

ТЕОРЕМА ВИТТА О СОКРАЩЕНИИ И ЕЁ ПРИМЕНЕНИЕ

Антон А. Клячко

Более подробное (и более правильное) изложение того, что я рассказывал (или хотел рассказать) на второй половине последней лекции весной 2025 года.

Теорема Витта о сокращении. Если невырожденные квадратичные формы

$$q(x_1, \dots, x_n) + r(y_1, \dots, y_m) \quad \text{и} \quad q(x_1, \dots, x_n) + s(y_1, \dots, y_m)$$

от $n + m$ переменных над полем характеристики не два эквивалентны (то есть получаются друг из друга линейными невырожденными заменами переменных), то квадратичные формы $r(y_1, \dots, y_m)$ и $s(y_1, \dots, y_m)$ эквивалентны.

Лемма. Если на векторном пространстве $V \ni u, v$ над полем характеристики не два имеется квадратичная форма f , причём $f(u) = f(v) \neq 0$ (то есть векторы u и v имеют одинаковую ненулевую «длину») то найдётся линейный оператор A , сохраняющий форму f (то есть $f(Aw) = f(w)$ для любого $w \in V$) такой, что $f(u) = f(v)$.

Доказательство. Рассмотрим билинейную симметрическую функцию $(\cdot, \cdot)$, соответствующую форме f , то есть $(a, b) \stackrel{\text{онп}}{=} \frac{1}{2}(f(a+b) - f(a) - f(b))$.

- Векторы $u+v$ и $u-v$ ортогональны: действительно, $(u+v, u-v) = f(u) - f(v) = 0$.
- Один из этих векторов имеет ненулевую «длину»: действительно, $u = \frac{1}{2}((u+v) + (u-v))$, в силу ортогональности это означает, что $f(u) = (u, u) = \frac{1}{4}(u+v, u+v) + \frac{1}{4}(u-v, u-v)$, но $f(u) \neq 0$ по условию.
- Пусть $f(u+v) \neq 0$. Тогда $V = \langle u+v \rangle \oplus \langle u+v \rangle^\perp$ (где $\langle u+v \rangle^\perp \stackrel{\text{онп}}{=} \{w \in V \mid (w, u+v) = 0\}$ — ортогональное дополнение относительно f к прямой $\langle u+v \rangle$). Возьмём тогда в качестве A отражение относительно прямой $\langle u+v \rangle$, то есть $A(\lambda(u+v) + w) \stackrel{\text{онп}}{=} \lambda(u+v) - w$ (где $w \in \langle u+v \rangle^\perp$). Этот оператор очевидно сохраняет форму f (проверьте!), а $Au = \frac{1}{2}(A(u+v) + A(u-v)) = \frac{1}{2}(u+v+v-u) = v$.
- Если же $f(u-v) \neq 0$, то $V = \langle u-v \rangle \oplus \langle u-v \rangle^\perp$. Тогда в качестве A можно взять отражение относительно подпространства $\langle u-v \rangle^\perp$, то есть $A(\lambda(u-v) + w) \stackrel{\text{онп}}{=} \lambda(u-v) + w$ (где $w \in \langle u-v \rangle^\perp$). Этот оператор, разумеется тоже сохраняет форму f (проверьте!), и опять $Au = \frac{1}{2}(A(u-v) + A(u+v)) = \frac{1}{2}(v-u+u+v) = v$.

Это завершает доказательство леммы.

Доказательство теоремы Витта. Поскольку квадратичная форма q над полем характеристики не два приводится к сумме квадратов с коэффициентами, дело сразу сводится к случаю $n = 1$ (понимаете?).

Задачу теперь можно переформулировать в бескоординатной форме так:

на векторном пространстве над полем характеристики не два имеется невырожденная квадратичная форма f и $f(e_0) = f(e'_0) \neq 0$ для некоторых векторов e_0 и e'_0 . Надо показать, что ограничения формы f на $\langle e_0 \rangle^\perp$ и на $\langle e'_0 \rangle^\perp$ эквивалентны, то есть существует линейное биективное отображение $B: \langle e_0 \rangle^\perp \rightarrow \langle e'_0 \rangle^\perp$, сохраняющее форму f (в естественном смысле: $f(Bw) = f(w)$ для всех $w \in \langle e_0 \rangle^\perp$).

Понимаете, почему это переформулировка (в случае $n = 1$)? (Просто исходные две квадратичные формы $ax_0^2 + r(y_1, \dots, y_m)$ и $ax_0^2 + s(y_1, \dots, y_m)$ можно считать одной формой f , записанной в разных базисах $\{e_0, \dots, e_m\}$ и $\{e'_0, \dots, e'_m\}$, поскольку дано, что формы эквивалентны; тогда форма r есть ограничение формы f на $\langle e_0 \rangle^\perp$, а s — это ограничение формы f на $\langle e'_0 \rangle^\perp$).

По лемме существует невырожденный оператор A , сохраняющий форму f и переводящий e_0 в e'_0 (и, следовательно, переводящий $\langle e_0 \rangle^\perp$ в $\langle e'_0 \rangle^\perp$). Теперь в качестве B достаточно взять ограничение оператора A на $\langle e_0 \rangle^\perp$. Это завершает доказательство.

Упражнение 1. Можно ли убрать условие невырожденности форм из теоремы Витта? (Я думаю, можно.)

Упражнение 2. Верна ли аналогичная теорема о сокращении для линейных операторов: следует ли из подобия блочно-диагональных матриц $\begin{pmatrix} A & 0 \\ 0 & B \end{pmatrix}$ и $\begin{pmatrix} A & 0 \\ 0 & C \end{pmatrix}$ подобие матриц B и C ?

