

Общематематические понятия

0.1. Математический язык и логическая символика.

1. **ВЫСКАЗЫВАНИЯ.** В математике мы часто встречаемся с высказываниями. Под высказыванием понимают повествовательное предложение, в котором что-то утверждается о каком-то объекте, и при этом можно судить, верно это утверждение или нет. Для высказывания вместо слова «верно» говорят также «истинно», «справедливо», «выполнено», «имеет место» и т. п., для не верного высказывания используют также термин «ложно».

Есть несколько способов получения новых высказываний на основе имеющихся. Ясно, что для каждого из них мы должны указать, в каких случаях новое высказывание будет истинным в зависимости от истинности имевшихся высказываний.

Можно получить новое высказывание, имея в распоряжении всего одно высказывание, а именно можно взять *отрицание* данного высказывания, т. е. если высказывание обозначить одной буквой, например P , то новое высказывание будет (не P). Символически высказывание (не P) или (неверно P) обозначают через $\neg P$.

Рассмотрим два высказывания, обозначим одно из них буквой P , а другое — буквой Q . Путем их соединения одним из союзов «и», «или» можно получить новые высказывания: $(P \text{ и } Q)$ и $(P \text{ или } Q)$ (в последнем вместо союза «или» используют иногда союз «либо»). Высказывание $(P \text{ и } Q)$ называют *конъюнкцией высказываний P и Q* и обозначают через $P \wedge Q$ или $P \& Q$. Высказывание $(P \text{ или } Q)$ называют *дизъюнкцией высказываний P и Q* и обозначают через $P \vee Q$.

Высказывание $(P \text{ и } Q)$ истинно, если оба высказывания P , Q истинны, и ложно при других раскладах. Высказывание $(P \text{ или } Q)$ истинно, если хотя бы одно из высказываний P , Q истинно, и ложно в одном случае — когда оба высказывания P , Q ложны.

Еще один способ получения нового высказывания из P , Q — это составление условного утверждения, а именно высказывания

из P следует Q

или, иначе говоря,

если (выполнено) P , то (выполнено) Q .

Это высказывание кратко обозначают так:

$$P \Rightarrow Q,$$

и называют *импликацией*.

Если $P \Rightarrow Q$ и $Q \Rightarrow P$, то высказывания P и Q называют *равносильными* или *эквивалентными* и используют при этом обозначение $P \iff Q$. Два высказывания равносильны в том и только в том случае, если они оба либо истинны, либо ложны.

Импликацию $P \Rightarrow Q$ можно переформулировать иным способом, привлекая рассуждение «от противного»: предположим, что Q не верно, если при этом P также оказывается ложным, то считается, что верно высказывание $P \Rightarrow Q$. Иначе говоря, высказывание $P \Rightarrow Q$ равносильно высказыванию $(\text{не } Q) \Rightarrow (\text{не } P)$.

Импликация $P \Rightarrow Q$ ложна только в одном случае, когда P истинно, а Q ложно, в остальных случаях она истинна. Считать $P \Rightarrow Q$ ложной, если P истинно, а Q ложно, достаточно естественно — вряд ли мы согласимся, что из истинного утверждения, рассуждая корректно, можно получить ложное. Также естественно считать ее истинной, если P и Q оба истинны. Некоторую настороженность могут вызвать случаи, когда P ложно. Попробуем договориться так. Если Q истинно, то нам все равно, откуда оно могло бы следовать, так что импликацию $P \Rightarrow Q$ в этом случае будем считать истинной, даже если P ложно.

Осталась ситуация, когда P , Q оба ложны. Однако ввиду того, что импликация $P \Rightarrow Q$ равносильна импликации $(\text{не } Q) \Rightarrow (\text{не } P)$, в случае, когда оба P и Q ложны, их отрицания $(\text{не } P)$ и $(\text{не } Q)$ истинны. В таком случае импликацию $(\text{не } Q) \Rightarrow (\text{не } P)$ истинна, а вместе с ней истинной считаем и импликацию $P \Rightarrow Q$.

Информацию об истинности высказываний, составленных с использованием конъюнкции, дизъюнкции, импликации и эквивалентности, в зависимости от истинности их составляющих удобно представить в виде

таблицы, в которой 1 означает, что высказывание истинно, а 0 — что оно ложно:

P	Q	$P \wedge Q$	$P \vee Q$	$P \Rightarrow Q$	$P \Leftrightarrow Q$
1	1	1	1	1	1
1	0	0	1	0	0
0	1	0	1	1	0
0	0	0	0	1	1

О равносильности составных высказываний можно судить по их таблицам истинности. А именно отличительной особенностью равносильных высказываний служит совпадение их таблиц истинности. Стало быть, если оказалось, что таблицы совпали, то такие высказывания равносильны.

2. ОТРИЦАНИЕ СОСТАВНЫХ ВЫСКАЗЫВАНИЙ. Рассмотрим высказывание

не $(P$ и $Q)$.

В нем говорится о том, что не выполнены утверждения P и Q одновременно, а это указывает на то, что по крайней мере одно из них не выполнено, т. е. не $(P$ и $Q)$ равносильно высказыванию $((\text{не } P) \text{ или } (\text{не } Q))$. Символически можно это записать так:

$$\neg(P \wedge Q) \iff (\neg P) \vee (\neg Q).$$

Рассмотрим высказывание

не $(P$ или $Q)$.

