

Сергей Николенко

СПбГУ — Санкт-Петербург

17 ноября 2025 г.

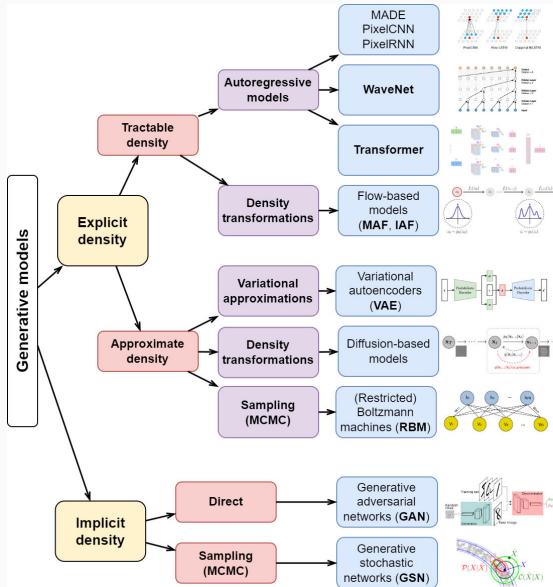
Random facts:



- 17 ноября — Международный день студентов (в память об аресте >1200 пражских студентов нацистами 17 ноября 1939 года); кроме того, 17 ноября — Международный день недоношенных детей, а в России — День участковых уполномоченных полиции
- 17 ноября 1875 г. Елена Блаватская и Генри Стил Олкотт основали в Нью-Йорке Теософское общество, в которое входили Томас Эдисон, отец Джавахарлала Неру Мотилал, Уильям Батлер Йейтс и другие; общество пришло в упадок после смерти Блаватской в 1891 г., но созданный в 1907 г. российский филиал процветал до 1918-го
- 17 ноября 1959 г. южноафриканская компания De Beers начала производство искусственных алмазов
- 17 ноября 1853 г. произошло первое в истории сражение пароходофрегат: русский «Владимир» победил турецкий «Перваз-Бахри»
- 17 ноября 2018 г. американский миссионер Джон Чау погиб от рук аборигенов Андаманских островов; Чау считал их последним оплотом Сатаны на Земле, но вскоре после высадки был застрелен из лука, посмертно получив за это премию Дарвина
- 17 ноября 2019 г. в Китае выявили первого инфицированного COVID-19

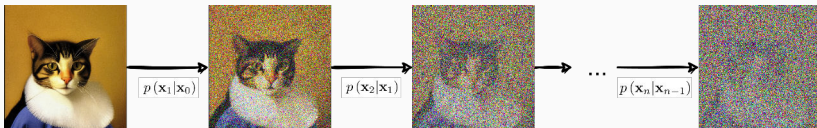
Порождающие модели в глубоком обучении

ГЛУБОКИЕ ПОРОЖДАЮЩИЕ МОДЕЛИ



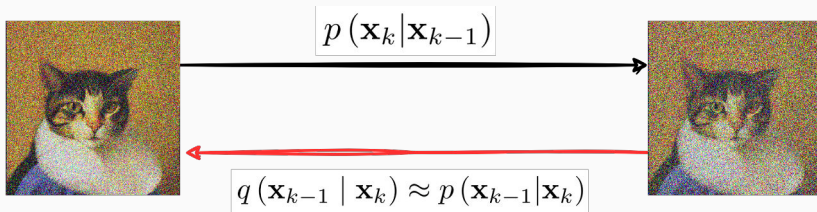
ДИФФУЗИОННЫЕ МОДЕЛИ

- (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:

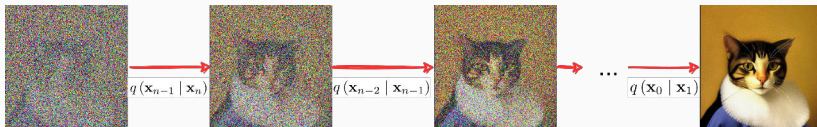


ДИФфуЗИОННЫЕ МОДЕЛИ

- А потом попробуем обратить каждый шаг:



