound $\Lambda_\xi(w_s^t, z) \leq z_t^{\delta_+(z) + o(1)}$, with $o(1)$ uniform in z .

807 On the other hand, fix $\eta > 0$. The set $S_\eta(z) := \{\xi \in \text{supp}(\xi \mid Z = z) : -\hat{\mathbf{w}}_s \cdot \xi > \delta_+(z) - \eta\}$ has
808 positive $(\xi \mid Z = z)$ -probability by definition of $\delta_+(z)$. Let $p_\eta(z) := P(\xi \in S_\eta(z) \mid Z = z) > 0$.
809 Then:

$$\Lambda_\xi(w_s^t, z) \geq p_\eta(z) \cdot \exp((\delta_+(z) - \eta) \log z_t + O(\log \log z_t)) \geq z_t^{\delta_+(z)} \cdot p_\eta(z) z_t^{-\eta} O((\log z_t)^{-C})$$

810 where the $O(\dots)$ is uniform in η . Sending $\eta \rightarrow 0$ gives the lower bound $\Lambda_\xi(w_s^t, z) \geq z_t^{\delta_+(z) + o(1)}$
811 with $o(1)$ uniform in z . $\square$

812 (iii) ODE (63) reformulation.

Lemma F.5. *For all group $g \in \{\text{min}, \text{maj}\}$, the group- g drift function becomes*

$$\mathcal{I}_g^{\text{eff}}(a_t^g, w_s^t) := \mathbb{E}_{\mathcal{G}_g}[Z \sigma_{\text{eff}}^t(-a_t^g Z, Z)] = \Lambda_\xi(w_s^t, z_g^*) \cdot \frac{\tilde{\gamma}_g \lambda_g(\tilde{\gamma}_g) e^{-\tilde{\gamma}_g a_t^g}}{a_t^g} (1 + O(1/a_t^g)), \quad (120)$$

for some t -dependent $z_g^ \in (\tilde{\gamma}_g, K)$.*

Similarly, the classification error integral becomes:

$$\mathcal{J}_g^{\text{eff}}(a_t^g, w_s^t) := \mathbb{E}_{\mathcal{G}_g}[\sigma_{\text{eff}}^t(-a_t^g Z, Z)] = \Lambda_\xi(w_s^t, z_g^*) \cdot \frac{\lambda_g(\tilde{\gamma}_g) e^{-\tilde{\gamma}_g a_t^g}}{a_t^g} (1 + O(1/a_t^g)). \quad (121)$$

814

815 *Proof.* By Eq (116):

$$\begin{aligned} \mathcal{I}_g^{\text{eff}}(a_t^g) &= \int_{\tilde{\gamma}_g}^K z \sigma_{\text{eff}}^t(-a_t^g z, z) \lambda_g(z) dz \\ &\stackrel{(1)}{=} \int_{\tilde{\gamma}_g}^K z \Lambda_\xi(w_s^t, z) e^{-a_t^g z} \lambda_g(z) dz (1 + O(z_t^{-c'_0})) \\ &\stackrel{(2)}{=} \Lambda_\xi(w_s^t, z_g^*) \int_{\tilde{\gamma}_g}^K z e^{-a_t^g z} \lambda_g(z) dz (1 + O(z_t^{-c'_0})) \end{aligned} \quad (122)$$

816 where the factorisation (1) used the fact that $O(z_t^{-c'_0})$ is uniform in $Z \mid g$, and equality (2) used the
 817 mean value theorem for integrals with $z_g^* \in [\tilde{\gamma}_g, K]$. Using Lemma H.2 and more precisely Eq (178)
 818 on the integral gives the stated result, since $z_t^{-c'_0} = o(1/a_t^g)$ for $c'_0 > 0$ (because $a_t^g = \Theta(\log z_t)$). $\square$

819 Two direct consequences of Lemma F.5 are

$$\mathcal{J}_g^{\text{eff}} = \mathcal{I}_g^{\text{eff}} / \tilde{\gamma}_g \cdot (1 + O(1/a_t^g)), \quad (123)$$

820 and that

$$\frac{\mathcal{I}_{\text{maj}}^{\text{eff}}(a_t)}{\mathcal{I}_{\text{min}}^{\text{eff}}(b_t)} = \frac{\Lambda_\xi(w_s^t, z_{\text{maj}}^*)}{\Lambda_\xi(w_s^t, z_{\text{min}}^*)} \cdot \frac{\tilde{\gamma}_{\text{maj}} \lambda_{\text{maj}}(\tilde{\gamma}_{\text{maj}})}{\tilde{\gamma}_{\text{min}} \lambda_{\text{min}}(\tilde{\gamma}_{\text{min}})} \cdot \frac{b_t}{a_t} e^{-\tilde{\gamma}_{\text{maj}} a_t + \tilde{\gamma}_{\text{min}} b_t} (1 + O(1/b_t)). \quad (124)$$

821 By Lemma F.4, the Λ_ξ ratio in Eq (124) satisfies:

$$\frac{\Lambda_\xi(w_s^t, z_{\text{maj}}^*)}{\Lambda_\xi(w_s^t, z_{\text{min}}^*)} = z_t^{\delta_+(z_{\text{maj}}^*) - \delta_+(z_{\text{min}}^*) + o(1)}. \quad (125)$$

822 Since the Laplace weight $e^{-a_t^g z}$ concentrates all mass on the boundary $z = \tilde{\gamma}_g$, $z_g^* \rightarrow \tilde{\gamma}_g$ as $t \rightarrow \infty$.
 823 Therefore, with δ_+ being continuous and bounded, Eq (125) can be rewritten as:

$$\frac{\Lambda_\xi(w_s^t, z_{\text{maj}}^*)}{\Lambda_\xi(w_s^t, z_{\text{min}}^*)} = z_t^{\delta_{\text{maj}} - \delta_{\text{min}} + o(1)}. \quad (126)$$

824 Regarding the scalar update for a_t^g , Eqs (53)-(54) acquire a noise dependent term, so that Eqs (57)-(58)
 825 become:

$$\begin{aligned} a_{t+1}^g - a_t^g &= h_t \left(\varepsilon_g (1 + \mu_g) \mathcal{I}_g^{\text{eff}}(a_t^g) + \varepsilon_{\bar{g}} (1 + \mu) \mathcal{I}_{\bar{g}}^{\text{eff}}(a_t^{\bar{g}}) \right) \\ &\quad + h_t (A_g \mathbf{v}) \cdot \left(\varepsilon_g \mathbb{E}_g[q^t \xi] + \varepsilon_{\bar{g}} \mathbb{E}_{\bar{g}}[q^t \xi] \right), \end{aligned} \quad (127)$$

826 so that the gap function Δ (see Eq (65)) becomes $\Delta_t := (a_t - \tilde{\gamma}_{\text{min}} b_t) - (\delta_{\text{maj}} - \delta_{\text{min}}) \log z_t - d_t$
 827 with the new synchronisation term d_t defined by the following:

$$\begin{aligned}
\Delta_{t+1} - \Delta_t &= h_t \left((1 - \varepsilon)(1 + \mu_A)(1 - \alpha) \mathcal{I}_{\text{maj}}^{\text{eff}}(a_t) - \varepsilon(\tilde{\gamma}_{\min}(1 + \mu_B) - (1 + \mu)) \mathcal{I}_{\min}^{\text{eff}}(b_t) \right) \\
&\quad + h_t (\mathbf{A}\mathbf{v} - \tilde{\gamma}_{\min} B\mathbf{v}) \mathbb{E}[q^t \xi] + (\delta_{\text{maj}} - \delta_{\min}) \log(z_{t+1}/z_t) + (d_{t+1} - d_t) \\
&= \varepsilon(\tilde{\gamma}_{\min}(1 + \mu_B) - (1 + \mu)) \mathcal{I}_{\min}^{\text{eff}}(b_t) \left(\frac{(1 - \varepsilon)(1 + \mu_A)(1 - \alpha)}{\varepsilon(\tilde{\gamma}_{\min}(1 + \mu_B) - (1 + \mu))} \frac{\mathcal{I}_{\text{maj}}^{\text{eff}}(a_t)}{\mathcal{I}_{\min}^{\text{eff}}(b_t)} - 1 \right) \\
&\quad + O(h_t \mathbb{E}[q^t \xi]) + \Theta(h_t/z_t) + (d_{t+1} - d_t) \\
&= \varepsilon(\tilde{\gamma}_{\min}(1 + \mu_B) - (1 + \mu)) \mathcal{I}_{\min}^{\text{eff}}(b_t) \left(\frac{(1 - \varepsilon)(1 + \mu_A)(1 - \alpha)}{\varepsilon(\tilde{\gamma}_{\min}(1 + \mu_B) - (1 + \mu))} \times \right. \\
&\quad \left. \frac{\lambda_{\text{maj}}(1)}{\tilde{\gamma}_{\min} \lambda_{\min}(\tilde{\gamma}_{\min})} \cdot \frac{b_t}{\tilde{\gamma}_{\min} b_t + (\delta_{\text{maj}} - \delta_{\min}) \log z_t + d_t + \Delta_t} \times \right. \\
&\quad \left. \underbrace{e^{-a_t + \tilde{\gamma}_{\min} b_t + (\delta_{\text{maj}} - \delta_{\min} + o(1)) \log z_t}}_{= e^{-\Delta_t - d_t + o(\log z_t)}} (1 + O(1/b_t)) - 1 \right) \\
&\quad + O(h_t \mathbb{E}[q^t \xi]) + \Theta(h_t/z_t) + (d_{t+1} - d_t), \tag{128}
\end{aligned}$$

so that in order to get $\Delta_t \rightarrow 0$, we necessarily need

$$\begin{aligned}
e^{-d_t + o(\log z_t)} &\sim \frac{\varepsilon(\tilde{\gamma}_{\min}(1 + \mu_B) - (1 + \mu))}{(1 - \varepsilon)(1 + \mu_A)(1 - \alpha)} \cdot \frac{\tilde{\gamma}_{\min} \lambda_{\min}(\tilde{\gamma}_{\min})}{\lambda_{\text{maj}}(1)} \cdot \frac{\tilde{\gamma}_{\min} b_t + (\delta_{\text{maj}} - \delta_{\min}) \log z_t + d_t}{b_t} \\
&\sim \frac{\varepsilon(\tilde{\gamma}_{\min}(1 + \mu_B) - (1 + \mu))}{(1 - \varepsilon)(1 + \mu_A)(1 - \alpha)} \cdot \frac{\tilde{\gamma}_{\min} \lambda_{\min}(\tilde{\gamma}_{\min})}{\lambda_{\text{maj}}(1)} \cdot \frac{\tilde{\gamma}_{\min} \hat{\phi} \cdot \mathbf{v} + \delta_{\text{maj}} - \delta_{\min}}{\hat{\phi} \cdot \mathbf{v}}, \tag{129}
\end{aligned}$$

the analogous of the free-noise case condition (68), and where we used the b_t expression from Eq (113).

We now show that the $\Theta(1/z_t)$ error rate from Section E.4 is preserved under feature noise, for both groups.

(iii) *Preservation of the rates decay (case $\alpha < 1$).*

When $\alpha < 1$, the synchronisation argument (Lemma E.5 extended to the noisy setting) carries over with the same Lyapunov structure, ensuring $\Delta_t \rightarrow 0$, and therefore $\mathcal{I}_{\text{maj}}^{\text{eff}}(a_t)/\mathcal{I}_{\min}^{\text{eff}}(b_t) = R_0(1 + o(1))$ with the same R_0 as in the noise free case Eq (76). Therefore, the b -recurrence (127) reads

$$\begin{aligned}
b_{t+1} - b_t &= h_t \tilde{c}_t \mathcal{I}_{\min}^{\text{eff}}(b_t) + h_t O(C_\xi \mathbb{E}[q^t]) \\
&\stackrel{(1)}{=} h_t \tilde{c}_t \Lambda_\xi(w_s^t, z_{\min}^\star) \frac{\tilde{\gamma}_{\min} \lambda_{\min}(\tilde{\gamma}_{\min}) e^{-\tilde{\gamma}_{\min} b_t}}{b_t} (1 + O(1/b_t)) + h_t O(C_\xi \mathbb{E}[q^t]) \\
&\stackrel{(2)}{=} h_t \tilde{c}_t z_t^{\delta_{\min} + o(1)} \frac{\tilde{\gamma}_{\min} \lambda_{\min}(\tilde{\gamma}_{\min}) e^{-\tilde{\gamma}_{\min} b_t}}{b_t} (1 + O(1/b_t)) + h_t O(C_\xi \mathbb{E}[q^t]), \tag{130}
\end{aligned}$$

where $C_\xi := \|B\mathbf{v}\| R$, we used Lemma F.5 for line (1), Lemma F.4 and $z_{\min}^\star \rightarrow \tilde{\gamma}_{\min}$ for line (2), and where $\tilde{c}_t = c_0 + o(1)$ is exactly the same as in the noise-free setup Eq (87).

Under synchronisation, $\mathcal{I}_{\text{maj}}^{\text{eff}} = R_0 \mathcal{I}_{\min}^{\text{eff}}(1 + o(1))$, and combined with Eq (123), we have $\mathbb{E}[q^t] = (1 - \varepsilon) \mathcal{J}_{\text{maj}}^{\text{eff}}(a_t, w_s^t) + \varepsilon \mathcal{J}_{\min}^{\text{eff}}(b_t, w_s^t) \leq C' \mathcal{I}_{\min}^{\text{eff}}(b_t, w_s^t)$ for some $C' > 0$ and all t large enough. Therefore (130) can be rewritten as

$$b_{t+1} - b_t = h_t (\tilde{c}_t + O(C_\xi)) z_t^{\delta_{\min} + o(1)} \frac{\tilde{\gamma}_{\min} \lambda_{\min}(\tilde{\gamma}_{\min}) e^{-\tilde{\gamma}_{\min} b_t}}{b_t} (1 + O(1/b_t)), \tag{131}$$

with the term $O(C_\xi)$ dependent on the noise and on t and bounded uniformly in ξ . This is the exact prerequisite of Proposition H.5, which yields

$$\tilde{\gamma}_{\min} b(z) = (1 + \delta_{\min} + o(1)) \ln z - \ln \ln z + O(1), \quad (132)$$

$$z^{\delta_{\min} + o(1)} \mathcal{I}_{\min}(b(z)) = \frac{1}{(\tilde{\gamma}_{\min} c_0 + O(C_\xi))(1 + \delta_{\min})} \cdot \frac{e^{O(1)}}{z}, \quad (133)$$

844 with the bounded terms $O(\dots)$ independent on c_0 .

845 Combining with Lemma H.1 and Lemma H.2: $\mathcal{I}_{\min}/\mathcal{J}_{\min} = \tilde{\gamma}_{\min} (1 + O(1/b_t))$, we have:

$$\begin{aligned} z^{\delta_{\min} + o(1)} \mathcal{J}_{\min}(b(z)) &= \frac{1}{\tilde{\gamma}_{\min}} z^{\delta_{\min} + o(1)} \mathcal{I}_{\min}(b(z)) (1 + O(1/b)) \\ &= \frac{1}{(\tilde{\gamma}_{\min}^2 c_0 + O(C_\xi))(1 + \delta_{\min})} \cdot \frac{e^{O(1)}}{z} \end{aligned} \quad (134)$$

846 Now, from Eq (121):

$$\mathcal{J}_{\min}^{\text{eff}}(b_t) = \Lambda_\xi(w_s^t, z_{\min}^*) \cdot \mathcal{J}_{\min}(b_t) (1 + o(1)) = z_t^{\delta_{\min} + o(1)} \cdot \mathcal{J}_{\min}(b_t) (1 + o(1)),$$

847 and combining with Eq (134):

Proposition F.6 (Rate decay preserved, $\alpha < 1$). *When $\alpha < 1$*

$$\mathbb{E}_{\mathcal{G}_{\min}}[1 - p_y^t] = \mathcal{J}_{\min}^{\text{eff}}(b_t) = \frac{e^{O(1)}}{(\tilde{\gamma}_{\min}^2 c_0 + O(C_\xi))(1 + \delta_{\min})} \cdot \frac{1}{z_t} \quad (135)$$

the exact same rate as Eq (46) in the free-noise case (Theorem D.4), with some $O(1)$ and $e^{O(1)}$ extra-terms who depend on the noise. Especially this confirms the ε -dependency on the rates prefactors.

