

КАК РАБОТАЮТ ДИФFUЗИОННЫЕ МОДЕЛИ

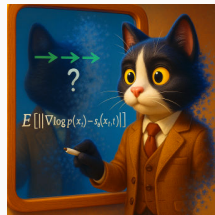
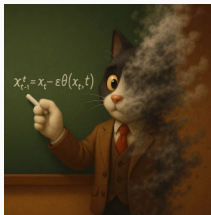
Сергей Николенко

СПбГУ — Санкт-Петербург

20 ноября 2025 г.

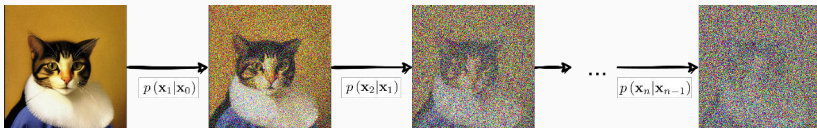
Random facts:

- 20 ноября в ООН — Всемирный день ребёнка; в этот день в 1959 г. была принята «Декларация прав ребёнка», а в 1989 — «Конвенция прав ребёнка»
- 20 ноября 1805 г. венский мясник Иоганн Ланер изобрёл сосиски
- 20 ноября 1910 г. Франсиско Мадеро опубликовал «план Сан-Луис-Потоси», в котором объявлял результаты президентских выборов недействительными и призывал к борьбе с режимом Порфирио Диаса, правившего к тому времени 34 года, чем запустил Мексиканскую революцию; а в доме Ивана Озолина, начальника станции Астапово, в тот же день умер Лев Толстой
- 20 ноября 1947 г. в Вестминстерском аббатстве принцесса Елизавета вышла замуж за лейтенанта Филипа Маунтбаттена, который при этом стал герцогом Эдинбургским
- 20 ноября 1979 г. у стен Каабы появился человек и объявил, что настало время для исполнения древнего пророчества, во имя которого в пределах Запретной Мечети была пролита кровь; захват заложников длился до 4 декабря, погибло 255 человек
- 20 ноября 1985 г. вышла первая графическая операционная система от Microsoft, Windows 1.0



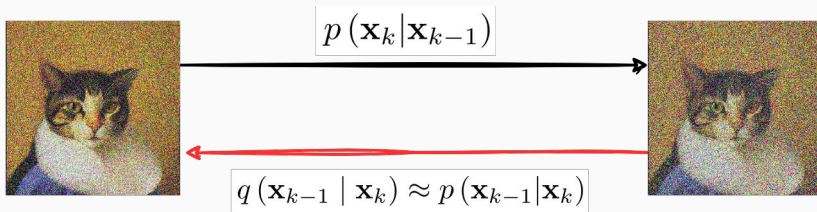
ДИФФУЗИОННЫЕ МОДЕЛИ

- (Sohl-Dickstein et al., 2015): Deep Unsupervised Learning using Nonequilibrium Thermodynamics
- Идея из статистической физики: давайте обучать марковскую цепь, которая постепенно сделает из случайного шума нужное распределение
- Для этого начнём с цепи, которая сделает из нужного распределения случайный шум:

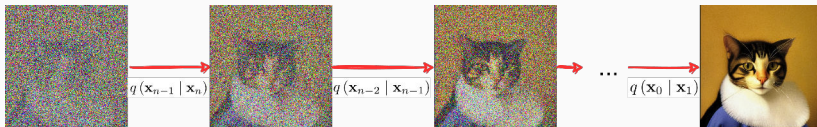


ДИФфуЗИОННЫЕ МОДЕЛИ

- А потом попробуем обратить каждый шаг:



- Если получится, то потом мы можем начать со случайного шума и двигаться обратно:



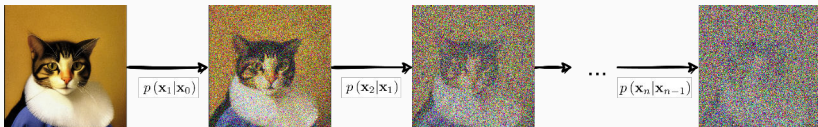
ДИФфуЗИОННЫЕ МОДЕЛИ

- Forward diffusion: начнём с $p(\mathbf{x}_0)$ и будем добавлять гауссовский шум:

$$q(\mathbf{x}_t | \mathbf{x}_{t-1}) = N(\mathbf{x}_t | \sqrt{1 - \beta_t} \mathbf{x}_{t-1}, \beta_t \mathbf{I}).$$

- Тогда постепенно $\mathbf{x}_0$ превращается в шум, скажем, за T шагов:

$$q(\mathbf{x}_{1:T} | \mathbf{x}_0) = \prod_{t=1}^T q(\mathbf{x}_t | \mathbf{x}_{t-1}).$$



- И для каждого t мы можем просэмплировать $\mathbf{x}_t$ в замкнутой форме через reparametrization trick: пусть $\epsilon \sim N(0, \mathbf{I})$ и обозначим $\alpha_t = 1 - \beta_t$, $A_t = \prod_{i=1}^T \alpha_i$; тогда

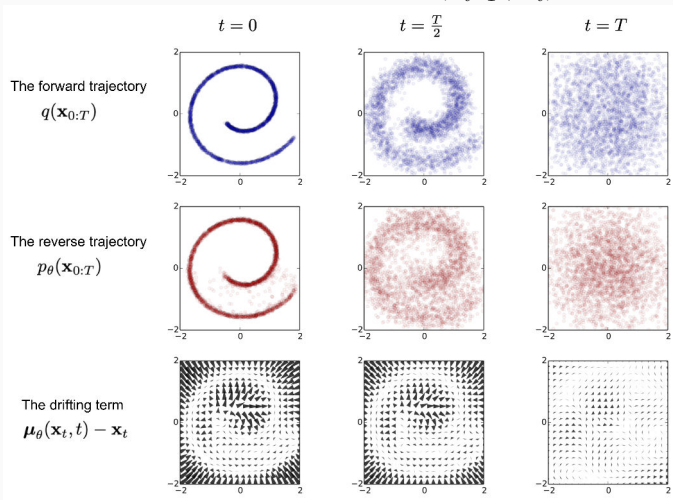
$$\begin{aligned}\mathbf{x}_t &= \sqrt{\alpha_t} \mathbf{x}_{t-1} + \sqrt{1 - \alpha_t} \epsilon = \sqrt{\alpha_t \alpha_{t-1}} \mathbf{x}_{t-2} + \sqrt{1 - \alpha_t \alpha_{t-1}} \epsilon = \dots \\ \dots &= \sqrt{A_t} \mathbf{x}_0 + \sqrt{1 - A_t} \epsilon, \text{ то есть } q(\mathbf{x}_t | \mathbf{x}_0) = N(\mathbf{x}_t | \sqrt{A_t} \mathbf{x}_0, (1 - A_t) \mathbf{I})\end{aligned}$$

потому что когда мы складываем два гауссиана $N(0, \sigma_1^2 \mathbf{I})$ и $N(0, \sigma_2^2 \mathbf{I})$, получается гауссиан $N(0, (\sigma_1^2 + \sigma_2^2) \mathbf{I})$, и у нас это

$$(1 - \alpha_t) + \alpha_t (1 - \alpha_t) = 1 - \alpha_t \alpha_{t-1}.$$

ДИФфуЗИОННЫЕ МОДЕЛИ

- А теперь задача, казалось бы, в том, чтобы обратить этот процесс, научиться сэмплировать из $q(\mathbf{x}_{t-1} \mid \mathbf{x}_t)$



- Для маленьких β_t обратные распределения тоже будут гауссовскими, но параметры их теперь уже зависят от точек данных.
- Reverse diffusion: теперь уже надо обучать распределения

$$p_{\theta}(\mathbf{x}_{0:T}) = p_{\theta}(\mathbf{x}_T) \prod_{t=1}^T p_{\theta}(\mathbf{x}_{t-1} \mid \mathbf{x}_t), \quad \text{где}$$

$$p_{\theta}(\mathbf{x}_{t-1} \mid \mathbf{x}_t) = N(\mathbf{x}_{t-1} \mid \mu_{\theta}(\mathbf{x}_t, t), \Sigma_{\theta}(\mathbf{x}_t, t)).$$

- Если бы мы знали $\mathbf{x}_0$, то обратить было бы можно аналитически, кстати:

$$q(\mathbf{x}_{t-1} | \mathbf{x}_t, \mathbf{x}_0) = q(\mathbf{x}_t | \mathbf{x}_{t-1}, \mathbf{x}_0) \frac{q(\mathbf{x}_{t-1} | \mathbf{x}_0)}{q(\mathbf{x}_t | \mathbf{x}_0)} \propto \\ \propto e^{-\frac{1}{2} \left(\frac{(\mathbf{x}_t - \sqrt{\alpha_t} \mathbf{x}_{t-1})^2}{\beta_t} + \frac{(\mathbf{x}_{t-1} - \sqrt{A_{t-1}} \mathbf{x}_0)^2}{1 - A_{t-1}} + \frac{(\mathbf{x}_t - \sqrt{A_t} \mathbf{x}_0)^2}{1 - A_t} \right)},$$

и в этой формуле теперь надо выделить полный квадрат по $\mathbf{x}_{t-1} \dots$

- Получится

$$q(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0) = N(\mathbf{x}_{t-1} \mid \tilde{\mu}(\mathbf{x}_t, \mathbf{x}_0), \tilde{\beta}_t \mathbf{I}), \quad \text{где}$$

$$\tilde{\beta}_t = \frac{1 - A_{t-1}}{1 - A_t} \cdot \beta_t,$$

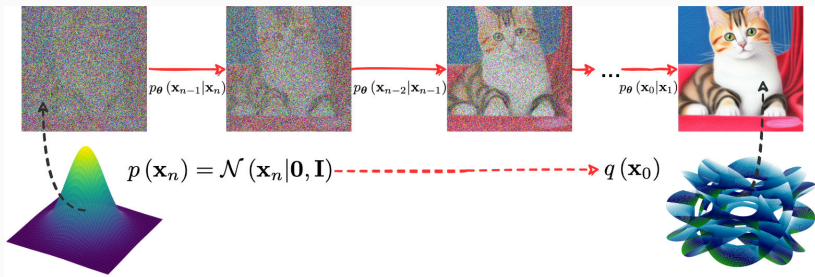
$$\tilde{\mu}(\mathbf{x}_t, \mathbf{x}_0) = \frac{\sqrt{\alpha_t}(1 - A_{t-1})}{1 - A_t} \mathbf{x}_t + \frac{\sqrt{A_{t-1}}\beta_t}{1 - A_t} \mathbf{x}_0,$$