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



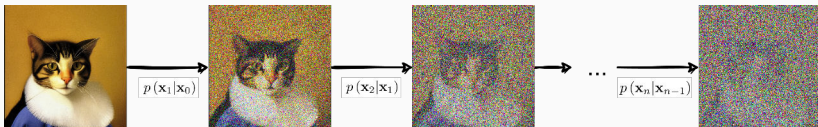
ДИФфуЗИОННЫЕ МОДЕЛИ

- 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

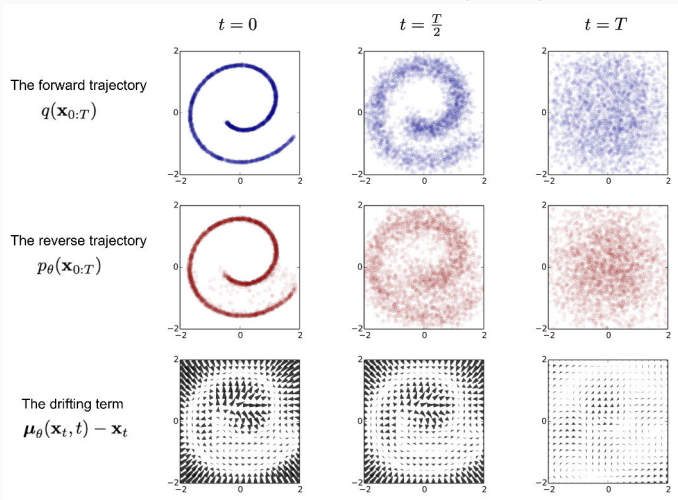
$$\begin{aligned}\mathbf{x}_t &= \sqrt{\alpha_t} \mathbf{x}_{t-1} + \sqrt{1 - \alpha_t} \epsilon = \sqrt{\alpha_t \alpha_{t-1}} \mathbf{x}_{t-2} + \sqrt{1 - \alpha_t \alpha_{t-1}} \epsilon = \dots \\ \dots &= \sqrt{A_t} \mathbf{x}_0 + \sqrt{1 - A_t} \epsilon, \text{ то есть } q(\mathbf{x}_t | \mathbf{x}_0)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

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

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

ДИФFUЗИОННЫЕ 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

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

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

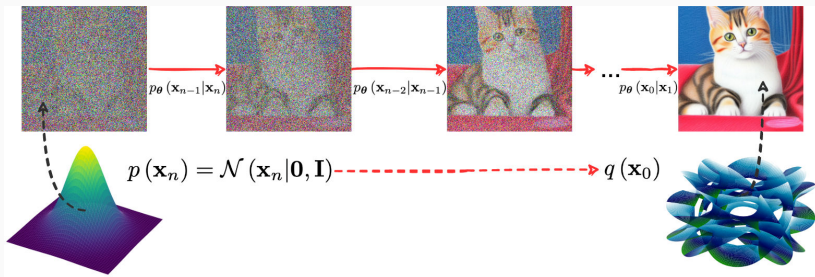
$$\tilde{\mu}(\mathbf{x}_t, \mathbf{x}_0) = \frac{\sqrt{\alpha_t}(1 - A_{t-1})}{1 - A_t} \mathbf{x}_t + \frac{\sqrt{A_{t-1}}\beta_t}{1 - A_t} \mathbf{x}_0,$$