В нем говорится: не верно, что выполнено P или выполнено Q , т. е. оба высказывания P и Q не верны. Таким образом, не $(P$ или $Q)$ равносильно тому, что $((\text{не } P) \text{ и } (\text{не } Q))$. Символически:

$$\neg(P \vee Q) \iff (\neg P) \wedge (\neg Q).$$

Прежде чем заняться отрицанием импликации $P \Rightarrow Q$ сформируем равносильное ей высказывание с использованием дизъюнкции и отрицания и запишем отрицание импликации как отрицание такого высказывания. Когда мы считаем высказывание «если P , то Q » истинным? В случае, когда P не выполнено, нам безразлична справедливость Q , но если P верно, то Q должно быть также верным. Это можно сказать так:

или P не верно, или Q верно, т. е. записать в виде дизъюнкции $\neg P \vee Q$. Легко убедиться в справедливости равносильности этой дизъюнкции и импликации $P \Rightarrow Q$, написав таблицу истинности:

P	Q	$\neg P \vee Q$	$P \Rightarrow Q$
1	1	1	1
1	0	0	0
0	1	1	1
0	0	1	1

Теперь легко сформировать отрицание (не верно $P \Rightarrow Q$). Заменяя $P \Rightarrow Q$ ему равносильным $\neg P \vee Q$ и взяв его отрицание, получим высказывание $P \wedge \neg Q$. Словесно его можно оформить так: неверно, что из P следует Q , означает, что P выполнено, а при этом Q — нет.

3. ВЫСКАЗЫВАНИЯ С ПЕРЕМЕННЫМИ. КВАНТОРЫ. Высказывания могут быть такими, что их истинность зависит от выбора какого-то объекта (или нескольких объектов) из заданной совокупности объектов. В таком случае высказывание включает в себя некоторую букву (или несколько букв), на место которой (которых) можно подставлять разные объекты из данной совокупности, и тогда говорят, что рассматривается высказывание с переменной (переменными). При этом в результате подстановки каждого конкретного объекта мы должны получать высказывание без переменной, т. е. каждый раз должны иметь возможность судить, истинно высказывание для конкретного объекта или ложно.

С высказываниями с переменными связан важный способ образования новых высказываний (без переменных). Пусть дано высказывание $P(x)$, истинность которого зависит от объектов, обозначенных здесь буквой x и выбираемых из некоторого множества X . Из этого высказывания с переменной можно сделать новые высказывания, сообщив, при каких обстоятельствах относительно x будем рассматривать данное высказывание.

Одно из этих обстоятельств выражается фразами «для любого», «для каждого», «для всех» и приводит к новому высказыванию вида «(для любого $x \in X$) $P(x)$ », которое читается обычно так: для любого (для каждого, для всех) x из X выполнено $P(x)$. Ясно, что словами эти фразы

писать не всегда удобно, поэтому для их обозначения используют символ $\forall$, называемый *квантором общности*. О высказывании $(\forall x \in X) P(x)$ говорят, что переменная $x \in X$ в нем *связана квантором общности*.

Второе обстоятельство выражается словами «найдется», «существует», «можно подобрать» и т. п. и приводит к высказыванию вида «(существует $x \in X$) $P(x)$ », читается так: существует x из X , для которого выполнено $P(x)$. Для обозначения слов «найдется», «существует» и т. п. используется символ $\exists$, который называют *квантором существования*. О высказывании $(\exists x \in X) P(x)$ говорят, что переменная $x \in X$ в нем *связывается квантором существования*.

Заметим, что при таком способе построения высказываний множество X , из которого берется x , всегда конкретно определено, однако иногда не указывается, если это множество ясно из контекста. Легко убедиться, что связывание переменной разными кванторами, вообще говоря, приводит к разным результатам. Например, высказывание

$$\exists x \in \mathbb{R} \quad x > 0$$

истинно, а

$$\forall x \in \mathbb{R} \quad x > 0$$

уже ложно.

Все переменные в высказывании должны быть связаны!

4. ПОРЯДОК КВАНТОРОВ. Имея предложение с двумя или более переменными, мы можем превратить его в высказывание, связав каждую переменную своим квантором. Если переменные связаны одинаковыми кванторами, то кванторные приставки могут идти в произвольном порядке, истинность высказывания от этого не зависит. Так, следующие высказывания эквивалентны:

$$\begin{aligned} \forall x \in A \forall y \in B \quad P(x, y), \\ \forall y \in B \forall x \in A \quad P(x, y). \end{aligned}$$

Если же кванторы разные, порядок важен. Так, высказывания

$$\forall x \in A \exists y \in B \quad P(x, y)$$

и

$$\exists y \in B \forall x \in A \quad P(x, y)$$

различны и никак между собой не связаны. Также нельзя менять кванторы местами. Например, высказывания

$$\forall x \in A \exists y \in B \quad P(x, y)$$

и

$$\exists x \in A \forall y \in B \quad P(x, y)$$

различны и, как и выше, между собой не связаны.

5. ЗАПИСЬ ПРЕДЛОЖЕНИЙ С КВАНТОРАМИ СРЕДСТВАМИ ЯЗЫКА. Говоря о разнообразии выражения высказываний с кванторами средствами языка, мы имеем в виду в первую очередь не замену одних слов им равнозначными, например, вместо слов «для всех» можно говорить «для любого» или же вычурное «каково бы ни было», а «найдется» заменять словом «существует». Существенное разнообразие записи высказываний с кванторами связано с положением кванторных приставок и даже употреблением по умолчанию, а также с использованием особых словосочетаний, которые заменяют кванторные приставки и даже их комбинации (последнее подробно обсуждается ниже в пункте *Макросы*.)

Рассмотрим особенности записи высказываний с квантором общности на примере следующих эквивалентных высказываний:

- (1) $\forall x \in A \quad f(x) \geq 0$;
- (2) $f(x) \geq 0, x \in A$;
- (3) $f(x) \geq 0$.