The same results are also true for $\mathcal{G}_{\text{maj}}$ with the same analysis.

848

849 (iv) *Preservation of the rates decay (case $\alpha > 1$).*

850 From Eq (127), the raw gap $G_t := a_t - \tilde{\gamma}_{\min} b_t$ evolves as

$$\begin{aligned} G_{t+1} - G_t &= h_t \left[(1 - \varepsilon)(1 + \mu_A)(1 - \alpha) \mathcal{I}_{\text{maj}}^{\text{eff}}(a_t) - \varepsilon(\tilde{\gamma}_{\min}(1 + \mu_B) - (1 + \mu)) \mathcal{I}_{\min}^{\text{eff}}(b_t) \right] \\ &\quad + h_t (A\mathbf{v} - \tilde{\gamma}_{\min} B\mathbf{v}) \cdot \mathbb{E}[q^t \xi]. \end{aligned} \quad (136)$$

851 Since $|\mathbb{E}_g[q^t \xi]| \leq R \mathcal{J}_g^{\text{eff}}(a_t^g) \leq \frac{R}{\tilde{\gamma}_g} \mathcal{I}_g^{\text{eff}}(a_t^g) (1 + O(1/a_t^g))$ for t large (by Eq (123)), $\mathbb{E}[q^t] =$
 852 $(1 - \varepsilon) \mathcal{J}_{\text{maj}}^{\text{eff}}(a_t) + \varepsilon \mathcal{J}_{\min}^{\text{eff}}(b_t) \leq C' (\mathcal{I}_{\text{maj}}^{\text{eff}}(a_t) + \mathcal{I}_{\min}^{\text{eff}}(b_t))$ for some structural constant $C' > 0$. The noise
 853 contribution to (136) is therefore bounded by $C_\xi (\mathcal{I}_{\text{maj}}^{\text{eff}} + \mathcal{I}_{\min}^{\text{eff}})$, where $C_\xi := R |A\mathbf{v} - \tilde{\gamma}_{\min} B\mathbf{v}| C'$.

854 Both coefficients of $\mathcal{I}_{\text{maj}}^{\text{eff}}$, and $\mathcal{I}_{\min}^{\text{eff}}$ in Eq (136) are non-positive (because $\alpha, \tilde{\gamma}_{\min} \geq 1$), therefore the
 855 same argument as in Section E.5 shows $G_t \rightarrow -\infty$, since the noise contribution is dominated by the
 856 non-positive terms for reasonably small R or C_ξ .

857 Furthermore, on one hand, from the continuous version of Eq (136), we have

$$\dot{G} \leq [(1 - \varepsilon)(1 + \mu_A)(1 - \alpha) + C_\xi] \mathcal{I}_{\text{maj}}^{\text{eff}}(a), \quad (137)$$

858 and on another hand, from Eq (127), we also have

$$\dot{a} \geq [(1 - \varepsilon)(1 + \mu_A) - C_\xi] \mathcal{I}_{\text{maj}}^{\text{eff}}(a). \quad (138)$$

859 Dividing (137) by (138):

$$\frac{\dot{G}}{\dot{a}} \leq \frac{(1 - \varepsilon)(1 + \mu_A)(1 - \alpha) + C_\xi}{(1 - \varepsilon)(1 + \mu_A) - C_\xi} =: 1 - \alpha_\xi < 0, \quad (139)$$

860 where

$$\alpha_\xi := \frac{(1 - \varepsilon)(1 + \mu_A) \alpha - 2C_\xi}{(1 - \varepsilon)(1 + \mu_A) - C_\xi} = \alpha + O(C_\xi) \xrightarrow{R \rightarrow 0} \alpha \quad \text{uniformly in } z. \quad (140)$$

861 Integrating:

$$G(z) \leq (1 - \alpha_\xi) a(z) + O(1). \quad (141)$$

862 Finally, since (138) is a scalar inequality of the form $\dot{a} \geq c_- z^{\delta_{\text{maj}}+o(1)} \mathcal{I}_{\text{maj}}(a)$ with $c_- :=$
 863 $(1-\varepsilon)(1+\mu_A) - C_\xi > 0$ (using $\mathcal{I}_g^{\text{eff}} = z^{\delta_g+o(1)} \mathcal{I}_g$ from Lemma F.5 and F.4), Proposition H.5
 864 gives the lower bound:

$$a(z) \geq (1+\delta_{\text{maj}}+o(1)) \ln z - \ln \ln z + O(1). \quad (142)$$

865 Defining the effective drift ratio

$$R^{\text{eff}}(z) := \frac{\mathcal{I}_{\text{min}}^{\text{eff}}(b)}{\mathcal{I}_{\text{maj}}^{\text{eff}}(a)}, \quad (143)$$

866 Eq (126) and Lemma F.5 give

$$R^{\text{eff}}(z) = z^{\delta_{\text{min}}-\delta_{\text{maj}}+o(1)} \frac{\tilde{\gamma}_{\text{min}} \lambda_{\text{min}}(\tilde{\gamma}_{\text{min}})}{\lambda_{\text{maj}}(1)} \frac{a}{b} e^G (1+o(1)). \quad (144)$$

867 Since $G \leq 0$ for z large ($G \rightarrow -\infty$), $a/b = O(1)$. Substituting in Eq (144) along with Eq (141) and
 868 Eq (142):

$$\begin{aligned} R^{\text{eff}}(z) &\leq z^{\delta_{\text{min}}-\delta_{\text{maj}}+o(1)} \cdot O(1) \cdot e^{(1-\alpha_\xi) a(z)+O(1)} \\ &\leq z^{\delta_{\text{min}}-\delta_{\text{maj}}+(1-\alpha_\xi)(1+\delta_{\text{maj}})+o(1)} \cdot O\left((\ln z)^{-(1-\alpha_\xi)}\right) \\ &= z^{1+\delta_{\text{min}}-\alpha_\xi(1+\delta_{\text{maj}})+o(1)} \cdot O\left((\ln z)^{-(1-\alpha_\xi)}\right). \end{aligned} \quad (145)$$

869 When the noise satisfies $\alpha_\xi > \frac{1+\delta_{\text{min}}}{1+\delta_{\text{maj}}} \geq 1$, the bound forces $R^{\text{eff}} \rightarrow 0$ and the dynamics decouple. In
 870 this case, the a -equation (Eq (127)) reads

$$a_{t+1} - a_t = h_t \left[(1-\varepsilon)(1+\mu_A) + O(1) \right] z_t^{\delta_{\text{maj}}+o(1)} \frac{\lambda_{\text{maj}}(1) e^{-a_t}}{a_t} (1 + O(1/a_t)), \quad (146)$$

871 where the $O(1)$ absorbs $\varepsilon(1+\mu) R^{\text{eff}} + O(\mathbb{E}[q^t]/\mathcal{I}_{\text{maj}}^{\text{eff}})$. This is the exact prerequisite of Proposi-
 872 tion H.5 which yields

$$a(z) = (1+\delta_{\text{maj}}+o(1)) \ln z - \ln \ln z + O(1), \quad (147)$$

$$z^{\delta_{\text{maj}}+o(1)} \mathcal{I}_{\text{maj}}(a(z)) = \frac{1}{((1-\varepsilon)(1+\mu_A) + O(1))(1+\delta_{\text{maj}})} \cdot \frac{e^{O(1)}}{z} = \Theta(1/z). \quad (148)$$

873 leading to

$$\mathbb{E}_{\mathcal{G}_{\text{maj}}}[1-p_y^t] = \mathcal{J}_{\text{maj}}^{\text{eff}}(a_t) = \frac{e^{O(1)}}{((1-\varepsilon)(1+\mu_A) + O(1))(1+\delta_{\text{maj}})} \cdot \frac{1}{z_t}. \quad (149)$$

874 Regarding the b -dynamic, the ratio $\dot{b}/\dot{a}$ (Eq (127)), with $R^{\text{eff}} \rightarrow 0$ and $\mathbb{E}[q^t] \leq C' \mathcal{I}_{\text{maj}}^{\text{eff}}(a_t) (1+o(1))$:

$$\begin{aligned} \frac{\dot{b}}{\dot{a}} &= \frac{(1-\varepsilon)(1+\mu) + O(R^{\text{eff}}) + O(\mathbb{E}[q^t]/\mathcal{I}_{\text{maj}}^{\text{eff}})}{(1-\varepsilon)(1+\mu_A) + O(R^{\text{eff}}) + O(\mathbb{E}[q^t]/\mathcal{I}_{\text{maj}}^{\text{eff}})} \\ &= \frac{(1-\varepsilon)(1+\mu) + O(C_\xi)}{(1-\varepsilon)(1+\mu_A) + O(C_\xi)} + O(R^{\text{eff}}) \\ &= \frac{\alpha}{\tilde{\gamma}_{\text{min}}} + O(C_\xi) + O(R^{\text{eff}}), \end{aligned} \quad (150)$$

875 where the $O(C_\xi)$ terms absorb the noise contribution $O(\mathbb{E}[q^t]/\mathcal{I}_{\text{maj}}^{\text{eff}}) = O(R |A_g \mathbf{v}|/\tilde{\gamma}_{\text{maj}}) = O(C_\xi)$
 876 and converge to 0 as $R \rightarrow 0$ uniformly in t .

877 From (145) with $\alpha_\xi > \frac{1+\delta_{\text{min}}}{1+\delta_{\text{maj}}}$: $R^{\text{eff}} = O(z^{-p})$ for some $p > 0$, so $\dot{a} R^{\text{eff}} = O(z^{-p-1})$. Therefore,
 878 integrating Eq (150):

$$\tilde{\gamma}_{\min} b(z) = (\alpha + O(C_\xi)) a(z) + O(1) = \alpha_\xi a(z) + O(1), \quad (151)$$

879 with α_ξ defined in Eq (140). Substituting (147):

$$\tilde{\gamma}_{\min} b(z) = \alpha_\xi (1 + \delta_{\text{maj}} + o(1)) \ln z - \alpha_\xi \ln \ln z + O(1), \quad (152)$$

$$e^{-\tilde{\gamma}_{\min} b_t} = z_t^{-\alpha_\xi (1 + \delta_{\text{maj}}) + o(1)} \cdot (\ln z_t)^{\alpha_\xi} \cdot e^{O(1)} \quad (153)$$

880 Mixing everything with Eq (121):

$$\mathbb{E}_{\mathcal{G}_{\min}}[1 - p_y^t] = \mathcal{J}_{\min}^{\text{eff}}(b_t) = z_t^{\delta_{\min} + o(1)} \frac{\lambda_{\min}(\tilde{\gamma}_{\min}) e^{-\tilde{\gamma}_{\min} b_t}}{b_t} \quad (154)$$

$$\begin{aligned} &= z_t^{\delta_{\min} - \alpha_\xi (1 + \delta_{\text{maj}}) + o(1)} (\ln z_t)^{\alpha_\xi - 1} \cdot \Theta(1) \\ &= \Theta\left(\frac{(\ln z_t)^{\alpha_\xi - 1}}{z_t^{\alpha_\xi (1 + \delta_{\text{maj}}) - \delta_{\min} + o(1)}}\right). \end{aligned} \quad (155)$$

Proposition F.7 (Rate decay, $\alpha_\xi > \frac{1 + \delta_{\min}}{1 + \delta_{\text{maj}}}$). When $\alpha_\xi > \frac{1 + \delta_{\min} + o(1)}{1 + \delta_{\text{maj}}} \geq 1$ with α_ξ defined in Eq (140), we obtain the rates Eq (149) and Eq (155). They present the same rates as in the noise-free case (Theorem D.4)-part(ii), with the $O(1)$ and $o(1)$ extra-terms depending on the noise ξ . Especially when $R = 0$, we recover the exact noise-free rates.

881

882 G Proof of Corollary 7.1

883 We first recall the general accuracy ceiling under correlation shift Corollary 7.1.

Corollary G.1 (Accuracy ceiling under correlation shift). *Assume the infimum in (17) is achieved at the boundary $Z = \tilde{\gamma}_{\text{test}}$:*

$$\gamma_{\text{test}} = \gamma_{\text{test}}^{(r)} - \delta_{\text{test}}, \quad \text{with} \quad \gamma_{\text{test}}^{(r)} := \text{ess inf}_{\mathbf{r}|Z=\tilde{\gamma}_{\text{test}}} \{\hat{\chi} \cdot \mathbf{r}\}. \quad (156)$$

We have the following cases:

(i) *If $\gamma_{\text{test}}(C) > 0$ (test group still separated by $\hat{w}$), then*

$$\mathbb{E}_{\mathcal{D}_{\text{test}}}[1 - p_y(\mathbf{x}, w^t)] = \Theta\left(\frac{(\log z_t)^{\gamma_{\text{test}} + \delta_{\text{test}} + \tilde{\delta}_{\text{test}} - 1}}{z_t^{\gamma_{\text{test}} + o(1)}}\right), \quad (157)$$

where the $o(1)$ is uniformly noise dependent. The test group error still vanishes, but at a rate governed by $\gamma_{\text{test}}(C)$.

(ii) *If $\gamma_{\text{test}}(C) < 0$ (test group not separated by $\hat{w}$):*

$$\lim_{t \rightarrow \infty} \mathbb{E}_{\mathcal{D}_{\text{test}}}[1 - p_y(\mathbf{x}, w^t)] = P_{\mathcal{D}_{\text{test}}}(\hat{\chi} \cdot \mathbf{r} + \hat{w}_s \cdot \xi < 0) > 0. \quad (158)$$

884

885 Part (ii) reveals a non-vanishing error floor caused by the correlation shift.

886 *Proof.* The test logit at time t is: $\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi$ with $\chi_t := w_r^t + C^\top w_s^t$ and by Proposition 5.4,

$$\chi_t = (\log z_t - \log \log z_t) \hat{\chi} + O(1), \quad w_s^t = (\log z_t - \log \log z_t) \hat{\mathbf{w}}_s + O(1). \quad (159)$$

887 The test error is:

$$\begin{aligned} \mathbb{E}_{\mathcal{D}_{\text{test}}}[1 - p_y^t] &= \mathbb{E}[\sigma(-(\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi))] = \mathbb{E}_Z[\mathbb{E}_{\mathbf{r}_\perp, \xi|Z}[\sigma(-\chi_t \cdot \mathbf{r} - w_s^t \cdot \xi)]] \\ &= \int_{\tilde{\gamma}_{\text{test}}}^K \mathbb{E}_{\mathbf{r}_\perp, \xi|Z=z}[\sigma(-\chi_t \cdot \mathbf{r} - w_s^t \cdot \xi)] \lambda_{\text{test}}(z) dz \end{aligned} \quad (160)$$

888 *Proof of (i).* Assume $\gamma_{\text{test}} > 0$.