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

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

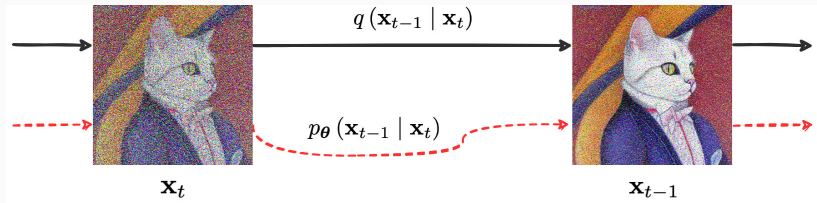
ДИФфуЗИОННЫЕ МОДЕЛИ

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



ДИФфуЗИОННЫЕ МОДЕЛИ

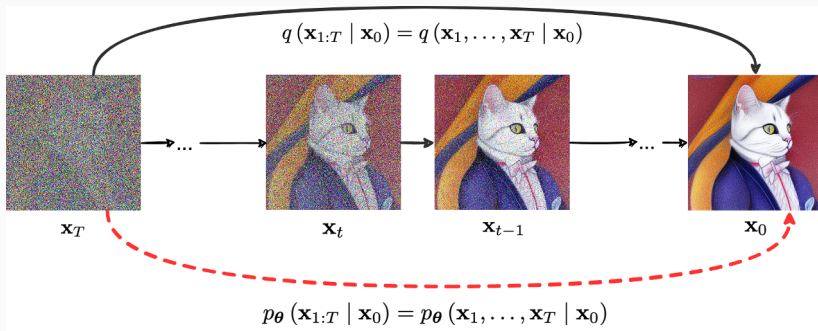
- Что делать? Как всегда в байесовском выводе, приближать!
- На каждом шаге модель должна стать хорошим приближением к $q(\mathbf{x}_t | \mathbf{x}_{t-1})$, без условия на неизвестное $\mathbf{x}_0$:



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 $q(\mathbf{z})$:

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

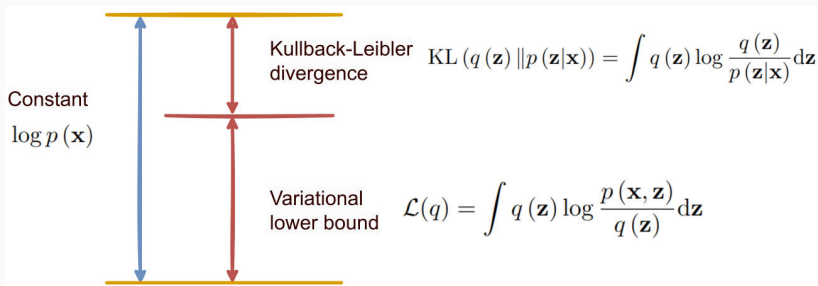
а затем сделаем стандартные преобразования:

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

- 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:

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

- 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,

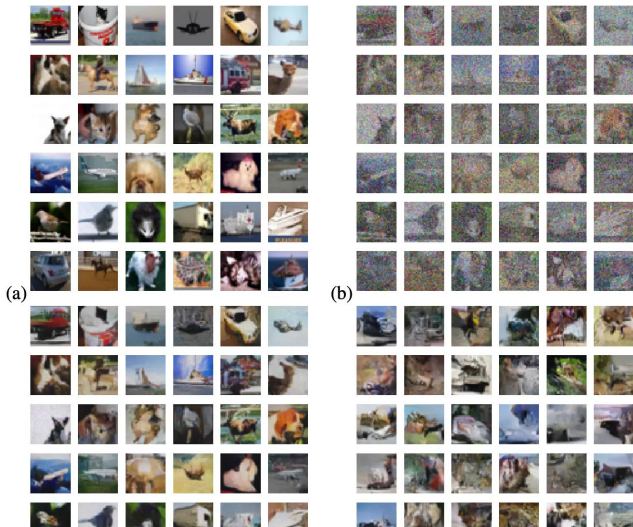
$$\mu_{\theta}(\mathbf{x}_t, t) \approx \tilde{\mu}_t(\mathbf{x}_t, \mathbf{x}_0) = \frac{1}{\sqrt{A_t}} \left(\mathbf{x}_t - \frac{1 - \alpha_t}{\sqrt{1 - A_t}} \epsilon_t \right),$$

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

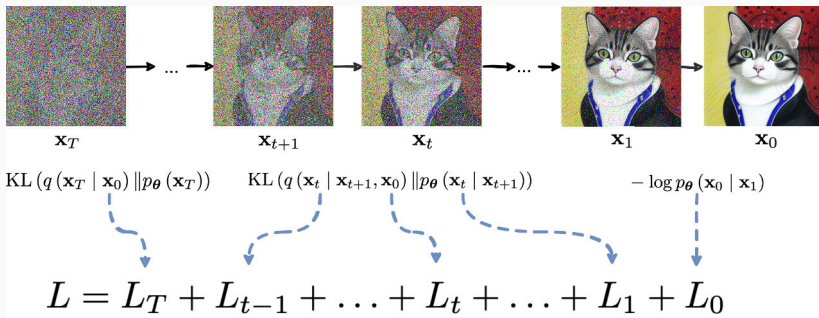
$$p_{\theta}(\mathbf{x}_{t-1} | \mathbf{x}_t) = N \left(\mathbf{x}_{t-1} \left| \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{x}_t - \frac{1 - \alpha_t}{\sqrt{1 - A_t}} \epsilon_{\theta}(\mathbf{x}_t, t) \right), \Sigma_{\theta}(\mathbf{x}_t, t) \right) \right).$$