Аналогично высказывания с квантором $\exists$ также допускают различную запись:

- (4) $\exists(a, b) \subset \mathbb{R} \quad \forall x \in (a, b) \quad f(x) \geq 0$;
- (5) $\exists(a, b) \subset \mathbb{R} \quad f(x) \text{ на } (a, b)$;
- (6) $f(x) \geq 0$ на некотором интервале (a, b) .

В (2) кванторная приставка уходит в конец, и происходит это по той причине, что мы преимущественно интересуемся высказыванием, понимая или подразумевая, как в нем связана переменная, а в (3) пропадает вовсе — подразумевается, что она легко восстанавливается из контекста. В (5) кванторная приставка $\forall x \in (a, b)$ отправляется в конец, превращаясь в лаконичное «на (a, b) ». После этого и первая приставка $\exists(a, b) \subset \mathbb{R}$ убегает в конец, занимая свое место между «на» и (a, b) в виде слова «некотором». Данные примеры показывают основную тенденцию, имеющую место при записи сложных высказываний с кванторами, — кванторные приставки иногда уходят в конец и принимают изящные лаконичные формы. Математики стремятся украсить и разнообразить свою

речь. Действительно, записи типа (3) и (6) по сравнению с формальными развернутыми формами (1) и (4) легче читаются, однако чаще остаются непонятыми.

Иногда переход кванторной приставки в конец приводит к неоднозначной трактовке. Например, как понимать предложение

Для всех $x \in A$ верно $P(x, y)$ для некоторого $y \in B$?

Как

$$\forall x \in A \exists y \in B \quad P(x, y)$$

или

$$\exists y \in B \forall x \in A \quad P(x, y)?$$

Ясно, что ситуации, приводящие к неоднозначности, недопустимы.

Работая с высказываниями, будем придерживаться следующих принципов.

- (1) Если важна структура высказывания, например, при выводе одного высказывания с кванторами из другого, будем использовать развернутую форму типа (1), (4) (см. также *Дано–надо*).
- (2) Все переменные должны либо иметь конкретные значения, либо быть связаны кванторами. Однако кванторы общности, стоящие в начале, можно опускать: если переменная явно никак не связана, то подразумевается, что она связана квантором $\forall$, как в случае (3). (Квантор $\exists$ опускается в исключительных ситуациях, например, «в окрестности точки» означает «в некоторой окрестности точки».) Вообще опусканием кванторов следует пользоваться с осторожностью.

6. МАКРОСЫ. В программировании макросами называют короткие команды, которые заменяют большие часто повторяющиеся части текста. В математике также встречаются макросы. Под *макросом* будем понимать некоторое компактное словосочетание, которое разворачивается в какое-то детальное высказывание. Приведем несколько примеров.

$P(n)$ выполняется начиная с некоторого номера

означает

$$\exists n_0 \in \mathbb{N} \forall n \geq n_0 \quad P(n);$$

$P(x)$ выполняется для достаточно больших x

означает

$$\exists x_0 \in \mathbb{R} \quad \forall x \geq x_0 \quad P(x);$$

$f(x)$ локально ограничена

означает

$$\forall [a, b] \subset \text{dom } f \quad \exists C > 0 \quad \forall x \in [a, b] \quad |f(x)| \leq C;$$

$f(x)$ локально обратима

означает

для любого $x \in \text{dom } f$ существует интервал (a, b) , содержащий x , такой, что сужение $f(x)$ на (a, b) обратимо.

В некотором смысле многие понятия в анализе, такие как предел и производная, — это макросы.

7. ОТРИЦАНИЕ ВЫСКАЗЫВАНИЯ С КВАНТОРАМИ. Разберемся в том, как брать отрицание высказываний с кванторами, Рассмотрим утверждение вида

$$(\exists x \in X) P(x).$$

Его отрицание, т. е. утверждение

$$\text{не верно } ((\exists x \in X) P(x)),$$

можно проговорить так: не верно, что существует $x \in X$, для которого выполнено $P(x)$. Но если не существует такого $x \in X$, для которого выполнено $P(x)$, то для любого $x \in X$ должно выполняться отрицание высказывания $P(x)$, т. е. должно быть

$$(\text{для любого } x \in X) \quad \text{не верно } P(x)$$

или, символически,

$$(\text{не } (\exists x \in X) P(x)) \iff (\forall x \in X) \text{ не } P(x).$$

Что произошло? Квантор существования сменился квантором общности, а отрицание проникло вглубь утверждения и разместилось за квантором.

Теперь рассмотрим утверждение вида

$$(\forall x \in X) Q(x).$$

Его отрицание, т. е. утверждение

$$\text{не верно } ((\forall x \in X) Q(x)),$$

можно проговорить так: не верно, что для любого $x \in X$ выполнено $Q(x)$, т. е. не для любого $x \in X$ выполнено $Q(x)$, т. е. существует такое $x \in X$, для которого не выполнено $Q(x)$. Символически:

$$\text{не } ((\forall x \in X) Q(x)) \iff (\exists x \in X) \text{ не } Q(x).$$

Здесь квантор общности сменился квантором существования, а отрицание стало после квантора.

Последовательно применяя эти простейшие шаги к более сложным высказываниям с кванторами, можно легко получать их отрицания. Так, например,

$$\begin{aligned} \text{не } ((\exists C \forall x \in X) f(x) \leq C) &\iff \\ &\iff \forall C \text{ не } ((\forall x \in X) f(x) \leq C) \iff \\ &\iff (\forall C \exists x \in X) \text{ не верно } f(x) \leq C \iff \\ &\iff (\forall C \exists x \in X) f(x) > C, \end{aligned}$$

или, опуская промежуточные шаги,

$$\text{не } ((\exists C \forall x \in X) f(x) \leq C) \iff (\forall C \exists x \in X) f(x) > C.$$

Знатоки могут отметить, что мы сначала указали свойство ограниченности сверху функции f на множестве X и в виде отрицания получили свойство неограниченности сверху функции f на множестве X .