889 By Eq (17), $\hat{\chi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi \geq \gamma_{\text{test}} > 0$ a.s., so with Eq (159):

$$\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi = (\hat{\chi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi)(\log z_t - \log \log z_t) + O(1) \geq \gamma_{\text{test}}(\log z_t - \log \log z_t) + O(1) > 0 \quad (161)$$

890 a.s. for t large. Using the identity $\sigma(-v) = e^{-v}/(1 + e^{-v})$ for $v > 0$, we write pointwise:

$$\sigma(-(\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi)) = e^{-(\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi)} \cdot \frac{1}{1 + e^{-(\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi)}} = e^{-\chi_t \cdot \mathbf{r}} e^{-w_s^t \cdot \xi} \tilde{h}_t(\mathbf{r}, \xi), \quad (162)$$

891 where $\tilde{h}_t(\mathbf{r}, \xi) := 1/(1 + e^{-(\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi)}) \in (1/2, 1)$ since $\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi > 0$. Therefore

$$\sigma(-(\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi)) = \Theta(e^{-\chi_t \cdot \mathbf{r}} e^{-w_s^t \cdot \xi}) \quad \text{a.s. for } t \text{ large, uniformly in } \mathbf{r}, \xi. \quad (163)$$

892 By Eq (163), the integrand in Eq (160) satisfies

$$\begin{aligned} \mathbb{E}_{\mathbf{r}_\perp, \xi|Z=z}[\sigma(-\chi_t \cdot \mathbf{r} - w_s^t \cdot \xi)] \lambda_{\text{test}}(z) &= \lambda_{\text{test}}(z) \mathbb{E}_{\mathbf{r}_\perp, \xi|Z=z}[\Theta(e^{-\chi_t \cdot \mathbf{r}} e^{-w_s^t \cdot \xi})] \\ &= \Lambda_\xi(w_s^t, z) \lambda_{\text{test}}(z) \mathbb{E}_{\mathbf{r}_\perp|Z=z}[\Theta(e^{-\chi_t \cdot \mathbf{r}})] \quad (\text{see Eq (105)}) \\ &= e^{-z \chi_t \cdot \mathbf{v}} \Lambda_\xi(w_s^t, z) \lambda_{\text{test}}(z) \Theta(\mathbb{E}_{\mathbf{r}_\perp|Z=z}[e^{-\chi_t \cdot \mathbf{r}_\perp}]) \quad (\text{all } \Theta(\dots) \text{ are uniform in } \mathbf{r}, z) \\ &= \Theta(e^{-z \chi_t \cdot \mathbf{v}} \Lambda_\xi(w_s^t, z) \Lambda_{\mathbf{r}_\perp}(\chi_t, z) \lambda_{\text{test}}(z)) \quad (\text{with Eq (159)}), \end{aligned} \quad (164)$$

893 where

$$\Lambda_{\mathbf{r}_\perp}(w_r, z) := \mathbb{E}_{\mathbf{r}_\perp|Z=z}[e^{-w_r \cdot \mathbf{r}_\perp}]. \quad (165)$$

894 Therefore, Eq (160) rewrites thanks to the mean-value theorem:

$$\begin{aligned}
\mathbb{E}_{\mathcal{D}_{\text{test}}}[1 - p_y^t] &= \Theta(1) \cdot \int_{\tilde{\gamma}_{\text{test}}}^K e^{-z \chi_t \cdot \mathbf{v}} \Lambda_{\xi}(w_s^t, z) \Lambda_{\mathbf{r}_{\perp}}(\chi^t, z) \lambda_{\text{test}}(z) dz \\
&= \Theta(1) \Lambda_{\xi}(w_s^t, z_{\text{test}}^*) \Lambda_{\mathbf{r}_{\perp}}(\chi^t, z_{\text{test}}^*) \cdot \int_{\tilde{\gamma}_{\text{test}}}^K e^{-z \chi_t \cdot \mathbf{v}} \lambda_{\text{test}}(z) dz, \tag{166}
\end{aligned}$$

895 for some t -dependent $z_{\text{test}}^* \in [\tilde{\gamma}_{\text{test}}, K]$. By Lemma F.4:

$$\Lambda_{\xi}(w_s^t, z_{\text{test}}^*) = z_t^{\delta_+(z_{\text{test}}^*) + o(1)}, \quad \Lambda_{\mathbf{r}_{\perp}}(\chi^t, z_{\text{test}}^*) = z_t^{\tilde{\delta}_+(z_{\text{test}}^*) + o(1)}, \tag{167}$$

896 where $\tilde{\delta}_+$ is defined exactly as δ_+ in Eq (117) and has the same properties (due to separability
897 $\gamma_{\text{test}} > 0$):

$$\tilde{\delta}_+(z) := \sup_{\mathbf{r}_{\perp} \in \text{supp}(\mathbf{r}_{\perp} | Z=z)} -\hat{\chi} \cdot \mathbf{r}_{\perp} \tag{168}$$

898 Since the Laplace weight $e^{-z \chi_t \cdot \mathbf{v}} = e^{-z(\log z_t - \log \log z_t) \hat{\chi} \cdot \mathbf{v} + O(1)}$ concentrates at $z = \tilde{\gamma}_{\text{test}}$, the mean
899 value theorem points satisfy $z_{\text{test}}^* \rightarrow \tilde{\gamma}_{\text{test}}$, hence $\delta_+(z_{\text{test}}^*) \rightarrow \delta_{\text{test}}$ and $\tilde{\delta}_+(z_{\text{test}}^*) \rightarrow \tilde{\delta}_+(\tilde{\gamma}_{\text{test}}) = \tilde{\delta}_{\text{test}}$.

900 For the remaining integral, a similar Watson Lemma argument as in Lemma H.1 gives:

$$\begin{aligned}
\int_{\tilde{\gamma}_{\text{test}}}^K e^{-z \chi_t \cdot \mathbf{v}} \lambda_{\text{test}}(z) dz &\sim \frac{\lambda_{\text{test}}(\tilde{\gamma}_{\text{test}}) e^{-\tilde{\gamma}_{\text{test}} \chi_t \cdot \mathbf{v}}}{\chi_t \cdot \mathbf{v}} \\
&\sim \frac{\lambda_{\text{test}}(\tilde{\gamma}_{\text{test}}) z_t^{-\tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v}} (\log z_t)^{\tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v}} e^{O(1)}}{\hat{\chi} \cdot \mathbf{v} \log z_t} \\
&= \Theta\left(z_t^{-\tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v}} (\log z_t)^{\tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v} - 1}\right) \tag{169}
\end{aligned}$$

901 But the Assumption (156) reads,

$$\begin{aligned}
\gamma_{\text{test}}^{(r)} &= \text{ess inf}_{\mathbf{r} | Z=\tilde{\gamma}_{\text{test}}} \{\hat{\chi} \cdot \mathbf{r}\} = \tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v} + \text{ess inf}_{\mathbf{r}_{\perp} | Z=\tilde{\gamma}_{\text{test}}} \{\hat{\chi} \cdot \mathbf{r}_{\perp}\} = \tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v} - \tilde{\delta}_{\text{test}}, \\
\gamma_{\text{test}} &= \gamma_{\text{test}}^{(r)} - \delta_{\text{test}} = \tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v} - \tilde{\delta}_{\text{test}} - \delta_{\text{test}}. \tag{170}
\end{aligned}$$

902 Therefore combining all factors:

$$\begin{aligned}
\mathbb{E}_{\mathcal{D}_{\text{test}}}[1 - p_y^t] &= \Theta(1) \cdot z_t^{\delta_{\text{test}} + \tilde{\delta}_{\text{test}} + o(1)} \cdot z_t^{-\tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v}} (\log z_t)^{\tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v} - 1} \\
&= \Theta(1) \cdot z_t^{-\gamma_{\text{test}} + o(1)} \cdot (\log z_t)^{\tilde{\gamma}_{\text{test}} \hat{\chi} \cdot \mathbf{v} - 1} \\
&= \Theta\left(\frac{(\log z_t)^{\gamma_{\text{test}} + \delta_{\text{test}} + \tilde{\delta}_{\text{test}} - 1}}{z_t^{\gamma_{\text{test}} + o(1)}}\right). \tag{171}
\end{aligned}$$

903 *Proof of (ii).* Assume $\gamma_{\text{test}} < 0$.

904 Writing the test-time logit $\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi = (\hat{\chi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi)(\log z_t - \log \log z_t) + O(1)$, we have:

905 On $\{\hat{\chi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi > 0\}$: $\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi \rightarrow +\infty$, hence $\sigma(-(\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi)) \rightarrow 0$.

906 On $\{\hat{\chi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi < 0\}$: $\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi \rightarrow -\infty$, hence $\sigma(-(\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi)) \rightarrow 1$.

907 Since $\gamma_{\text{test}} = \text{ess inf} \hat{\chi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi < 0$, the set $\{\hat{\chi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi < 0\}$ has positive probability. By
908 dominated convergence:

$$\lim_{t \rightarrow \infty} \mathbb{E}[\sigma(-(\chi_t \cdot \mathbf{r} + w_s^t \cdot \xi))] = P(\hat{\chi} \cdot \mathbf{r} + \hat{\mathbf{w}}_s \cdot \xi < 0) > 0. \tag{172}$$

909 □

910 G.1 Application to the Proof of Theorem 5.1

911 A natural application of the analysis presented above is the proof of Theorem 5.1.

912 The isotropic regime in Theorem 5.3 was mathematically load-bearing because it pushed the pa-
913 rameter updates to remain confined to $\text{Span}\{\mathbf{v}\}$ (Proposition E.2), facilitating the reduction of the

914 problem into a 2D-ODE (63). When this condition fails, the gradient continuously injects a perpen-
 915 dicular component into ψ_t, ϕ_t , physically rotating the parameters out of $\text{Span}\{\mathbf{v}\}$. This removes the
 916 possibility to reduce the problem to its projection on $\text{Span}\{\mathbf{v}\}$ utilized in Proposition E.3.
 917 Recall the theorem describing learning dynamics in this general regime:

Theorem G.2 (General Decay Rates). *Assume the infimums in (25) are achieved at the bound-
 aries $Z = \tilde{\gamma}_{\text{maj}}$, and $Z = \tilde{\gamma}_{\text{min}}$, then*

$$\mathbb{E}_{\mathcal{G}_{\text{maj}}}[1 - p_y(\mathbf{x}, w^t)] = \Theta\left(\frac{(\log z_t)^{\gamma_{\text{maj}} + \delta_{\text{maj}} + \tilde{\delta}_{\text{maj}} - 1}}{z_t^{\gamma_{\text{maj}} + o(1)}}\right), \quad (173)$$

$$\mathbb{E}_{\mathcal{G}_{\text{min}}}[1 - p_y(\mathbf{x}, w^t)] = \Theta\left(\frac{(\log z_t)^{\gamma_{\text{min}} + \delta_{\text{min}} + \tilde{\delta}_{\text{min}} - 1}}{z_t^{\gamma_{\text{min}} + o(1)}}\right), \quad (174)$$

where γ_g are the group-specific hard-margin defined in Eq (25) and the $o(1)$ are uniformly
 noise-dependent.

918

919 *Proof.* The proof follows directly from part (i) of Corollary 7.1 with $\mathcal{D}_{\text{test}} = \mathcal{G}_{\text{maj}}$ and $\mathcal{D}_{\text{test}} = \mathcal{G}_{\text{min}}$
 920 respectively, and with the $\mathbf{r}_\perp$ -boundary exponents for each group:

$$\tilde{\delta}_g := \tilde{\delta}_+(\tilde{\gamma}_g), \quad g \in \{\text{maj}, \text{min}\}. \quad (175)$$

921

□

922 H Toolbox

923 H.1 Asymptotic Behavior of $\mathcal{J}$ and $\mathcal{I}$

Lemma H.1. *Let λ be a continuous probability density function on $[\tilde{\gamma}, K]$ with $\lambda(\tilde{\gamma}) > 0$, and let $\mathcal{J}(x) := \int_{\tilde{\gamma}}^K \frac{\lambda(z)}{1+e^{xz}} dz$. Then,*

$$\mathcal{J}(x) \sim \frac{\lambda(\tilde{\gamma}) e^{-\tilde{\gamma}x}}{x} \quad \text{and} \quad \mathcal{J}'(x) \sim -\frac{\tilde{\gamma} \lambda(\tilde{\gamma}) e^{-\tilde{\gamma}x}}{x} \quad \text{as } x \rightarrow +\infty. \quad (176)$$

924 *Proof.* We proceed in two steps. For $x > 0$,

$$\begin{aligned} \left| \mathcal{J}(x) - \int_{\tilde{\gamma}}^K e^{-xz} \lambda(z) dz \right| &= \left| - \int_{\tilde{\gamma}}^K \lambda(z) \frac{e^{-xz}}{1+e^{xz}} dz \right| \\ &\leq e^{-\tilde{\gamma}x} \int_{\tilde{\gamma}}^K \frac{\lambda(z)}{1+e^{xz}} dz \leq e^{-\tilde{\gamma}x} \int_{\tilde{\gamma}}^K e^{-xz} \lambda(z) dz \\ &= o\left(\int_{\tilde{\gamma}}^K e^{-xz} \lambda(z) dz \right) \quad \text{as } x \rightarrow +\infty, \end{aligned}$$

926 Hence $\mathcal{J}(x) \sim \int_{\tilde{\gamma}}^K e^{-xz} \lambda(z) dz$. By Watson's Lemma Wong (2001, Chapter 1):

$$\int_{\tilde{\gamma}}^K e^{-xz} \lambda(z) dz \sim \frac{\lambda(\tilde{\gamma}) e^{-\tilde{\gamma}x}}{x} \quad \text{as } x \rightarrow +\infty.$$

927 Combining everything yields $\mathcal{J}(x) \sim \lambda(\tilde{\gamma}) e^{-\tilde{\gamma}x}/x$.

928 We use the same arguments for $\mathcal{J}'(x) = - \int_{\tilde{\gamma}}^K \frac{z \lambda(z) e^{xz}}{(1+e^{xz})^2} dz \sim - \int_{\tilde{\gamma}}^K z \lambda(z) e^{-xz} dz$, which yields

929 $\mathcal{J}'(x) \sim -\tilde{\gamma} \lambda(\tilde{\gamma}) e^{-\tilde{\gamma}x}/x$. □

Lemma H.2. *Assume $\lambda \in C^1([\gamma, K])$ with $\lambda(\gamma) > 0$. Define $\mathcal{I}(x) := \int_{\gamma}^K \frac{z \lambda(z)}{1+e^{xz}} dz$, then*

$$\mathcal{I}(x) = \frac{\gamma \lambda(\gamma) e^{-\gamma x}}{x} \left[1 + \frac{\nu}{x} + O\left(\frac{1}{x^2}\right) \right] \quad \text{as } x \rightarrow +\infty, \quad (177)$$

where $\nu := 1/\gamma + \lambda'(\gamma)/\lambda(\gamma)$.

In particular, writing $r(x) := \mathcal{I}(x)/(\gamma \lambda(\gamma) e^{-\gamma x}/x) - 1$, one has $|r(x)| = O(1/x)$.

930 *Proof.* We reuse the same approach as in Lemma H.1:

$$\begin{aligned} \left| \mathcal{I}(x) - \int_{\gamma}^K z \lambda(z) e^{-xz} dz \right| &= \left| \int_{\gamma}^K z \lambda(z) \frac{e^{-xz}}{1+e^{xz}} dz \right| \\ &\leq e^{-2\gamma x} \int_{\gamma}^K z \lambda(z) dz = O(e^{-2\gamma x}) \quad \text{as } x \rightarrow +\infty. \end{aligned}$$

932 Hence as $x \rightarrow +\infty$,

$$\mathcal{I}(x) = \int_{\gamma}^K z \lambda(z) e^{-xz} dz + O(e^{-2\gamma x}) = \int_{\gamma}^K z \lambda(z) e^{-xz} dz + O\left(\frac{e^{-\gamma x}}{|x|^3}\right) \quad (178)$$

933 .