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:



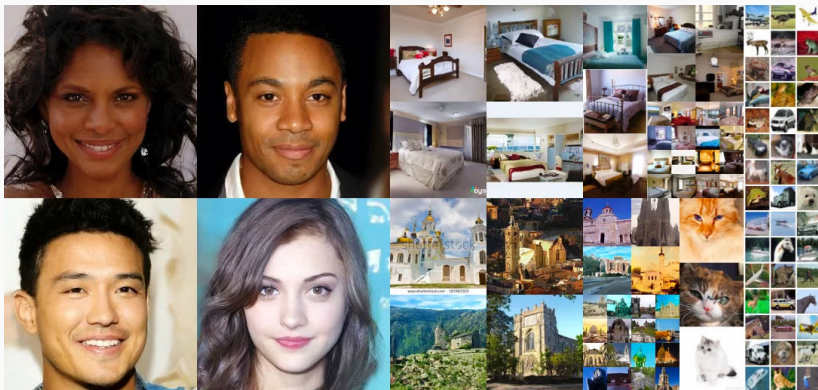
- Но дальше DDPM делает следующие замечания:
 - дисперсии прямой диффузии β_t — постоянные гиперпараметры, их не обучают, так что все q не обучаются; т.к. $p_\theta(\mathbf{x}_T)$ — фиксированное распределение, из которого мы будем сэмплировать, L_T является константой;
 - для промежуточных шагов дисперсии 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 $x_{0,i}$

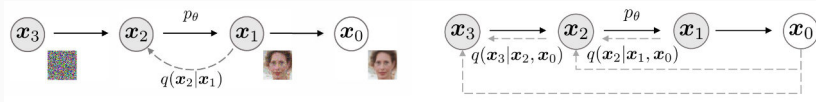
- Т.е. можно подставить другую модель на последнем шаге и использовать $\mu_\theta(\mathbf{x}_1)$ без шума во время сэмплирования

- 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:



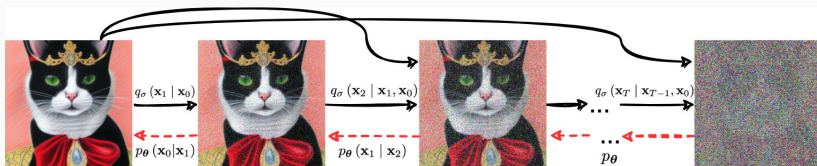
- Для DDIM мы определим диффузионный процесс в терминах его обратной формы:

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

- Теперь можно выразить распределения прямой диффузии по теореме Байеса:

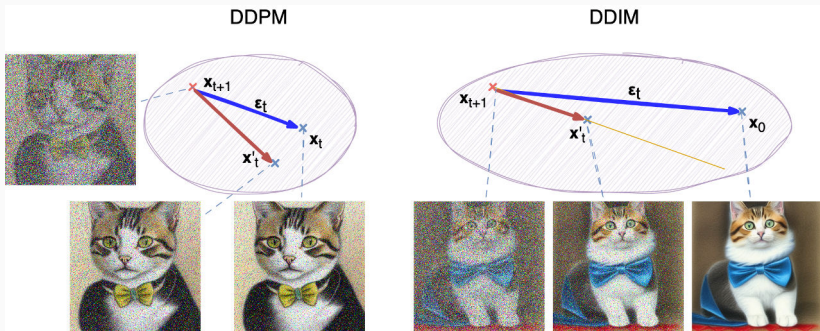
$$q_{\sigma}(\mathbf{x}_t \mid \mathbf{x}_{t-1}, \mathbf{x}_0) = \frac{q_{\sigma}(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0) q_{\sigma}(\mathbf{x}_t \mid \mathbf{x}_0)}{q_{\sigma}(\mathbf{x}_{t-1} \mid \mathbf{x}_0)}.$$