и поскольку мы знаем, что $\mathbf{x}_t = \sqrt{A_t}\mathbf{x}_0 + \sqrt{1 - A_t}\epsilon$,

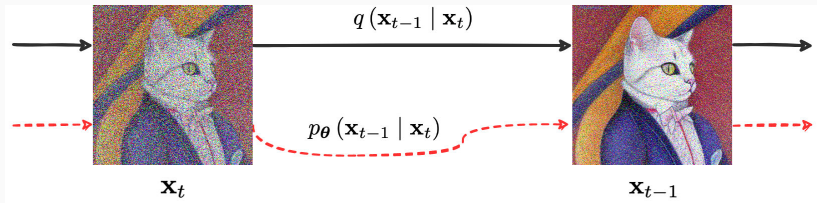
$$\tilde{\mu}(\mathbf{x}_t, \mathbf{x}_0) = \frac{1}{\sqrt{A_t}} \left(\mathbf{x}_t - \frac{1 - \alpha_t}{\sqrt{1 - A_t}} \epsilon_t \right).$$

ДИФфуЗИОННЫЕ МОДЕЛИ

- В конечном счёте хочется, чтобы обратная диффузия реконструировала $q(\mathbf{x}_0)$ из стандартного входа в $\mathbf{x}_n$; примерно так:



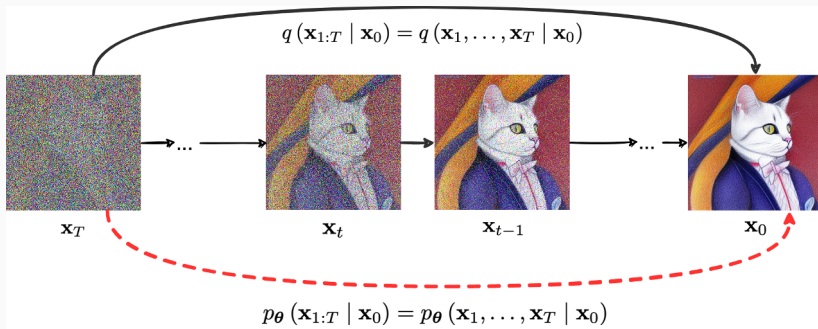
- Что делать? Как всегда в байесовском выводе, приближать!
- На каждом шаге модель должна стать хорошим приближением к $q(\mathbf{x}_t | \mathbf{x}_{t-1})$, без условия на неизвестное $\mathbf{x}_0$:



ДИФфуЗИОННЫЕ МОДЕЛИ

- Чтобы получить эту аппроксимацию, нужна вариационная нижняя оценка, похожая на ту, которая была в VAE и DALL-E
- Начинаем с оценки на глобальное распределение

$$q(\mathbf{x}_{1:T} \mid \mathbf{x}_0) = q(\mathbf{x}_1, \dots, \mathbf{x}_T \mid \mathbf{x}_0):$$



- А потом она разложится на оценки для отдельных шагов диффузионного процесса

- Вспомним идею вариационных приближений:

$$p(\mathbf{x}, \mathbf{z}) = p(\mathbf{x}) p(\mathbf{z}|\mathbf{x}),$$

$$\log p(\mathbf{x}) = \log p(\mathbf{x}, \mathbf{z}) - \log p(\mathbf{z}|\mathbf{x}),$$

возьмём ожидание по $q(\mathbf{z})$:

$$\mathbb{E}_{q(\mathbf{z})} [\log p(\mathbf{x})] = \mathbb{E}_{q(\mathbf{z})} [\log p(\mathbf{x}, \mathbf{z})] - \mathbb{E}_{q(\mathbf{z})} [\log p(\mathbf{z}|\mathbf{x})],$$

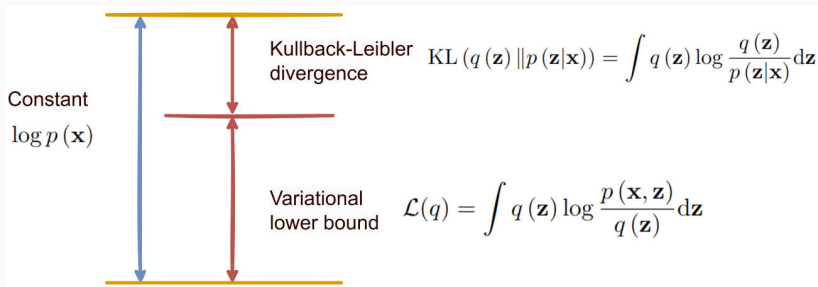
а затем сделаем стандартные преобразования:

$$\begin{aligned} \log p(\mathbf{x}) &= \mathbb{E}_{q(\mathbf{z})} [\log p(\mathbf{x}, \mathbf{z})] - \mathbb{E}_{q(\mathbf{z})} [\log q(\mathbf{z})] + \\ &\quad + \mathbb{E}_{q(\mathbf{z})} [\log q(\mathbf{z})] - \mathbb{E}_{q(\mathbf{z})} [\log p(\mathbf{z}|\mathbf{x})], \\ \log p(\mathbf{x}) &= \int q(\mathbf{z}) \log \frac{p(\mathbf{x}, \mathbf{z})}{q(\mathbf{z})} d\mathbf{z} + \int q(\mathbf{z}) \log \frac{q(\mathbf{z})}{p(\mathbf{z}|\mathbf{x})} d\mathbf{z}. \end{aligned}$$

- Итого получается, что

$$\log p(\mathbf{x}) = L(q) + \text{KL}(q(\mathbf{z}) \| p(\mathbf{z}|\mathbf{x})), \quad \text{где}$$

$$L(q) = \int q(\mathbf{z}) \log \frac{p(\mathbf{x}, \mathbf{z})}{q(\mathbf{z})} d\mathbf{z}.$$



- Давайте для диффузионной модели выведем вариационную оценку снова из первых принципов:

$$\log p_{\theta}(\mathbf{x}_0) = \log p_{\theta}(\mathbf{x}_0, \dots, \mathbf{x}_T) - \log p_{\theta}(\mathbf{x}_1, \dots, \mathbf{x}_T \mid \mathbf{x}_0),$$

$$\begin{aligned}\mathbb{E}_{q(\mathbf{x}_0)} [\log p_{\theta}(\mathbf{x}_0)] &= \mathbb{E}_{q(\mathbf{x}_{0:T})} [\log p_{\theta}(\mathbf{x}_{0:T})] - \mathbb{E}_{q(\mathbf{x}_{0:T})} [\log p_{\theta}(\mathbf{x}_{1:T} \mid \mathbf{x}_0)] = \\ &= \mathbb{E}_{q(\mathbf{x}_{0:T})} \left[\log \frac{p_{\theta}(\mathbf{x}_{0:T})}{q(\mathbf{x}_{1:T} \mid \mathbf{x}_0)} \right] + \mathbb{E}_{q(\mathbf{x}_{0:T})} \left[\log \frac{q(\mathbf{x}_{1:T} \mid \mathbf{x}_0)}{p_{\theta}(\mathbf{x}_{1:T} \mid \mathbf{x}_0)} \right].\end{aligned}$$

- Здесь второе слагаемое — это KL-дивергенция, и мы получаем функцию ошибки как минус вариационную нижнюю оценку:

$$L = \mathbb{E}_{q(\mathbf{x}_{0:T})} \left[\log \frac{q(\mathbf{x}_{1:T} \mid \mathbf{x}_0)}{p_{\theta}(\mathbf{x}_{0:T})} \right] \geq -\mathbb{E}_{q(\mathbf{x}_0)} [\log p_{\theta}(\mathbf{x}_0)].$$

- Эту оценку можно подсчитать как сумму:

$$\begin{aligned}
 L &= \mathbb{E}_q \left[\log \frac{q(\mathbf{x}_{1:T} | \mathbf{x}_0)}{p_\theta(\mathbf{x}_{0:T})} \right] = \mathbb{E}_q \left[\log \frac{\prod_{t=1}^T q(\mathbf{x}_t | \mathbf{x}_{t-1})}{p_\theta(\mathbf{x}_T) \prod_{t=1}^T p_\theta(\mathbf{x}_{t-1} | \mathbf{x}_t)} \right] \\
 &= \mathbb{E}_q \left[-\log p_\theta(\mathbf{x}_T) + \sum_{t=1}^T \log \frac{q(\mathbf{x}_t | \mathbf{x}_{t-1})}{p_\theta(\mathbf{x}_{t-1} | \mathbf{x}_t)} \right] \\
 &= \mathbb{E}_q \left[-\log p_\theta(\mathbf{x}_T) + \sum_{t=2}^T \log \frac{q(\mathbf{x}_t | \mathbf{x}_{t-1})}{p_\theta(\mathbf{x}_{t-1} | \mathbf{x}_t)} + \log \frac{q(\mathbf{x}_1 | \mathbf{x}_0)}{p_\theta(\mathbf{x}_0 | \mathbf{x}_1)} \right] \\
 &= \mathbb{E}_q \left[-\log p_\theta(\mathbf{x}_T) + \sum_{t=2}^T \log \left(\frac{q(\mathbf{x}_{t-1} | \mathbf{x}_t, \mathbf{x}_0)}{p_\theta(\mathbf{x}_{t-1} | \mathbf{x}_t)} \frac{q(\mathbf{x}_t | \mathbf{x}_0)}{q(\mathbf{x}_{t-1} | \mathbf{x}_0)} \right) + \log \frac{q(\mathbf{x}_1 | \mathbf{x}_0)}{p_\theta(\mathbf{x}_0 | \mathbf{x}_1)} \right] \\
 &= \mathbb{E}_q \left[-\log p_\theta(\mathbf{x}_T) + \sum_{t=2}^T \log \frac{q(\mathbf{x}_{t-1} | \mathbf{x}_t, \mathbf{x}_0)}{p_\theta(\mathbf{x}_{t-1} | \mathbf{x}_t)} + \log \frac{q(\mathbf{x}_T | \mathbf{x}_0)}{q(\mathbf{x}_1 | \mathbf{x}_0)} + \log \frac{q(\mathbf{x}_1 | \mathbf{x}_0)}{p_\theta(\mathbf{x}_0 | \mathbf{x}_1)} \right] \\
 &= \mathbb{E}_q \left[\log \frac{q(\mathbf{x}_T | \mathbf{x}_0)}{p_\theta(\mathbf{x}_T)} + \sum_{t=2}^T \log \frac{q(\mathbf{x}_{t-1} | \mathbf{x}_t, \mathbf{x}_0)}{p_\theta(\mathbf{x}_{t-1} | \mathbf{x}_t)} - \log p_\theta(\mathbf{x}_0 | \mathbf{x}_1) \right].
 \end{aligned}$$