934 Now, substitute $z = \gamma + t/x$ in $\int_{\gamma}^K \phi(z) e^{-xz} dz$ with $\phi(z) := z \lambda(z)$, giving

$$\int_{\gamma}^K \phi(z) e^{-xz} dz = \frac{e^{-\gamma x}}{x} \int_0^{x(K-\gamma)} e^{-t} \phi\left(\gamma + \frac{t}{x}\right) dt.$$

935 Taylor-expanding $\phi(\gamma + s) = \phi(\gamma) + \phi'(\gamma)s + \epsilon(s)s^2$ with ϵ bounded on $[0, K - \gamma]$ by $\|\epsilon\|_{\infty}$.
 936 Integrating term by term then yields

$$\begin{aligned} \int_0^{x(K-\gamma)} e^{-t} \phi(\gamma) dt &= \phi(\gamma) + O(e^{-x(K-\gamma)}) = \phi(\gamma) + O\left(\frac{1}{x^2}\right) \\ \int_0^{x(K-\gamma)} e^{-t} \frac{t}{x} \phi'(\gamma) dt &= \frac{\phi'(\gamma)}{x} + O(e^{-x(K-\gamma)}) = \frac{\phi'(\gamma)}{x} + O\left(\frac{1}{x^2}\right) \\ \left| \int_0^{x(K-\gamma)} e^{-t} \epsilon\left(\frac{t}{x}\right) \frac{t^2}{x^2} dt \right| &\leq \frac{\|\epsilon\|_{\infty}}{x^2} + O(e^{-x(K-\gamma)}) = O\left(\frac{1}{x^2}\right). \end{aligned}$$

937 Finally,

$$\int_{\gamma}^K \phi(z) e^{-xz} dz = \frac{e^{-\gamma x}}{x} \left(\phi(\gamma) + \frac{\phi'(\gamma)}{x} + O\left(\frac{1}{x^2}\right) \right),$$

938 with $\phi(\gamma) = \gamma \lambda(\gamma)$, $\phi'(\gamma) = \lambda(\gamma) + \gamma \lambda'(\gamma)$. We then get

$$\mathcal{I}(x) = \frac{\gamma \lambda(\gamma) e^{-\gamma x}}{x} + \frac{(\lambda(\gamma) + \gamma \lambda'(\gamma)) e^{-\gamma x}}{x^2} + O\left(\frac{e^{-\gamma x}}{x^3}\right),$$

939 which rearranges into (177). □

940 H.2 Asymptotic Behavior of the Solution of 1D ODEs

Proposition H.3 (Non-autonomous ODE). *Let $\mathcal{I} \in C(\mathbb{R}^+, \mathbb{R}_+^*)$ satisfy $\mathcal{I}(x) \sim \gamma \lambda_0 e^{-\gamma x}/x$ as $x \rightarrow +\infty$ for some $\lambda_0 > 0$ and $\gamma > 0$. Let $c : [z_0, \infty) \rightarrow \mathbb{R}$ be continuous and assume there exist $0 < c_- \leq c_+ < \infty$ such that $c_- \leq c(z) \leq c_+$ for all $z \geq z_0$. Then any solution of the Cauchy problem*

$$\dot{v}(z) = c(z) \mathcal{I}(v(z)), \quad v(z_0) \in \mathbb{R}^+,$$

with $v(z_0)$ sufficiently large, satisfies $v(z) \rightarrow +\infty$ monotonically and

$$v(z) = \frac{1}{\gamma} \left(\ln z - \ln \ln z + O(1) \right) \quad \text{as } z \rightarrow +\infty. \quad (179)$$

Consequently $\mathcal{I}(v(z)) = \Theta(1/z)$: there exist constants $0 < M_- \leq M_+ < \infty$ and $z_1 \geq z_0$ such that

$$\frac{M_-}{z} \leq \mathcal{I}(v(z)) \leq \frac{M_+}{z} \quad \text{for all } z \geq z_1.$$

941

942 *Proof.* Since $\mathcal{I}(v) > 0$ and $c(z) > 0$, $\dot{v}(z) > 0$, hence $v(\cdot)$ is strictly increasing and $v(z) \rightarrow +\infty$
 943 (monotonicity rules out any finite limit because $\mathcal{I}(l) > 0$ at any finite $l > 0$).

944 Let $F(x)$ be a primitive of $1/\mathcal{I}$ on $\mathbb{R}^+$. Separating variables,

$$F(v(z)) - F(v(z_0)) = \int_{z_0}^z c(s) ds. \quad (180)$$

945 Given that, $1/\mathcal{I}(x) \sim C x e^{\gamma x}$ as $x \rightarrow +\infty$ with $C = 1/(\gamma \lambda_0)$, and since $x e^{\gamma x}$ is non-integrable on
 946 $(0, \infty)$ the theorem of integration of asymptotic equivalents (Jourdain, 2018, *Theorem 4.11*) yields

$$F(x) \sim C x e^{\gamma x} / \gamma \quad \text{as } x \rightarrow +\infty. \quad (181)$$

947 Meanwhile, the hypothesis $c_- \leq c \leq c_+$ gives that the right-hand side of (180) is a $\Theta(z)$. Combining
 948 with (181),

$$C v(z) e^{\gamma v(z)} (1/\gamma + o(1)) = \Theta(z), \quad \text{i.e.} \quad \gamma v(z) e^{\gamma v(z)} = \Theta(z).$$

949 This means $\gamma v e^{\gamma v} \in [\alpha z, \beta z]$ (for some $0 < \alpha \leq \beta$ and z large). We apply the principal real
 950 branch of the Lambert W function (Visser, 2018, Eq. A2): $\gamma v(z) \in [W(\alpha z), W(\beta z)]$ with $W(y) =$
 951 $\ln y - \ln \ln y + O(1)$ as $y \rightarrow +\infty$. Hence

$$v(z) = \frac{1}{\gamma} \left(\ln z - \ln \ln z + O(1) \right),$$

952 which proves (179).

953 Plugging into $\mathcal{I}(x) = \gamma \lambda_0 e^{-\gamma x} x^{-1} (1 + o(1))$ gives

$$\begin{aligned} \mathcal{I}(v(z)) &= \frac{\gamma \lambda_0 e^{-\gamma v(z)}}{v(z)} (1 + o(1)) \\ &= \frac{\gamma^2 \lambda_0 e^{-\ln z + \ln \ln z + O(1)}}{\ln z (1 + o(1))} (1 + o(1)) \\ &= \frac{\gamma^2 \lambda_0}{z} e^{O(1)} (1 + o(1)) = \Theta(1/z), \end{aligned}$$

954 which gives the stated $M_{\pm}$ bounds. □

Proposition H.4 (Perturbed autonomous ODE). *Let $c_0 > 0$, let $\mathcal{I}$ be a continuous function satisfying $\mathcal{I}(x) \sim \gamma \lambda_0 e^{-\gamma x}/x$ as $x \rightarrow +\infty$ for some $\lambda_0 > 0$, $\gamma > 0$, and let $\delta: [0, \infty) \rightarrow \mathbb{R}$ be a continuous function with $\delta(z) \rightarrow 0$ as $z \rightarrow \infty$. The solution v of the perturbed Cauchy problem*

$$v'(z) = c_0 \mathcal{I}(v(z)) (1 + \delta(z)), \quad v(0) = v_0, \quad (182)$$

satisfies the asymptotic expansion:

$$\gamma v(z) = \ln z - \ln \ln z + \ln(\gamma^3 c_0 \lambda_0) + o(1) \quad \text{as } z \rightarrow +\infty. \quad (183)$$

955

956 *Proof.* We proceed step by step.

957 Dividing both sides of (182) by $c_0 \mathcal{I}(v)$ and integrating from 0 to z :

$$F(v(z)) - F(v_0) = z + \underbrace{\int_0^z \delta(s) ds}_{=: D(z)}, \quad (184)$$

958 where F is any primitive of $1/(c_0 \mathcal{I})$. Since $\delta(z) \rightarrow 0$ as $z \rightarrow \infty$, the Cesàro mean converges to the
 959 same limit:

$$\frac{D(z)}{z} = \frac{1}{z} \int_0^z \delta(s) ds \xrightarrow{z \rightarrow \infty} 0,$$

960 i.e. $D(z) = o(z)$, and Eq (184) becomes

$$F(v(z)) = z (1 + o(1)) + F(v_0). \quad (185)$$

961 By the hypothesis $\mathcal{I}(x) \sim \gamma \lambda_0 e^{-\gamma x}/x$ and the integration-of-equivalents theorem (Jourdain, 2018,
 962 *Theorem 4.11*) (the integrand $x e^{\gamma x}$ is not integrable on $[0, +\infty)$, so the theorem applies):

$$F(x) = \int^x \frac{du}{c_0 \mathcal{I}(u)} \sim C \int^x u e^{\gamma u} du = \frac{C}{\gamma^2} e^{\gamma x} (\gamma x - 1), \quad C := \frac{1}{c_0 \gamma \lambda_0}, \quad (186)$$

963 Since $v(z) \rightarrow +\infty$ (by the same argument as in Proposition H.3), combining Eq (185) with Eq (186):

$$\begin{aligned} \frac{C}{\gamma^2} \gamma v e^{\gamma v} (1 + o(1)) &= z (1 + o(1)) \\ \gamma v e^{\gamma v} &= \frac{\gamma^2 z}{C} (1 + o(1)) = \gamma^3 c_0 \lambda_0 z (1 + o(1)). \end{aligned} \quad (187)$$

964 Using the principal real branch of the Lambert W function (Visser, 2018, Eq. A2):

$$\begin{aligned}\gamma v(z) &= W(\gamma^3 c_0 \lambda_0 z (1 + o(1))) \\ &= \ln(\gamma^3 c_0 \lambda_0 z) - \ln \ln(\gamma^3 c_0 \lambda_0 z) + o(1) \\ &= \ln z - \ln \ln z + \ln(\gamma^3 c_0 \lambda_0) + o(1).\end{aligned}$$

965

□

Proposition H.5 (Non-autonomous ODE with polynomial factor). *Let $\mathcal{I} \in C(\mathbb{R}^+, \mathbb{R}_+^*)$ satisfy $\mathcal{I}(x) \sim \gamma \lambda_0 e^{-\gamma x}/x$ as $x \rightarrow +\infty$ for some $\lambda_0 > 0$ and $\gamma > 0$. Let $\delta \geq 0$, let $\epsilon(z) = o(1)$, and let $c : [z_0, \infty) \rightarrow \mathbb{R}$ be continuous with $c(z) = c_0 + O(1)$, where $0 < c_- \leq c(z) \leq c_+$ for all $z \geq z_0$. Then the solution of*

$$\dot{v}(z) = c(z) z^{\delta+\epsilon(z)} \mathcal{I}(v(z)), \quad v(z_0) \in \mathbb{R}^+,$$

with $v(z_0)$ large enough, satisfies

$$v(z) = \frac{1}{\gamma} \left((1+\delta+\epsilon(z)) \ln z - \ln \ln z + \ln(\gamma^3 \lambda_0 c_0 + \textcolor{blue}{O}(1)) + O(1) \right) \quad \text{as } z \rightarrow +\infty, \quad (188)$$

where both $O(1)$ are independent on c_0 , $\textcolor{blue}{O}(1)$ is proportional to the input $O(1) = c(z) - c_0$, and

$$z^{\delta+\epsilon(z)} \mathcal{I}(v(z)) = \frac{1}{(\gamma c_0 + \textcolor{blue}{O}(1))(1+\delta)} \cdot \frac{e^{O(1)}}{z} = \Theta(1/z), \quad (189)$$

with both $O(1)$ independent on c_0 , and $\textcolor{blue}{O}(1)$ is proportional to the input $O(1) = c(z) - c_0$.

966

967 *Proof.* Let $G(z) := \int_{z_0}^z c(s) s^{\delta+\epsilon(s)} ds$. Separating variables: $F(v(z)) - F(v(z_0)) = G(z)$ with F
968 as in Proposition H.3.

969 Since $c_- \leq c \leq c_+$:

$$c_- \int_{z_0}^z s^{\delta+\epsilon(s)} ds \leq G(z) \leq c_+ \int_{z_0}^z s^{\delta+\epsilon(s)} ds.$$

970 By Karamata's theorem (Feller et al., 1971, Chapter VIII, Section 9, Theorem 1) (applicable since
971 $\delta \geq 0 > -1$ and with ϵ decreasing fast enough):

$$\int_{z_0}^z s^{\delta+\epsilon(s)} ds \sim \frac{z^{1+\delta+\epsilon(z)}}{1+\delta} \quad \text{as } z \rightarrow +\infty. \quad (190)$$

972 Therefore $G(z) = z^{1+\delta+\epsilon(z)} (c_0 + \textcolor{blue}{O}(1))$. From $F(v) \sim v e^{\gamma v}/(\gamma^2 \lambda_0)$ (see Proposition H.3) and
973 $F(v(z)) \sim G(z)$:

$$v(z) e^{\gamma v(z)} = \gamma^2 \lambda_0 G(z) (1+o(1)) = z^{1+\delta+\epsilon(z)} (\gamma^2 \lambda_0 c_0 + \textcolor{blue}{O}(1)).$$

974 Applying the Lambert W function:

$$\begin{aligned}\gamma v(z) &= W\left(z^{1+\delta+\epsilon(z)} (\gamma^3 \lambda_0 c_0 + \textcolor{blue}{O}(1))\right) \\ &= \ln\left(z^{1+\delta+\epsilon(z)} (\gamma^3 \lambda_0 c_0 + \textcolor{blue}{O}(1))\right) - \ln \ln\left(z^{1+\delta+\epsilon(z)} (\gamma^3 \lambda_0 c_0 + \textcolor{blue}{O}(1))\right) + O(1) \\ &= (1+\delta+\epsilon(z)) \ln z + \ln(\gamma^3 \lambda_0 c_0 + \textcolor{blue}{O}(1)) - \ln\left((1+\delta+\epsilon(z)) \ln z + O(1)\right) + O(1) \\ &= (1+\delta+\epsilon(z)) \ln z - \ln \ln z + \ln(\gamma^3 \lambda_0 c_0 + \textcolor{blue}{O}(1)) + O(1),\end{aligned}$$

975 giving Eq (188).

976 Plug (188) into $\mathcal{I}(x) = \gamma \lambda_0 e^{-\gamma x} x^{-1} (1+o(1))$:

$$\begin{aligned}
z^{\delta+\epsilon(z)} \mathcal{I}(v(z)) &= z^{\delta+\epsilon(z)} \cdot \frac{\gamma \lambda_0 e^{-(1+\delta+\epsilon(z)) \ln z + \ln \ln z - \ln(\gamma^3 \lambda_0 c_0 + \mathcal{O}(1))} + O(1)}{v(z)} (1+o(1)) \\
&= z^{\delta+\epsilon(z)} \cdot \frac{\gamma \lambda_0 z^{-(1+\delta+\epsilon(z))} \ln z \cdot e^{O(1)}}{(\gamma^3 \lambda_0 c_0 + \mathcal{O}(1))^{\frac{1+\delta+\epsilon(z)}{\gamma}} \ln z (1+o(1))} (1+o(1)) \\
&= \frac{1}{(\gamma c_0 + \mathcal{O}(1))(1+\delta)} \cdot \frac{e^{O(1)}}{z} (1+o(1)) = \Theta(1/z).
\end{aligned}$$

977

□

978 H.3 Euler Scheme Error Bound

Lemma H.6. Let $\mathcal{I} \in C^1(\mathbb{R}^+, \mathbb{R}_+^*)$ one of the function defined in Eq (55). Let $(c_t)_{t \geq 0}$ be a sequence satisfying $c_t \in [c_-, c_+]$ for some $0 < c_- \leq c_+ < \infty$ and all t large enough. Consider the scalar recurrence

$$b_{t+1} = b_t + h_t c_t \mathcal{I}(b_t), \quad b_0 \geq 1, \quad (191)$$

where $(h_t)_{t \geq 0}$ satisfies Condition (21), and set $z_t := \sum_{k=0}^{t-1} h_k$.