8. ВОСПРИЯТИЕ КВАНТОРОВ В ЗАВИСИМОСТИ ОТ СИТУАЦИЙ С ИХ ИСПОЛЬЗОВАНИЕМ. Поскольку встречаться с кванторами приходится на протяжении всего времени общения с математикой, важно хорошо понимать, как работать с утверждениями, содержащими кванторы.

Отношение к кванторам общности и существования неоднозначно и зависит от статуса утверждения с их участием, а именно от того, надо ли такое утверждение доказывать или оно уже дано и его можно использовать.

Ситуация 1. Квантор общности участвует в утверждении, которое надо доказывать. Тогда квантор «для любого» можно представлять себе как внешнее требование, не зависящее от нас, которое мы должны удовлетворить. Если в утверждении есть слова «для любого», представьте себе, что кто-то будет предлагать вам какие-то значения величины, сопровождаемой словами «для любого», а вам предстоит при этом выполнить все действия, которые следуют за этими словами. Пусть для определенности величина, к которой относятся слова «для любого», обозначается буквой ε , и для еще большей определенности (так чаще всего бывает) добавим требование: $\varepsilon > 0$. Тогда выражение «для любого $\varepsilon > 0$ » надо понимать так: кто-то будет давать какие-то положительные числа в качестве ε . Мы, в свою очередь, должны будем обеспечить справедливость какого-то утверждения (обозначим его через $P(\varepsilon)$), связанного с выбранным ε . С чего начать обеспечение требуемого утверждения? Во-первых, надо посмотреть, что от нас потребуется в утверждении $P(\varepsilon)$, и задать-ся вопросом: откуда можно получить выполнение $P(\varepsilon)$? Иначе говоря, пойти с конца, пытаясь понять, какие действия надо предпринять для обеспечения $P(\varepsilon)$. Значение ε нам неизвестно, поэтому надо заботиться о том, чтобы выработать **правило** обеспечения справедливости $P(\varepsilon)$, зависящее, естественно, от ε . Если такое правило выработать удастся, то можно гарантировать, что мы обеспечим выполнение утверждения

$$(\forall \varepsilon > 0) \quad P(\varepsilon).$$

Ситуация 2. Квантор общности встретился в утверждении, которое дано, и мы имеем право это утверждение использовать. Тогда выбор той величины, к которой относится этот квантор, в наших руках, мы вправе придать ей то значение, которое нам требуется, и использовать с этим значением всё, что сообщается в утверждении.

Очень важно не смешивать эти две ситуации!

Иначе следует относиться к квантору существования.

Ситуация 3. Квантор существования стоит в утверждении, которое мы доказываем, то это наши возможности, мы можем подбирать величину, о существовании которой говорится, исходя из всех имеющихся ресурсов. Это наш регулятор, которым мы управляем. Конечно, в каждом конкретном случае гарантия существования осуществляется по-разному, однако в некоторых ситуациях есть наблюдения, носящие рекомендательный характер.

Ситуация 4. Квантор существования находится в утверждении, которое известно, нам дано, и мы можем его использовать. В этом случае квантор существования — это пришедшая извне (из имеющегося утверждения) информация, которую мы можем использовать в наших целях.

Формирование корректного отношения к кванторам — процесс длительный, и он эффективно реализуется на простейших высказываниях с кванторами, заложенных в определениях базовых понятий, связанных с функциями. Целесообразно располагать появление таких понятий в порядке возрастания трудности восприятия кванторов. Самые простые понятия связаны с одним квантором существования, это область определения и множество значений, затем можно перейти к понятиям, основанным на одном кванторе общности, это четность и нечетность. Следующий шаг — два квантора общности, это монотонность. Затем можно перейти к определениям, основанным на двух кванторах, сначала существования, потом общности, это периодичность и ограниченность, и сначала общности, потом существования, это неограниченность и точные границы. Следующий по трудности восприятия шаг — определение предела последовательности. На восприятии кванторов именно в этой конфигурации, а именно $\forall\exists\forall$, мы подробно остановимся ниже ввиду традиционного наличия трудностей в изучении понятия предела последовательности.

9. АКСИОМЫ. УТВЕРЖДЕНИЯ. ДОКАЗАТЕЛЬСТВА. Математика состоит из высказываний, истинность которых декларируется, и такие высказывания называют *аксиомами*, и высказываний, доказываемых на основе аксиом и уже доказанных высказываний с помощью логических рассуждений. Доказываемые высказывания называют *утверждениями*, и в зависимости от дополнительных обстоятельств используют другие слова для обозначения утверждений. Например, утверждение, носящее вспомогательный для данного контекста характер, называют леммой, а утверждение, в данном контексте фундаментальное, важное — теоремой, и т. д.

В математических утверждениях не встречается свободных, не связанных кванторами объектов. При этом есть договоренность, что если не указано явно, каким из кванторов связывается переменная, то считается, что она связана квантором общности.

Формулировки условных утверждений организуют обычно так. Сначала сообщают обстановку, в рамках которой будет происходить действие, т. е. указывают, какие объекты и при каких обстоятельствах рассматри-

ваются. Затем дают условие, начиная фразу, например, словом «если». После того как условие записано, сообщают результат (вывод). Эта часть утверждения начинается, например, словом «то». Описание обстановки можно вынести за рамки утверждения, однако при изучении и использовании данного утверждения всегда надо иметь в виду обстановку, в которой оно формулируется. Такая организация в явном виде встречается не всегда, есть и другие способы выражения условного утверждения, но эта наиболее прозрачна и требует минимума усилий при раскодировании содержания утверждения.