Применять теорему Витта можно следующим образом. Пусть мы хотим выяснить эквивалентны ли квадратичные формы f и g над полем характеристики не два.

- для начала перепишем форму g в виде $g = \sum b_i x_i^2$ (этого всегда можно добиться заменой переменных, как нам известно).

- Если элемент $b_1 \neq 0$ не является значением квадратичной формы f (то есть $f(c_1, c_2, \dots) \neq b_1$ ни для каких c_i из поля), то формы очевидно не эквивалентны (поскольку $g(1, 0, 0, \dots) = b_1$, а у эквивалентных форм множество значений одинаковое).
- Если же элемент b_1 является значением формы f , то есть $f(c_1, c_2, \dots) = b_1$ то мы выберем вектор $(c_1, c_2, \dots)^T$ в качестве первого базисного, остальные базисные векторы выберем в его ортогональном дополнении и приведём тем самым форму f к виду $f = b_1 x_1^2 + \tilde{f}(x_2, x_3, \dots)$.
- Теорема Витта теперь сводит наш вопрос к вопросу об эквивалентности форм $\tilde{f}$ и $b_2 x_2^2 + b_3 x_3^2 + \dots$ от меньшего числа переменных.

Таким образом, задача об эквивалентности квадратичных форм сводится к задаче о множестве значений квадратичной формы. Этот вопрос может оказаться непростым (например, для поля рациональных чисел), но для конечного поля всё оказывается удивительно просто. (Когда-если вы доучитесь до второго курса, вы узнаете все детали о строении конечных полей; но в нашей задаче никакие детали не нужны.)

Теорема о значениях квадратичной формы над конечным полем. *Над конечным полем характеристики не два множество значений любой квадратичной формы ранга большего единицы — это всё поле.*

Доказательство. Сперва надо понять, как устроено множество квадратов элементов конечного поля F характеристики не два.

- 1) Ровно половина ненулевых элементов поля F является квадратами.
- 2) Если элемент $\varepsilon \in F$ не является квадратом, то $\{\text{неквадраты}\} = \varepsilon \cdot \{\text{квадраты}\} \setminus \{0\}$.
- 3) Произведение двух неквадратов — это квадрат (и, очевидно, произведение двух квадратов — квадрат, а произведение ненулевого квадрата на неквадрат есть неквадрат).

Таким образом, в этом смысле, конечные поля характеристики не два, можно сказать, похожи на вещественные числа (разница только в том, что минус единица может и являться квадратом, например, в $\mathbb{Z}_5$ это так, как вы понимаете). Объяснение у этих фактов 1), 2), 3) простое:

- 1) возведение в квадрат представляет собой сюръективное отображение $F \setminus \{0\} \rightarrow \{\text{квадраты}\} \setminus \{0\}$, при котором у каждого элемента есть ровно два прообраза ($a^2 = b^2$ тогда и только тогда, когда $a = \pm b$);
- 2) включение $\supseteq$ в 2) очевидно (поскольку частное двух квадратов является квадратом); а обратное включение теперь следует из совпадения мощностей (по пункту 1));
- 3) это немедленно вытекает из 2): $(\varepsilon a^2) \cdot (\varepsilon b^2) = (\varepsilon ab)^2$;

А как ведут себя квадраты по отношению к сложению? Здесь аналогия с вещественными числами терпит полный крах:

любой элемент конечного поля характеристики не два раскладывается в сумму двух квадратов. (*)

Объяснение простое: для каждого $c \in F$ мощность множества $\{c - a^2 \mid a \in F\}$ очевидно равна мощности множества всех квадратов, что больше мощности множества всех неквадратов по пункту 1) (и поскольку ноль является квадратом). Стало быть, $c - a^2 = b^2$ при подходящих a и b , что и требовалось.

Приступим теперь собственно к доказательству теоремы о значениях квадратичной формы над конечным полем. Ясно, что её достаточно доказать для форм ранга два. (Понимаете?) Форма ранга два заменой переменных приводится, как известно, к виду $ax^2 + by^2$ (где $a \neq 0 \neq b$) и далее, согласно 2), (заменами $x \mapsto \alpha x$, $y \mapsto \beta y$ и, возможно, $x \leftrightarrow y$) к одному из трёх видов

$$f = x^2 + y^2, \quad f = \varepsilon x^2 + \varepsilon y^2 \quad \text{или} \quad f = x^2 + \varepsilon y^2, \quad \text{где } \varepsilon \text{ — произвольный фиксированный неквадрат.}$$

В первых двух случаях утверждение теоремы это просто утверждение (*), а в третьем случае

$$\{f(a, 0) \mid a \in F\} = \{\text{квадраты}\} \quad \text{и} \quad \{f(0, a) \mid a \in F\} = \{\text{неквадраты}\} \cup \{0\} \quad (\text{по пункту 2)).}$$

Это завершает доказательство теоремы о значениях квадратичной формы над конечным полем.

Вместе с теоремой Витта это даёт следующий неожиданный факт.

Классификации квадратичных форм над конечным полем. *Невырожденные квадратичные формы над конечным полем характеристики не два эквивалентны тогда и только тогда, когда определители их матриц одновременно являются или не являются квадратами.*

(В одну сторону это вытекает из того, что при замене переменных определитель квадратичной формы умножается, как известно, на квадрат определителя матрицы перехода; а в нетривиальную сторону утверждение вытекает из рассуждений выше.)

Упражнение 3. Запишите подробно доказательство этой классификации (и убедитесь, что теорема Витта здесь реально не используется вовсе :)

Упражнение 4. Как классифицировать произвольные (вырожденные) квадратичные формы над конечным полем характеристики не два?