Let $b(\cdot)$ be the solution of the perturbed ODE $\dot{b}(z) = c(z) \mathcal{I}(b(z))$ with coefficient $c(z) \in [c_-, c_+]$ and the same initial condition $b(0) = b_0$. Then:

- (i) $b_t \rightarrow +\infty$ monotonically.
- (ii) $b_t = \ln z_t - \ln \ln z_t + O(1)$ as $t \rightarrow \infty$.
- (iii) If additionally $c_t = c_0 + \delta_t$ with $\delta_t \rightarrow 0$, then

$$b_t = b(z_t) + o(1) \quad \text{as } t \rightarrow \infty. \quad (192)$$

979

980 *Proof.* (i) Since $\mathcal{I}(b_t) > 0$, $c_t > 0$, and $h_t > 0$, the increment $h_t c_t \mathcal{I}(b_t) > 0$, so (b_t) is strictly
981 increasing. If it were upper bounded by some $l < \infty$, then since $\mathcal{I}$ is decreasing, $\infty > \sum_t b_{t+1} - b_t \geq$
982 $c_- \mathcal{I}(l) \sum h_t = +\infty$, contradicting boundedness. Hence $b_t \rightarrow +\infty$.

983 (ii) Define the primitive $G(x) := \int_{b_0}^x \frac{du}{\mathcal{I}(u)}$.

984 By Taylor's theorem with the increment $\eta_t := b_{t+1} - b_t = h_t c_t \mathcal{I}(b_t)$:

$$\begin{aligned}
G(b_{t+1}) - G(b_t) &= G'(b_t) \eta_t + \frac{1}{2} G''(\xi_t) \eta_t^2 \\
&= \frac{h_t c_t \mathcal{I}(b_t)}{\mathcal{I}(b_t)} + \frac{1}{2} G''(\xi_t) \eta_t^2 \\
&= h_t c_t + \frac{1}{2} G''(\xi_t) \eta_t^2,
\end{aligned} \quad (193)$$

985 for some $\xi_t \in [b_t, b_{t+1}]$. Since $G''(x) = -\mathcal{I}'(x)/\mathcal{I}(x)^2$, and by Lemma H.1 (with $z\lambda(z)$ instead of
986 $\lambda(z)$):

$$\mathcal{I}(x) \sim \frac{\gamma \lambda(\gamma) e^{-\gamma x}}{x}, \quad \mathcal{I}'(x) \sim -\frac{\gamma^2 \lambda(\gamma) e^{-\gamma x}}{x}, \quad G''(x) = -\frac{\mathcal{I}'(x)}{\mathcal{I}(x)^2} \sim \frac{x e^{\gamma x}}{\lambda(\gamma)} \quad (194)$$

987 Therefore

$$G''(\xi_t) \eta_t^2 \sim \frac{\xi_t e^{\gamma \xi_t}}{\lambda(\gamma)} h_t^2 c_t^2 \mathcal{I}(b_t)^2 \sim \frac{\xi_t e^{\gamma(\xi_t - 2b_t)}}{b_t^2} \gamma^2 \lambda(\gamma) h_t^2 c_t^2 = o(h_t^2),$$

988 where we used that $\xi_t - 2b_t \leq \eta_t - b_t \rightarrow -\infty$ because $\eta_t = h_t c_t \mathcal{I}(b_t)$ is bounded.

989 Summing (193) from 0 to $t-1$:

$$\begin{aligned}
G(b_t) &= \underbrace{G(b_0)}_{=0} + \sum_{k=0}^{t-1} h_k c_k + \sum_{k=0}^{t-1} o(h_k^2) \\
&= \Theta(z_t) + O(1) = \Theta(z_t).
\end{aligned} \quad (195)$$

990 With the asymptotics $G(x) \sim \frac{x e^{\gamma x}}{\gamma^2 \lambda(\gamma)}$ (derived in the proof of Proposition H.3 or H.4), Eq (195)
 991 becomes $\gamma b_t e^{\gamma b_t} = \Theta(z_t)$. We then derive $\gamma b_t = \ln z_t - \ln \ln z_t + O(1)$ by the same Lambert W
 992 asymptotics (see Proposition H.3 or H.4).

993 (iii) Assume $c_t = c_0 + \delta_t$ with $\delta_t \rightarrow 0$. Resumming (193) gives

$$G(b_t) = c_0 z_t + \sum_{k=0}^{t-1} h_k \delta_k + O(1). \quad (196)$$

994 By the Cesàro, since $\delta_k \rightarrow 0$ and $\sum h_k = +\infty$:

$$\frac{\sum_{k=0}^{t-1} h_k \delta_k}{\sum_{k=0}^{t-1} h_k} \xrightarrow{t \rightarrow \infty} 0, \quad \text{i.e.} \quad \sum_{k=0}^{t-1} h_k \delta_k = o(z_t).$$

995 Therefore $G(b_t) = c_0 z_t (1 + o(1))$. The ODE trajectory satisfies the same relation: $G(b(z_t)) =$
 996 $c_0 z_t (1 + o(1))$, by (185) in the proof of Proposition H.4. Hence $G(b_t) - G(b(z_t)) = o(z_t)$. By the
 997 mean-value theorem,

$$b_t - b(z_t) = \frac{G(b_t) - G(b(z_t))}{G'(\xi_t)}$$

998 for some ξ_t between b_t and $b(z_t)$. Both satisfy $\gamma \xi_t = \ln z_t - \ln \ln z_t + O(1)$, so

$$G'(\xi_t) = \frac{1}{\mathcal{I}(\xi_t)} \sim \frac{\xi_t e^{\gamma \xi_t}}{\gamma \lambda(\gamma)} = \Theta(z_t).$$

999 Therefore

$$|b_t - b(z_t)| = \frac{o(z_t)}{\Theta(z_t)} = o(1). \quad \square$$

1000 H.4 Synchronisation Method Lemmas

Lemma H.7. *Let α , and β strictly positive. As $z \rightarrow +\infty$, we have the asymptotic*

$$\int_{z_c}^z \exp\left(-\int_s^z \frac{\alpha}{r} dr\right) \frac{B}{s \ln s} ds = O\left(\frac{1}{\ln z}\right). \quad (197)$$

1001

Proof.

$$\int_{z_c}^z \exp\left(-\int_s^z \frac{\alpha}{r} dr\right) \frac{B}{s \ln s} ds = B z^{-\alpha} \int_{z_c}^z \frac{s^{\alpha-1}}{\ln s} ds.$$

1002 Integration by parts with $u = 1/\ln s$, $dv = s^{\alpha-1} ds$, so that $du = -ds/(s(\ln s)^2)$ and $v = s^\alpha/\alpha$,
 1003 gives

$$\int_{z_c}^z \frac{s^{\alpha-1}}{\ln s} ds = \frac{z^\alpha}{\alpha \ln z} - \frac{z_c^\alpha}{\alpha \ln z_c} + \frac{1}{\alpha} \int_{z_c}^z \frac{s^{\alpha-1}}{(\ln s)^2} ds.$$

1004 Using the theorem of integration of asymptotic equivalents (Jourdain, 2018, *Theorem 4.10*) gives
 1005 $\int_{z_c}^z s^{\alpha-1}/(\ln s)^2 ds = O(z^\alpha/(\ln z)^2)$. Therefore

$$\int_{z_c}^z \frac{s^{\alpha-1}}{\ln s} ds = \frac{z^\alpha}{\alpha \ln z} (1 + O(1/\ln z)),$$

1006 Combining everything:

$$\int_{z_c}^z \exp\left(-\int_s^z \frac{\alpha}{r} dr\right) \frac{B}{s \ln s} ds = B z^{-\alpha} \frac{z^\alpha}{\alpha \ln z} (1 + O(1/\ln z)) = O(1/\ln z). \quad \square$$

1007

Lemma H.8. Consider a sequence $(u_t)_t$ a positive sequence satisfying:

$$u_{t+1} \leq \left(1 - \frac{h_t \gamma_-}{2z_t}\right) u_t + \frac{h_t B}{z_t \ln z_t}, \quad t \geq t_0,$$

where γ_- and B are strictly positive, $t_0 > 0$, $(h_t)_t$ a sequence satisfying condition (21), and $z_t := \sum_{k=0}^{t-1} h_k$. Then, $u_t = O(1/\ln z_t)$.

1008

1009 *Proof.* A quick recurrence gives:

$$u_t \leq u_{t_0} \prod_{j=t_0}^{t-1} \left(1 - \frac{h_j \gamma_-}{2z_j}\right) + \sum_{k=t_0}^{t-1} \frac{h_k B}{z_k \ln z_k} \prod_{j=k+1}^{t-1} \left(1 - \frac{h_j \gamma_-}{2z_j}\right).$$

1010 Noticing that for j large, $0 < h_j \gamma_- / (2z_j) < 1$, and $\ln(1 - h_j \gamma_- / (2z_j)) \leq -h_j \gamma_- / (2z_j)$, we have

$$\prod_{j=k+1}^{t-1} \left(1 - \frac{h_j \gamma_-}{2z_j}\right) \leq \exp\left(-\sum_{j=k+1}^{t-1} \frac{h_j \gamma_-}{2z_j}\right).$$

1011 Now $\sum_{j=k+1}^{t-1} h_j / z_j \geq \int_{z_{k+1}}^{z_t} ds/s$ (by the integral comparison, using $h_j = z_{j+1} - z_j$), giving

1012 $\sum_{j=k+1}^{t-1} \frac{h_j}{z_j} \geq \ln \frac{z_t}{z_{k+1}}$. Therefore

$$\prod_{j=k+1}^{t-1} \left(1 - \frac{h_j \gamma_-}{2z_j}\right) \leq \left(\frac{z_{k+1}}{z_t}\right)^{\gamma_-/2}. \quad (198)$$

1013 In particular, $\sum h_j / z_j = +\infty$ (since $z_t \rightarrow \infty$), so the homogeneous term $u_{t_1} \prod(\dots)$ decays
1014 polynomially in z_t , hence tends to 0. For the remaining term, setting $\alpha := \gamma_-/2 > 0$:

$$\sum_{k=t_0}^{t-1} \frac{h_k B}{z_k \ln z_k} \prod_{j=k+1}^{t-1} \left(1 - \frac{h_j \gamma_-}{2z_j}\right) \leq \frac{B}{z_t^\alpha} \sum_{k=t_0}^{t-1} \frac{h_k z_{k+1}^\alpha}{z_k \ln z_k} \leq \frac{2^\alpha B}{z_t^\alpha} \sum_{k=t_0}^{t-1} \frac{h_k z_k^{\alpha-1}}{\ln z_k}, \quad (199)$$

1015 using $z_{k+1}^\alpha \leq 2^\alpha z_k^\alpha$ for k large (since $z_{k+1}/z_k \rightarrow 1$). To bound the remaining sum, we use Abel's

1016 summation formula. Define $A_k := \sum_{j=t_0}^k h_j$ so that $A_{t-1} = z_t - z_{t_0}$ and $A_k - A_{k-1} = h_k$. Writing

1017 $f_k := z_k^{\alpha-1} / \ln z_k = \varphi(z_k)$:

$$\sum_{k=t_1}^{t-1} h_k f_k = A_{t-1} f_{t-1} - \sum_{k=t_1}^{t-2} A_k (f_{k+1} - f_k).$$

1018 $A_{t-1} = z_t - z_{t_1} \leq z_t$ and $f_{t-1} \sim z_t^{\alpha-1} / \ln z_t$, so $A_{t-1} f_{t-1} = O(z_t^\alpha / \ln z_t)$.

1019 By the mean-value theorem, $f_{k+1} - f_k = h_k \varphi'(\xi_k)$ for some $\xi_k \in [z_k, z_{k+1}]$, with $\varphi'(s) =$
1020 $s^{\alpha-2}((\alpha-1) \ln s - 1) / (\ln s)^2$. For all $\alpha > 0$ and s large, $|\varphi'(s)| \leq C_\varphi s^{\alpha-2}$ for some $C_\varphi > 0$.

1021 Therefore $|f_{k+1} - f_k| \leq C_\varphi h_k z_k^{\alpha-2}$ and

$$\sum_{k=t_1}^{t-2} A_k |f_{k+1} - f_k| \leq C_\varphi \sum_{k=t_1}^{t-2} z_k h_k z_k^{\alpha-2} h_k = C_\varphi \sum_{k=t_1}^{t-2} h_k^2 z_k^{\alpha-1}.$$

1022 Since $z_k^{\alpha-1} \leq z_k^{\max(\alpha-1, 0)} \leq \max(1, z_t^{\alpha-1})$ and $\sum h_k^2 < \infty$ by (21), this sum is $O(z_t^{\max(\alpha-1, 0)}) =$

1023 $O(z_t^\alpha / \ln z_t)$. Combining: $\sum_{k=t_1}^{t-1} h_k \varphi(z_k) = O(z_t^\alpha / \ln z_t)$. Substituting back into (199):

$$\sum_{k=t_0}^{t-1} \frac{h_k B}{z_k \ln z_k} \prod_{j=k+1}^{t-1} \left(1 - \frac{h_j \gamma_-}{2z_j}\right) \leq \frac{2^\alpha B}{z_t^\alpha} \cdot \frac{C}{\alpha \ln z_t} = O(1/\ln z_t).$$

1024 Therefore the forcing contribution is $O(1/\ln z_t)$.

1025 Combining: $u_t = O(z_t^{-\alpha}) + O(1/\ln z_t) = O(1/\ln z_t)$ □

1026 H.5 Discrete Synchronisation

1027 The synchronisation Lemma E.5 establishes $a(z) - b(z) - d \rightarrow 0$ for the ODE trajectory. To close
1028 the proof, we need the analogous statement for the *discrete* iterates (a_t, b_t) .

Lemma H.9 (Discrete synchronisation). *Under Condition 10 and Condition (21), if there exists $t_0 > 0$ such that the discrete iterates of Proposition E.3 satisfy*

$$|a_{t_0} - \tilde{\gamma}_{\min} b_{t_0} - d| \leq \eta_*, \quad b_{t_0} \geq b_*, \quad \text{and} \quad h_t \text{ sufficiently small for all } t \geq t_0,$$

with η_ and b_* as described in Lemma E.5. Then*

$$|a_t - \tilde{\gamma}_{\min} b_t - d| \leq \eta_* \quad \text{for all } t \geq t_0.$$

Moreover, the discrete gap converges: $a_t - \tilde{\gamma}_{\min} b_t \rightarrow d$.

1029

1030 *Proof.* The proof mirrors the previous proof on the continuous case.