- Итого получается

$$L = L_T + L_{T-1} + \dots + L_0, \quad \text{where}$$

$$L_T = \text{KL} (q (\mathbf{x}_T \mid \mathbf{x}_0) \parallel p_\theta (\mathbf{x}_T)),$$

$$L_t = \text{KL} (q (\mathbf{x}_t \mid \mathbf{x}_{t+1}, \mathbf{x}_0) \parallel p_\theta (\mathbf{x}_t \mid \mathbf{x}_{t+1})), \quad t = 1, \dots, T-1,$$

$$L_0 = -\log p_\theta (\mathbf{x}_0 \mid \mathbf{x}_1).$$

ДИФфуЗИОННЫЕ МОДЕЛИ

- Всё это KL-дивергенции между гауссианами, их можно подсчитать и подставить в функцию ошибки
- Например, в L_t мы используем гауссовскую параметризацию

$$p_{\theta}(\mathbf{x}_{t-1} | \mathbf{x}_t) = N(\mathbf{x}_{t-1} | \mu_{\theta}(\mathbf{x}_t, t), \Sigma_{\theta}(\mathbf{x}_t, t))$$

и пытаемся её параметры совместить с $q(\mathbf{x}_t | \mathbf{x}_{t-1}, \mathbf{x}_0)$. Для среднего, например,

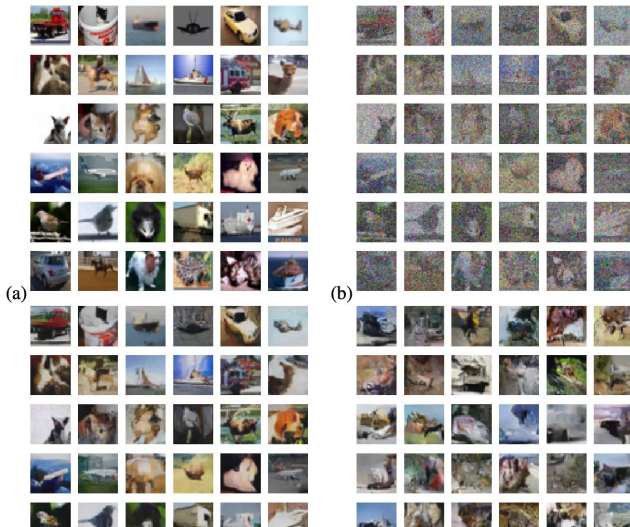
$$\mu_{\theta}(\mathbf{x}_t, t) \approx \tilde{\mu}_t(\mathbf{x}_t, \mathbf{x}_0) = \frac{1}{\sqrt{A_t}} \left(\mathbf{x}_t - \frac{1 - \alpha_t}{\sqrt{1 - A_t}} \epsilon_t \right),$$

и поскольку мы знаем $\mathbf{x}_t$ при обучении, мы можем параметризовать шум напрямую:

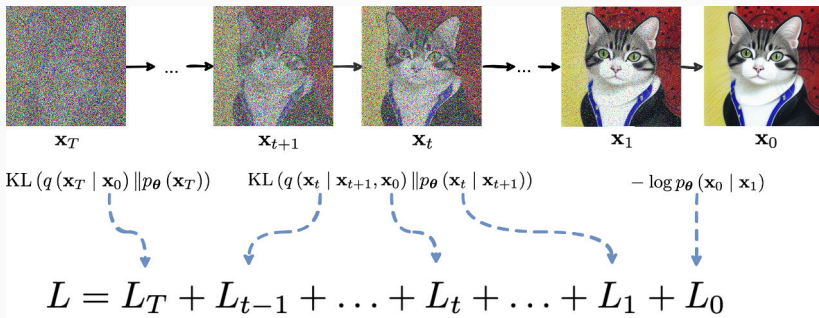
$$p_{\theta}(\mathbf{x}_{t-1} | \mathbf{x}_t) = N \left(\mathbf{x}_{t-1} \left| \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{x}_t - \frac{1 - \alpha_t}{\sqrt{1 - A_t}} \epsilon_{\theta}(\mathbf{x}_t, t) \right), \Sigma_{\theta}(\mathbf{x}_t, t) \right) \right).$$

ДИФфуЗИОННЫЕ МОДЕЛИ

- (Sohl-Dickstein et al., 2015) само по себе не слишком хорошо работало, разве что на CIFAR:



- Следующий шаг: “Denoising Diffusion Probabilistic Models” (DDPM; Ho et al., 2020)
- Та же основная идея, та же структура функции ошибки:



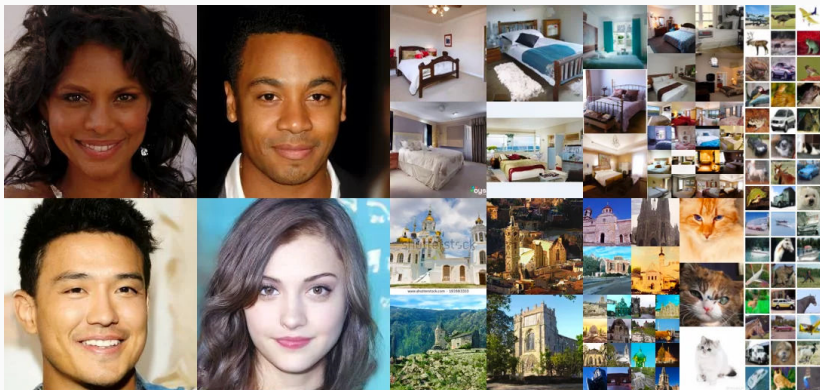
- Но дальше DDPM делает следующие замечания:
 - дисперсии прямой диффузии β_t — постоянные гиперпараметры, их не обучают, так что все q не обучаются; т.к. $p_\theta(\mathbf{x}_T)$ — фиксированное распределение, из которого мы будем сэмплировать, L_T является константой;
 - для промежуточных шагов дисперсии в $p_\theta(\mathbf{x}_t | \mathbf{x}_{t+1})$ тоже не обучаются, и можно найти простую замкнутую форму для L_t ;
 - главное — давайте введём отдельный дискретный декодер для L_0 ; если данные состоят из целых чисел от 0 до 255, отмасштабированных в $[-1, 1]$ (это естественное представление для картинок), DDPM моделирует

$$p_\theta(\mathbf{x}_0 | \mathbf{x}_1) = \prod_{i=1}^D \int_{\delta_-(x_0,i)}^{\delta_+(x_0,i)} N(x | \mu_{\theta,i}(\mathbf{x}_1), \sigma_1^2) dx,$$

где i идёт по пикселям, $\mu_\theta(\mathbf{x}_1)$ — декодер, пределы интегрирования $\frac{1}{255}$ в каждую сторону от $x_{0,i}$

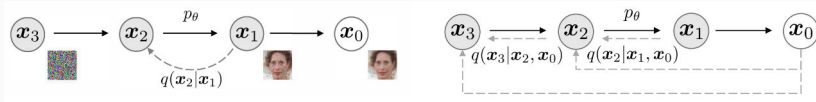
- Т.е. можно подставить другую модель на последнем шаге и использовать $\mu_\theta(\mathbf{x}_1)$ без шума во время сэмплирования

- DDPM уже давала отличное порождение, сравнимое с лучшими GAN'ами того времени:



- Следующий шаг пришёл очень скоро: “Denoising Diffusion Implicit Models” (DDIM; Song et al., 2020)
- Важный недостаток всех диффузионных моделей пока что в том, что они *очень медленные*
- Порождение должно пройти по всем шагам диффузии, а их буквально тысячи, и они идут последовательно, не параллелизуются
- Song et al. упоминают, что сэмплирование из обученного GAN примерно в 1000 раз быстрее, чем сэмплирование из DDPM для того же размера картинок

- Как можно ускорить диффузионные модели?
- Song et al. обобщают диффузионные модели и в частности DDPM
- Функция ошибки L не зависит напрямую от совместного распределения $q(\mathbf{x}_{1:T} | \mathbf{x}_0)$, а только от маргиналов $q(\mathbf{x}_t | \mathbf{x}_0)$
- Это значит, что мы можем переиспользовать ту же функцию ошибки для другого совместного распределения, если у него те же маргиналы
- В частности, обобщить можно даже на немарковские диффузионные процессы, если для них получится найти обратную марковскую цепь:



- Зададим вероятностную модель для данных $\mathbf{x}_0$ (из распределения $q(\mathbf{x}_0)$) через латентные переменные $\mathbf{x}_1, \dots, \mathbf{x}_T$
- Сначала прямой проход (noising) $q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0)$, в котором к данным постепенно добавляется шум:

$$q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0) = q_\sigma(\mathbf{x}_T \mid \mathbf{x}_0) \prod_{t=2}^T q_\sigma(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0).$$