- И теперь можно показать, что у полученного процесса ровно те же маргиналы, так что обратную диффузию можно обучать с той же функцией ошибки, и она будет марковской цепью:



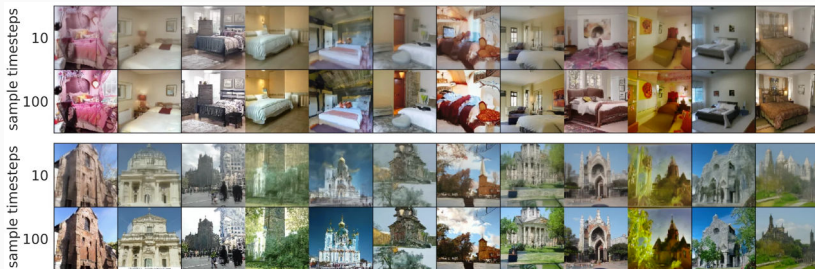
- Пока это не очень помогает: мы расширили класс распределений для прямой диффузии, но сэмплирование новой картинки всё равно должно пройти через все шаги обратной диффузии
- Но теперь, вместо того чтобы приближать случайный шум ϵ_t , который ведёт от $\mathbf{x}_t$ к $\mathbf{x}_{t+1}$, можно приближать случайный шум ϵ_t , который приведёт сразу от $\mathbf{x}_0$ к $\mathbf{x}_{t+1}$!
- И когда мы идём в обратном направлении, мы приближаем направление не к $\mathbf{x}_t$, а сразу к $\mathbf{x}_0$, и двигаться в эту сторону можно быстрее

- Иллюстрация:

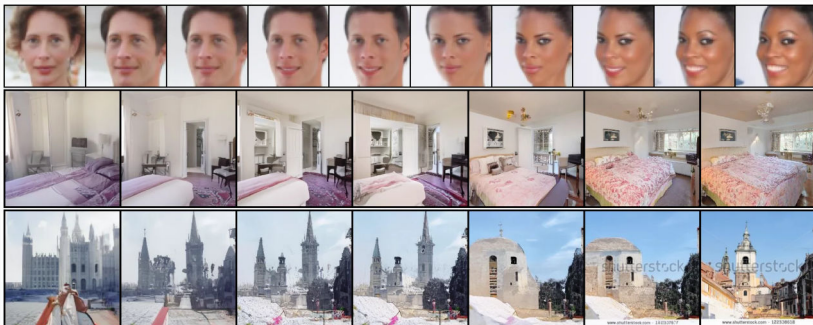


- Сложно сказать, какое приближение лучше, но теперь мы разделили зависимости ϵ_t от x_0 , и мы можем перепрыгнуть на несколько шагов сразу, двигаясь от x_t к x_{t+k} за один шаг с увеличенным ϵ !

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



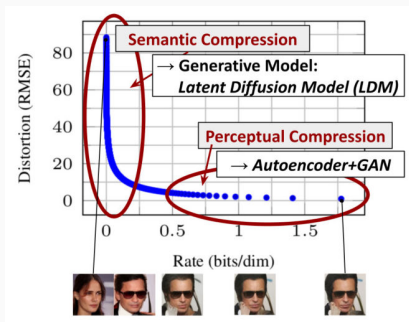
- Порождение в DDIMs тоже не обязано быть стохастическим; можно дисперсию обратной диффузии установить в ноль при порождении
- И теперь латентный код в пространстве $\mathbf{x}_T$ соответствует ровно одной картинке, а DDIM — хорошая модель для латентных представлений; вот, например, интерполяции в латентном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



STABLE DIFFUSION

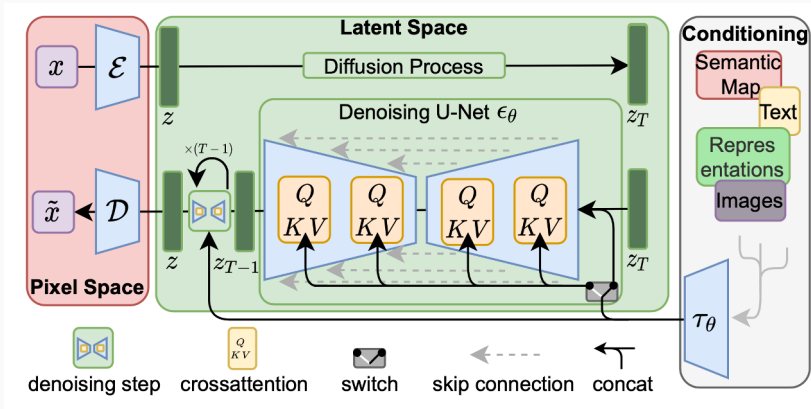
STABLE DIFFUSION

- 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):