В простейшем случае можно считать, что теорема — это утверждение вида $A \rightarrow B$, которое остается, если из длинной цепочки

$$\underbrace{A}_{\text{Тема}} \rightarrow \underbrace{A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_n}_{\text{Доказательство}} \rightarrow \underbrace{B}_{\text{Рема}}$$

вполне понятных переходов выбросить среднюю часть, которая называется *доказательством*. Утверждение A в формулировке теоремы называется в шутку *темой*, а B — *ремой*. При формулировке теоремы важно понимать, где тема, а где — рема, и никогда не путать их. Напомним, что истинность обратного утверждения, которое получается, если поменять местами тему и рему, никак не связана с истинностью исходного утверждения $A \rightarrow B$. Подчеркнем, что изображенная схема построения утверждения и доказательства соответствует простейшим случаям, обычно в доказательстве используются известные факты, лежащие за пределами данного утверждения, так что нередко схема получения из условия требуемого результата сложнее.

Доказать теорему — значит восстановить часть цепочки

$$\rightarrow A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_n \rightarrow ,$$

которую опускают и заменяют стрелкой $\rightarrow$ при формулировке теоремы. Если наша теорема — критерий, т. е. имеет вид $A \leftrightarrow B$, то восстанавливать нужно две цепочки

$$\begin{aligned} A &\rightarrow A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_n \rightarrow B, \\ B &\rightarrow B_1 \rightarrow B_2 \rightarrow \dots \rightarrow B_m \rightarrow A \end{aligned}$$

или сразу строить цепочку

$$A \leftrightarrow A_1 \leftrightarrow A_2 \leftrightarrow \dots \leftrightarrow A_n \leftrightarrow B.$$

В первом случае доказательство разбивается на две части, именуемые *необходимостью* и *достаточностью*.

Общего правила, по которому строятся переходы $A_i \rightarrow A_{i+1}$, нет, в их построении и состоит искусство математика. В следующих двух пунктах мы обсудим моменты, связанные с построением таких переходов, которые неизменно вызывают трудности у начинающих.

10. ДАНО-НАДО. Доказывая переход $A_i \rightarrow A_{i+1}$, нужно четко осознавать, в чем состоят утверждения A_i и A_{i+1} , особенно если это утверждения с кванторами. Дело в том, что значения кванторной приставки меняются в зависимости от того, стоит она в A_i или A_{i+1} (см. пункт о восприятии кванторов). Приставка $\forall x \in X$ в A_i означает, что мы можем удобным для нас образом выбрать x и для этого x будет выполняться оставшаяся часть утверждения A_i . Если же $\forall x \in X$ стоит в A_{i+1} , то это означает, что нам будут задавать различные значения x и для каждого из них мы должны обеспечить выполнение части A_{i+1} , стоящей после этой кванторной приставки. Как это сделать, подробно разбирается в доказательствах первых же теорем. Квантор $\exists$ должен вызывать следующие ассоциации. Если $\exists y \in Y$ стоит в A_i , то существование y с нужным свойством гарантировано и мы можем использовать его в наших дальнейших выкладках. Появление же $\exists y \in Y$ в A_{i+1} означает, что нам достаточно предъявить хотя бы один y , для которого выполняется оставшаяся часть A_{i+1} .

В сложных случаях, когда кванторов много, мы постараемся подробно расписывать содержания A_i и A_{i+1} предваряя первое словом **дано**, а второе — словом **надо** (обратите внимание, что эти слова состоят из одних и тех же букв!).

11. МЕТОД ДОКАЗАТЕЛЬСТВА ОТ ПРОТИВНОГО. Легко проверить, что утверждения

$$A_i \rightarrow A_{i+1} \quad \text{и} \quad \neg A_{i+1} \rightarrow \neg A_i$$

равносильны, т. е. их таблицы истинности совпадают. При этом доказательство второго утверждения может оказаться проще, чем доказательство первого. В этом случае переход $A_i \rightarrow A_{i+1}$ в доказательстве

$$\rightarrow A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_i \rightarrow A_{i+1} \rightarrow \dots \rightarrow A_n \rightarrow$$

можно мысленно взять в скобки и заменить его переходом $\neg A_{i+1} \rightarrow \neg A_i$:

$$\rightarrow A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_i \{ \neg A_{i+1} \rightarrow \neg A_i \} A_{i+1} \rightarrow \dots \rightarrow A_n \rightarrow .$$

В этом состоит метод доказательства от противного. Иногда такая замена происходит на самом верхнем уровне, т. е. вместо $A \rightarrow B$ сразу доказывают $\neg B \rightarrow \neg A$, а иногда этим приемом пользуются лишь для части цепочки. В этом случае важно понимать, какая именно часть цепочки доказывается методом от противного, т. е. где стоят скобки $\{$ и $\}$. Мы будем выделять эти места, используя следующую конструкцию. Докажем $A_i \rightarrow A_{i+1}$ методом от противного. Предположим, что A_{i+1} неверно. Далее идет цепочка рассуждений $\neg A_{i+1} \rightarrow \neg A_i$. В итоге, дойдя до $\neg A_i$, будем говорить: **противоречие**. (Противоречие с тем, что, на самом деле, A_i верно, как следует из предшествующей части цепочки $\rightarrow A_1 \rightarrow A_2 \rightarrow \dots \rightarrow A_i$.) Используется также краткая конструкция: Если A_{i+1} неверно, то $\dots$ имеет место $\neg A_i$, что невозможно. Мы будем использовать разные конструкции, когда способ доказательства от противного встречается внутри самого перехода $\neg A_{i+1} \rightarrow \neg A_i$.