1031 (a) *Removing z from Lemma E.5 setting.* Define the discrete gap $\eta_t := a_t - \tilde{\gamma}_{\min} b_t - d$. From
1032 Eqs (57)–(58):

$$\eta_{t+1} = \eta_t + h_t \Phi(\eta_t, b_t), \quad (200)$$

1033 where

$$\Phi(\eta, b) := (1 - \varepsilon)(1 + \mu_A)(1 - \alpha) \mathcal{I}_{\text{maj}}(\tilde{\gamma}_{\min} b + d + \eta) - \varepsilon \left(\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu) \right) \mathcal{I}_{\min}(b). \quad (201)$$

1034 Noticing that s in Eq (71) depends on z only through $b(z)$, we can rewrite Φ like:

$$\Phi(\eta, b) = \varepsilon \left(\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu) \right) \mathcal{I}_{\min}(b) \left[e^{-\eta} (1 + s(b, \eta)) - 1 \right], \quad (202)$$

1035 with s rewritten as:

$$s(b, \eta) := \frac{\tilde{\gamma}_{\min} b}{\tilde{\gamma}_{\min} b + d + \eta} \cdot \frac{1 + r_{\text{maj}}(\tilde{\gamma}_{\min} b + d + \eta)}{1 + r_{\min}(b)} - 1. \quad (203)$$

1036 The estimates (72)–(73) remain valid for any $b \leq b_*$ and $|\eta| \leq 1$:

$$|s(b, \eta)| \leq \frac{K_s}{|b|}, \quad |s(b, \eta) - s(b, 0)| \leq \frac{K'_s |\eta|}{|b|}. \quad (204)$$

1037 (b) *Discrete tube invariance.* The tube argument is identical to step (i) of Lemma E.5: for $b \geq b_*$,
1038 the same computation as in (67) gives $\Phi(\eta_*, b) < 0$ and $\Phi(-\eta_*, b) > 0$. Therefore on the boundary
1039 $\eta = \eta_*$ for instance, $\eta_{t+1} = \eta_* + h_t \Phi(\eta_*, b_t) \in [-\eta_*, \eta_*]$ provided $h_t |\Phi(\eta_*, b_t)| < 2\eta_*$, which
1040 holds for all $t > t_0$ with h_t sufficiently small and given $|\Phi|$ is bounded on $\mathbb{R}_{\geq 1}^2$.

1041 Since $b_t \rightarrow +\infty$ monotonically (Remark E.4), b_t eventually remains above b_* . Choosing t_0 as the
1042 first time both h_t is small enough and $b_t \geq b_*$, the tube $\{|\eta_t| \leq \eta_*\}$ is forward invariant for $t \geq t_0$.

1043 (c) *Convergence.* Like in step (a), noticing that γ, β and R depends on z only through b and η , the
1044 linearisation Eq (80) generalizes as well from Eq (202) into:

$$\Phi(\eta, b) = -\hat{\gamma}(b) \eta + \hat{\beta}(b) + \hat{R}(b, \eta)$$

1045 with

$$\hat{\gamma}(b) := \varepsilon \left(\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu) \right) \mathcal{I}_{\min}(b), \quad \hat{\beta}(b) := \varepsilon \left(\tilde{\gamma}_{\min} (1 + \mu_B) - (1 + \mu) \right) \mathcal{I}_{\min}(b) s(b, 0), \quad (205)$$

1046

$$|\hat{R}(b, \eta)| \leq \hat{K} |\mathcal{I}_{\min}(b)| \left(\eta^2 + \frac{|\eta|}{b} \right) \quad (206)$$

1047 with $\hat{K}$ an explicit constant. By the tube invariance (step (b)), $|\eta_t| \leq \eta_*$ for $t \geq t_0$, so the quantity

$$\tilde{c}_t := \varepsilon(1 + \mu_B) + (1 - \varepsilon)(1 + \mu) \frac{\mathcal{I}_{\text{maj}}(a_t)}{\mathcal{I}_{\text{min}}(b_t)}; \quad \text{s.t.} \quad b_{t+1} - b_t = \tilde{c}_t \mathcal{I}_{\text{min}}(b_t) \quad (\text{see Eq (58)})$$

1048 satisfies $\tilde{c}_t \in [\tilde{c}_-, \tilde{c}_+] \subset (0, \infty)$ for $t \geq t_1$, by the same bounded-ratio argument as in (77).
 1049 Lemma H.6, part (ii), then gives

$$\gamma b_t = \ln z_t - \ln \ln z_t + O(1), \quad \text{s.t.} \quad \mathcal{I}_{\text{min}}(b_t) = \Theta(1/z_t) \quad (\text{as in Eq (79)}). \quad (207)$$

1050 Defining $\gamma_t := \hat{\gamma}(b_t)$, $\beta_t := \hat{\beta}(b_t)$, $\mathcal{R}_t := \hat{R}(b_t, \eta_t)$, we also derive the analogous constants and
 1051 bounds:

1052 (i) $\frac{M_-}{z_t} \leq \mathcal{I}_{\text{min}}(b_t) \leq \frac{M_+}{z_t}$, with $0 < M_- \leq M_+$ and $\frac{\gamma_-}{z_t} \leq \gamma_t \leq \frac{\gamma_+}{z_t}$,

$$\gamma_{\pm} := \varepsilon(\tilde{\gamma}_{\text{min}}(1 + \mu_B) - (1 + \mu)) M_{\pm},$$

1053 (ii) $|\beta_t| \leq B/(z_t \ln z_t)$ with $B > 0$ by Eq (204),

1054 (iii) from (206) with $(b, \eta) = (b_t, \eta_t)$ and using $b_t \geq c_b \ln z_t$ for some $c_b > 0$:

$$|\mathcal{R}_t| \leq \hat{K} |\mathcal{I}_{\text{min}}(b_t)| \left(\eta_t^2 + \frac{|\eta_t|}{b_t} \right) \leq \frac{\hat{K} M_+}{\min(1, c_b) z_t} \left(\eta_t^2 + \frac{|\eta_t|}{\ln z_t} \right) = \frac{\tilde{K}}{z_t} \left(\eta_t^2 + \frac{|\eta_t|}{\ln z_t} \right),$$

1055 where we wrote $\tilde{K}$ for $\hat{K} M_+ / \min(1, c_b)$ to lighten notation). The recurrence (200) with (205) reads

$$\eta_{t+1} = (1 - h_t \gamma_t) \eta_t + h_t \beta_t + h_t \mathcal{R}_t.$$

1056 Setting $u_t := |\eta_t|$ and using $|\eta_t| \leq \eta_*$ with η_* small enough and t_0 large enough that $\tilde{K} \eta_* \leq \gamma_-/4$
 1057 and $\tilde{K} / \ln z_{t_0} \leq \gamma_-/4$ (exactly as we did to get Eq (84), we absorb $\mathcal{R}_t$ into γ_t :

$$u_{t+1} \leq \left(1 - \frac{h_t \gamma_-}{2 z_t} \right) u_t + \frac{h_t B}{z_t \ln z_t}, \quad t \geq t_0, \quad (208)$$

1058 which is the exact hypothesis of Lemma H.8. We can therefore conclude that $u_t = O(1/\ln z_t)$ and
 1059 $\eta_t = a_t - \tilde{\gamma}_{\text{min}} b_t - d \rightarrow 0$. $\square$

I Generalizing Soudry et al. (2018, Theorem 9)

I.1 Setup and assumptions

We consider a data distribution $\mathcal{D}$ over $(\mathbf{x}, y)$ with $\mathbf{x} \in \mathbb{R}^d$ and binary labels $y \in \{-1, +1\}$. Following the relabelling convention of Soudry et al. (2018), we absorb the label into the feature: redefine $\mathbf{x} \leftarrow y\mathbf{x}$ so that we may assume $y = +1$ for all samples without loss of generality. The population expected risk (in contrast to empirical risk) loss is

$$\mathcal{L}(\mathbf{w}) = \mathbb{E}[\ell(\mathbf{w}^\top \mathbf{x})], \quad (209)$$

where ℓ is any loss function satisfying Soudry et al. (2018, Assumptions 2 and 3). Without loss of generality we will consider $\ell(u) = e^{-u}$ for simplicity.

We perform *population (full-batch) gradient descent* with fixed learning rate $\eta > 0$:

$$\mathbf{w}(t+1) = \mathbf{w}(t) - \eta \nabla \mathcal{L}(\mathbf{w}(t)) = \mathbf{w}(t) + \eta \mathbb{E}[e^{-\mathbf{w}(t)^\top \mathbf{x}} \mathbf{x}]. \quad (210)$$

We impose the following standard assumptions on $\mathcal{D}$ inspired from 2.2 and D.1.

Assumption I.1 (Linear separability). *There exists $\mathbf{w}_* \in \mathbb{R}^d$ such that $\mathbf{w}_*^\top \mathbf{x} \geq 1$ almost surely.*

Under Assumption I.1, the population L_2 max-margin (analogous to hard-margin SVM) solution is well-defined:

$$\hat{\mathbf{w}} = \underset{\mathbf{w} \in \mathbb{R}^d}{\operatorname{argmin}} \|\mathbf{w}\|^2 \quad \text{s.t.} \quad \mathbf{w}^\top \mathbf{x} \geq 1 \text{ a.s.} \quad (211)$$

Assumption I.2 (Bounded features). *There exists $K > 0$ such that $\|\mathbf{x}\| \leq K$ almost surely. Combined with Assumptions I.1 this assumption gives $\mathcal{L}$ is β -smooth with $\beta = K^2 \mathcal{L}(\mathbf{w}(0))$. For technical reasons in the proof below we set $\beta := \max(\mathcal{L}(\mathbf{w}(0)), 8) K^2$ instead.*

Assumption I.3 (Margin density). *$\mathcal{D}$ has continuous density and if we define the margin variable*

$$\mathbf{m} := \hat{\mathbf{w}}^\top \mathbf{x}, \quad (212)$$

then $\mathbf{m}$ has a continuous strictly positive density p_m on $[1, M]$ where $M = K\|\hat{\mathbf{w}}\|$.

Assumption I.4. *Let $\mathbf{x}_\perp := \bar{\Pi}_{\hat{\mathbf{w}}} \mathbf{x}$ (see (22) for the notation), then*

$$\mathbb{E}[\mathbf{x}_\perp \mid \mathbf{m}] = 0. \quad (213)$$

This assumption mirrors Assumption D.1 and similarly could be weakened by requiring that for every direction $\hat{\boldsymbol{\delta}}$ such that $\hat{\boldsymbol{\delta}}^\top \hat{\mathbf{w}} = 0$, $\mathbf{x}_\perp \mid \mathbf{m}=1$ never puts all its mass on one single side of the hyperplane $(\hat{\boldsymbol{\delta}})^\perp$ (non-degeneracy F.1):

$$\Pr(\hat{\boldsymbol{\delta}}^\top \mathbf{x}_\perp < 0 \mid \mathbf{m}=1) > 0. \quad (214)$$

In geometric terms, the origin lies in the interior of the convex hull of the support of $\mathbf{x}_\perp \mid \mathbf{m}=1$ in the subspace $(\hat{\mathbf{w}})^\perp$. However, this would require considerable more complex mathematical development, therefore we opt for this simpler and relatively general setting.

In Soudry et al. (2018) finite-sample framework, the decomposition $\mathbf{w}(t) = \hat{\mathbf{w}} \log t + \boldsymbol{\rho}(t)$ with *bounded* $\boldsymbol{\rho}$ holds for “almost all datasets” (their Theorem 9); the degenerate case where $\|\boldsymbol{\rho}\| = O(\log \log t)$ is relegated to Theorem 13. Although our continuous-distribution study does not observe such degeneracy, we will show that a $O(\log \log t)$ correction that is absent in the discrete case is nevertheless introduced.

We now state the following proposition generalizing the theorem to our continuous distribution setting.

Proposition I.5 (Implicit bias of GD - continuous-distribution setting). *Under Assumptions I.1-I.4, let $\mathbf{w}(t)$ denote the iterates of population gradient descent (210) with any initial point $\mathbf{w}(0)$ and learning rate η small enough (specified in the proof). Then, as $t \rightarrow \infty$:*

$$\mathbf{w}(t) = \hat{\mathbf{w}} \log t - \hat{\mathbf{w}} \log \log t + O(1). \quad (215)$$

In particular, the direction of $\mathbf{w}(t)$ still converges to the max-margin direction:

$$\lim_{t \rightarrow \infty} \frac{\mathbf{w}(t)}{\|\mathbf{w}(t)\|} = \frac{\hat{\mathbf{w}}}{\|\hat{\mathbf{w}}\|}.$$

I.2 Proof of Proposition 5.4

The proof proceeds in four steps inspired from the original discrete proof.

(i) Loss convergence and norm divergence

Lemma I.6. For η small enough: $\mathcal{L}(\mathbf{w}(t)) \rightarrow 0$, and $\|\mathbf{w}(t)\| \rightarrow \infty$ as $t \rightarrow \infty$.

Proof. The argument follows Soudry et al. (2018, Lemma 1) textually. For any $\mathbf{w}$:

$$\mathbf{w}_*^\top \nabla \mathcal{L}(\mathbf{w}) = -\mathbb{E} \left[e^{-\mathbf{w}^\top \mathbf{x}} \mathbf{w}_*^\top \mathbf{x} \right] < 0 \quad (216)$$

since $\mathbf{w}_*^\top \mathbf{x} > 0$ a.s. Therefore there's no finite critical point $\mathbf{w}$ such that $\nabla \mathcal{L}(\mathbf{w}) = 0$. Since $\mathcal{L}$ is β -smooth, gradient descent with $\eta < 2/\beta$ converges to a critical point: $\nabla \mathcal{L}(\mathbf{w}(t)) \rightarrow 0$ (see Soudry et al. (2018, Lemma 10)). The absence of finite critical points forces $\|\mathbf{w}(t)\| \rightarrow \infty$. Finally, using $\mathbf{w}_*^\top \mathbf{x} \geq 1$ a.s, we have the following bound

$$0 \leq \mathcal{L}(\mathbf{w}(t)) = \mathbb{E} \left[e^{-\mathbf{w}^\top \mathbf{x}} \right] \leq \mathbb{E} \left[e^{-\mathbf{w}^\top \mathbf{x}} \mathbf{w}_*^\top \mathbf{x} \right] = \left| \mathbf{w}_*^\top \nabla \mathcal{L}(\mathbf{w}(t)) \right| \leq \|\mathbf{w}_*\| \|\nabla \mathcal{L}(\mathbf{w}(t))\| \rightarrow 0,$$

hence $\mathcal{L}(\mathbf{w}(t)) \rightarrow 0$. $\square$

Decomposing $\mathbf{w}(t)$ like this:

$$\mathbf{w}(t) = f(t) \hat{\mathbf{w}} + \boldsymbol{\delta}(t), \quad f(t) := \frac{\mathbf{w}(t)^\top \hat{\mathbf{w}}}{\|\hat{\mathbf{w}}\|^2}, \quad \hat{\mathbf{w}}^\top \boldsymbol{\delta}(t) = 0, \quad (217)$$

we now show $f(t) \rightarrow \infty$.

(ii) $f(t)$ diverges.

Lemma I.7 (Coarse f growth rate). $f(t) \rightarrow +\infty$, and more precisely, $f(t) = \Theta(\log t)$.