- Здесь мы по определению задаём, как связаны $\mathbf{x}_1, \dots, \mathbf{x}_T$ с $\mathbf{x}_0$
- Важный момент: переходы $q_\sigma(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0)$ зависят и от $\mathbf{x}_t$, и от $\mathbf{x}_0$, т.е. процесс $\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_T$ не марковский! это осознанное решение

- Мы выбираем следующие распределения (с тем же маргиналом, что и в DDPM):

$$q_{\sigma}(\mathbf{x}_T | \mathbf{x}_0) = \mathcal{N}(\sqrt{A_T} \mathbf{x}_0, (1 - A_T) \mathbf{I}),$$

$$q_{\sigma}(\mathbf{x}_{t-1} | \mathbf{x}_t, \mathbf{x}_0) = \mathcal{N}\left(\sqrt{A_{t-1}} \mathbf{x}_0 + \sqrt{1 - A_{t-1} - \sigma_t^2} \frac{\mathbf{x}_t - \sqrt{A_t} \mathbf{x}_0}{\sqrt{1 - A_t}}, \sigma_t^2 \mathbf{I}\right).$$

- σ_t^2 – дисперсия добавочного шума в обратном переходе; форма среднего такая для того, чтобы маргиналы $q_{\sigma}(\mathbf{x}_t | \mathbf{x}_0)$ оставались гауссианами и совпадали с DDPM

- Ключевое утверждение: для любого t

$$q_{\sigma}(\mathbf{x}_t \mid \mathbf{x}_0) = \mathcal{N}(\sqrt{A_t} \mathbf{x}_0, (1 - A_t) \mathbf{I}).$$

- Эквивалентно можно записать в виде «одноступенчатой» формулы:

$$\mathbf{x}_t = \sqrt{A_t} \mathbf{x}_0 + \sqrt{1 - A_t} \epsilon_t, \quad \epsilon_t \sim \mathcal{N}(\mathbf{0}, \mathbf{I}).$$

- Это ключевая формула DDPM/DDIM: любое зашумлённое состояние $\mathbf{x}_t$ можно получить напрямую из $\mathbf{x}_0$, не прогоняя цепочку шаг за шагом

- Доказательство по индукции. Предположим, что

$$q_{\sigma}(\mathbf{x}_t \mid \mathbf{x}_0) = \mathcal{N}(\sqrt{A_t} \mathbf{x}_0, (1 - A_t) \mathbf{I}), \quad \text{то есть}$$

$$\mathbf{x}_t = \sqrt{A_t} \mathbf{x}_0 + \sqrt{1 - A_t} \epsilon_t, \quad \epsilon_t \sim \mathcal{N}(\mathbf{0}, \mathbf{I}).$$

- Подставим это в $q_{\sigma}(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0)$, получим

$$q_{\sigma}(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0) = \mathcal{N}\left(\sqrt{A_{t-1}} \mathbf{x}_0 + \sqrt{1 - A_{t-1} - \sigma_t^2} \epsilon_t, \sigma_t^2 \mathbf{I}\right).$$

- Мы просто заменили $\frac{\mathbf{x}_t - \sqrt{A_t} \mathbf{x}_0}{\sqrt{1 - A_t}}$ на ϵ_t , используя выражение для $\mathbf{x}_t$ через $\mathbf{x}_0$ и ϵ_t ; среднее стало линейной функцией $\mathbf{x}_0$ и ϵ_t

- Используем факт про гауссианы: если

$$p(\mathbf{x}) = \mathcal{N}(\mathbf{x} \mid \mu, \Lambda^{-1}),$$

$$p(\mathbf{y} \mid \mathbf{x}) = \mathcal{N}(\mathbf{y} \mid A\mathbf{x} + \mathbf{b}, L^{-1}), \quad \text{то}$$

$$p(\mathbf{y}) = \mathcal{N}(\mathbf{y} \mid A\mu + \mathbf{b}, L^{-1} + A\Lambda^{-1}A^\top).$$

- В нашем случае $\mathbf{x}$ – это ϵ_t , распределённый как $\mathcal{N}(\mathbf{0}, \mathbf{I})$; $\mathbf{y}$ – это $\mathbf{x}_{t-1}$; матрица A и сдвиг $\mathbf{b}$ задаются коэффициентами при ϵ_t и $\mathbf{x}_0$

- Сопоставим параметры:

$$\mu \equiv \mathbf{0}, \quad \Lambda^{-1} \equiv \mathbf{I},$$

$$\mathbf{A} \equiv \sqrt{1 - A_{t-1} - \sigma_t^2} \mathbf{I}, \quad \mathbf{b} \equiv \sqrt{A_{t-1}} \mathbf{x}_0, \quad L^{-1} \equiv \sigma_t^2 \mathbf{I}.$$

- Подставив в формулу, получим

$$q_\sigma(\mathbf{x}_{t-1} \mid \mathbf{x}_0) = \mathcal{N}(A\mu + \mathbf{b}, L^{-1} + A\Lambda^{-1}A^\top) = \mathcal{N}(\sqrt{A_{t-1}}\mathbf{x}_0, (1 - A_{t-1})\mathbf{I}),$$

потому что

$$A\mu + \mathbf{b} = \sqrt{A_{t-1}} \mathbf{x}_0,$$

$$L^{-1} + A\Lambda^{-1}A^\top = \sigma_t^2 \mathbf{I} + (1 - A_{t-1} - \sigma_t^2) \mathbf{I} = (1 - A_{t-1}) \mathbf{I}.$$

- Из полученного выражения

$$\mathbf{x}_t = \sqrt{A_t} \mathbf{x}_0 + \sqrt{1 - A_t} \epsilon_t$$

можно выразить $\mathbf{x}_0$ через $\mathbf{x}_t$ и ϵ_t :

$$\mathbf{x}_0 = \frac{\mathbf{x}_t - \sqrt{1 - A_t} \epsilon_t}{\sqrt{A_t}}.$$

- Эта формула важна для порождения, т.к. в обратном процессе мы будем оценивать ϵ_t нейросетью и тем самым получать оценку $\mathbf{x}_0$ из $\mathbf{x}_t$

- Поскольку истинный шум ϵ_t нам неизвестен, мы вводим аппроксимацию

$$\epsilon_{\theta}^{(t)}(\mathbf{x}_t) \approx \epsilon_t,$$

где $\epsilon_{\theta}^{(t)}$ – нейросеть, зависящая от $\mathbf{x}_t$ и (неявно) от шага t

- Тогда оценка $\mathbf{x}_0$ на шаге t задаётся как

$$f_{\theta}^{(t)}(\mathbf{x}_t) \equiv \frac{\mathbf{x}_t - \sqrt{1 - A_t} \epsilon_{\theta}^{(t)}(\mathbf{x}_t)}{\sqrt{A_t}}.$$

- Обратное выражение:

$$\epsilon_{\theta}^{(t)}(\mathbf{x}_t) = \frac{\mathbf{x}_t - \sqrt{A_t} f_{\theta}^{(t)}(\mathbf{x}_t)}{\sqrt{1 - A_t}}.$$

- Теперь переходим к обратному (порождающему) процессу
- Совместное распределение модели задаётся как

$$p_{\theta}(\mathbf{x}_{0:T}) = p_{\theta}(\mathbf{x}_T) \prod_{t=1}^T p_{\theta}^{(t)}(\mathbf{x}_{t-1} \mid \mathbf{x}_t),$$

где обычно берут

$$p_{\theta}(\mathbf{x}_T) = \mathcal{N}(\mathbf{0}, \mathbf{I}),$$

а переходы $p_{\theta}^{(t)}$ параметризуются через нейросеть

- Правдоподобие данных:

$$\log p(\mathbf{x}_0) = \log \int p_{\theta}(\mathbf{x}_{0:T}) d\mathbf{x}_{1:T}$$

- Введём вспомогательное распределение – наш прямой процесс $q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0)$, умножим и поделим на него внутри интеграла и применим неравенство Йенсена:

$$\begin{aligned}\log p(\mathbf{x}_0) &= \log \int p_\theta(\mathbf{x}_{0:T}) d\mathbf{x}_{1:T} = \log \int \frac{p_\theta(\mathbf{x}_{0:T})}{q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0)} q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0) d\mathbf{x}_{1:T} \\ &= \log \mathbb{E}_{q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0)} \left[\frac{p_\theta(\mathbf{x}_{0:T})}{q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0)} \right] \\ &\geq \mathbb{E}_{q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0)} [\log p_\theta(\mathbf{x}_{0:T}) - \log q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0)].\end{aligned}$$

- Обозначим $J_\sigma = -\mathbb{E}_{q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0)} [\log p_\theta(\mathbf{x}_{0:T}) - \log q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0)]$, подставим факторизации p_θ и q_σ :

$$\begin{aligned}q_\sigma(\mathbf{x}_{1:T} \mid \mathbf{x}_0) &= q_\sigma(\mathbf{x}_T \mid \mathbf{x}_0) \prod_{t=2}^T q_\sigma(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0), \\ p_\theta(\mathbf{x}_{0:T}) &= p_\theta(\mathbf{x}_T) \prod_{t=1}^T p_\theta^{(t)}(\mathbf{x}_{t-1} \mid \mathbf{x}_t).\end{aligned}$$

- Тогда после раскрытия логарифмов получаем

$$J_{\sigma} = \mathbb{E}_{q_{\sigma}(\mathbf{x}_{1:T}|\mathbf{x}_0)} \left[\log q_{\sigma}(\mathbf{x}_T | \mathbf{x}_0) + \sum_{t=2}^T \log q_{\sigma}(\mathbf{x}_{t-1} | \mathbf{x}_t, \mathbf{x}_0) - \sum_{t=1}^T \log p_{\theta}^{(t)}(\mathbf{x}_{t-1} | \mathbf{x}_t) - \log p_{\theta}(\mathbf{x}_T) \right].$$

- Следующий шаг – перегруппировать всё это так, чтобы получилось что-то похожее на сумму KL-дивергенций, плюс отдельное слагаемое для реконструкции $\mathbf{x}_0$