12. ЗАПИСЬ ВЫСКАЗЫВАНИЙ ПРИ ПОМОЩИ ЯЗЫКА ОБЩЕНИЯ. Язык общения гораздо богаче математического языка символов (хотя при этом менее строг) и предоставляет разные способы записи одного и того же факта (высказывания). Особенно это разнообразие проявляется в случае импликации. Обсудим некоторые способы записи импликации $P \Rightarrow Q$:

- (1) пусть P , тогда верно Q ;
- (2) если P , то Q ;
- (3) из P следует Q .

Это понятные и нейтральные формы записи. Мы постараемся пользоваться именно такими формами.

Однако есть и другие способы записи $P \rightarrow Q$:

- (4) для того чтобы выполнялось P , необходимо выполнение Q ;
- (5) Q — необходимое условие для выполнения P ;
- (6) P выполнено только тогда, когда выполнено Q ;
- (7) P выполнено только в том случае, если выполнено Q ;
- (8) P выполнено, только если выполнено Q .

Это уже более тонкие формы, несущие в себе некоторую окраску. Например, формы (4) и (5) подчеркивают, что при выполнении P с неизбежностью (с необходимостью) выполняется и Q , или что P не может выполняться без Q . Это несколько непривычное выражение можно легко понять, если заметить, что слово «необходимо» здесь означает «не обходимо», т. е. «нельзя обойти», т. е. «обязательно будет выполнено», т. е. если выполнено утверждение P , то нельзя обойти выполнение утвержде-

ния Q , оно также будет выполнено, т. е. из P следует Q . Такие формы встречаются в формулировках теорем в случае, когда P — интересующий нас факт, а Q — (относительно) легко проверяемое условие. Такая теорема означает, что нечего пытаться доказать P , если Q не верно (ср. *Доказательство от противного* ниже).

Если записать ту же импликацию в другом порядке: $Q \leftarrow P$, то можно описать ее словами так:

- (9) Q следует из P .
- (10) Для выполнения Q достаточно выполнения P .
- (11) P — достаточное условие для выполнения Q .
- (12) Q выполнено, когда выполнено P .
- (13) Q выполнено в том случае, если выполнено P .

И здесь (9) нейтральная, а (10)–(13) — окрашенные формы. Они применяются в случае, когда Q — интересующий нас факт, а P — легко проверяемое условие, и подчеркивают, что если мы хотим установить Q , нам достаточно доказать, что выполняется P .

13. ЗНАЧЕНИЕ НЕОБХОДИМЫХ УСЛОВИЙ. Важно понимать суть и значение необходимых условий. Допустим, что мы интересуемся утверждением P и хотим узнать, справедливо ли оно, но пока располагаем лишь утверждением типа «из P следует Q ». Можно ли исходя из этого утверждения выяснить, выполнено ли P ? Конечно, нет. Дело в том, что в утверждении «из P следует Q » интересующее нас утверждение P стоит в условии, на месте предположения, а не в результате, не на месте вывода, т. е. нам известно, что в **предположении справедливости** P будет выполнено какое-то заключение, обозначенное нами символом Q . Исходя из этого невозможно сделать какое-либо заключение о справедливости P .

Тем не менее утверждение типа «из P следует Q », или «для выполнения P необходимо выполнение Q », какую-то информацию об утверждении P сообщает. Действительно, перепишем его в виде «если не выполнено Q , то не выполнено P » и заметим, что мы тем самым получаем возможность выяснить, когда же P **не выполнено**, т. е. при каких обстоятельствах не надо интересоваться справедливостью P — это произойдет в том случае, если не выполнено Q . Отсюда можно сделать вывод о том, что ожидать выполнения P можно лишь среди объектов, для которых выполнено Q . Те объекты, на которых Q не выполняется, можно не рассматривать, ибо на них P заведомо не выполняется.

В отличие от необходимых, достаточные условия гарантируют вы-

полнение интересующего нас свойства — их и называют достаточными потому, что их проверки вполне хватает для выполнения требуемого свойства. Однако прежде чем применить достаточные условия, сначала следует учесть необходимые условия с тем, чтобы достаточные применять только там, где выполнены условия необходимые.

Рассмотрим пример. Допустим, мы интересуемся, будет ли точка x из интервала $(a, b) \subset \mathbb{R}$ точкой экстремума дифференцируемой на (a, b) функции f . Если никакой информации о возможности экстремума в данной точке нет, то задача неподъемная — не исследовать же все точки из (a, b) , их слишком много, времени не хватит. Помогает необходимое условие: для того чтобы точка x была точкой экстремума функции f , необходимо, чтобы ее производная в этой точке обращалась в нуль, т. е. чтобы было выполнено равенство $f'(x) = 0$. Иначе это можно сформулировать так: если x — точка экстремума функции f на (a, b) , то $f'(x) = 0$. Это утверждение, сформулированное в виде: если $f'(x) \neq 0$, то x — не точка экстремума, подсказывает нам, что если $f'(x) \neq 0$, то на предмет наличия экстремума такую точку исследовать не надо — в ней экстремума нет! Следовательно, исследованию подлежат только те точки, где $f'(x) = 0$, а их, как правило, не очень много, и задача становится обозримой. Какими средствами изучать те точки, где $f'(x) = 0$, это уже другой вопрос, как правило, к ним применяют достаточное условие экстремума.