Proof. Projecting (210) onto $\hat{\mathbf{w}}$ and using Assumption I.1:

$$f(t+1) - f(t) = \frac{\eta}{\|\hat{\mathbf{w}}\|^2} \mathbb{E} [e^{-\mathbf{w}(t)^\top \mathbf{x}} \hat{\mathbf{w}}^\top \mathbf{x}] = \frac{\eta}{\|\hat{\mathbf{w}}\|^2} \mathbb{E} [e^{-\mathbf{w}(t)^\top \mathbf{x}} \mathbf{m}] \quad (218)$$

$$\geq \frac{\eta}{\|\hat{\mathbf{w}}\|^2} \mathbb{E} [e^{-\mathbf{w}(t)^\top \mathbf{x}}] = \frac{\eta}{\|\hat{\mathbf{w}}\|^2} \mathcal{L}(\mathbf{w}(t)) > 0. \quad (219)$$

So f is strictly increasing for all t and summing (219):

$$f(t) - f(0) \geq \frac{\eta}{\|\hat{\mathbf{w}}\|^2} \sum_{t=0}^{t-1} \mathcal{L}(\mathbf{w}(t)). \quad (220)$$

It remains to show $\sum_t \mathcal{L}(\mathbf{w}(t)) = +\infty$. As in Soudry et al. (2018, Lemma 10), using $\mathcal{L}$'s β -smoothness, we have:

$$\mathcal{L}(\mathbf{w}(t)) - \mathcal{L}(\mathbf{w}(t+1)) \geq \eta \left(1 - \frac{\eta\beta}{2}\right) \|\nabla \mathcal{L}(\mathbf{w}(t))\|^2 > 0, \quad (221)$$

$$\mathcal{L}(\mathbf{w}(t)) - \mathcal{L}(\mathbf{w}(t+1)) \leq \eta \left(1 + \frac{\eta\beta}{2}\right) \|\nabla \mathcal{L}(\mathbf{w}(t))\|^2 \leq \eta K^2 \left(1 + \frac{\eta\beta}{2}\right) \mathcal{L}(\mathbf{w}(t))^2, \quad (222)$$

1112 where the last inequality comes from $\|\nabla \mathcal{L}\| = \|\mathbb{E}[e^{-\mathbf{w}^\top \mathbf{x}} \mathbf{x}]\| \leq K \mathcal{L}(\mathbf{w})$.

1113 Dividing both sides by $\mathcal{L}(\mathbf{w}(t)) \mathcal{L}(\mathbf{w}(t+1))$ and using $\mathcal{L}(\mathbf{w}(t+1)) \leq \mathcal{L}(\mathbf{w}(t))$ (by Eq (221)):

$$\begin{aligned} \frac{1}{\mathcal{L}(\mathbf{w}(t+1))} - \frac{1}{\mathcal{L}(\mathbf{w}(t))} &= \frac{\mathcal{L}(\mathbf{w}(t)) - \mathcal{L}(\mathbf{w}(t+1))}{\mathcal{L}(\mathbf{w}(t)) \mathcal{L}(\mathbf{w}(t+1))} \\ &\leq \left(1 + \frac{\eta\beta}{2}\right) \frac{\eta K^2 \mathcal{L}(\mathbf{w}(t))^2}{\mathcal{L}(\mathbf{w}(t)) \mathcal{L}(\mathbf{w}(t+1))} \\ &\leq \left(1 + \frac{\eta\beta}{2}\right) \frac{\eta K^2 \mathcal{L}(\mathbf{w}(t))}{\mathcal{L}(\mathbf{w}(t+1))} \leq 4\eta K^2, \end{aligned} \quad (223)$$

1114 where the last step uses $\eta < 2/\beta$ and the fact that for t large enough, $\mathcal{L}(\mathbf{w}(t))/\mathcal{L}(\mathbf{w}(t+1)) \leq 2$.

1115 In fact, for t large enough, $\mathcal{L}(\mathbf{w}(t)) < 1$ (by Lemma I.6) which gives by combining to Eq (222):

$$\begin{aligned} \mathcal{L}(\mathbf{w}(t)) - \mathcal{L}(\mathbf{w}(t+1)) &\leq \eta K^2 \left(1 + \frac{\eta\beta}{2}\right) \mathcal{L}(\mathbf{w}(t)) \\ &= \frac{\eta\beta}{\max(\mathcal{L}(\mathbf{w}(0)), 8)} \left(1 + \frac{\eta\beta}{2}\right) \mathcal{L}(\mathbf{w}(t)) \quad \text{by Assumption I.2} \\ &= \frac{\mathcal{L}(\mathbf{w}(t))}{2} \quad \text{since } \eta < 2/\beta, \end{aligned}$$

1116 so that $\mathcal{L}(\mathbf{w}(t+1)) \geq \mathcal{L}(\mathbf{w}(t))/2$. Summing Eq (223): $1/\mathcal{L}(\mathbf{w}(t)) \leq 1/\mathcal{L}(\mathbf{w}(0)) + 4\eta K^2 t$, hence

$$\mathcal{L}(\mathbf{w}(t)) \geq \frac{c'}{t} \quad \text{for } c' := \frac{1}{4\eta K^2 + 1/\mathcal{L}(\mathbf{w}(0))} > 0. \quad (224)$$

1117 Therefore $\sum_t \mathcal{L}(\mathbf{w}(t)) = +\infty$ and (220) gives $f(t) \rightarrow +\infty$.

1118 We can also similarly derive an upper-bound for the loss as follows. By Eq (216), $\hat{\mathbf{w}}^\top \nabla \mathcal{L} = -\mathbb{E}[e^{-\mathbf{w}^\top \mathbf{x}} \mathbf{m}] \leq -\mathcal{L}(\mathbf{w})$ (since $\mathbf{m} \geq 1$ a.s.), so $\|\nabla \mathcal{L}\| \geq |\hat{\mathbf{w}}^\top \nabla \mathcal{L}|/\|\hat{\mathbf{w}}\| \geq \mathcal{L}/\|\hat{\mathbf{w}}\|$. Therefore
1119 with Eq (221):
1120

$$\mathcal{L}(\mathbf{w}(t)) - \mathcal{L}(\mathbf{w}(t+1)) \geq \left(1 - \frac{\eta\beta}{2}\right) \frac{\eta}{\|\hat{\mathbf{w}}\|^2} \mathcal{L}(\mathbf{w}(t))^2. \quad (225)$$

1121 By the same telescoping argument as above: $\frac{1}{\mathcal{L}(\mathbf{w}(t+1))} - \frac{1}{\mathcal{L}(\mathbf{w}(t))} \geq \left(1 - \frac{\eta\beta}{2}\right) \frac{\eta}{\|\hat{\mathbf{w}}\|^2}$, giving

1122 $\mathcal{L}(\mathbf{w}(t)) \leq C/t$ for $C := \left(1 - \frac{\eta\beta}{2}\right)^{-1} \|\hat{\mathbf{w}}\|^2/\eta$. Combined with (224):

$$\mathcal{L}(\mathbf{w}(t)) = \Theta(1/t). \quad (226)$$

1123 A similar analysis as in Eq (219) gives the upper bound $f(t+1) - f(t) \leq \eta K \mathcal{L}(\mathbf{w}(t))/\|\hat{\mathbf{w}}\|$ (from
1124 $\mathbf{m} \leq K\|\hat{\mathbf{w}}\|$ a.s.). When combined with $\mathcal{L}(\mathbf{w}(t)) \leq C/t$ we get $f(t) \leq C' \log t + C''$. Together
1125 with $\mathcal{L} \geq c'/t$ and (219):

$$f(t) = \Theta(\log t). \quad (227)$$

1126 □

1127 **(iii) $\delta(t)$ is bounded.**

Lemma I.8. $\delta(t) = O(1)$ and therefore $\frac{\mathbf{w}(t)}{\|\mathbf{w}(t)\|} \xrightarrow{t \rightarrow \infty} \frac{\hat{\mathbf{w}}}{\|\hat{\mathbf{w}}\|}$.

1128

1129 *Proof.* Projecting (210) onto $\text{Span}\{\hat{\mathbf{w}}\}^\perp$:

$$\delta(t+1) - \delta(t) = \eta \mathbb{E}[e^{-f(t)\mathbf{m} - \delta(t)^\top \mathbf{x}_\perp} \mathbf{x}_\perp], \quad (228)$$

1130 so that

$$\|\delta(t+1) - \delta(t)\| \leq \eta K \mathcal{L}(\mathbf{w}(t)) \leq \eta K \frac{C}{t}. \quad (229)$$

1131 Therefore:

$$\sum_{t=1}^{\infty} \|\delta(t+1) - \delta(t)\|^2 \leq \eta^2 K^2 C^2 \sum_{t=1}^{\infty} \frac{1}{t^2} < \infty. \quad (230)$$

1132 Now let's write

$$\begin{aligned} \|\delta(t+1)\|^2 - \|\delta(t)\|^2 &= 2\delta(t)^\top [\delta(t+1) - \delta(t)] + \|\delta(t+1) - \delta(t)\|^2 \\ &= 2\eta \mathbb{E}[e^{-f\mathbf{m} - \delta(t)^\top \mathbf{x}_\perp} \delta(t)^\top \mathbf{x}_\perp] + \|\delta(t+1) - \delta(t)\|^2 \quad (\text{by Eq (228)}) \\ &= 2\eta \underbrace{\mathbb{E}[e^{-f\mathbf{m}} (e^{-z} - 1) z]}_{=: R(t) \leq 0} + 2\eta \underbrace{\mathbb{E}[e^{-f\mathbf{m}} z]}_{=: L(t)} + \|\delta(t+1) - \delta(t)\|^2, \end{aligned} \quad (231)$$

1133 where we use $e^{-z} = (e^{-z} - 1) + 1$ with $z = \delta(t)^\top \mathbf{x}_\perp$ along with the pointwise inequality $z(e^{-z} - 1) \leq 0$
1134 for all $z \in \mathbb{R}$.

1135 Under Assumption I.4 and the total expectation property, we have:

$$L(t) = \delta(t)^\top \mathbb{E}[e^{-f\mathbf{m}} \mathbf{x}_\perp] = \delta(t)^\top \mathbb{E}_{\mathbf{m}}[e^{-f\mathbf{m}} \mathbb{E}[\mathbf{x}_\perp | \mathbf{m}]] = 0. \quad (232)$$

1136 Substituting $L(t) = 0$ into (231):

$$\|\delta(t+1)\|^2 - \|\delta(t)\|^2 = 2\eta R(t) + \|\delta(t+1) - \delta(t)\|^2 \leq \|\delta(t+1) - \delta(t)\|^2. \quad (233)$$

1137 Summing from 0 to $t-1$ and using Eq (230):

$$\|\delta(t)\|^2 \leq \|\delta(0)\|^2 + \sum_{t=0}^{T-1} \|\delta(t+1) - \delta(t)\|^2 \leq \|\delta(0)\|^2 + \eta^2 K^2 C^2 \sum_{t=1}^{\infty} t^{-2}, \quad (234)$$

1138 hence $\|\delta(t)\| = O(1)$. Combining with $f(t) \rightarrow +\infty$:

$$\frac{\|\delta(t)\|}{f(t)} \rightarrow 0, \quad \text{and therefore} \quad \frac{\mathbf{w}(t)}{\|\mathbf{w}(t)\|} \xrightarrow{t \rightarrow \infty} \frac{\hat{\mathbf{w}}}{\|\hat{\mathbf{w}}\|}. \quad (235)$$

1139 □

1140 (iv) *Sharp growth rate of f .*

Lemma I.9. $f(t) = \log t - \log \log t + O(1)$.

1142 *Proof.* From (218):

$$f(t+1) - f(t) = \frac{\eta}{\|\hat{\mathbf{w}}\|^2} \mathbb{E}[e^{-f(t)\mathbf{m} - \delta(t)^\top \mathbf{x}_\perp} \mathbf{m}] = \frac{\eta}{\|\hat{\mathbf{w}}\|^2} \mathbb{E}_{\mathbf{m}}[e^{-f(t)\mathbf{m}} \mathbf{m} \mathbb{E}[e^{-\delta(t)^\top \mathbf{x}_\perp} | \mathbf{m}]].$$

1143 By the disintegration theorem Ambrosio et al. (2005, *Theorem 5.3.1*) and Assumption I.3, there exists
1144 a continuous function $H(m, t) := m \mathbb{E}_{\mathbf{x} | \mathbf{m}=m}[e^{-\delta(t)^\top \mathbf{x}_\perp} | \mathbf{m} = m] p_m(m)$ such that

$$\mathbb{E}_{\mathbf{m}}[e^{-f(t)\mathbf{m}} \mathbf{m} \mathbb{E}[e^{-\delta(t)^\top \mathbf{x}_\perp} | \mathbf{m}]] = \int_1^{K\|\hat{\mathbf{w}}\|} e^{-f(t)m} H(m, t) dm. \quad (236)$$

1145 Since $\|\delta\| = O(1)$ and $\|\mathbf{x}_\perp\| \leq K$ a.s., there exists a constant $D > 0$ such that for all t :

$$e^{-DK} m p_m(m) \leq H(m, t) \leq e^{DK} m p_m(m). \quad (237)$$

1146 Define the t -independent bounding function $h(m) := m p_m(m)$ (continuous with $h(1) = p_m(1) >$
1147 0). Eq (237) gives

$$e^{-DK} \int_1^M e^{-f m} h(m) dm \leq \int_1^M e^{-f m} H(m, t) dm \leq e^{DK} \int_1^M e^{-f m} h(m) dm.$$

1148 Watson's lemma Wong (2001, *Chapter 1*) for $f \rightarrow +\infty$ applied to the bounding integrals gives:

$$\int_1^M e^{-f m} h(m) dm = \frac{p_m(1)}{f} e^{-f} (1 + O(1/f)). \quad (238)$$

1149 Combining both results yields:

$$\eta c_- \frac{e^{-f(t)}}{f(t)} (1 + O(1/f)) \leq f(t+1) - f(t) \leq \eta c_+ \frac{e^{-f(t)}}{f(t)} (1 + O(1/f)), \quad (239)$$

1150 where $c_{\pm} := \frac{e^{\pm DK} p_m(1)}{\|\mathbf{w}\|^2} > 0$. We write this compactly as

$$f(t+1) - f(t) = \eta C(t) \frac{e^{-f(t)}}{f(t)}, \quad 0 < c_- \leq C(t) \leq c_+, \quad (240)$$

1151 potentially adjusting $c_{\pm}$ accordingly to absorb the $O(1/f)$.