- После некоторых мучений получим:

$$J_{\sigma} = \sum_{t=2}^T \mathbb{E}_{q(\mathbf{x}_t | \mathbf{x}_0)} \left[D_{\text{KL}}(q_{\sigma}(\mathbf{x}_{t-1} | \mathbf{x}_t, \mathbf{x}_0) \| p_{\theta}^{(t)}(\mathbf{x}_{t-1} | \mathbf{x}_t)) \right] \\ - \mathbb{E}_{q(\mathbf{x}_1 | \mathbf{x}_0)} [\log p_{\theta}^{(0)}(\mathbf{x}_0 | \mathbf{x}_1)].$$

- Здесь $p_{\theta}^{(0)}(\mathbf{x}_0 | \mathbf{x}_1)$ – отдельный шаг, как в DDPM; в реальных реализациях его часто опускают или сливают с общей ошибкой

- Параметризация обратного шага и KL-дивергенция: для $t > 1$ возьмём

$$p_{\theta}^{(t)}(\mathbf{x}_{t-1} \mid \mathbf{x}_t) = \mathcal{N}\left(\sqrt{A_{t-1}} f_{\theta}^{(t)}(\mathbf{x}_t) + \sqrt{1 - A_{t-1} - \sigma_t^2} \epsilon_{\theta}^{(t)}(\mathbf{x}_t), \sigma_t^2 \mathbf{I}\right).$$

- Это та же структура, что и у $q_{\sigma}(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0)$, только вместо истинного $\mathbf{x}_0$ и истинного шума ϵ_t мы подставляем оценки $f_{\theta}^{(t)}(\mathbf{x}_t)$ и $\epsilon_{\theta}^{(t)}(\mathbf{x}_t)$
- KL между двумя гауссианами с одинаковой ковариацией имеет особенно простую форму – это просто квадратичная функция разности средних

- У нас:

$$q_{\sigma}(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0) = \mathcal{N}\left(\sqrt{A_{t-1}}\mathbf{x}_0 + \sqrt{1 - A_{t-1} - \sigma_t^2}\epsilon_t, \sigma_t^2 \mathbf{I}\right),$$

$$p_{\theta}^{(t)}(\mathbf{x}_{t-1} \mid \mathbf{x}_t) = \mathcal{N}\left(\sqrt{A_{t-1}}f_{\theta}^{(t)}(\mathbf{x}_t) + \sqrt{1 - A_{t-1} - \sigma_t^2}\epsilon_{\theta}^{(t)}(\mathbf{x}_t), \sigma_t^2 \mathbf{I}\right).$$

- Для двух гауссианов с одинаковой ковариацией $\Sigma = \sigma_t^2 \mathbf{I}$ KL-дивергенция равна

$$\text{KL}(\mathcal{N}(\mu_q, \Sigma) \parallel \mathcal{N}(\mu_p, \Sigma)) = \frac{1}{2} (\mu_q - \mu_p)^{\top} \Sigma^{-1} (\mu_q - \mu_p).$$

- Подставляя наши средние и $\Sigma^{-1} = \sigma_t^{-2} \mathbf{I}$, получаем

$$\begin{aligned} D_{\text{KL}}(q_{\sigma}(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0) \parallel p_{\theta}^{(t)}(\mathbf{x}_{t-1} \mid \mathbf{x}_t)) &= \\ &= \frac{1}{2\sigma_t^2} \left\| \sqrt{A_{t-1}}(\mathbf{x}_0 - f_{\theta}^{(t)}(\mathbf{x}_t)) + \sqrt{1 - A_{t-1} - \sigma_t^2}(\epsilon_t - \epsilon_{\theta}^{(t)}(\mathbf{x}_t)) \right\|^2. \end{aligned}$$

- Используя связь

$$\mathbf{x}_0 = \frac{\mathbf{x}_t - \sqrt{1 - A_t} \epsilon_t}{\sqrt{A_t}}, \quad f_{\theta}^{(t)}(\mathbf{x}_t) = \frac{\mathbf{x}_t - \sqrt{1 - A_t} \epsilon_{\theta}^{(t)}(\mathbf{x}_t)}{\sqrt{A_t}},$$

можно показать (после упрощений), что KL-член пропорционален

$$D_{\text{KL}}(\cdot) = \frac{\gamma_t}{2\sigma_t^2} \|\epsilon_t - \epsilon_{\theta}^{(t)}(\mathbf{x}_t)\|^2 + \text{const},$$

где γ_t – некоторая известная функция A_t и σ_t^2 , которую можно заранее посчитать по расписанию шумов, как вес для шага t

- Главный вывод: оптимизация KL между истинным и параметризованным обратным переходом эквивалентна MSE между истинным шумом ϵ_t и предсказанным шумом $\epsilon_{\theta}^{(t)}(\mathbf{x}_t)$, с некоторым весом, зависящим от t

- Если собрать все шаги t и учесть ожидание по $\mathbf{x}_0 \sim q(\mathbf{x}_0)$ и $\epsilon_t \sim \mathcal{N}(0, \mathbf{I})$, то окончательная функция ошибки (с весами) имеет вид

$$\mathbb{E}_{t, \mathbf{x}_0, \epsilon_t} \left[\frac{\gamma_t}{2\sigma_t^2} \|\epsilon_t - \epsilon_{\theta}^{(t)}(\mathbf{x}_t)\|^2 \right].$$

- На практике, как и в DDPM, обычно игнорируют веса $\frac{\gamma_t}{2\sigma_t^2}$, беря более простой loss

$$L_{\text{simple}} = \mathbb{E}_{t, \mathbf{x}_0, \epsilon} [\|\epsilon - \epsilon_{\theta}^{(t)}(\mathbf{x}_t)\|^2],$$

где t выборочно сэмплируется из $\{1, \dots, T\}$

Вывод DDIM

- Итого обучение идёт точно как в DDPM: берём $\mathbf{x}_0$, сэмплируем $\mathbf{x}_t$, обновляем параметры сети по ошибке $L = \|\epsilon - \epsilon_\theta^{(t)}(\mathbf{x}_t)\|$
- А вот inference после обучения имеет вид

$$p_\theta^{(t)}(\mathbf{x}_{t-1} \mid \mathbf{x}_t) = \mathcal{N}(\mu_\theta(\mathbf{x}_t, t), \sigma_t^2 \mathbf{I}),$$

где среднее

$$\mu_\theta(\mathbf{x}_t, t) = \sqrt{A_{t-1}} f_\theta^{(t)}(\mathbf{x}_t) + \sqrt{1 - A_{t-1} - \sigma_t^2} \epsilon_\theta^{(t)}(\mathbf{x}_t)$$

- Если мы явно подставим $f_\theta^{(t)}(\mathbf{x}_t)$ через $\epsilon_\theta^{(t)}(\mathbf{x}_t)$, то получится более удобная форма:

$$\mathbf{x}_{t-1} = \sqrt{A_{t-1}} \left(\frac{\mathbf{x}_t - \sqrt{1 - A_t} \epsilon_\theta^{(t)}(\mathbf{x}_t)}{\sqrt{A_t}} \right) + \sqrt{1 - A_{t-1} - \sigma_t^2} \epsilon_\theta^{(t)}(\mathbf{x}_t) + \sigma_t \epsilon_t$$

где $\epsilon_t \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ – новый случайный шум

- Часто это записывают как

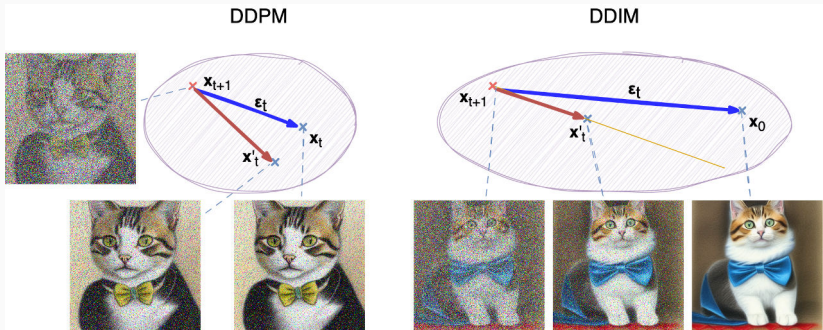
$$\mathbf{x}_{t-1} = \underbrace{\sqrt{A_{t-1}} \hat{\mathbf{x}}_0^{(t)}}_{\text{проекция предсказанного } \mathbf{x}_0} + \underbrace{\sqrt{1 - A_{t-1} - \sigma_t^2} \epsilon_\theta^{(t)}(\mathbf{x}_t)}_{\text{направление к текущему } \mathbf{x}_t} + \underbrace{\sigma_t \epsilon_t}_{\text{случайный шум}},$$

где $\hat{\mathbf{x}}_0^{(t)} = f_\theta^{(t)}(\mathbf{x}_t)$.