14. СВОЙСТВО, ПРИЗНАК, КРИТЕРИЙ. ПРЯМОЕ И ОБРАТНОЕ УТВЕРЖДЕНИЯ. Для условного утверждения «если выполнено P , то выполнено Q » в зависимости от ситуации употребляются следующие названия. Допустим, что мы интересуемся справедливостью утверждения Q , тогда P называют *признаком* выполнения Q . Если же у нас есть P и мы интересуемся, что из этого можно получить, то высказывание Q называют *свойством* P . Например, в утверждении «в равнобедренном треугольнике углы при основании равны» высказывание Q = «углы при основании равны» является свойством высказывания P = «данный треугольник равнобедренный». Если же мы интересуемся высказыванием Q = «данный треугольник равнобедренный» и имеем в распоряжении высказывание P = «два угла треугольника равны между собой», то оно служит признаком равнобедренности, так как имеет место утверждение «если два угла треугольника равны между собой, то треугольник равнобедренный»,

Если у нас есть некое высказывание P и мы установили эквивалентность « P равносильно Q », то Q называют *критерием* P . Например, если

посмотреть на приведенный пример с точки зрения такой формулировки: «равнобедренность треугольника равносильна тому, что два его угла равны между собой», то равенство углов оказывается критерием равнобедренности треугольника.

Пусть даны утверждения P и Q . Из них можно составить два условных утверждения, а именно «если верно P , то верно Q » и «если верно Q , то верно P » или, иными словами, «из P следует Q » и «из Q следует P ». Считая одно из них данным, говорят «прямым», другое называют к нему *обратным*. Например, если утверждение

из P следует Q

данное (прямое), то утверждение

из Q следует P

к нему обратное.

Обратим внимание на то, как оформляется в доказательстве равносильности каких-то утверждений переход к доказательству обратного утверждения. Это в некоторой мере зависит от того, в какой манере дана формулировка. Так, если в утверждении написано, что для выполнения P необходимо и достаточно выполнения Q , и доказана необходимость, т. е. «если выполнено P , то выполнено Q », далее говорят: «докажем достаточность», имея в виду, что теперь будет доказываться импликация «из Q следует P ».

Если утверждение сформулировано так: «высказывание P выполняется тогда и только тогда, когда выполнено Q », и доказано прямое (в данном контексте) утверждение, т. е. что из P следует Q , и теперь переходим к доказательству обратного утверждения, то этот момент оформляется так: «обратно, пусть выполнено Q , докажем выполнение P ». Используемое в этой ситуации слово «обратно» как признак начала доказательства обратного утверждения иногда заменяют словом «наоборот», здесь совершенно неуместным: слово «обратно» указывает на начало доказательства обратного, а не «наоборотного» утверждения.

Может возникнуть вопрос: как связаны между собой прямое и обратное утверждения в плане их справедливости? Ответ простой: никак. Поэтому если вместо требуемого утверждения предлагают обратное к нему, это грубая ошибка, встречающаяся, в частности, при формулировке теоремы о необходимых условиях экстремума, особенно в стиле, использующем именно слово «необходимо».

15. ОПРЕДЕЛЕНИЯ. Когда та или иная ситуация, объект или свойство встречаются часто, математики дают им имена или вводят новые понятия. Так возникают предел, производная, интеграл и т. д. С другой стороны, любое понятие в математике (в отличие от некоторых других наук) имеет четко определенное содержание, т. е. может быть более подробно расписано через другие, уже определенные понятия и не допускает произвольного (интуитивного) толкования. Это означает, что, формулируя и доказывая теоремы, мы должны понимать смысл всех слов. Проблема заключается в том, что часто новые понятия называются словами обыденного языка, имеющими общепринятые интуитивно понятные значения. Это вызывает иллюзию понимания математического текста. Встречая такое слово, начинающий иногда не задумывается над тем, что в математике это слово имеет строго определенное значение. Так, например, происходит со словами предел, монотонность, граница, максимум, последовательность. Напротив, слова инфимум, интеграл, резольвента, не встречающиеся в обыденном языке, невольно заставляют задуматься о своем значении. Итак, подчеркнем еще раз, что в математике мы должны знать точное математическое значение всех произносимых слов, а не довольствоваться интуитивным пониманием.

В заключение заметим, что слово *если*, часто используемое в определениях, в данном контексте имеет смысл двусторонней импликации или равносильности.

0.2. Множества.

Имея намерение определять все понятия, мы на этот раз его нарушим. Исключение будет сделано для понятия множества, которое будет для нас синонимом слов совокупность, набор, система. Опыт показывает, что в рамках курса математического анализа такой интуитивный подход (именуемый наивной теорией множеств) не подводит. Более того, строгое определение понятия множества на этом этапе не приводит к более полному пониманию предмета.

При работе с множествами всегда подразумевается, что состав любого множества вполне определен, т. е. для любого объекта x и любого множества M либо $x \in M$, либо $x \notin M$, и в случае $x \in M$ говорят, что x *принадлежит* M или что x — *элемент множества* M . Иногда бывает сложно сказать, какой именно из указанных двух случаев реализуется, однако из них один и только один имеет место. Множества считаются равными если они имеют один состав, иначе говоря, $A = B$ означает, что

если $x \in A$, то $x \in B$, и если $x \in B$, то $x \in A$.

Множество можно задать просто перечислением его элементов. Другой способ — выделить в уже имеющемся множестве M множество элементов, обладающих некоторым свойством $P(x)$:

$$A = \{x \in M : \text{верно } P(x)\}$$

(читается «множество элементов $x \in M$ таких, что верно $P(x)$ »). Вместо двоеточия в этой записи используют также черту:

$$A = \{x \in M \mid \text{верно } P(x)\}$$

Часто множество M понятно из контекста, в этом случае можно просто писать $A = \{x : \text{верно } P(x)\}$. Двоеточие (иногда на этом месте ставят черту) является сокращением слов «таких, что», «для которых» и т. п.