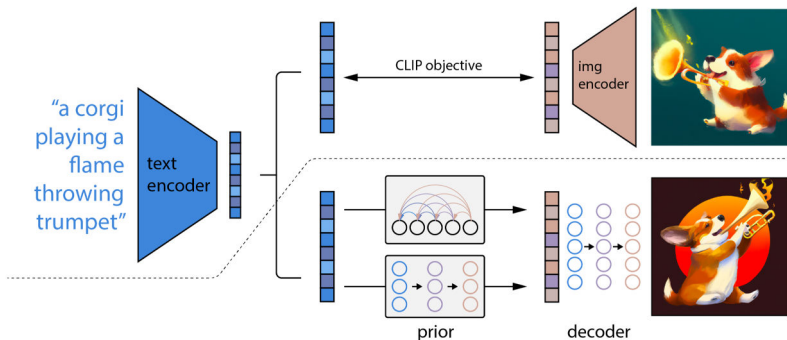


STABLE DIFFUSION

- 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:



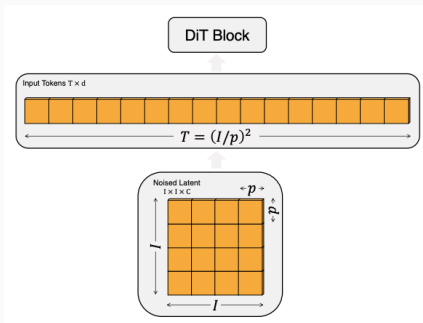




DIFFUSION TRANSFORMERS

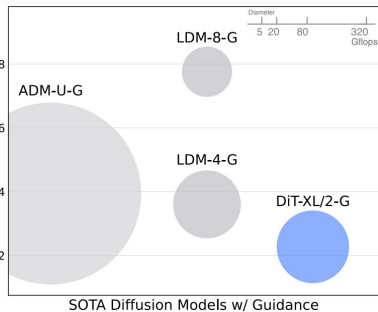
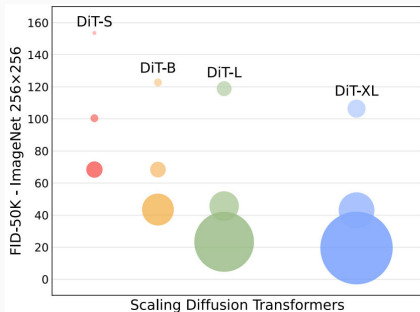
DIFFUSION TRANSFORMERS

- 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



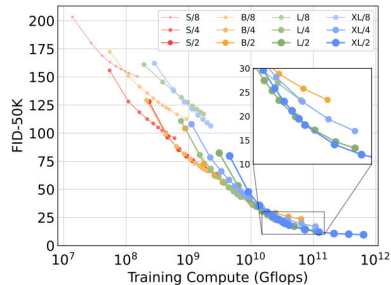
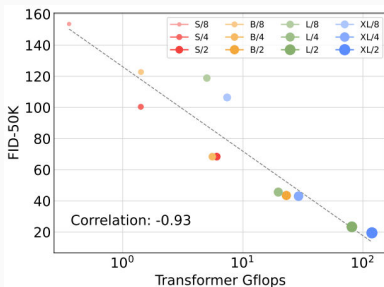
DIFFUSION TRANSFORMERS

- Получается, что DiT куда более вычислительно эффективны:



DIFFUSION TRANSFORMERS

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

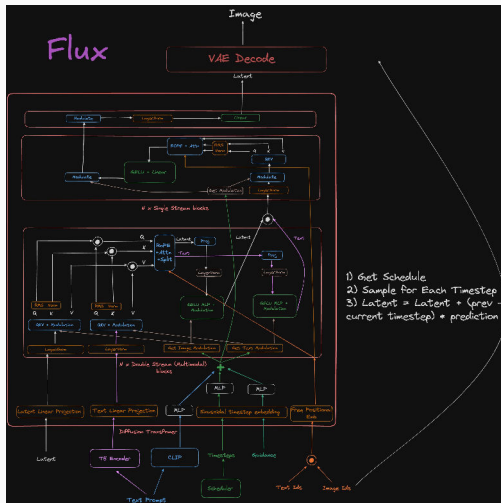


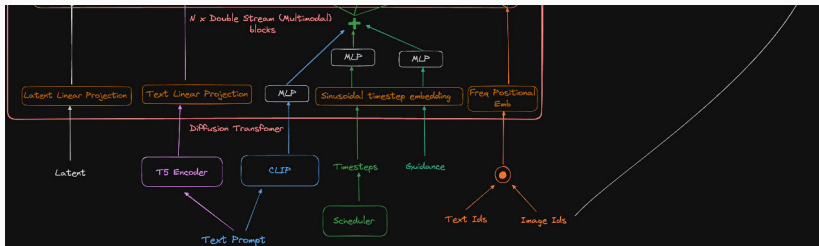
DIFFUSION TRANSFORMERS

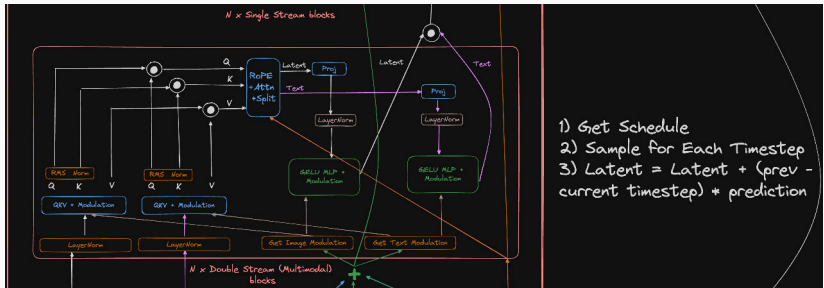
- Хорошие сэмплы:

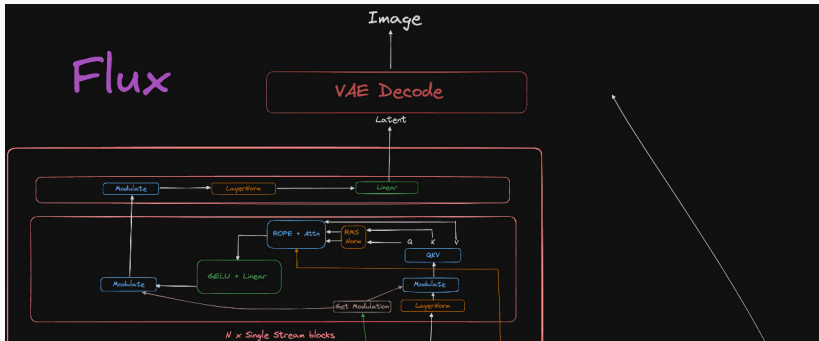


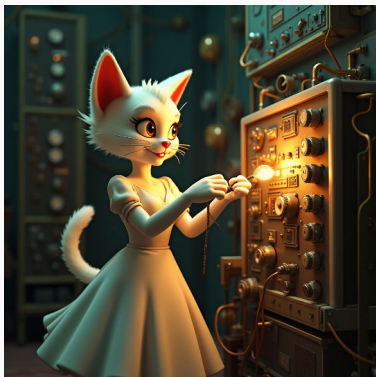
FLUX

















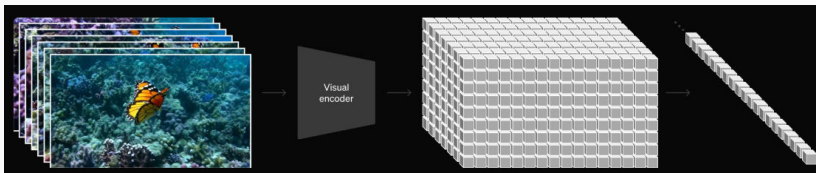


OPENAI SORA

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



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



СПАСИБО!

Спасибо за внимание!