- Получается, что первый член возвращает нас к предсказанной версии $\mathbf{x}_0$, второй корректирует направление в пространстве к текущему $\mathbf{x}_t$, третий добавляет случайность
- В DDIM можно взять $\sigma_t = 0$ и получить детерминированный обратный процесс; это тоже очень важное преимущество

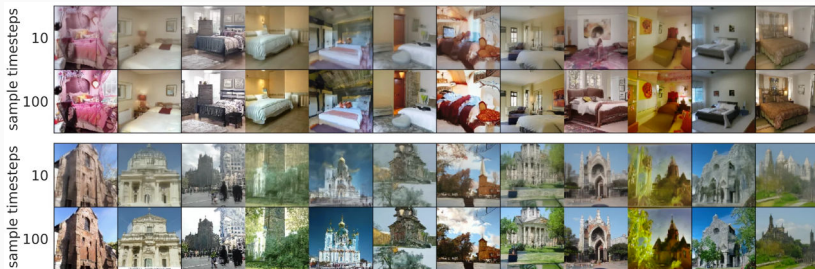
- Но главное, конечно, в том, что теперь можно перепрыгивать через некоторые шаги!
- В DDIM можно рассмотреть подмножество шагов $\tau = \{\tau_1, \dots, \tau_S\} \subset \{1, \dots, T\}$, по которым мы действительно делаем обратные переходы, и делать всё то же, но по τ
- Интуитивно говоря, мы сохраняем маргинал $q(\mathbf{x}_t \mid \mathbf{x}_0)$ для всех t , но обратный процесс будем реализовывать только на подмножестве τ , перепрыгивая через остальные шаги
- Таким образом можно тысячу шагов превратить в примерно пятьдесят

- Иллюстрация:

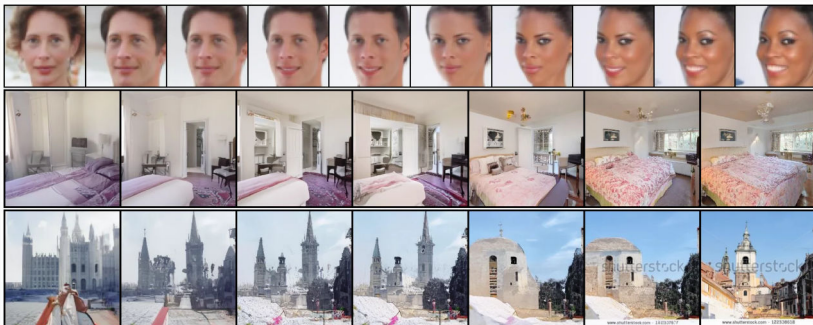


- Сложно сказать, какое приближение лучше, но теперь мы разделили зависимости ϵ_t от x_0 , и мы можем перепрыгнуть на несколько шагов сразу, двигаясь от x_t к x_{t+k} за один шаг с увеличенным ϵ !

- Это привело к 10х–100х ускорению по сравнению с DDPM без потери качества:



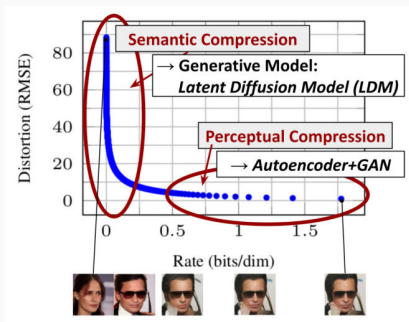
- Порождение в DDIMs тоже не обязано быть стохастическим; можно дисперсию обратной диффузии установить в ноль при порождении, $\sigma \rightarrow 0$
- И теперь латентный код в пространстве $\mathbf{x}_T$ соответствует ровно одной картинке, а DDIM — хорошая модель для латентных представлений; вот, например, интерполяции в латентном пространстве:



STABLE DIFFUSION

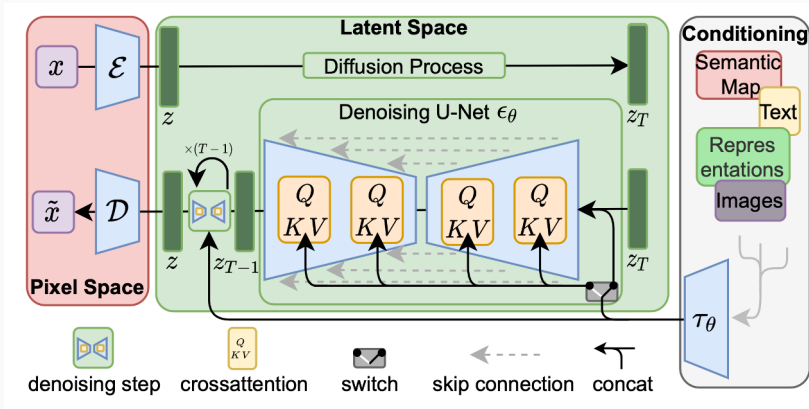
STABLE DIFFUSION

- Ещё один прорыв пришёл тогда, когда стали использовать этот диффузионный процесс в *латентном пространстве*
- Stable Diffusion (Rombach, Blattmann et al., 2022):



STABLE DIFFUSION

- Stable Diffusion (Rombach, Blattmann et al., 2022):

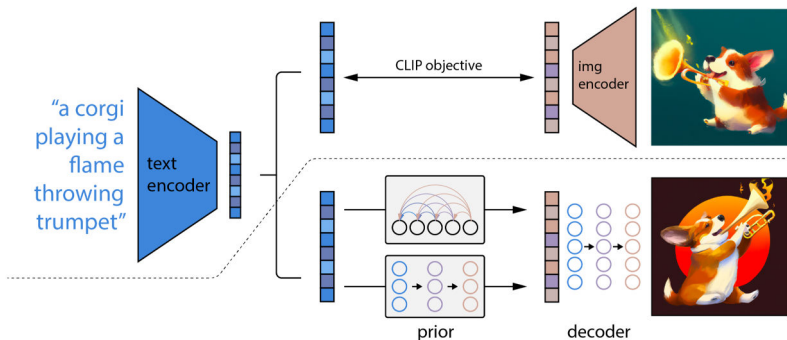


STABLE DIFFUSION

- unCLIP, он же DALL-E 2 (Ramesh et al., 2022), обучает

$$p(\mathbf{x}|y) = p(\mathbf{x}|\mathbf{c}, y) p(\mathbf{c}|y),$$

где prior model $p(\mathbf{c}|y)$ строит CLIP embedding по тексту, а в $p(\mathbf{x}|\mathbf{c}, y)$ сначала порождается «шум» (prior) для картинки, а потом по ней диффузионная модель рисует саму картинку:



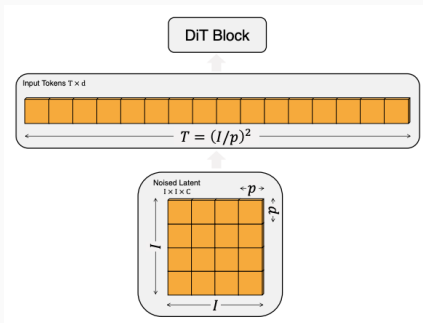




DIFFUSION TRANSFORMERS

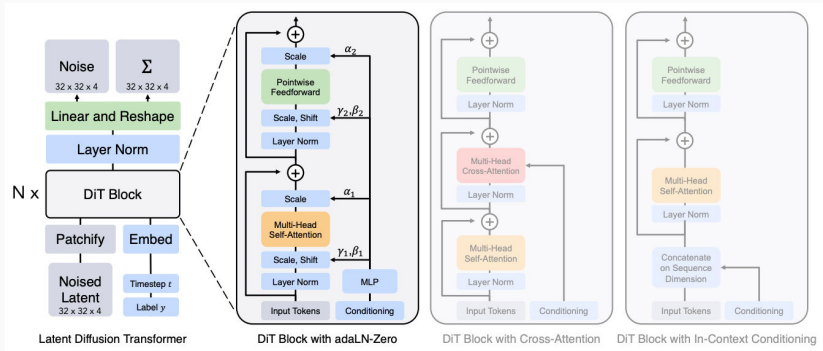
DIFFUSION TRANSFORMERS

- Уже вполне современная архитектура – DiT (Peebles, Xie, 2022)
- Это модель латентной диффузии, но вместо U-Net-подобной архитектуры для латентных кодов используют трансформер
- Вход трансформера – тензор $I \times I \times C$, "patchified" в последовательность $p \times p$ патчей длины $T = (I/p)^2$; p – гиперпараметр, квадратично влияющий на сложность



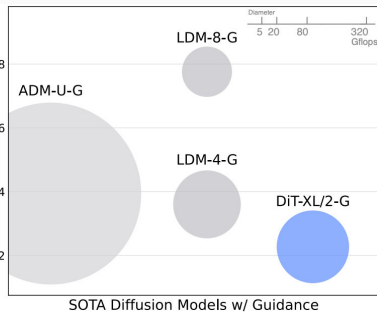
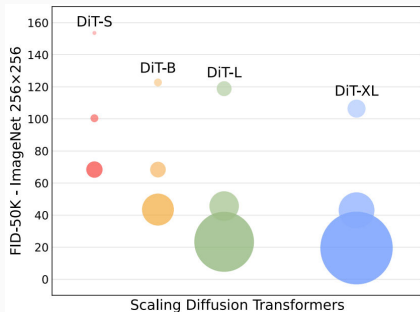
DIFFUSION TRANSFORMERS

- Затем его обрабатывают блоками DiT; интересная часть здесь – adaptive layer norm с параметрами, взятыми из условия



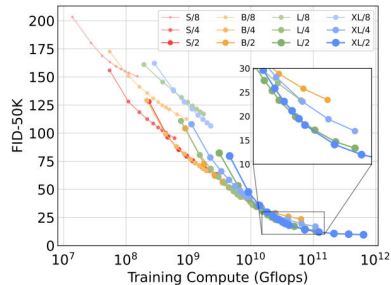
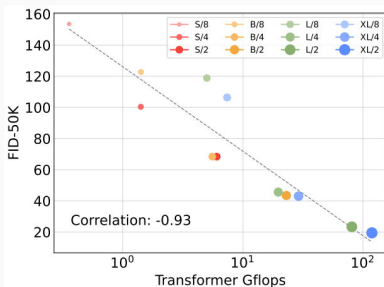
DIFFUSION TRANSFORMERS

- Получается, что DiT куда более вычислительно эффективны:



DIFFUSION TRANSFORMERS

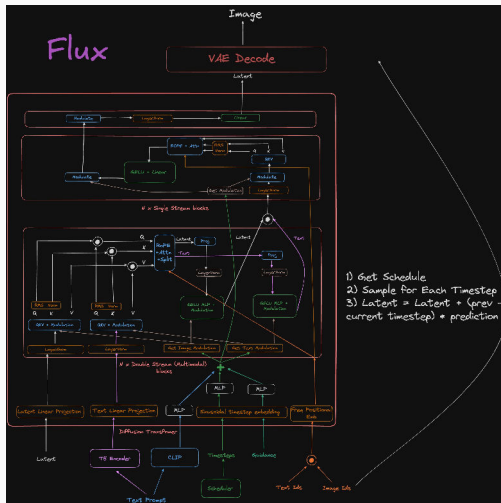
- Transformer GFlops сильно коррелированы с качеством, а большие модели DiT используют вычисления более эффективно:

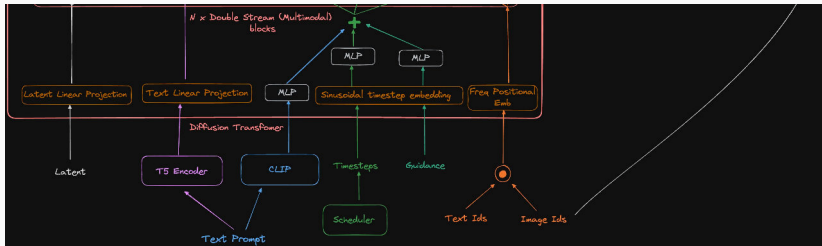


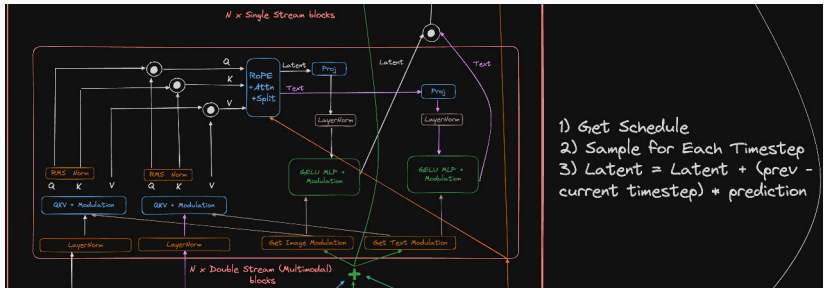
DIFFUSION TRANSFORMERS

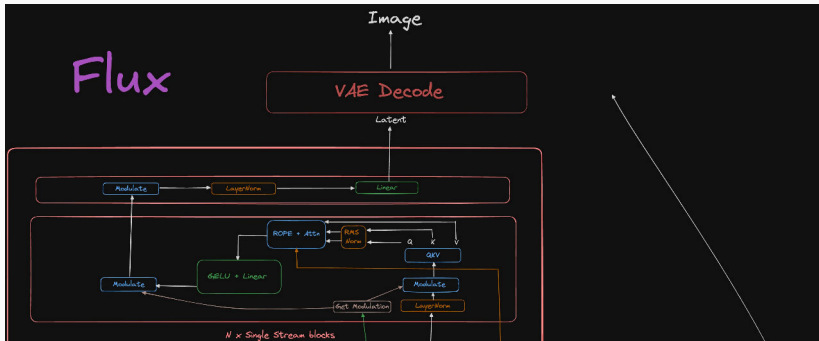
- Хорошие сэмплы:

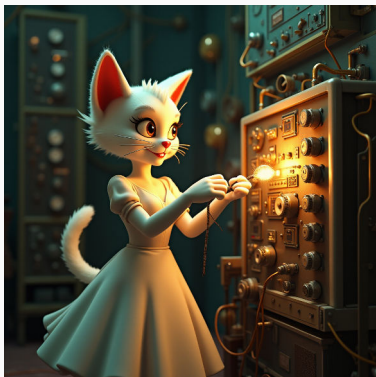
















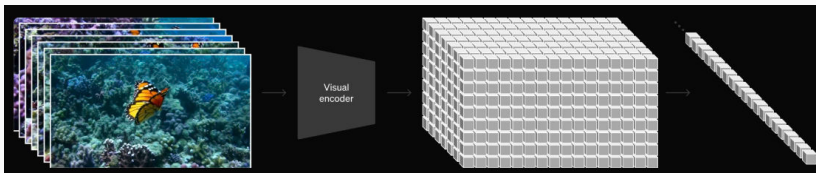




- Текущее применение DiT: OpenAI Sora (Feb 2024, и вот ещё совсем недавно)



- OpenAI всех деталей не раскрывает, но там точно есть DiT:



ОТ ДИСКРЕТНЫХ ПРОЦЕССОВ К НЕПРЕРЫВНЫМ

- Начнём с идеи DDIM: вместо шага $\mathbf{x}_t \rightarrow \mathbf{x}_{t-1}$ мы приближаем направление сразу к $\mathbf{x}_0$.
- Пусть

$$\alpha_t = 1 - \beta_t, \quad A_t = \prod_{i=1}^t \alpha_i, \quad \hat{\mathbf{x}}_0(\mathbf{x}_t, t) = \frac{\mathbf{x}_t - \sqrt{1 - A_t} \epsilon_\theta(\mathbf{x}_t, t)}{\sqrt{A_t}}.$$

- Один шаг DDIM с параметром $\eta = 0$:

$$\mathbf{x}_{t-1} = \sqrt{A_{t-1}} \hat{\mathbf{x}}_0(x_t, t) + \sqrt{1 - A_{t-1}} \hat{\epsilon}(\mathbf{x}_t, t).$$

- Интуитивно это значит, что мы движемся по «детерминированной» траектории от $\mathbf{x}_t$ к $\mathbf{x}_0$.

- Далее переходим к *непрерывному времени*: время $t \in [0, T]$, шаги $\Delta t \rightarrow 0$.
- Будем описывать прямую диффузию как стохастическое дифференциальное уравнение (SDE); обратный процесс — как *обратную* SDE.
- Зашумление — это стохастическое дифференциальное уравнение

$$d\mathbf{x} = \mathbf{f}(\mathbf{x}, t) dt + g(t) d\mathbf{w}_t, \quad \mathbf{f}(\mathbf{x}, t) = -\frac{1}{2}\beta(t) \mathbf{x}, \quad g(t) = \sqrt{\beta(t)},$$

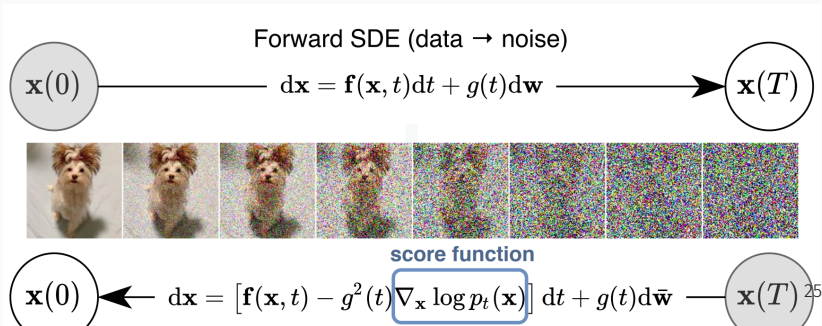
где $\mathbf{w}_t$ – винеровский процесс (броуновское движение), $\mathbf{f}$ – детерминированный дрейф, g – интенсивность шума; этот процесс постепенно превращает данные в шум (обычно стандартный гауссиан); пусть $p_t(\mathbf{x})$ – это плотность $\mathbf{x}(t)$

От DDIM к непрерывному времени

- Самый главный здесь результат — Anderson (1982): для такого процесса обратный тоже является диффузионным процессом и задаётся уравнением

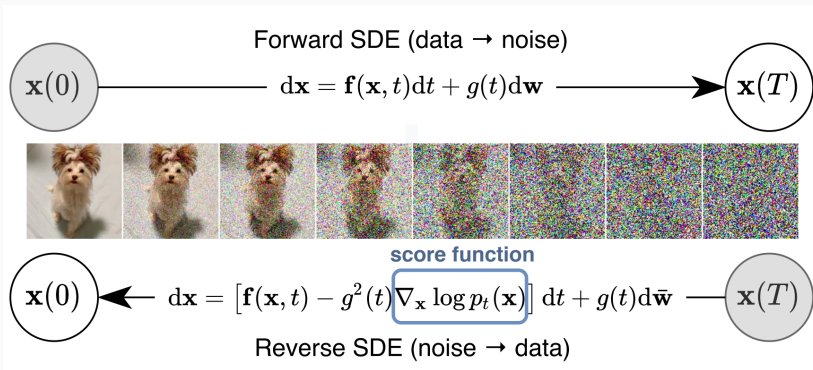
$$d\mathbf{x} = [\mathbf{f}(\mathbf{x}, t) - g(t)^2 \nabla_{\mathbf{x}} \log p_t(\mathbf{x})] dt + g(t) d\bar{\mathbf{w}},$$

где $\bar{\mathbf{w}}$ – тоже стандартный винеровский процесс, но для времени от T до 0, а dt теперь отрицательное



От DDIM к непрерывному времени

- Функцию $s_*(\mathbf{x}, t) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x})$ в этой науке называют score; если score известен для каждого t , то процесс можно обратить
- Так что наш вопрос сводится к тому, откуда взять score

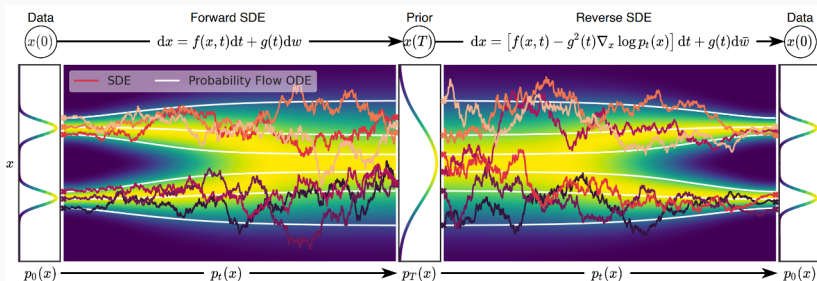


От DDIM к НЕПРЕРЫВНОМУ ВРЕМЕНИ

- А score, конечно же, можно и нужно обучать через score matching: ищем s_θ с параметрами θ , минимизируя

$$\mathbb{E}_t \left[\lambda(t) \mathbb{E}_{\mathbf{x}(0)} \left[\mathbb{E}_{\mathbf{x}(t)|\mathbf{x}(0)} \left[\left\| s_\theta(\mathbf{x}(t), t) - \nabla_{\mathbf{x}(t)} \log p_t(\mathbf{x}(t)|\mathbf{x}(0)) \right\|_2^2 \right] \right] \right],$$

где λ – веса, а $p_t(\mathbf{x}(t)|\mathbf{x}(0))$ обычно можно так или иначе подсчитать (не будем уж в подробностях)



- И теперь DDPM – это дискретизация такого процесса: наша марковская цепь $\mathbf{x}_t = \sqrt{1 - \beta_t} \mathbf{x}_{t-1} + \sqrt{\beta_t} \epsilon$ сходится в пределе к стохастическому дифференциальному уравнению

$$d\mathbf{x} = -\frac{1}{2}\beta(t)\mathbf{x} dt + \sqrt{\beta(t)} d\mathbf{w},$$

т.е. как раз к такому процессу, который мы рассматривали

- Там есть два варианта: variance preserving (VP) как выше (дисперсия сохраняется единичной) и variance exploding (VE), в котором она всегда растёт, но это нам сейчас не важно
- Важно, что DDIM – это просто более грубая дискретизация того же процесса!