Еще один способ построения множеств — взятие прямого произведения. Имея два элемента a и b (вообще говоря, разных множеств), можно назвать a первым, а b — вторым и образовать новый объект (a, b) , называемый *упорядоченной парой*. При этом упорядоченные пары (a, b) и (c, d) равны тогда и только тогда, когда $a = c$ и $b = d$, так что, вообще говоря, $(a, b) \neq (b, a)$ (в то время как $\{a, b\} = \{b, a\}$).

Пусть A и B — два множества, их *прямым произведением* называют множество

$$A \times B = \{(a, b) : a \in A, b \in B\}$$

Если множества A и B в прямом произведении равны то их прямое произведение $A^2 = A \times A = \{(a, b) : a \in A, b \in A\}$ называют *квадратом* A . Аналогично определяют упорядоченные наборы (n -ки) $(a_1, \dots, a_n)$, прямые произведения n множеств, а также n -ю степень множества. Например,

$$\mathbb{R}^3 = \mathbb{R} \times \mathbb{R} \times \mathbb{R} = \{(x, y, z) : x \in \mathbb{R}, y \in \mathbb{R}, z \in \mathbb{R}\}$$

— трехмерное координатное пространство.

Напомним, что для пары множеств A, B определены их объединение $A \cup B$, пересечение $A \cap B$ и разность $A \setminus B$. Напомним также, что важно не путать знаки принадлежности $\in$ и включения $\subset$. Первый ставится между элементом и множеством, а второй — между двумя множествами.

0.3. Краткие обозначения для суммы и произведения. Познакомимся с простыми обозначениями, с которыми будем часто встречаться в дальнейшем.

Рассмотрим конечный набор чисел $x_1, x_2, \dots, x_n$, $n \in \mathbb{N}$, занумерованных последовательными натуральными числами от 1 до n . Для суммы $x_1 + x_2 + \dots + x_n$ используют обозначение $\sum_{k=1}^n x_k$, в котором букву k называют *индексом суммирования*. Она является внутренней переменной в том смысле, что ее использование ограничивается данным выражением и за его пределами можно снова использовать эту же букву. Важно не то, какой буквой обозначен индекс суммирования, а то, в каких конфигурациях эта буква участвует. Иначе говоря, выражение $\sum_{k=1}^n x_k$ представляет собой краткую форму записи следующей процедуры: на то место, где встречается индекс k в выражении x_k , подставляются последовательно числа $1, 2, \dots, n$ и полученные значения суммируются. Например,

$$\sum_{k=1}^n \frac{1}{k^2} = 1 + \frac{1}{2^2} + \frac{1}{3^2} + \dots + \frac{1}{n^2},$$

$$\begin{aligned} \sum_{i=1}^m \frac{1}{(2i-1)^2} &= 1 + \frac{1}{(2 \cdot 2 - 1)^2} + \frac{1}{(2 \cdot 3 - 1)^2} + \dots + \frac{1}{(2 \cdot m - 1)^2} \\ &= 1 + \frac{1}{3^2} + \frac{1}{5^2} + \frac{1}{(2m-1)^2}. \end{aligned}$$

Аналогичное обозначение используют для обозначения произведения конечной последовательности x_k , $k = 1, \dots, n$:

$$\prod_{k=1}^n = x_1 \cdot x_2 \cdot \dots \cdot x_n.$$

Произведение $1 \cdot 2 \cdot 3 \cdot \dots \cdot n$ первых n натуральных чисел обозначают символом $n!$ называемым *факториалом числа n* . Для удобства по определению полагают $0! = 1$.

Полезно иметь в виду свойство сочетаемости сумм, состоящее в следующем. Рассмотрим конечный набор чисел, занумерованный посредством двух натуральных чисел, например x_{ij} , где i может принимать значения от 1 до m , а j — от 1 до n , т. е. набор чисел

$$x_{11}, \dots, x_{1n}, x_{21}, \dots, x_{2n}, \dots, x_{m1}, \dots, x_{mn}.$$

Ясно, что сумма всех этих чисел не зависит от того, в каком порядке их складывают. Например, при каждом фиксированном $j = 1, \dots, n$ можно взять сумму по всем индексам $i = 1, \dots, m$, а затем все результаты

просуммировать, и это в кратком виде запишется так: $\sum_{j=1}^n \sum_{i=1}^m x_{ij}$, а можно поступить наоборот, т. е. сначала просуммировать по j при каждом фиксированном i , а затем собрать результаты. Это простое наблюдение выразится таким равенством:

$$\sum_{i=1}^m \sum_{j=1}^n x_{ij} = \sum_{j=1}^n \sum_{i=1}^m x_{ij}.$$

0.4. Метод математической индукции. Строение множества натуральных чисел порождает метод математической индукции доказательства утверждений, истинность которых зависит от натуральных чисел. Он состоит в следующем.

Предположим, что некоторое утверждение $P(n)$, справедливость которого зависит от натурального n , верно для первого натурального числа $n = 1$ и из того, что верно $P(n)$, следует, что верно $P(n + 1)$. Принцип математической индукции гласит, что в таком случае утверждение $P(n)$ верно для любого $n \in \mathbb{N}$.

Для использования метода математической индукции с целью доказательства некоторого утверждения $P(n)$, $n \in \mathbb{N}$, надо проверить $P(1)$, т. е. справедливость утверждения для первого натурального числа, а затем обосновать условное утверждение, так называемый индуктивный переход. А именно, предполагая справедливость утверждения для некоторого натурального числа n , доказать его выполнение для следующего натурального числа $n + 1$. Разумеется, в процессе доказательства надо активно использовать индуктивное предположение, т. е. допущение справедливости $P(n)$. Если утверждение $P(n)$ надо доказывать не для всех натуральных n , а только начиная с некоторого натурального n_0 , то в качестве базы индукции проверяется утверждение $P(n_0)$ и индуктивный переход обосновывается для $n \geq n_0$.