1152 Let Φ the function $f \mapsto (f-1)e^f$. By Taylor's theorem applied to Φ at $f(t)$ with increment
1153 $a_t := f(t+1) - f(t) \geq 0$:

$$\begin{aligned} \Phi(f(t+1)) - \Phi(f(t)) &= \Phi'(f(t)) a_t + \frac{1}{2} \Phi''(\xi_t) a_t^2 \\ &= \eta C(t) + \frac{1}{2} \Phi''(\xi_t) a_t^2 \quad (\text{by Eq (240)}) \end{aligned} \quad (241)$$

1154 for some $\xi_t \in (f(t), f(t+1))$. Since $f(t)e^{f(t)} \leq \Phi''(\xi_t) \leq 2f(t)e^{f(t)}$ for t large:

$$0 \leq \Phi''(\xi_t) a_t^2 \leq 2f(t)e^{f(t)} \cdot \eta^2 C(t)^2 \frac{e^{-2f(t)}}{f(t)^2} = 2\eta^2 C^2(t) \frac{e^{-f(t)}}{f(t)} \leq 2\eta c_+ a_t.$$

1155 Combining and summing from t_0 to $t-1$ for sufficiently large t :

$$\begin{aligned} \Phi(f(t)) &\leq \Phi(f(t_0)) + \eta \sum_{s=t_0}^{t-1} C(s) + 2\eta c_+ \sum_{s=t_0}^{t-1} a_s \\ &\leq \Phi(f(t_0)) + \eta c_+ (t - t_0) + 2\eta c_+ (f(t) - f(t_0)) \\ &\leq \Phi(f(t_0)) + \eta c_+ t + 2\eta c_+ f(t) \end{aligned} \quad (242)$$

$$\begin{aligned} \Phi(f(t)) &\geq \Phi(f(t_0)) + \eta \sum_{s=t_0}^{t-1} C(s) \\ &\geq \Phi(f(t_0)) + \eta c_- (t - t_0) \\ &\geq \Phi(f(t_0)) + \frac{\eta c_-}{2} t. \end{aligned} \quad (243)$$

1156 Using the Lambert W function which is increasing, and noticing that $\Phi(f) = (f-1)e^f = b$ implies
1157 $f-1 = W(b/e)$, these inequalities yield:

$$W\left(\frac{\Phi(f(t_0))}{e} + \frac{\eta c_-}{2e} t\right) \leq f(t) - 1 \leq W\left(\frac{\Phi(f(t_0))}{e} + \frac{\eta c_+}{e} t + \frac{2\eta c_+}{e} f(t)\right). \quad (244)$$

1158 By Lambert W asymptotics Visser (2018, *Eq A2*), and Eq (227):

$$\begin{aligned} W\left(\frac{\Phi(f(t_0))}{e} + \frac{\eta c_+}{e} t + \frac{2\eta c_+}{e} f(t)\right) &= \log(t + \Theta(\log t)) - \log \log(t + \Theta(\log t)) + O(1) \\ &= \log t - \log \log t + O(1), \\ W\left(\frac{\Phi(f(t_0))}{e} + \frac{\eta c_-}{2e} t\right) &= \log t - \log \log t + O(1). \end{aligned}$$

1159 Hence

$$f(t) = \log t - \log \log t + O(1). \quad (245)$$

1160 $\square$

1161 **(v) Conclusion** From (234) and (245) (Step 2):

$$\mathbf{w}(t) = f(t) \hat{\mathbf{w}} + \delta(t) = \hat{\mathbf{w}} (\log t - \log \log t + O(1)) + O(1) = \hat{\mathbf{w}} \log t - \hat{\mathbf{w}} \log \log t + O(1).$$

1162 This completes the proof of Proposition I.5.

1163 I.3 Generalisation to variable step-size

1164 The argument extends to step-sizes (h_t) satisfying condition (21). We propose a version of Proposi-
1165 tion I.5 and how to adapt the proof.

Proposition I.10 (Variable step-size). *Under Assumptions I.1–I.4, let $\mathbf{w}(t)$ denote the iterates of population gradient descent (210) with any initial point $\mathbf{w}(0)$, and assume our scheduler $(\eta_t)_t$ satisfies*

$$2/\beta > \eta_t > 0, \quad \sum \eta_t = +\infty, \quad \sum \eta_t^2 < \infty,$$

with β introduced in Assumption I.2. Then, as $t \rightarrow \infty$:

$$\mathbf{w}(t) = \hat{\mathbf{w}} \log z_t - \hat{\mathbf{w}} \log \log z_t + O(1), \quad (246)$$

with $z_t := \sum_{k=0}^{t-1} \eta_k$.

1166

1167 *Proof.* The proof of Lemma I.6 is still mostly valid with a little rearrangement of the deductions.
1168 Eq (221) is still valid with η_t instead of η , this means that $\mathcal{L}(\mathbf{w}(t))$ decreases to a limit $\mathcal{L}_\infty$. By
1169 Eq (216), $\|\nabla \mathcal{L}\| \geq \mathcal{L}/\|\hat{\mathbf{w}}\|$, so $\|\nabla \mathcal{L}(\mathbf{w}(t))\| \geq \mathcal{L}_\infty/(2\|\hat{\mathbf{w}}\|)$ for $t \geq t_0$. From the proof of Soudry
1170 et al. (2018, Lemma 10), we have this updated result:

$$\sum_{t=0}^{\infty} \eta_t \left(1 - \frac{\eta_t \beta}{2}\right) \|\nabla \mathcal{L}(\mathbf{w}(t))\|^2 < \infty, \quad (247)$$

1171 and given that $\sum_{t=0}^{\infty} \eta_t = \infty$, this implies $\mathcal{L}_\infty = 0$ and $\|\nabla \mathcal{L}(\mathbf{w}(t))\| \rightarrow 0$.

1172 The proof of Lemma I.7 undergoes the following adjustments. Eq (223) becomes $\frac{1}{\mathcal{L}(\mathbf{w}(t+1))} -$
1173 $\frac{1}{\mathcal{L}(\mathbf{w}(t))} \leq 4\eta_t K^2$, summing it gives $\mathcal{L}(\mathbf{w}(t)) \geq c'/z_t$ for $c' := \frac{1}{4K^2+1/\mathcal{L}(\mathbf{w}(0))} > 0$, instead of
1174 Eq (224). A similar argument gives $\mathcal{L}(\mathbf{w}(t)) \leq C/z_t$ with an explicit constant C . Collecting
1175 everything yields $\mathcal{L}(\mathbf{w}(t)) = \Theta(1/z_t)$ and $f(t) = \Theta(\log z_t)$ by the same integral comparison
1176 argument.

1177 The proof of Lemma I.8 is valid textually by using the fact that $(\eta_t)_t$ is a bounded sequence.

1178 Regarding Lemma I.9, the proof needs the following adjustment. The sums Eqs (242) and (243)
1179 become

$$\Phi(f(t)) \leq \Phi(f(t_0)) + c_+ z_t + 2\|\eta\|_\infty c_+ f(t), \quad (248)$$

$$\Phi(f(t)) \geq \Phi(f(t_0)) + \frac{c_-}{2} z_t, \quad (249)$$

1180 where $\|\eta\|_\infty := \sup_{t \geq 0} \eta_t$. The same Lambert W asymptotics yields the same analogous result
1181 $f(t) = \log z_t - \log \log z_t + O(1)$, which concludes the proof. $\square$

1182 NeurIPS Paper Checklist

1183 1. Claims

1184 Question: Do the main claims made in the abstract and introduction accurately reflect the
1185 paper’s contributions and scope?

1186 Answer: [Yes]

1187 Justification: All the claims and contributions are cited in the Introduction and discussed in
1188 the Discussion sections.

1189 Guidelines:

- 1190 • The answer [N/A] means that the abstract and introduction do not include the claims
1191 made in the paper.
- 1192 • The abstract and/or introduction should clearly state the claims made, including the
1193 contributions made in the paper and important assumptions and limitations. A [No] or
1194 [N/A] answer to this question will not be perceived well by the reviewers.
- 1195 • The claims made should match theoretical and experimental results, and reflect how
1196 much the results can be expected to generalize to other settings.
- 1197 • It is fine to include aspirational goals as motivation as long as it is clear that these goals
1198 are not attained by the paper.

1199 2. Limitations

1200 Question: Does the paper discuss the limitations of the work performed by the authors?

1201 Answer: [Yes]

1202 Justification: The limitations are listed in the Conclusion section and discussed along the
1203 paper as well.

1204 Guidelines:

- 1205 • The answer [N/A] means that the paper has no limitation while the answer [No] means
1206 that the paper has limitations, but those are not discussed in the paper.
- 1207 • The authors are encouraged to create a separate “Limitations” section in their paper.
- 1208 • The paper should point out any strong assumptions and how robust the results are to
1209 violations of these assumptions (e.g., independence assumptions, noiseless settings,
1210 model well-specification, asymptotic approximations only holding locally). The authors
1211 should reflect on how these assumptions might be violated in practice and what the
1212 implications would be.
- 1213 • The authors should reflect on the scope of the claims made, e.g., if the approach was
1214 only tested on a few datasets or with a few runs. In general, empirical results often
1215 depend on implicit assumptions, which should be articulated.
- 1216 • The authors should reflect on the factors that influence the performance of the approach.
1217 For example, a facial recognition algorithm may perform poorly when image resolution
1218 is low or images are taken in low lighting. Or a speech-to-text system might not be
1219 used reliably to provide closed captions for online lectures because it fails to handle
1220 technical jargon.
- 1221 • The authors should discuss the computational efficiency of the proposed algorithms
1222 and how they scale with dataset size.
- 1223 • If applicable, the authors should discuss possible limitations of their approach to
1224 address problems of privacy and fairness.
- 1225 • While the authors might fear that complete honesty about limitations might be used by
1226 reviewers as grounds for rejection, a worse outcome might be that reviewers discover
1227 limitations that aren’t acknowledged in the paper. The authors should use their best
1228 judgment and recognize that individual actions in favor of transparency play an impor-
1229 tant role in developing norms that preserve the integrity of the community. Reviewers
1230 will be specifically instructed to not penalize honesty concerning limitations.

1231 3. Theory assumptions and proofs

1232 Question: For each theoretical result, does the paper provide the full set of assumptions and
1233 a complete (and correct) proof?

Answer: [Yes]

Justification: All the core Assumptions are listed in the Problem Setup section and they are recalled in the Appendix with all the detailed proofs.

Guidelines:

- The answer [N/A] means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: All the codes are provided and can be easily reexecuted

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- If the paper includes experiments, a [No] answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

1287 Question: Does the paper provide open access to the data and code, with sufficient instruc-
1288 tions to faithfully reproduce the main experimental results, as described in supplemental
1289 material?

1290 Answer: [Yes]

1291 Justification: A readme with complete guidelines for reproducibility are provided

1292 Guidelines:

- 1293 • The answer [N/A] means that paper does not include experiments requiring code.
- 1294 • Please see the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- 1295 • While we encourage the release of code and data, we understand that this might not
1296 be possible, so [No] is an acceptable answer. Papers cannot be rejected simply for not
1297 including code, unless this is central to the contribution (e.g., for a new open-source
1298 benchmark).
- 1299 • The instructions should contain the exact command and environment needed to run to
1300 reproduce the results. See the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- 1301 • The authors should provide instructions on data access and preparation, including how
1302 to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- 1303 • The authors should provide scripts to reproduce all experimental results for the new
1304 proposed method and baselines. If only a subset of experiments are reproducible, they
1305 should state which ones are omitted from the script and why.
- 1306 • At submission time, to preserve anonymity, the authors should release anonymized
1307 versions (if applicable).
- 1308 • Providing as much information as possible in supplemental material (appended to the
1309 paper) is recommended, but including URLs to data and code is permitted.
- 1310
- 1311

1312 6. Experimental setting/details

1313 Question: Does the paper specify all the training and test details (e.g., data splits, hyperpa-
1314 rameters, how they were chosen, type of optimizer) necessary to understand the results?

1315 Answer: [Yes]

1316 Justification: The experiments are always introduced with all the settings

1317 Guidelines:

- 1318 • The answer [N/A] means that the paper does not include experiments.
- 1319 • The experimental setting should be presented in the core of the paper to a level of detail
1320 that is necessary to appreciate the results and make sense of them.
- 1321 • The full details can be provided either with the code, in appendix, or as supplemental
1322 material.

1323 7. Experiment statistical significance

1324 Question: Does the paper report error bars suitably and correctly defined or other appropriate
1325 information about the statistical significance of the experiments?

1326 Answer: [Yes]

1327 Justification: Each configuration is repeated over 3 random seeds; with shaded bands in all
1328 plots denote ± 1 standard deviation across seeds

1329 Guidelines:

- 1330 • The answer [N/A] means that the paper does not include experiments.
- 1331 • The authors should answer [Yes] if the results are accompanied by error bars, confidence
1332 intervals, or statistical significance tests, at least for the experiments that support the
1333 main claims of the paper.
- 1334 • The factors of variability that the error bars are capturing should be clearly stated (for
1335 example, train/test split, initialization, random drawing of some parameter, or overall
1336 run with given experimental conditions).
- 1337 • The method for calculating the error bars should be explained (closed form formula,
1338 call to a library function, bootstrap, etc.)

- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g., negative error rates).
- If error bars are reported in tables or plots, the authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: This information is provided in the readme of the codes

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We carefully checked the author handbook

Guidelines:

- The answer [N/A] means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer [No], they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [N/A]

Justification: This manuscript does not have societal impact

Guidelines:

- The answer [N/A] means that there is no societal impact of the work performed.
- If the authors answer [N/A] or [No], they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate Deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pre-trained language models, image generators, or scraped datasets)?

Answer: [N/A]

Justification: The current manuscript does not pose such risks

Guidelines:

- The answer [N/A] means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: -

Guidelines:

- The answer [N/A] means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- 1442 • If this information is not available online, the authors are encouraged to reach out to
1443 the asset’s creators.

1444 **13. New assets**

1445 Question: Are new assets introduced in the paper well documented and is the documentation
1446 provided alongside the assets?

1447 Answer: [N/A]

1448 Justification: [N/A]

1449 Guidelines:

- 1450 • The answer [N/A] means that the paper does not release new assets.
1451 • Researchers should communicate the details of the dataset/code/model as part of their
1452 submissions via structured templates. This includes details about training, license,
1453 limitations, etc.
1454 • The paper should discuss whether and how consent was obtained from people whose
1455 asset is used.
1456 • At submission time, remember to anonymize your assets (if applicable). You can either
1457 create an anonymized URL or include an anonymized zip file.

1458 **14. Crowdsourcing and research with human subjects**

1459 Question: For crowdsourcing experiments and research with human subjects, does the paper
1460 include the full text of instructions given to participants and screenshots, if applicable, as
1461 well as details about compensation (if any)?

1462 Answer: [N/A]

1463 Justification: [N/A]

1464 Guidelines:

- 1465 • The answer [N/A] means that the paper does not involve crowdsourcing nor research
1466 with human subjects.
1467 • Including this information in the supplemental material is fine, but if the main contribu-
1468 tion of the paper involves human subjects, then as much detail as possible should be
1469 included in the main paper.
1470 • According to the NeurIPS Code of Ethics, workers involved in data collection, curation,
1471 or other labor should be paid at least the minimum wage in the country of the data
1472 collector.

1473 **15. Institutional review board (IRB) approvals or equivalent for research with human
1474 subjects**

1475 Question: Does the paper describe potential risks incurred by study participants, whether
1476 such risks were disclosed to the subjects, and whether Institutional Review Board (IRB)
1477 approvals (or an equivalent approval/review based on the requirements of your country or
1478 institution) were obtained?

1479 Answer: [N/A]

1480 Justification: [N/A]

1481 Guidelines:

- 1482 • The answer [N/A] means that the paper does not involve crowdsourcing nor research
1483 with human subjects.
1484 • Depending on the country in which research is conducted, IRB approval (or equivalent)
1485 may be required for any human subjects research. If you obtained IRB approval, you
1486 should clearly state this in the paper.
1487 • We recognize that the procedures for this may vary significantly between institutions
1488 and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the
1489 guidelines for their institution.
1490 • For initial submissions, do not include any information that would break anonymity (if
1491 applicable), such as the institution conducting the review.

1492 **16. Declaration of LLM usage**

1493 Question: Does the paper describe the usage of LLMs if it is an important, original, or
1494 non-standard component of the core methods in this research? Note that if the LLM is used
1495 only for writing, editing, or formatting purposes and does *not* impact the core methodology,
1496 scientific rigor, or originality of the research, declaration is not required.

1497 Answer: [N/A]

1498 Justification: LLM were used in a standard way for drafting and editing

1499 Guidelines:

- 1500 • The answer [N/A] means that the core method development in this research does not
- 1501 involve LLMs as any important, original, or non-standard components.
- 1502 • Please refer to our LLM policy in the NeurIPS handbook for what should or should not
- 1503 be described.