- И даже более того – удивительно, но для любого диффузионного процесса есть и *детерминированный* процесс с теми же маргиналами $p_t(\mathbf{x})$! Для $s(\mathbf{x}, t) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x})$ это

$$d\mathbf{x} = [\mathbf{f}(\mathbf{x}, t) - \frac{1}{2}g(t)^2 s(\mathbf{x}, t)] dt$$

- Например, для DDPM как выше это будет

$$d\mathbf{x} = -\frac{1}{2}\beta(t)\mathbf{x} - \frac{1}{2}\beta(t)s(\mathbf{x}, t)$$

- Интуитивно говоря, в SDE вероятностная масса течёт и размазывается, а в таком процессе течёт без шума, но так, что плотность p_t по времени остаётся той же

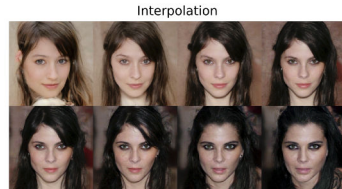
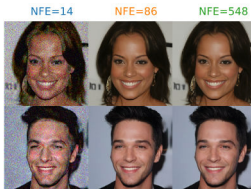
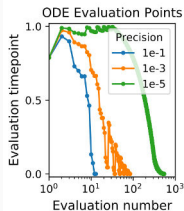
- И если теперь заменить s_* на s_θ , мы получим практическую формулу для порождения:

$$\frac{d\mathbf{x}}{dt} \approx -\frac{1}{2}\beta(t)\left(\mathbf{x} + s_\theta(\mathbf{x}, t)\right).$$

- Это и есть непрерывная версия DDIM: движение к $\mathbf{x}_0$ задаётся ОДУ, а не стохастическим процессом

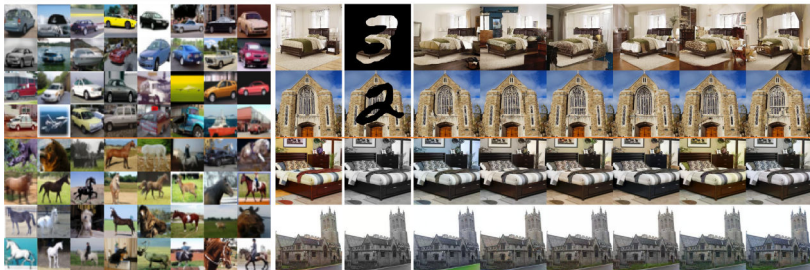
DDPM и DDIM как дискретизации SDE

- Как обучать на практике: берём $\mathbf{x}_0$ из данных, получаем $\mathbf{x}_t$ из замкнутой формы распределения, минимизируем квадрат отклонения $s_\theta(\mathbf{x}, t)$ от истинного $\log p_t(\mathbf{x}(t)|\mathbf{x}(0))$
- Благодаря тому, что процесс можно считать детерминированным, появляются полезные эффекты:
 - удобные латентные коды и интерполяции;
 - пропуски шагов, как в DDIM: можно брать разреженную сетку по t и ускорять процесс;
 - траектория к $\mathbf{x}_0$ «более прямая», чем в стохастической обратной SDE.



DDPM и DDIM как дискретизации SDE

- А ещё, кстати, всё это работает и для управляемого порождения: если нам известно $p_t(\mathbf{y}|\mathbf{x}(t))$ для некоторого условия $\mathbf{y}$ (например, метки класса или частично заданного изображения), то мы можем так же сэмплировать из $p_0(\mathbf{x}(0)|\mathbf{y})$

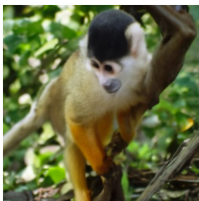


Что дальше?

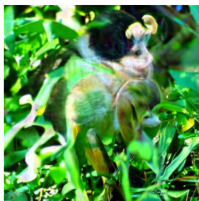
- Что произошло после DiT в диффузионных моделях?
- DPM-Solver++ (Lu et al., 2022): как дальше ускорить DDIM?
- Как мы видели выше, DDIM – это шаг первого порядка для интегрирования некоторого ОДУ
- Но ведь есть большая наука о том, как решать ОДУ лучше!
- Давайте выберем какую-нибудь многошаговую схему второго или третьего порядка, она будет понижать локальную ошибку, и в итоге можно будет шагать ещё шире, и шагов понадобится ещё меньше

ПОСЛЕДНИЕ НОВОСТИ ДИФфуЗИОННЫХ МОДЕЛЕЙ

- Например, DPM-Solver++ действительно стабильнее и лучше работает, чем другие варианты (о которых не будем):



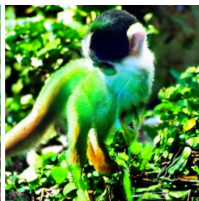
DDIM (order = 1)
(Song et al., 2021a)



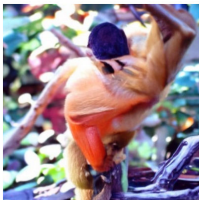
PNDM (order = 2)
(Liu et al., 2022b)



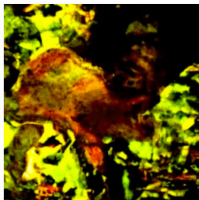
DEIS-1 (order = 2)
(Zhang & Chen, 2022)



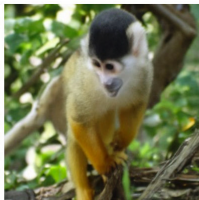
DEIS-2 (order = 3)
(Zhang & Chen, 2022)



DPM-Solver (order = 2)
(Lu et al., 2022)



DPM-Solver-3 (order = 3)
(Lu et al., 2022)



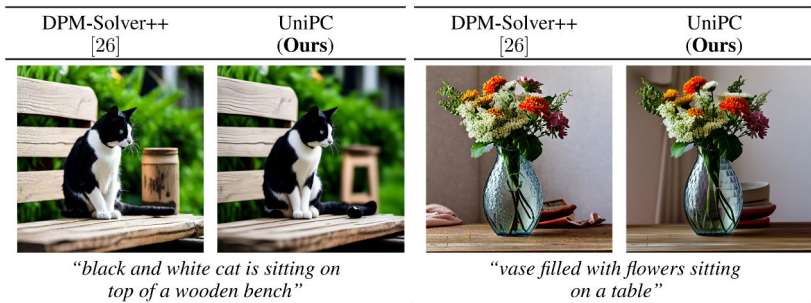
†DDIM (thresholding)
(Saharia et al., 2022b)



DPM-Solver++ (order = 2)
(ours)

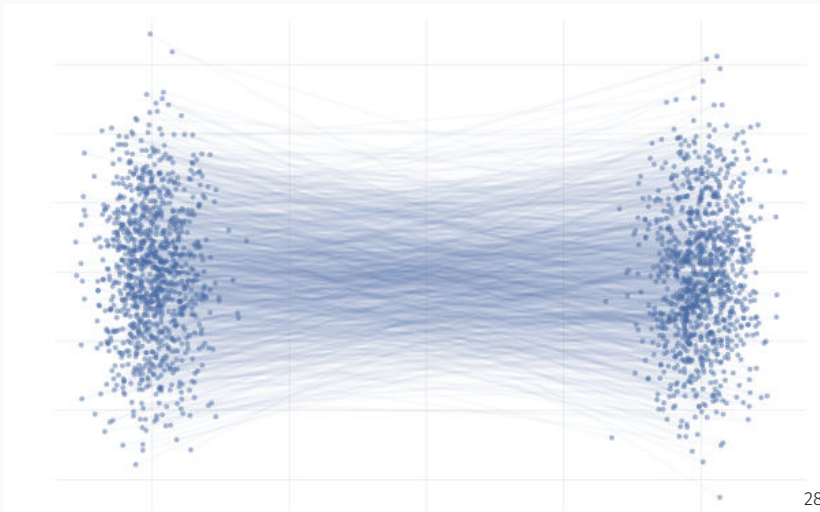
ПОСЛЕДНИЕ НОВОСТИ ДИФфуЗИОННЫХ МОДЕЛЕЙ

- UniPC (Zhao et al., 2023): метод типа «предиктор-корректор»
- Выводят универсальный корректирующий шаг (UniC), который можно добавить после любого базового предиктора, чтобы увеличить порядок точности без дополнительных вызовов модели, и выводят сопряжённый UniP-предиктор



Последние новости диффузионных моделей

- А самое интересное дальнейшее развитие событий – это, конечно, flow matching:



Последние новости диффузионных моделей

- На нём работают многие современные порождающие модели:

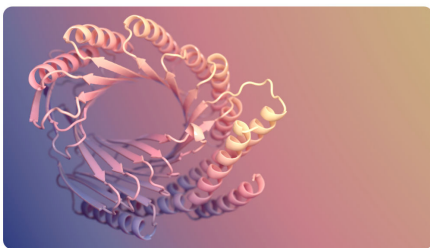


Figure 1: Protein generated by RFDiffusion (Watson et al., 2023).



Figure 2: Image from DALL-E 3 (Betker et al., 2023).

- Но это уже совсем другая история...

СПАСИБО!

Спасибо за внимание!

